

中国人民大学法学院 数字法学教研月报

2025 年第 8 期 (总第 20 期)

2025 年 8 月 25 日



本期看点

【数字法治大事件】

国家多部门聚焦“自媒体”医疗科普行为划定合规边界，以互联网新闻信息服务“持证亮牌”工程筑牢行业规范基石；“人工智能+”行动意见的出台为技术赋能实体经济指明方向，全国一体化算力网、数据基础设施等技术文件的意见征集，则推动数字基础设施建设向标准化、安全化迈进。地方层面同样动作频频，北京加速公共数据资源开发利用，河北立体推进公共数据登记，厦门出台数据产业高质量发展措施，更有企业服务“0 证明”示范场景落地赋能民营经济，形成上下联动的数字发展新格局。

多位专家围绕深圳高质量数据集建设、数据资源体系完善及信息化赋能路径的深度解读，也为数字领域高质量发展提供了理论支撑与实践参考，全方位勾勒出我国数字领域在法治护航下创新前行的生动图景。

【研究动态】

基础理论领域，多位学者围绕数字法学学科独立性、数字人权、打击网络恐怖主义等问题展开讨论，**天津大学法学院教授田野发表于《中国法学》的论文《数智时代的过劳问题及其法律因应》**针对“科技异化”背景下劳动者“过度劳动”问题进行了鞭辟入里的分析，并提出了解决方案。该篇论文全文转载于本期“数字法评”栏目。

数据确权领域，学者从民商事法律、刑事法律等多角度出发，探讨数据权益的多角度、多层次保护；人工智能法学围绕大模型训练版权困境、算法量刑证据化改造、AGI 国际权力重塑等前沿议题形成系列成果；数字行政与司法领域，学者聚焦于“电子证据”、“智能化背景下的政治权利保障”等议题展开讨论。

【教研活动】

受中国人民大学法学院、中国人民大学未来法治研究院、中国人民大学涉外法治研究院邀请：

美国斯坦福大学史波格利国际研究所研究者 Graham Webster，耶鲁大学蔡中曾中国中心高级研究员 Samm Sacks 围绕“中美人工智能与数据治理”展开了精彩的主题讲座。

美国哥伦比亚大学法学院副院长、中国法研究中心主任、罗伯特·立夫 (Robert L. Lief) 讲席教授本杰明·李本 (Benjamin L. Liebman) 围绕“人工智能与中美法院的数据使用”展开了精彩的主题讲座。

【数字法评】

《构思决定论视野下生成式人工智能生成物的版权保护》，《法学研究》2025 年第 4 期，作者：张金平

《数智时代的过劳问题及其法律因应》，《中国法学》2025 年第 3 期，作者：田野

本期目录

数字法治大事件	3	研究动态	22
关于规范“自媒体”医疗科普行为的通知.....	3	基础理论.....	22
国家网信办持续深入推进互联网新闻信息服务“持证亮牌”工程.....	4	数据确权与流通.....	26
国务院关于深入实施“人工智能+”行动的意见5		人工智能.....	29
全国一体化算力网算力池化、算网安全相关技术文件公开征求意见.....	8	平台治理.....	33
关于征求数据基础设施3项技术文件、可信数据空间3项技术文件意见的通知.....	9	数字行政与司法.....	34
地方动态 关于加快北京市公共数据资源开发利用的实施意见.....	9	教研活动	36
地方动态 河北立体推进公共数据资源登记工作12		讲座回顾 斯坦福大学 Graham Webster & 耶鲁大学 Samm Sacks: 中美人工智能与数据治理 .36	
地方动态 “国家公共数据‘跑起来’示范场景——企业服务‘0证明’赋能民营经济创新应用”启动活动成功举办.....	12	讲座回顾 哥大法学院李本教授: 人工智能与中美法院的数据使用.....	38
地方动态 厦门市人民政府办公厅关于印发促进数据产业高质量发展若干措施的通知.....	13	第四届“中英法治圆桌会议”在英成功举办..39	
专家解读 深圳市高质量数据集建设探索与实践15		中国政法大学与联合国教科文组织联合举办WAIC2025“人工智能规则全球对话与协同发展”论坛.....	40
专家解读 夯实高质量数据集底座: 完善数据资源体系, 助力“人工智能+”创新发展.....	18	“人工智能立法前沿与实践研讨会”在上海召开.....	41
专家解读 以“智”提质, 向“新”而行, 擘画信息化赋能新图.....	19	“人工智能伦理与法治”论坛在上海圆满收官42	
		数字法评	43
		构思决定论视野下生成式人工智能生成物的版权保护.....	43
		数智时代的过劳问题及其法律因应.....	56

学术顾问: 王利明

编委会: 张新宝 丁晓东 王莹 张吉豫

编辑部: 阮神裕 卞龙 艾薇 邓语鑫 何芮 梁因格 李佳丽 林诗敏 麻卓妍 乔彩霞 王昊 朱恬馨

联系方式: RUCdigitallaw@163.com

本期编辑: 王昊

数字法治大事件

导言：在数字化浪潮深度融入经济社会各领域的当下，规范与创新并重成为行业发展的核心主线。近期，从国家层面到地方实践，从顶层设计到技术落地，数字法治建设、行业规范完善、技术创新推进及数据资源开发利用等领域密集释放关键信号：国家多部门聚焦“自媒体”医疗科普行为划定合规边界，以互联网新闻信息服务“持证亮牌”工程筑牢行业规范基石；“人工智能+”行动意见的出台为技术赋能实体经济指明方向，全国一体化算力网、数据基础设施等技术文件的意见征集，则推动数字基础设施建设向标准化、安全化迈进。地方层面同样动作频频，北京加速公共数据资源开发利用，河北立体推进公共数据登记，厦门出台数据产业高质量发展措施，更有企业服务“0 证明”示范场景落地赋能民营经济，形成上下联动的数字发展新格局。此外，专家围绕深圳高质量数据集建设、数据资源体系完善及信息化赋能路径的深度解读，也为数字领域高质量发展提供了理论支撑与实践参考，全方位勾勒出我国数字领域在法治护航下创新前行的生动图景。

关于规范“自媒体”医疗科普行为的通知

原载：“网信中国”微信公众号

各省、自治区、直辖市党委网信办、卫生健康委、市场监管局（厅、委）、中医药管理局，新疆生产建设兵团党委网信办、卫生健康委、市场监管局、中医药管理局：

医疗科普行为是以提高公众的医学健康科学素养为目的，通过文字、图片、音视频等形式，生产传播预防或治疗疾病等相关知识的行为。为进一步压实网站平台信息内容管理主体责任，规范“自媒体”医疗科普信息发布传播行为，提升“自媒体”规范开展医疗科普行为意识，支持专业医疗科普内容生产传播，防范虚假医疗科普信息误导公众，维护人民群众合法权益，现就有关工作通知如下：

1.分类核查认证账号资质。网站平台应进一步完善医疗“自媒体”账号资质认证工作，对申请相关资质认证的账号，区分医疗机构从业人员、医学

院校人员、医药研发机构等不同医疗领域人员类型，分类开展账号资质核查。对医疗机构从业人员，划分医师（中西医）、护士、药学技术人员、医技人员、科研人员等类型，分类核查医师资格证、医师执业证、护士执业证、卫生专业技术资格证（药学）等资质证件，医疗机构出具的在职证明（具体到科室）等材料。对医学院校、医药研发机构人员，核查学校学院出具的在职证明、组织机构代码证、机构部门出具的在职证明等材料。其他医疗领域从业人员，结合账号名称、简介等情况，核查相应医疗资质证明材料。

2.清晰展示账号资质信息。网站平台应在资质核验基础上，在账号主页显著位置强化认证信息展示，主要包括：医师资格证、医师执业证、护士执业证、卫生专业技术资格证（药学）等相关资质证件的名称、专业范围等主要信息，医院科室/学校学院/研发机构部门等准确名称，是否在职，签约的MCN 机构名称等。保障用户清晰知晓账号运营者专业背景、从业经历等情况，以便用户全面评估“自媒体”账号发布医疗科普内容的权威性、专业性。网站平台应依法加强个人信息保护，对相关账号资质信息采取适当保护措施。

3.严格标注医疗科普信息来源。网站平台应明确要求提供医疗科普内容的“自媒体”账号，对发布转发医疗科普信息的真实性、科学性负责。引用转载专业医疗科普内容、引用医疗安全等旧闻旧事、结合医疗领域素材摆拍剧情、借助人工智能技术生成合成医疗科普信息、分享传播真实健康经历的，需严格标注信息来源或生成合成内容标识。不得恶意编造、散布虚假信息，不得随意拼接发布不准确信息。

4.认真做好资质核验工作。网站平台应进一步强化医疗科普领域认证材料的真实性审核，对证件、单位证明的时间期限、落款部门等明显可核查的情况，需严格比对核查。对医生、护士等执业信息，需通过国家卫生健康委官方信息查询渠道进行查验比对。如发现存在虚假不实证件信息，或同一科室/学院专业/部门认证人数明显超过常规人数设置

等显著存疑情形，应依法依规对账号采取措施并进一步核验，严防虚假认证。

5.严禁无资质账号生产发布专业医疗科普内容。网络平台应进一步完善用户协议，依法依规加强专业医疗科普行为管理，按照“复审存量、严管新增”原则，严禁无资质账号生产发布专业医疗科普内容。严格复审存量医疗科普“自媒体”账号，对未认证的账号通知其2个月内完成资质认证。对新注册的“自媒体”账号，不经医疗资质认证，不得生产发布专业医疗科普内容。

6.强化网络行为规范。网络平台应进一步健全激励机制，鼓励权威优质、科学专业的医疗科普信息生产传播。强化对医疗科普领域账号网络行为的规范，加强违法违规医疗科普行为发现处置，从严管控违规提供线上付费诊疗服务，认证账号直播时认证医生本人不出镜，将用户引流至社交平台或无医疗认证资质账号进行互动、交易等行为。

7.严禁违规变相发布广告。网络平台应明确告知“自媒体”账号不得以介绍健康、养生知识等形式，变相发布医疗、药品、医疗器械、保健食品、特殊医学用途配方食品广告。介绍健康、养生知识的，不得在同一页面或者同时出现相关医疗、药品、医疗器械、保健食品、特殊医学用途配方食品的商品经营者或者服务提供者地址、联系方式、购物链接等内容。

8.严处违法违规信息及账号。网络平台应坚决清理传授无底线蹭流量打造“网红医生”、借两性健康知识传播色情擦边内容、利用AI编造发布涉医领域同质化文案、编造健康故事售卖商品或药品、假冒医生身份开展科普、为售卖保健品鼓动拒绝就医等违法违规信息。对存在不按要求或虚假标注信息来源、无资质认证且持续生产发布专业医疗科普内容、违规发布广告、不遵守医疗科普行为规范等问题的账号，要依法依规采取取消互动功能、清理粉丝、取消营利权限、禁言、关闭等梯度措施。

各级网信部门应加强对网站平台的督促指导，强化医疗科普领域“自媒体”账号的规范管理，加强相关违法违规信息及账号的处置。各级卫生健康

部门、中医药主管部门应指导医院加强对中西医医务人员网络行为的规范管理，对存在不良网络行为的医生等，依法依规依纪予以处理。持续加大中西医相关医疗科普知识供给力度，将更多权威、专业的健康知识，以多样亲民的形式传达公众。各级市场监管部门应当依据职责，强化互联网商业营销活动监管，依法查处变相发布医疗广告等违法行为，保护消费者合法权益。

中央网信办秘书局

国家卫生健康委办公厅

市场监管总局办公厅

国家中医药管理局综合司

2025年7月28日

国家网信办持续深入推进互联网新闻信息服务“持证亮牌”工程

原载：“网信中国”微信公众号

为进一步规范网络传播秩序，提升互联网新闻信息服务辨识度，今年以来，国家网信办深入推进互联网新闻信息服务“持证亮牌”工程，部署指导地方网信办提升审批质效，督促腾讯、抖音、快手、微博等15家获批提供互联网新闻信息传播平台服务资质的网站平台优化完善系统功能，加强服务资质核验，对获批提供互联网新闻信息服务的公众账号统一增加红“V”标识，基本实现图文、音频、直播、短视频等内容场景全覆盖。截至2025年7月25日，共有13516个公众账号加注红“V”标识并明示服务主体名称、许可证编号和服务类别，4401家网站、平台等服务形式明示许可信息。

下一步，国家网信办将制度化、程序化、动态化推进互联网新闻信息服务“持证亮牌”，同时将研究出台流量支持政策，助力提升红“V”账号影响力、传播力。此外，将大力整治编发虚假不实新闻信息、假冒仿冒新闻媒体、出租出借转让互联网新闻信息服务许可资质等问题，严厉打击借舆论监督之名谋取非法利益、破坏营商网络环境的行为，着力营造清朗网络空间。

国务院关于深入实施“人工智能+”行动的意见

原载：“国家数据局”微信公众号

国务院关于深入实施“人工智能+”行动的意见

国发〔2025〕11号

各省、自治区、直辖市人民政府，国务院各部委、各直属机构：

为深入实施“人工智能+”行动，推动人工智能与经济社会各行业各领域广泛深度融合，重塑人类生产生活范式，促进生产力革命性跃迁和生产关系深层次变革，加快形成人机协同、跨界融合、共创分享的智能经济和智能社会新形态，现提出如下意见。

一、总体要求

以习近平新时代中国特色社会主义思想为指导，完整准确全面贯彻新发展理念，坚持以人民为中心的发展思想，充分发挥我国数据资源丰富、产业体系完备、应用场景广阔等优势，强化前瞻谋划、系统布局、分业施策、开放共享、安全可控，以科技、产业、消费、民生、治理、全球合作等领域为重点，深入实施“人工智能+”行动，涌现一批新基础设施、新技术体系、新产业生态、新就业岗位等，加快培育发展新质生产力，使全体人民共享人工智能发展成果，更好服务中国式现代化建设。

到2027年，率先实现人工智能与6大重点领域广泛深度融合，新一代智能终端、智能体等应用普及率超70%，智能经济核心产业规模快速增长，人工智能在公共治理中的作用明显增强，人工智能开放合作体系不断完善。到2030年，我国人工智能全面赋能高质量发展，新一代智能终端、智能体等应用普及率超90%，智能经济成为我国经济发展的重要增长极，推动技术普惠和成果共享。到2035年，我国全面步入智能经济和智能社会发展新阶段，为基本实现社会主义现代化提供有力支撑。

二、加快实施重点行动

（一）“人工智能+”科学技术

1.加速科学发现进程。加快探索人工智能驱动的新型科研范式，加速“从0到1”重大科学发现进程。加快科学大模型建设应用，推动基础科研平台和重大科技基础设施智能化升级，打造开放共享的高质量科学数据集，提升跨模态复杂科学数据处理水平。强化人工智能跨学科牵引带动作用，推动多学科融合发展。

2.驱动技术研发模式创新和效能提升。推动人工智能驱动的技术研发、工程实现、产品落地一体化协同发展，加速“从1到N”技术落地和迭代突破，促进创新成果高效转化。支持智能化研发工具和平台推广应用，加强人工智能与生物制造、量子科技、第六代移动通信（6G）等领域技术协同创新，以新的科研成果支撑场景应用落地，以新的应用需求牵引科技创新突破。

3.创新哲学社会科学研究方法。推动哲学社会科学研究方法向人机协同模式转变，探索建立适应人工智能时代的新型哲学社会科学研究组织形式，拓展研究视野和观察视域。深入研究人工智能对人类认知判断、伦理规范等方面的深层次影响和作用机理，探索形成智能向善理论体系，促进人工智能更好造福人类。

（二）“人工智能+”产业发展

1.培育智能原生新模式新业态。鼓励有条件的企业将人工智能融入战略规划、组织架构、业务流程等，推动产业全要素智能化发展，助力传统产业改造升级，开辟战略性新兴产业和未来产业发展新赛道。大力发展智能原生技术、产品和服务体系，加快培育一批底层架构和运行逻辑基于人工智能的智能原生企业，探索全新商业模式，催生智能原生新业态。

2.推进工业全要素智能化发展。推动工业全要素智能联动，加快人工智能在设计、中试、生产、服务、运营全环节落地应用。着力提升全员人工智能素养与技能，推动各行业形成更多可复用的专家知识。加快工业软件创新突破，大力发展智能制造装备。推进工业供应链智能协同，加强自适应供需匹配。推广人工智能驱动的生产工艺优化方法。深

化人工智能与工业互联网融合应用,增强工业系统的智能感知与决策执行能力。

3.加快农业数智化转型升级。加快人工智能驱动的育种体系创新,支持种植、养殖等农业领域智能应用。大力发展智能农机、农业无人机、农业机器人等智能装备,提高农业生产和加工工具的智能感知、决策、控制、作业等能力,强化农机农具平台化、智能化管理。加强人工智能在农业生产管理、风险防范等领域应用,帮助农民提升生产经营能力和水平。

4.创新服务业发展新模式。加快服务业从数字赋能的互联网服务向智能驱动的新型服务方式演进,拓展经营范围,推动现代服务业向智向新发展。探索无人服务与人工服务相结合的新模式。在软件、信息、金融、商务、法律、交通、物流、商贸等领域,推动新一代智能终端、智能体等广泛应用。

(三) “人工智能+”消费提质

1.拓展服务消费新场景。培育覆盖更广、内容更丰富的智能服务业态,加快发展提效型、陪伴型等智能原生应用,支持开辟智能助理等服务新入口。加强智能消费基础设施建设,提升文娱、电商、家政、物业、出行、养老、托育等生活服务品质,拓展体验消费、个性消费、认知和情感消费等服务消费新场景。

2.培育产品消费新业态。推动智能终端“万物智联”,培育智能产品生态,大力发展智能网联汽车、人工智能手机和电脑、智能机器人、智能家居、智能穿戴等新一代智能终端,打造一体化全场景覆盖的智能交互环境。加快人工智能与元宇宙、低空飞行、增材制造、脑机接口等技术融合和产品创新,探索智能产品新形态。

(四) “人工智能+”民生福祉

1.创造更加智能的工作方式。积极发挥人工智能在创造新岗位和赋能传统岗位方面的作用,探索人机协同的新型组织架构和管理模式,培育发展智能代理等创新型工作形态,推动在劳动力紧缺、环境高危等岗位应用。大力支持开展人工智能技能培训,激发人工智能创新创业和再就业活力。加强人

工智能应用就业风险评估,引导创新资源向创造就业潜力大的方向倾斜,减少对就业的冲击。

2.推行更富成效的学习方式。把人工智能融入教育教学全要素、全过程,创新智能学伴、智能教师等人机协同教育教学新模式,推动育人从知识传授为重向能力提升为本转变,加快实现大规模因材施教,提高教育质量,促进教育公平。构建智能化情景交互学习模式,推动开展方式更灵活、资源更丰富的自主学习。鼓励和支持全民积极学习人工智能新知识、新技术。

3.打造更有品质的美好生活。探索推广人人可享的高水平居民健康助手,有序推动人工智能在辅助诊疗、健康管理、医保服务等场景的应用,大幅提高基层医疗健康服务能力和效率。推动人工智能在繁荣文化生产、增强文化传播、促进文化交流中展现更大作为,利用人工智能辅助创作更多具有中华文化元素和标识的文化内容,壮大文化产业。充分发挥人工智能对织密人际关系、精神慰藉陪伴、养老托育助残、推进全民健身等方面的重要作用,拓展人工智能在“好房子”全生命周期的应用,积极构建更有温度的智能社会。

(五) “人工智能+”治理能力

1.开创社会治理人机共生新图景。有序推动市政基础设施智能化改造升级,探索面向新一代智能终端发展的城市规划、建设与治理,提升城市运行智能化水平。加快人工智能产品和服务向乡村延伸,推动城乡智能普惠。深入开展人工智能社会实验。安全稳妥有序推进人工智能在政务领域应用,打造精准识别需求、主动规划服务、全程智能办理的政务服务新模式。加快人工智能在各类公共资源招标投标活动中的应用,提升智能交易服务和监管水平。

2.打造安全治理多元共治新格局。推动构建面向自然人、数字人、智能机器人等多元一体的公共安全治理体系,加强人工智能在安全生产监管、防灾减灾救灾、公共安全预警、社会治安管理等方面的应用,提升监测预警、监管执法、指挥决策、现场救援、社会动员等工作水平,增强应用人工智能维护和塑造国家安全的能力。加快推动人工智能赋

能网络空间治理,强化信息精准识别、态势主动研判、风险实时处置等能力。

3.共绘美丽中国生态治理新画卷。提高空天地海一体化动态感知和国土空间智慧规划水平,强化资源要素优化配置。围绕大气、水、海洋、土壤、生物等多要素生态环境系统和全国碳市场建设等,提升人工智能驱动的监测预测、模拟推演、问题处置等能力,推动构建智能协同的精准治理模式。

(六) “人工智能+”全球合作

1.推动人工智能普惠共享。把人工智能作为造福人类的国际公共产品,打造平权、互信、多元、共赢的人工智能能力建设开放生态。深化人工智能领域高水平开放,推动人工智能技术开源可及,强化算力、数据、人才等领域国际合作,帮助全球南方国家加强人工智能能力建设,助力各国平等参与智能化发展进程,弥合全球智能鸿沟。

2.共建人工智能全球治理体系。支持联合国在人工智能全球治理中发挥主渠道作用,探索形成各国广泛参与的治理框架,共同应对全球性挑战。深化与国际组织、专业机构等交流合作,加强治理规则、技术标准等对接协调。共同研判、积极应对人工智能应用风险,确保人工智能发展安全、可靠、可控。

三、强化基础支撑能力

(七) 提升模型基础能力。加强人工智能基础理论研究,支持多路径技术探索和模型基础架构创新。加快研究更加高效的模型训练和推理方法,积极推动理论创新、技术创新、工程创新协同发展。探索模型应用新形态,提升复杂任务处理能力,优化交互体验。建立健全模型能力评估体系,促进模型能力有效迭代提升。

(八) 加强数据供给创新。以应用为导向,持续加强人工智能高质量数据集建设。完善适配人工智能发展的数据产权和版权制度,推动公共财政资助项目形成的版权内容依法合规开放。鼓励探索基于价值贡献度的数据成本补偿、收益分成等方式,加强数据供给激励。支持发展数据标注、数据合成等技术,培育壮大数据服务和数据服务产业。

(九) 强化智能算力统筹。支持人工智能芯片攻坚创新与使能软件生态培育,加快超大规模智算集群技术突破和工程落地。优化国家智算资源布局,完善全国一体化算力网,充分发挥“东数西算”国家枢纽作用,加大数、算、电、网等资源协同。加强智能算力互联互通和供需匹配,创新智能算力基础设施运营模式,鼓励发展标准化、可扩展的算力云服务,推动智能算力供给普惠易用、经济高效、绿色安全。

(十) 优化应用发展环境。布局建设一批国家人工智能应用中试基地,搭建行业应用共性平台。推动软件信息服务企业智能化转型,重构产品形态和服务模式。培育人工智能应用服务商,发展“模型即服务”、“智能体即服务”等,打造人工智能应用服务链。健全人工智能应用场景建设指引、开放度评价与激励政策,完善应用试错容错管理制度。加强知识产权保护、转化与协同应用。加快重点领域人工智能标准研制,推进跨行业、跨领域、国际化标准联动。

(十一) 促进开源生态繁荣。支持人工智能开源社区建设,促进模型、工具、数据集等汇聚开放,培育优质开源项目。建立健全人工智能开源贡献评价和激励机制,鼓励高校将开源贡献纳入学生学分认证和教师成果认定。支持企业、高校、科研机构等探索普惠高效的开源应用新模式。加快构建面向全球开放的开源技术体系和社区生态,发展具有国际影响力的开源项目和开发工具等。

(十二) 加强人才队伍建设。推进人工智能全学段教育和全社会通识教育,完善学科专业布局,加大高层次人才培养力度,超常规构建领军人才培养新模式,强化师资力量建设,推进产教融合、跨学科培养和国际合作。完善符合人工智能人才职业属性和岗位特点的多元化评价体系,更好发挥领军人才作用,给予青年人才更大施展空间,鼓励积极探索人工智能“无人区”。支持企业规范用好股权、期权等中长期激励方式引才留才育才。

(十三) 强化政策法规保障。健全国有资本投资人工智能领域考核评价和风险监控等制度。加大

人工智能领域金融和财政支持力度，发展壮大长期资本、耐心资本、战略资本，完善风险分担和投资退出机制，充分发挥财政资金、政府采购等政策作用。完善人工智能法律法规、伦理准则等，推进人工智能健康发展相关立法工作。优化人工智能相关安全评估和备案管理制度。

(十四) 提升安全能力水平。推动模型算法、数据资源、基础设施、应用系统等安全能力建设，防范模型的黑箱、幻觉、算法歧视等带来的风险，加强前瞻评估和监测处置，推动人工智能应用合规、透明、可信赖。建立健全人工智能技术监测、风险预警、应急响应体系，强化政府引导、行业自律，坚持包容审慎、分类分级，加快形成动态敏捷、多元协同的人工智能治理格局。

四、组织实施

坚持把党的领导贯彻到“人工智能+”行动全过程。国家发展改革委要加强统筹协调，推动形成工作合力。各地区各部门要紧密结合实际，因地制宜抓好贯彻落实，确保落地见效。要强化示范引领，适时总结推广经验做法。要加强宣传引导，广泛凝聚社会共识，营造全社会共同参与的良好氛围。

国务院

2025年8月21日

全国一体化算力网算力池化、算网安全相关技术文件公开征求意见

原载：“国家数据局”微信公众号

近日，全国数据标准化技术委员会（以下简称“全国数标委”）秘书处面向社会公开征求《全国一体化算力网 智算中心算力池化技术要求》《全国一体化算力网 安全保护要求》2项技术文件意见。至此，全国一体化算力网9项技术文件已全部发布，全国一体化算力网标准体系建设基本完善。

前期，为深入贯彻落实《关于深入实施“东数西算”工程 加快构建全国一体化算力网的实施意见》等政策文件，国家数据局指导全国数标委秘书处完成全国一体化算力网算力并网、数据资源管理

与调度、算力多量纲计费、算力算效衡量、算力运营服务、算力监测接口、算力中心能力等7项技术文件的研制并公开征求意见。新公开的《全国一体化算力网 智算中心算力池化技术要求》《全国一体化算力网 安全保护要求》2项技术文件，进一步提高全国一体化算力网算力调度的效能和安全，定义了算力资源抽象、资源池化管理、任务式资源申请与调度、业务编排等功能，提出了构建算力网多层次、全方位、可持续的安全保护能力要求。这9项技术文件的研制，标志着全国一体化算力网建设从谋划布局阶段进入到落地应用阶段，为构建全国算力资源“一盘棋”、监测“一本账”、调度“一张网”新发展格局，实现算力的普惠化、市政化、公共化服务奠定技术基础。

下一步，国家数据局将指导全国数标委秘书处组织做好技术文件意见收集及修改完善工作，加快系列技术文件的落地实施，推动构建联网调度、普惠易用、绿色安全的全国一体化算力网，助力网络强国、数字中国建设，着力打造全国统一大市场建设的数字基座。

附件：关于征求全国一体化算力网2项技术文件意见的通知

关于征求全国一体化算力网2项技术文件意见的通知

各有关单位：

贯彻落实《关于深入实施“东数西算”工程 加快构建全国一体化算力网的实施意见》《国家数据基础设施建设指引》等政策文件要求，充分发挥标准在加快算力池化先进技术部署应用及统筹算力发展与安全保障等方面的规范和引导作用，在国家数据局指导下，根据《全国一体化算力网 监测调度平台建设指南》体系框架中关于算力资源和算力安全部分规划，全国数据标准化技术委员会（以下简称“全国数标委”）秘书处牵头研制了《全国一体化算力网 智算中心算力池化技术要求》《全国一体化算力网 安全保护要求》等2项技术文件（见附件1和2），经研究讨论、修改完善，形成技术文件征求意见稿。

现面向社会公开征求技术文件意见,如有意见或建议,请于2025年8月15日前将附件3反馈至tc609@cesi.cn。

联系人及联系方式:邓老师 18910277158

附件:

1. 技术文件《全国一体化算力网 智算中心算力池化技术要求(征求意见稿)》

2. 技术文件《全国一体化算力网 安全保护要求(征求意见稿)》

3. 意见汇总处理表

全国数据标准化技术委员会秘书处

2025年8月1日

关于征求数据基础设施3项技术文件、可信数据空间3项技术文件意见的通知

原载:“国家数据局”微信公众号

各有关单位:

为贯彻落实《国家数据基础设施建设指引》《可信数据空间发展行动计划(2024—2028年)》等政策文件要求,促进地方、行业、领域、企业开展数据基础设施和可信数据空间的规划、建设、运营、管理,在国家数据局指导下,全国数据标准化技术委员会秘书处牵头研制了数据基础设施3项技术文件、可信数据空间3项技术文件(见附件1-6),经研究讨论、修改完善,形成技术文件征求意见稿。

现面向社会公开征求技术文件意见,如有意见或建议,请于2025年8月24日前将附件7反馈至tc609@cesi.cn。

联系人及联系方式:

许老师 15901508962(数据基础设施)

王老师 13621080102(可信数据空间)

全国数据标准化技术委员会秘书处

2025年8月11日

地方动态|关于加快北京市公共数据资源开发利用的实施意见

原载:“国家数据局”微信公众号

为深入贯彻《中共中央办公厅、国务院办公厅关于加快公共数据资源开发利用的意见》,建立健全本市公共数据资源开发利用制度机制,持续推进公共数据高质量供给、高效率流通、高水平应用,充分释放公共数据要素潜能,加快数据要素市场培育,赋能数据要素市场化配置改革。经市委、市政府同意,结合本市实际,现提出如下意见。

一、夯实公共数据资源开发利用基础

(一)完善公共数据目录。本市实行统一目录管理,动态更新数据的共享、开放属性。市数据管理部门统筹指导各有关行业主管部门组织本行业公共机构全量梳理公共数据目录,并由行业主管部门开展目录上链和数据汇聚。开展授权运营的授权主体指导运营机构建立公共数据授权运营产品和服务清单,纳入目录管理。各区(含北京经济技术开发区,下同)负责组织本区公共数据目录管理和数据汇聚,纳入全市目录体系调度。

(二)提升公共数据质量。各区各部门强化公共数据源头治理,保障数据供给质量,压实数据解读责任。遵循“一数一源一标准”原则,开展人、企业、城市部件等基础数据治理,形成城市基础数据底座。完善数据异议处理反馈机制,建立数据问题核查、反馈和修正闭环。加强供给质量评估和监督检查,全面提升公共数据质量。

(三)开展公共数据资源登记。制定公共数据资源登记制度,对纳入授权运营范围的公共数据资源实行登记管理。建设公共数据资源登记平台,提供规范化、标准化、便利化的登记服务。加强公共数据资源登记专业服务队伍力量。鼓励各部门对持有但未纳入授权运营范围的公共数据资源进行登记,鼓励企业对直接持有或管理的公共数据资源及形成的产品和服务进行登记。

二、畅通公共数据资源开发利用渠道

(四)高效开展政务数据共享。市数据管理部门统筹全市政务数据共享,各部门原则上不得单独建设数据共享通道。升级市大数据平台,通过数据、模型、接口等形式开展共享服务,提升目录区块链、数据探针、隐私保护计算、动态知识图谱等能力。

围绕“高效办成一件事”，组织有关部门形成数据共享责任清单，开展面向应用场景的数据共享。针对区、街道（乡镇）、社区（村）数据需求，依托市大数据平台探索数据集中管理、分级使用，促进市级数据下沉与基层数据反哺。各区各部门共享获取的数据应用于依法履职需要，不得用于其他目的。

（五）有序推动公共数据开放。健全公共数据开放管理制度，在维护国家数据安全、保护个人信息和商业秘密的前提下，依法依规有序开放公共数据。鼓励公共机构开放高价值数据。发布年度公共数据开放计划，动态更新开放目录。优先在与民生紧密相关、社会需求迫切的行业和领域推动全量数据目录和样例数据开放。升级公共数据开放平台，探索以政企合作方式开展开放平台基础设施建设和运营，提升数据产品开发效率和价值转化能力。建立开放数据接诉即办机制，响应社会主体数据需求和应用反馈。各区负责组织本区公共数据开放，并在公共数据开放平台发布。

（六）规范管理公共数据授权运营。修订完善公共数据授权运营管理规定。探索建立公共数据分类分级授权机制，加快建设一批公共数据专区，积极推动分领域授权、依场景授权，逐步探索整体授权运营。市数据管理部门统筹协调管理全市公共数据授权运营，指导监督各区各部门组织实施，各区各部门未经批准不得擅自开展授权运营。公共数据授权运营纳入“三重一大”决策范围，明确授权条件、运营模式、运营期限、退出机制和安全管理责任，授权符合条件的运营机构开展公共数据资源开发、产品经营和技术服务。市数据管理部门要履行行业监管职责，指导监督运营机构依法依规经营，运营机构要面向市场公平提供服务，严禁未经授权超范围使用数据。市财政部门会同市数据管理部门统筹数据资产全过程管理，研究建立收益分配路径。探索建设公共数据授权运营管理服务平台，实现授权运营数据的合规使用和安全管控。鼓励其他经营主体对运营机构交付的公共数据产品和服务再开发，融合多源数据进行创新。

三、加强公共数据资源开发利用服务能力

（七）建立授权运营价格形成机制。授权主体指导运营机构建立数据产品和服务清单，用于公共治理、公益事业的有条件无偿使用；用于产业发展、行业发展的，可收取公共数据运营服务费。公共数据运营服务费实行政府指导价管理，由发展改革部门会同市数据管理部门，按照“补偿成本、合理盈利”的原则核定运营机构最高准许收入；在最高准许收入范围内，由授权主体制定各类产品和服务的上限收费标准，并书面报告市发展改革部门、市数据管理部门；运营机构在不高于上限收费标准的范围内，确定具体收费标准。

（八）强化监督管理。市发展改革部门、市数据管理部门会同授权主体指导运营机构建立健全内部价格管理制度，单独核算并准确记录公共数据资源授权运营经营成本和收入等情况。定期对本市公共数据资源授权运营情况进行披露，授权主体应按规定公开授权对象、内容、范围和时限等情况，运营机构应公开公共数据使用情况、产品和服务清单及相关收费标准，接受社会监督。有关部门依法查处不遵守行业管理有关规定、不执行政府指导价、价格欺诈以及不按规定明码标价等行为。运营机构应依法依规在授权范围内开展业务，不得实施与其他经营主体达成垄断协议或滥用市场支配地位等垄断行为，不得实施不正当竞争行为。各区各部门政企合作项目涉及公共数据资源开发利用的，应按照国家信息化项目评审要求规范管理。

（九）布局新型数据基础设施，强化协同高效的数据服务能力建设，为社会主体提供“原始数据不出域，数据可用不可见”的公共数据资源开发利用创新环境。打造公共数据训练基地，推动语料库共建共享，赋能人工智能应用。建设智慧城市协同创新仿真实验平台，加大场景开放和数据开放，赋能创新资源聚合和价值释放。建设数据流通利用增值协作网络，促进本市公共数据产品和服务融入全国一体化数据市场建设。发挥公共数据引领作用，加快在文化、金融、医疗健康、教育、国土空间、气象、交通等领域形成开放互联、高效流通的城市可信数据空间。

四、释放数据要素市场创新活力

(十) **丰富数据应用场景。**实施“数据要素×”行动,推进应用场景落地,推出一批满足群众办事和民生服务需求的公共数据产品和服务。鼓励水、电、气、热、公共交通等企业对其持有的公共数据资源开发利用。鼓励和支持企事业单位和社会组织有条件无偿使用公共数据开发公益产品,提供便民利民服务。通过建立市场供需撮合机制,打造高质量数据集定制开发、迭代更新与开放使用的环境,形成一批行业领域高质量公共数据集,支持大模型开发、训练和应用。

(十一) **加强央地协同和区域合作发展。**推动与国家部委、央企总部、平台企业等建立数据合作机制,探索央地共建服务载体等方式,打造国家数据资源中心。积极争取国家部委数据授权使用,形成权责清晰的部市协同授权运营模式。探索建立公共数据资源开发利用区域合作和利益调节机制,完善京津冀蒙一体化算力协同机制,推动京津冀数据协同发展战略合作,建立京津冀公共数据共享协调机制,制定京津冀公共数据资源目录和共享清单。

(十二) **促进公共数据产品流通交易。**充分发挥北京国际大数据交易所枢纽作用,建立与数据流通利用增值协作网络、公共数据专区的联动机制,优先推动金融、医疗健康、交通等领域数据价值高、社会需求大的公共数据产品场内交易,引领高价值、高频次数据流通交易。推进京津冀数据流通交易协同,助力全国一体化数据市场建设。

(十三) **繁荣数据产业发展生态。**促进数据产业发展,完善数据产业链条布局,支持数据采集标注、分析挖掘、流通使用、数据安全等技术创新应用,鼓励开发数据模型、数据核验、评价指数等数据产品。鼓励运营机构以公共数据聚集行业数据,构建新业态新模式。培育一批创新活跃的数据资源企业、数据技术企业、数据服务企业、数据应用企业、数据安全企业和数据基础设施企业等高水平数据要素型企业。落实研发费用加计扣除、高新技术企业税收优惠等政策。支持数据行业协会、学会等社会团体和产业联盟发展,凝聚行业共识,加强行

业自律,推动行业发展。

五、统筹发展和安全

(十四) **加大创新激励。**明确公共数据管理和运营的责任边界,围绕强化管理职责优化机构编制资源配置。按照管运适度分离的原则,探索在保障政务应用和公共服务的前提下,承担数据运营职责的事业单位可按照有关规定转企改制。试点成立行业性、区域性运营机构,并按照国有资产有关法律法规进行管理,符合要求的纳入经营性国有资产集中统一监管。研究制定支持运营机构发展的激励政策。鼓励各区结合区域特色开展公共数据资源开发利用,积极参与公共数据分领域授权的建设运营,推动形成产业集聚的规模效应。

(十五) **强化安全管理。**强化数据安全和个人信息保护,网信、公安、数据、密码、保密等部门按照职责加强对数据资源生产、加工使用、产品经营等开发利用全过程的监督管理。建立健全分类分级、风险评估、监测预警、应急处置等工作机制,开展公共数据流通利用的安全风险评估和数据应用方案的规范性审查。运营机构应依据有关法律法规和政策要求,履行数据安全主体责任,采取必要安全措施,保护公共数据安全。加强技术能力建设,提升数据汇聚关联风险识别和管控水平。依法依规予以保密的公共数据不予开放,严格管控未依法依规公开的原始公共数据直接进入市场。

(十六) **鼓励先行先试。**充分考虑数据领域未知变量,落实“三个区分开来”,鼓励和保护干部担当作为,营造鼓励创新、包容创新的干事创业氛围,在数据产权、市场交易、权益分配、利益保护等领域开展先行先试。充分认识数据规模利用的潜在风险,坚决防止以数谋私等“数据上的腐败”,坚持有错必纠、有过必改,对苗头性、倾向性问题早发现早纠正,对失误错误及时采取补救措施,维护公共数据安全。

六、健全工作机制

(十七) **加强组织领导。**坚持和加强党对数据工作的全面领导,按照市委、市政府部署要求,各区各部门要抓好工作落实。市数据管理部门加强工

作统筹,动态掌握全市公共数据资源开发利用情况,及时协调解决工作中的重大问题。重要情况及时按程序向市委、市政府请示报告。

(十八) 强化资金保障。加强政府性投资等各类资金资源对数据基础设施、平台建设、生态培育等的支持力度,统筹安排数据产品和服务采购经费。鼓励各类金融机构创新产品和服务,加大对数据要素型企业和数据基础设施企业的融资支持力度。引导社会资本有序参与公共数据资源开发利用活动。

(十九) 增强支撑能力。开展全市公共数据资源调查,摸清数据资源底数。推动数据资源标准化、规范化建设,制定公共数据分类分级、数据治理、数据流通交易等标准规范。建立数据领域科技创新和技术攻关机制,实现数据加密、可信流通、安全治理等技术突破。深化首席数据官制度,加强数据领域人才队伍建设,将提高做好数据工作的能力纳入干部教育培训内容。积极参与数据领域国际交流合作。

(二十) 加强评价监督。市审计部门会同市数据管理部门建立审计电子数据共享工作协同机制。开展公共数据资源开发利用成效评价和第三方评估。优选公共数据资源开发利用优秀企业和典型案例,加强经验总结和宣传推广。

地方动态|河北立体推进公共数据资源登记工作

原载:“国家数据局”微信公众号

河北省扎实推进公共数据资源登记工作,形成“制度健全、平台互通、审核严格、服务精准”立体化工作体系。截至8月20日,累计向国家公共数据资源登记平台推送登记信息3539条,登记总量暂居全国首位,有力支撑公共数据合规高效流通使用。

一、加强制度平台建设,构建全链条登记体系。以《公共数据资源登记管理暂行办法》为遵循,率先出台《河北省公共数据资源登记管理实施细则(试行)》,配套制定登记工作规范、操作手册、安全评估报告模板等10项制度文件,构建覆盖数

据资源登记全流程的标准化体系。同步建设的省级登记平台已于5月1日上线试运行,并与国家平台实现数据互通。省级登记平台试运行以来,累计已受理登记申请7033条。

二、严格质量审核监管,确保登记数据真实可靠。建立健全“二级审核+动态复查”机制,严把登记质量关。在国家登记培训会后,立即组织各地各部门开展公共数据资源登记,严格明确数据名称、安全风险评估报告、存证等关键信息审核要求,实行“一次性告知”制度,确保登记主体“只改一次”。

三、全面优化配套服务,支撑公共数据登记落地。建立“提前储备+精准指导”工作模式,结合省直部门大走访、数据赋能产业专项行动,提前梳理储备1582条数据资源。组织省直64个单位和14个市开展登记业务培训,通过“上门讲政策、一对一指导”方式,为石家庄、沧州等地及民政、医保等部门提供全流程服务,协助完成数据存证、应用场景描述等关键环节填报。

四、深化数据开发利用,释放公共数据要素价值。推动公共数据开放共享,数据目录开放2130项,政务数据共享9333亿次。探索公共数据授权运营,省级形成分析报告、API接口、数据集等数据产品,秦皇岛、承德等市利用信用、公积金等5.65亿条数据完成20单授权运营交易。

下一步,河北将按照国家数据局统一部署,进一步完善登记规范体系,加强登记机构建设,优化平台功能,持续推动公共数据资源登记,扩大登记范围、提升数据质量,为积极推进数据要素市场化、价值化打牢基础。

地方动态|“国家公共数据‘跑起来’示范场景—企业服务‘0证明’赋能民营经济创新应用”启动活动成功举办

原载:“国家数据局”微信公众号

8月21日,由福建省数据管理局、福建省农村信用社联合社、泉州市人民政府主办的“国家公共数据‘跑起来’示范场景—企业服务‘0证明’

赋能民营经济创新应用”启动活动在泉州成功举办。本次活动以“数智赋能·金融为民”为主题，旨在推进公共数据赋能金融服务“产品变场景、行业变生态”，充分发挥示范场景赋能民营经济发展。

“企业服务‘0证明’赋能民营经济创新应用”是省数据管理局、省农信联社、泉州市数据管理局联合人民银行泉州市分行、泉州金融监管分局、泉州市地方金融管理局等单位创新打造的电子证照在金融领域社会化应用的有力举措和探索数字化金融场景服务政策性、公益性和普惠性的重要实践。

活动现场，省农信联社解读“企业服务‘0证明’赋能民营经济创新应用”场景建设方案并发布福建农信企业服务“0证明”赋能民营经济创新应用”金融服务方案。本次活动还创新运用AI智能体形式，由“源小美”AI智能体发布泉州市数据管理局和省农信联社泉州办事处联合推出的支持泉州市数据产业高质量发展的若干措施，通过专项信贷资金支持、丰富“泉数”系列金融产品、打造数字主题银行、支持可信数据空间建设、推进行业高质量数据集建设、推动数据产业集聚发展、创新数据资产化服务模式、共建泉州市数字金融实验室等“十二项措施”，全面助力数字金融与实体经济深度融合，赋能泉州市数据产业高质量发展。

国家数据局数据资源司、省数据管理局、省农信联社、省大数据集团、泉州市有关领导，市有关单位领导，省农信联社有关部门领导，各县（市、区）领导，各县（市、区）数据管理部门、农信系统有关负责同志参加。

地方动态|厦门市人民政府办公厅关于印发促进数据产业高质量发展若干措施的通知

原载：“国家数据局”微信公众号

厦门市人民政府办公厅关于印发促进数据产业高质量发展若干措施的通知

厦府办规〔2025〕10号

各区人民政府，市直各委、办、局，各开发区管委会，各有关单位：

《厦门市促进数据产业高质量发展若干措施》已经第116次市政府常务会议研究通过，现印发给你们，请认真组织实施。

厦门市人民政府办公厅

2025年8月4日

厦门市促进数据产业高质量发展若干措施

为贯彻落实党中央、国务院关于推动数据要素市场化配置改革的决策部署，促进数据产业高质量发展，提速建设厦门数据港，构建以数据为关键要素的数字经济，为塑造发展新动能新优势提供有力支撑，结合本市实际，制定本措施。

一、构建数据基础设施支撑体系

（一）强化数据流通利用基础设施建设。支持企业开展可信数据空间、数场、数联网、数据元件、区块链、隐私保护计算等技术实践，带动上下游企业共建场景驱动、技术兼容、标准互通的数据流通利用基础设施，每年评选一批成效显著的项目，每个项目按不超过核定总投资的30%给予一次性补助，最高不超过1000万元。鼓励各类企业接入城市可信数据空间，加快打造以城市可信数据空间为核心，行业、企业可信数据空间为重点，资源集聚、价值共创、生态繁荣、互联互通的数据流通利用基础设施体系。

（二）大力推进行业高质量数据集建设。鼓励企业面向人工智能应用创新，围绕电子信息、机械装备、商贸物流、金融服务等重点行业领域建设高质量的训练、验证、测试、语料等数据集，提供跨企业、跨行业的数据服务。根据数据规模、数据质量和应用效果，每年评选一批成效显著的高质量数据集，每个项目按不超过核定总投资的30%给予一次性补助，最高不超过500万元。

二、加强数据产业规划布局

（三）培育多元数据经营主体。鼓励有条件的行业龙头企业、互联网平台企业设立数据业务独立经营主体。支持数据资源、数据技术、数据服务、数据应用、数据安全、数据基础设施等多元数据经营主体发展，构建大中小企业融通发展、产业链上下游协同创新的生态体系。围绕重点领域培育一批

生态带动和产业引领能力强的数据企业，首次获评的省级标杆数据企业，每家给予一次性最高不超过50万元奖励。

(四) 培育产业创新企业。支持数据企业联合上下游企业、科研机构 and 高校等建立创新联合体，优化产学研协作机制，加大对数据领域基础研究和前沿技术、原创性技术创新，面向人工智能发展，提升数据采集、治理、应用的智能化水平，强化数据标注、数据合成等核心技术攻关，加快科技成果转化和应用落地。首次入选省数字经济核心产业“独角兽”“未来独角兽”“瞪羚”的创新企业，分别给予一次性最高不超过100万元、40万元、20万元奖励，对于晋级企业给予一次性补差奖励。

(五) 打造数据要素创新创业载体。鼓励各单位为数据企业用数、用云、用能、用地和人才引进等提供便利政策。推动数据产业区域聚集，支持各区立足产业基础和资源禀赋，围绕资源汇聚、技术创新、应用牵引、算力支撑等方向，建设数据产业集聚区、数据要素产业园，形成协同互补、特色发展的格局，对获评省数据产业集聚区、省数据要素产业园的，分别给予最高不超过2000万元、500万元资金支持。对获评国家级数据产业集聚区、数据要素产业园的，分别给予最高不超过3000万元、750万元资金支持，已享受上级标准奖励的，给予差额奖励。

三、打造企业数据流通新范式

(六) 提升企业数据治理能力。鼓励企业建立首席数据官制度，健全数据资源管理机制。推动数据管理相关国家标准贯标，规范开展数据治理能力评估，强化企业数据治理和质量管理能力建设。支持企业开展数据管理能力成熟度(DCMM)评估认证，对首次通过DCMM三级、四级、五级评估认证的企业，分别给予一次性最高不超过30万元、40万元、50万元奖励，对已经享受较低等级认证奖励的企业再次通过更高级别给予差额奖励。对首次通过数据安全能力成熟度模型(DSMM)二级及以上评估认证的企业，给予一次性最高不超过20万元奖励。

(七) 推进企业数据资产登记。支持企业对合法合规持有的数据产品和服务进行数据资源持有、数据加工使用权、数据产品经营权确权登记，保护数据权益。企业数据产品和服务首次在符合规定的数据资产登记中心取得数据资产登记证书且通过数据交易机构挂牌的，给予一次性最高不超过10万元奖励。

(八) 培育规范化数据交易市场。鼓励企业贴近市场需求，开发数据产品和服务，合规高效开展数据流通交易活动。支持企业加强行业数据资源整合，扩大行业高质量数据产品和服务供给。支持企业通过数据交易机构开展数据交易，对年度数据交易额达100万元以上的数据供方企业，按照不超过年度数据交易额的5%给予奖励，同一主体年度奖励金额最高不超过200万元。

(九) 发展专业化数据服务生态。培育和引进一批综合型或专业型的第三方数据服务机构，在数据交易机构、数据资产登记平台等提供合规认证、质量评价、资产评估等数据专业服务的企业，年度专业服务累计达到100万元的，给予不超过实际服务合同金额10%的奖励，最高不超过200万元。支持法律、会计、审计、资产评估和金融等机构面向数据资产化需求，建立健全行业标准指引，提升专业化支撑服务水平。

(十) 支持数据跨境流动创新实践。支持企业规范数据跨境流动，企业首次通过国家网络安全监管部门数据出境评估、备案，实现数据合规出境的，给予一次性最高不超过50万元奖励。鼓励企业开展“来数加工”“游戏出海”、文化出口、跨境电商等跨境数据业务，对接国际规则开展数据领域国际合作，打造连接两岸、辐射东南亚、面向“一带一路”和金砖国家的厦门数据港。

四、推动数据应用创新

(十一) 扩大数据资源供给。鼓励行业龙头企业、互联网平台企业在落实数据分类分级保护制度要求和保障各方合法权益的前提下开放数据，支持企业、研究机构和行业组织开展合作，促进数据可信交互、高效协同和融合利用，推动跨行业跨领域

数据互通共享。支持在依法保护个人信息权益的前提下，加强个人数据开发利用。

(十二) 加快公共数据开发利用。鼓励企业参与政务数据应用创新，推动公共数据“便民利企”，提升公共服务水平。支持市、区各级行业主管部门积极供数并发动企业参与省市公共数据社会化开发利用，每年评选一批社会效益显著的应用场景，每个应用场景给予开发主体一次性最高不超过20万元奖励。

(十三) 推动企业数据价值释放。推动数据赋能城市全域数字化转型，支持通过数字化转型提升社会服务智能便捷水平和市民高品质生活水平的应用场景建设，每年评选不超过5个数字应用场景优秀案例，给予一次性最高不超过50万元奖励。聚焦工业、商贸、金融、交通、文旅等重点领域，每年评选不超过5个数据要素驱动的数字化转型示范项目，对入选项目按不超过核定总投资的30%给予一次性最高不超过200万元补助。

(十四) 激发数据应用创新活力。企业数据开发应用项目纳入国家数据主管部门评选的培育试点项目的，每个项目按不超过核定总投资的30%给予一次性最高不超过500万元补助；支持企业围绕工业制造、商贸流通、金融服务、绿色低碳等行业领域，打造一批“数据要素×”典型场景，在“数据要素×”大赛全国总决赛获得一等奖、二等奖、三等奖的单位，分别给予一次性最高不超过100万元、50万元、30万元奖励。

五、优化产业发展环境

(十五) 提高数据领域动态安全保障能力。鼓励数据安全企业拓展数据安全业务，丰富数据安全产品，培育壮大适应数据流通特征和人工智能应用的安全服务业态，促进基础设施安全、数据安全、应用安全协同发展。对创新发展数据分类分级、隐私保护、安全监测、应急处置等数据安全产品和服务的单位，给予一次性最高不超过200万元奖励。

(十六) 强化产业金融支持。在依法合规、风险可控的前提下，支持金融机构探索数据资产增信融资、数据资产质押融资、数据资产保险、数据信

托以及数据资产证券化等金融服务产品。建立数据企业培育库，入库企业进行数据资产质押融资的，融资利率固定为1.5%/年，融资利率与银行利率的差额由财政贴息，每家企业年贴息额度最高不超过200万元。将数据企业新建研发设施项目、数据企业培育库中的重点数据企业纳入技术创新基金白名单，享受项目固定资产融资、企业研发投入融资支持。符合条件的企业可纳入中小微企业融资增信基金予以支持。同一企业同时符合本措施中不同融资政策条件的，贷款贴息按“就高不重复”原则执行。培育数据领域专业性投资机构，鼓励产业创投引导基金对符合条件的数据产业项目给予支持。

六、其他事项

(一) 算力、模型服务、数据领域技术创新攻关、标准化建设、重大创新平台建设、人才培养等方面按本市现有政策给予支持。

(二) 本措施按照年度预算安排兑现和拨付。涉及的各项扶持资金，由市、区按照财政体制共担。本措施第(三)点、第(五)点奖励标准与省级标准不一致的，按省级要求执行。同一项目、同一场景同时符合本措施多项奖补政策的，或市级同类政策与本措施不一致的，按照“就高不重复”原则兑现。在实施过程中，如遇国家、省、市有关政策调整，以新出台政策为准。

(三) 本措施由市数据管理局会同相关部门负责解释。本措施所涉及到的具体评选标准、支持方式、操作流程由市数据管理局会同相关部门按照公开、公平、公正的原则另行制定，并向社会公布。

(四) 本措施自发布之日起施行，有效期至2027年12月31日，2025年1月1日至本措施发布之日前符合奖补条件的，参照执行。

专家解读|深圳市高质量数据集建设探索与实践

原载：“国家数据局”微信公众号

文|深圳市政务服务和数据管理局党组书记、局长周剑明

深圳市委、市政府深入学习领会习近平总书记

关于人工智能的系列重要论述,深刻把握其核心要义与实践要求,切实把学习成效转化为前瞻布局、创新突破的强大动力,大力开展“人工智能+”和“数据要素×”行动,奋力打造具有国际影响力的人工智能先锋城市。数据作为人工智能的核心要素,训练数据的规模质量决定了人工智能发展的高度,深圳以建设人工智能全域全时全场景应用先锋城市为牵引,积极探索高质量数据集建设的特色实践路径。

一、紧扣全市战略,统筹规划建设方向

(一) 强化统筹推动人工智能发展。深圳市从强化统筹、政策引导、科技创新、应用牵引、要素保障、生态培育、规范治理七个方面全面发力,推动人工智能建设。市政府层面持续开展“人工智能+”和“数据要素×”联动调度,确保人工智能发展与数据建设有机协同。组建实体化运作的市人工智能产业办公室,开展产业发展集中攻坚。出台《深圳市打造人工智能先锋城市的若干措施》《深圳市加快打造人工智能先锋城市行动计划(2025-2026年)》等文件,加强政策规划指引。

(二) 呼应人工智能产业数据需求。深圳人工智能产业链条完备,综合实力位居全国第一梯队,现有人工智能企业2600余家。相关产业链覆盖芯片、模型、硬件及应用等全环节,既有华为、腾讯等龙头企业稳步牵引,优必选、众擎等机器人产品不断突破,元戎启行、城市之光等全车智能技术持续进步,还有打造了全球首个工业多模态大模型的思谋科技、建立药物研发智能化新范式的晶泰科技等创新型企业。人工智能产业快速发展催生了对高质量训练数据的迫切需求,也成为深圳高质量数据集建设流通的驱动力。

(三) 紧跟全市战略谋划数据集建设。深圳高质量数据集建设贯彻全市“人工智能+”和“数据要素×”发展战略,对照形成“场景应用最开放、算力供给最普惠、产业生态最健全、创新创业最便捷”的产业发展环境、建成具有国际影响力的人工智能先锋城市目标要求,聚焦数据要素市场化配置改革主线,坚持应用导向,面向产业发展要素需求

供给,加强创新生态建设,全面构建高质量数据集供给、流通、应用体系。

二、聚焦数智融合,系统推动建设培育

(一) 加强政策制度标准引领。深圳市出台专项资金政策,明确2025-2026年每年发放最高5000万元“语料券”专项资金,资助语料数据采购和开放。同步发布“训力券”“模型券”以及人工智能行业应用和政务应用场景开放等资助政策。发布建设机制指引,结合深圳市中级人民法院等单位探索经验,编制《推进人工智能行业大模型落地路径探索指引》,指导政府部门大模型应用和语料建设。组织学习借鉴美国、德国、阿联酋等国家地区人工智能和数据协同建设经验,印发工作方案,加快推进相关领域建设和人工智能创新发展应用。建立公共数据语料制度标准,起草公共数据语料汇聚加工、质检运营、安全合规等全流程业务共10项机制规范、12份技术标准,规范语料数据建设使用。开展知识工程建设探索,面向大模型能力提升所需的法规政策、经验方法、思维链等高阶知识需求,深圳政务服务和数据管理局会同宝安、福田等区结合政务人工智能应用建设经验,探索形成知识工程方法论指导实践。

(二) 促进数据集建设与流通。开展公共数据语料建设,印发《深圳公共数据语料集(训练类)建设工作方案》,基于政务云搭建语料智能标注平台,已汇聚全市3.5PB公共数据语料开展加工标注,涵盖生态环境、交通运输、文化旅游、科学研究等领域。同时,罗湖区探索整合辖区内市区两级医院资源,建设医学语料库,支持病理、脑神经、GCP等领域开展深度合作,一期建设语料规模预计超过3PB。推动行业高质量数据集建设培育,根据国家数据局关于行业高质量数据集建设部署要求,起草全市建设推动方案,组织各区各部门、相关行业协会和代表性企业开展行业高质量数据集建设,形成29个市级储备项目加强培育。先后向国家数据局报送高质量数据集备选项目5个、典型案例10个、数据标注案例17个,其中交通领域高质量数据集项目获国家专项资金支持。深化高质量数据集流通

交易,依托深圳数据交易所、市数据要素流通服务中心等平台,探索建立数据集确权、登记、评估、交易、结算等机制,上架严选高质量数据集42个,持续深化语料数据撮合交易。深圳数据交易所建设跨境数据专区,深化与国际数据企业战略合作。探索提升数据集融合供给能力,结合公共数据资源授权运营和可信数据空间建设探索,支持高质量公共数据和企业数据等融合应用,已在征信金融、气象、商保理赔等领域开展试点,取得较好成效。

(三) 培育语料创新生态体系。加强行业整合性平台建设,深圳市成立人工智能语料联盟和开放算料专委会,发布1100多个多模态开源语料数据集,推动高质量数据集建设、开源开放和流通。市人工智能行业协会、产业协会、人工智能与机器人研究院等机构积极发挥平台整合优势,助力产业政策落实,推动行业高质量数据集供需对接和流通整合。打造具身智能数据采集实训基地,依托广东省具身智能创新中心,在深圳市宝安、龙华等区建设具身智能数据采集基地,构建具身智能开源开放平台和多模态训练开源数据集,依托深圳数字孪生先锋城市建设推动打造城市级仿真训练平台,加快赋能具身智能、自动驾驶等领域大模型实训。开展素养培训和开发者社群打造,举办11期集中培训班和12期专题讲座,超3000人次参加,有效提升政府部门领导干部和业务人员人工智能应用与语料建设素养。推动构建开发者社群,扩大开发者关系网络,开展技术、标准、认证等系列培训。

(四) 引导数字深圳联合创新。深圳市组建数字深圳联合创新中心,以“人工智能+”“数据要素×”赋能高质量发展为主线,坚持政府引导、市场主导,“先实验、后实战”,打造政府、企业、科研院所、投资机构、交易服务机构、专家智库等多方参与的联合创新生态格局。创新中心下设人工智能、数据要素、数字孪生等专业实验室,现已完成首批10个人工智能应用场景需求开放,结合公共数据语料和行业高质量数据集支撑,推动人工智能应用创新发展。

三、深化领域赋能,彰显建设实践成效

(一) 政务领域应用不断深入。深圳政务领域人工智能应用正实现从技术落地向效能转化的纵深推进,在多层级治理场景中涌现出诸多创新实践。深圳市中级人民法院建成全国首个人工智能辅助审判系统,正式启用首个司法审判垂直领域大模型,在最高人民法院提供的案例库、法答网、法信等权威知识体系支持下,形成万亿汉字规模语料库,法官办案所需的法律法规、条文释义、裁判规则、权威案例、法律观点等均实现覆盖。自2024年6月28日上线人工智能辅助审判系统,至2024年底,实现平均结案时间缩短38天,2025年第一季度又进一步下降19天;法官人均结案495件,同比增加74件。深圳市政务服务和数据管理局上线全国首个面向公众的实用型AI政务助手“深小i”。在充分收集国家、省、市关于政务服务的法律法规、政策文件、办事指南、常用问答等数据资料并开展精细化治理的基础上,围绕企业开办、社保、公积金等8个高频重点领域,梳理出超350万字精细化专业知识图谱。遴选优质基座大模型,采用“通专结合”的大模型组合技术路线,推动实现全市域、全领域、全智能的政策解答和办事指引,并围绕高频重点服务事项上线边聊边办、智能辅助申办等功能,协同“@深圳—民意速办”服务体系,实现“智能应答+人工客服”的协同应用,打造全周期、全流程、全闭环“AI+政务服务”工作体系。目前,“深小i”日均处理咨询量超2万件,在政务办事领域应答率超过97%,解答准确率超过94%。根据近4000份调查问卷反馈,超过92%的用户认为“深小i”解决了他们的问题。精准高效的智能服务降低了企业群众政策获取和办事成本,有效提高了服务效率。深圳市福田区基于DeepSeek升级打造福田区政务大模型2.0版,归集1.2亿条近十年各类政务数据,构建覆盖政策法规、办事指南、历史案例等专属知识图谱,在此基础上打造70名AI“数智员工”,业务覆盖会议纪要、公文写作、工单分拨等11大类278个政务场景。“数智员工”上岗后,实现公文格式修正准确率超95%,审核时间缩短90%,民生诉求分拨准确率从70%提升至95%。

企业招商分析筛选效率提升 30%。

(二) 行业创新发展持续涌现。高质量数据集真正成为驱动产业创新的基石性要素，持续释放多行业变革动能。深圳国家高新技术产业创新中心构建全国首个创新情报高质量数据集，拥有由超过 1.1 亿条机构数据、1.4 亿条专利数据、100 亿条舆情数据、3900 万条人才数据、3000 万条行业特色数据等构成的数据集，打造覆盖 34 类实体、56 种关系的多维关联知识图谱。面向政府部门研发产业监测、招商引资等应用矩阵，赋能产业基础摸底效率提升 50%、重点赛道筛选精准度提升 60%。中国联通深圳分公司建设联通 AI 数据集管理平台，构建覆盖多模态数据“采、洗、标、测、用、评”能力，平台全流程智能化占比 65%，数据可用率超 95%，匿名化准确率超 99%，打造自动驾驶、具身智能、通信、金融等八大行业数据集，赋能超百个模型训练调优。深圳华大基因股份有限公司基于 3500 多万例检测数据、超 100PB 测序数据基础，构建中国人群精准数据库，填补东亚人群遗传数据库空白。自研的 AI 标注系统致病变异注释准确率达 99% 以上，效率提升 3 倍，降低基因检测成本 50%，实现核心工具 100% 国产化。深圳市南山区联合北京大学深圳研究生院、深圳埃空间科技、鹏城实验室等单位，创新性融合数十万条高价值冷冻电镜专有数据、200TB 动态蛋白数据、百亿级蛋白质数据库条目，建成高质量蛋白质设计数据集，将早期药物发现周期从传统的 24 个月缩短至 5 个月、药物发现研发成本降低 60%。

当前，高质量数据集已成为“人工智能+”创新发展的核心引擎，深圳将在既有探索实践基础上，面向人工智能与数据要素产业生态协同发展深层需求，进一步推动政策法规创新，强化应用场景牵引，健全要素支撑体系，找准痛点、难点、发力点，深化高质量数据集建设应用，为人工智能创新和数据产业繁荣注入持续动能。

专家解读|夯实高质量数据集底座：完善数据资源体系，助力“人工智能+”创新发展

原载：“国家数据局”微信公众号

文|北京市人民政府副秘书长 北京市政务服务和数据管理局党组书记、局长 沈彬华

随着大模型和人工智能技术快速演进，人工智能产业范式从“以模型为中心”转向“以数据为中心”，高质量数据集成为人工智能能力提升和“人工智能+”场景落地的关键支撑。

早期，业界普遍认为“数据越多越好”，数据集的概念主要围绕数据的“大规模”特征，认为通过海量数据就可以训练出更好的模型。随着应用的深入，大规模低质量的数据集局限性逐步显现，“高质量”数据集成为影响大模型“智商”的核心因素，数据清洗、标注等工作受到重视。伴随着人工智能在工业制造、医疗健康、教育教学等领域的应用落地，通用高质量数据集难以满足细分场景训练需求，高应用价值、高知识密度和高技术含量的行业高质量数据集供给日趋关键。今年 5 月，国家数据局印发《数字中国建设 2025 年行动方案》，明确提出“加强交通、医疗、金融、制造、农业等重点领域数据标注，建设行业高质量数据集”，为相关工作指明了方向。

但在实际工作中，我们仍面临着诸多挑战，制约了行业高质量数据集的高效建设与应用。

一是数据采集标准与转化机制有待进一步完善。各级公共数据平台归集整合能力不断加强，企业数字化转型持续加速，但受数据标准不一、采集误差等影响，数据存在分布偏差、颗粒度不一、采集缺失等状况，导致大量数据沉淀且难以直接使用。同时，为更好支持数据资源向可供人工智能大模型使用的高质量数据集转化，还需进一步完善面向应用端的数据治理、标注、评估和开发利用机制。

二是数据治理技术融合创新有待提升。行业高质量数据集是数据资源和专业知识的融合产物，现阶段行业专识数据集主要依赖人工标注，亟需智能

化、自动化标注工具以及精准的数据合成技术支持,以提升数据集生产效率,满足专业场景对数据集“规模”“质量”“附加知识”的多重需求。

三是高质量数据集专项支持政策有待完善。高知识密度、高应用价值的数据集开发周期长、成本高、复用率低,数据价值转化路径不清,市场回报机制不明,缺乏专门针对行业专识数据集的投资或补贴政策。同时,高质量数据集价值实现面临流通慢、责任界定不清等问题,影响供需双方的积极性和规模化交易,需要进一步构建涵盖高质量数据集建设、流通交易、创新应用、运营收益的政策体系。

为应对上述挑战,需要多方协同发力,推动形成涵盖资源汇聚、流通、应用以及技术创新、制度建设的高质量数据集建设体系,有力支撑人工智能应用创新发展。

一是畅通高质量数据集流通交易渠道。将高质量公共数据集纳入公共数据管理,实现集中管理、高效调用、智能应用,提升在政府部门间的整体利用效能。进一步完善供给渠道,打造高质量数据集流通交易体系,一方面依托公共数据开放平台,打造高质量数据集开放专题,持续保障面向企业和社会公众的高质量数据集普惠供给;另一方面鼓励公共数据专区运营单位,结合本领域市场需求,定向开展高质量数据集融合建设,提升高质量数据集市场化供给能力。同时,支持相关开源平台等规范化、规模化运营,探索高质量数据集开源机制。支持数据交易机构加快构建人工智能行业高质量数据集供需对接能力,进一步整合外部资源力量,引入数据集清洗、标注、合成、质检等领域生态合作伙伴,形成高质量数据集开发治理、供需对接、评估定价等服务能力。

二是加大高质量数据集相关技术攻关力度。聚焦关键环节突破,加大科技研发投入,鼓励相关市场主体打造智能化、自动化的行业高质量数据集标注工具,强化人机协同能力,提升标注效率与精准度。组织攻关多源异构数据融合技术,建立统一跨行业数据格式标准,破解“数据孤岛”难题。推动数据合成等技术迭代,探索模拟稀缺高质量数据集

的有效路径,通过技术创新夯实数据集建设根基。

三是健全数据集建设保障制度。发挥数据标准化技术委员会作用,推动高质量数据集格式、质量、流通有关标准建设。探索原创数据集确权、价值评估、流通交易、收益分配等机制建设,培育可持续供给生态。鼓励各类社会主体共建数据要素创新安全可信环境,充分利用“数据要素×”竞赛活动等渠道,加强对高质量数据集评估测试和应用落地的全面支撑。

专家解读 | 以“智”提质,向“新”而行,擘画信息化赋能新图

原载:“网信中国”微信公众号

近日,国家互联网信息办公室发布《国家信息化发展报告(2024年)》(以下简称《报告》),系统展现了2024年各地区、各部门积极谋划改革创新举措,凝聚政策资源合力,推动信息化发展取得的显著成果。其中,以人工智能为代表的网络信息技术蓬勃发展,以“智”筑基,以“智”增效,推动生产力和生产关系发生深刻变革,驱动信息化迈向数字化、网络化、智能化全面跃升的新阶段,擘画了信息化赋能高质量发展的新图景。

一、以“智”筑基,锻造融合创新“新”能力

习近平总书记指出,“要推动人工智能科技创新与产业创新深度融合,构建企业主导的产学研用协同创新体系”。在多方力量驱动下,我国人工智能创新加速迭代,能力集中涌现,为信息化发展向智能化跃升注入强劲动能。

从技术创新看,大模型技术加速迭代发展。当前,我国人工智能大模型发展按下“提速键”,处于全球第一梯队。《报告》显示,人工智能技术和产品加速规模化应用,截至2024年底,共302款生成式人工智能服务完成备案,注册用户总数超过6亿。人工智能在各学科领域应用不断深化,文献情报、化工、海洋、气象等重点领域大模型有效提升科研效率,加速教育科研模式变革。

从设施支撑看,算力基础设施提档升级。近年来,我国算力基础设施发展成效显著,梯次优化的

算力供给体系初步构建,算力基础设施规模和水平位居全球前列。《报告》显示,截至2024年底,我国智能算力规模达493EFLOPS(FP16),移动物联网加快向“万物智联”发展,移动物联网(蜂窝)用户达26.56亿户。

从要素供给看,高质量数据集建设步伐加快。近年来,国家充分调动社会各方力量,积极推动高质量数据集建设,促进高质量数据集与人工智能结合,助力“人工智能+”行动落地见效。《报告》显示,2024年,我国年度数据生产量达41.06ZB,同比增长25%,累计数据存储量达2.09ZB。数据标注成为新兴产业,7个数据标注基地加快建设,数商企业数量超过100万家。

从人才培养看,智能化人才队伍加速培育。近年来,我国加速构建多层次人工智能人才培养体系,多个高校成立人工智能学院,地方积极实施人工智能人才计划,人工智能训练师、人机交互设计师等新职业、新工种持续涌现。部署“人工智能应用”就业育人项目,惠及学生超过620万人。《报告》显示,信息化人才培养体系更加健全,全国数字经济本科专业达227个。国家教育数字化战略行动深入实施,建成世界最大的国家智慧教育公共服务平台。

从产业创新看,人工智能产业生态加速构建。我国持续加强人工智能基础研究,以应用导向推动新技术向场景纵深渗透,已形成覆盖基础层、框架层、模型层、应用层的完整人工智能产业体系。《报告》显示,2024年,我国操作系统、数据库、人工智能等开源社区持续完善,人工智能等重点领域标准体系建设指南发布。

二、以“智”增效,激发高质量发展“新”动能

习近平总书记强调,“要充分发挥新型举国体制优势,坚持自立自强,突出应用导向,推动我国人工智能朝着有益、安全、公平方向健康有序发展。”当前,我国人工智能正迎来大规模应用窗口期,加速重构产业发展格局。

深化“人工智能+农业”,从“会种地”向“慧

种地”发展。2025年,中央一号文件提出运用人工智能等技术,建设现代化农业,铸造“农业新质生产力”。人工智能正通过多种技术路径赋能农业科技创新,推动农业生产向智能化、精准化、高效化转型。在各类政策的推动下,国家智慧农业创新应用项目加快建设,国产化智慧农业技术中试熟化、推广应用稳步开展。《报告》显示,2024年,我国印发实施《关于大力发展智慧农业的指导意见》,明确了未来5-10年智慧农业发展的总体思路和任务路径。

深化“人工智能+制造业”,从“制造”向“智造”转型。人工智能与制造业融合发展加速,智能产品装备、大模型等在电子、原材料、消费品等行业加快落地,在研发设计、中试验证、生产制造等环节得到应用,智能工厂、智能车间等新模式新业态竞相涌现。《报告》显示,截至2024年底,重点工业企业数字化研发设计工具普及率、关键工序数控化率分别达到82.7%、65.3%。

深化“人工智能+服务业”,从“传统服务”向“智能服务”升级。我国持续拓展人工智能技术在消费、出行、旅游、餐饮、物流等领域场景化应用,智能家居、智能网联汽车、沉浸式文旅、AI点餐、低空配送、即时零售等新业态、新模式持续涌现,为现代服务业繁荣发展注入新活力。《报告》显示,2024年,我国实施数字消费提升行动,数字消费规模超6万亿元,网上零售额达15.23万亿元。

深化“人工智能+政务服务”,从“能办”向“好办”提升。以生成式人工智能为代表的通用人工智能技术,在数字政府建设进程中展现出重要潜力。据不完全统计,阿里云、商汤科技等超过56家大模型厂商已部署政务相关大模型,应用场景覆盖政务问答、智能办公、城市治理等多个关键领域。《报告》显示,2024年,我国电子政务发展指数全球排名第35位。各地依托政务服务平台推进“高效办成一件事”落地,重点事项范围逐步扩大,人工智能等新技术在政务服务部署应用。“数字人大”“数字政协”“数字纪检监察”建设成效显著,数智化赋能司法公平正义。

深化“人工智能+惠民服务”，从“智慧图景”向“幸福实景”落地。人工智能技术广泛渗透医疗健康、社会保障、养老服务、文化体验等民生领域，AI手机、AI电脑、AI电视等新型智能终端产品不断走进千家万户，成为提升百姓生活品质的新引擎。据有关数据显示，我国生成式人工智能产品的用户规模达2.49亿人，占整体人口的17.7%。《报告》显示，根据网络问卷调查结果，88.7%的受访者表示使用过智能家居产品。超80%的受访者曾经使用生成式人工智能应用来查询资料或搜索信息、形成内容摘要或辅助分析、生成文稿。2024年，数智化技术应用丰富拓展公共服务，医疗、健康、社保、就业、养老等民生保障服务信息化水平不断提高，群众获得感幸福感安全感持续提升。

三、以“智”拓界，迈向智能化发展“新”时代

习近平总书记指出，“人工智能是引领新一轮科技革命和产业变革的重要驱动力，正深刻改变着人们的生产、生活、学习方式，推动人类社会迎来人机协同、跨界融合、共创分享的智能时代。”浪潮已至，未来已来，持续深入推进“人工智能+”行动，驱动经济社会智能蝶变，是时代赋予信息化发展的重要使命，也是塑造发展新动能新优势的必由之路。

“人工智能+”赋能科研范式变革。人工智能前瞻性、战略性、系统性科研布局将不断深化，“基础理论—核心算法—平台工具—场景应用”的创新体系将加速构建。人工智能与量子计算、生物科技、新能源等领域深度交织，引发多点突破、齐头并进的链式变革，形成智能、高效、协同的融通创新生态。此外，人工智能与科学研究深度融合，加快科研范式从人力驱动向智能驱动、从单点攻关到生态构建的跨越式变革，推动更多原创性、颠覆性创新

成果持续涌现。

“人工智能+”赋能现代产业体系建设。人工智能将在种植、养殖、渔业等领域横向铺开，以及生产加工、仓储物流、数字营销等环节纵向拓展，推动农业生产向智能化、精准化、高效化转型。“人工智能+制造”行动深入实施，企业将深化通用大模型、行业大模型和智能体在重点场景的应用，智能工厂加速培育，智能制造装备、工业软件和系统集成创新成果加速应用和迭代升级。人工智能还将重塑服务业体系，智能医疗诊断、个性化在线教育等新模式持续涌现。此外，人工智能将促进具身智能、脑机接口、无人驾驶等新兴产业发展壮大，创造大量新经济增长点。

“人工智能+”赋能群众美好生活。在医疗领域，人工智能在疾病诊断、治疗方案制定、健康管理等方面的应用将更加深入，为健康服务注入新动力。在交通领域，智能网联汽车、无人驾驶技术不断成熟，自动驾驶、智能导航系统逐步走进现实。在家居领域，语音助手、智能家电推陈出新，让百姓生活更加便捷。在教育领域，云端学校、智造空间探索建设，智能学伴、数字导师等人机协同教学新模式将得到广泛应用，推动教育事业向更公平、更优质、更具活力的方向迈进。

“人工智能+”赋能社会发展治理。人工智能大模型将以其强大的感知、预测、协同能力，不断提升城市治理、能源管理、环保监测的治理高效性和决策科学性，为公共安全保障、突发事件应对、韧性城市建设等方面提供坚实支撑，为国家治理带来数字化、智能化变革机遇，助力国家治理体系和治理能力现代化。

（技术编辑：何芮）

研究动态



基础理论

1. 为数字法学辩护

(韩旭至)

来源：《华东政法大学学报》2025 年第 4 期

自数字法学被提出以来，引发了一定的学术争鸣。有观点主张，数字法学是“传统法学的数字化”，杂志对数字法学的“偏爱”催生了“学术泡沫”。这不得不让人联想到 1996 年弗兰克（Frank Easterbrook）所提出的“马法之议”。他认为所谓网络法就像是“（关于）马的法律”，终究要回到既有的各法律部门寻求依据。时至今日，网络法的基本体系已经确立，以《电子商务法》《网络安全法》为代表的网络立法已经较为成熟，对网络法的质疑已几乎消失。

2. 走出法学的“马法悖论”：数字法学的反思之反思

(李学尧)

来源：《华东政法大学学报》2025 年第 4 期

自治主义、改良主义和整合主义三分是当前中国法学面对数字化挑战时所呈现的三种主要理论进路。通过对自治主义的反思性分析可知，兼具现实适配力与未来理论张力的双重立场才是可取的。

面对数字技术的指数式迭代，自治主义强调传统法学理论的稳定性和体系自治性；改良主义主张在保持法学传统结构的基础上进行局部修正与功能调适；而整合主义则运用技术驱动的逻辑，意图推动法学的系统性、结构性范式更新。数字法学在理论进路和政策建议上应保持审慎过渡的路径安排，以避免不成熟的概念突进破坏既有的法秩序。在这一背景下，数字法学的知识建构应坚持“自治为本，改良为径，整合为势”的理论定位，从而在保障法学规范基础稳定性的同时，为未来法学范式的演进保留充足的制度弹性与认知空间。

3. 数字法学的理论范式及独立性

(罗有成)

来源：《华东政法大学学报》2025 年第 4 期

现有关于数字法学的论争主要集中在数字法学是否独立存在以及它与传统法学的关系上。理性的回应是，摒弃对传统法学“内部知识”的依赖，转向本体论意义上的数字法学，聚焦数字法学的独特性与独立性。数字技术挑战了传统的规则治理模式，催生出一系列原生性数字问题与“私人化法律”机制，进而动摇了法律的深层结构，为数字法学奠

定了立论基础。与此相对应,法学研究范式正逐步从“规制主义”向“技术主义”转变。数字法学理论范畴的构建需超越学科边界,从法学与多学科知识的融合中展开。如基于人工认知系统的基本原理,可提炼出控制架构、人机对齐、相关关系与涌现自治四个理论范畴。由此可见,数字法学本体论层面的独立性体现为论域的全方位扩展、理论范畴的重构,以及研究范式的转换。它与传统法学是一种迭代升级的关系,而不是附庸于传统法学的“数字内容”。

4. 数字安全价值的生成逻辑及其法治保障

(郑智航)

来源:《华东政法大学学报》2025年第4期

数字安全是指人们在确保数字技术稳定可靠运行的同时,形成的一种具有技术信任和安全认同的主观心理状态。数字安全是数字社会的基础性价值、原发性价值和整体性价值。数字安全价值的目的在于强化数字技术运用过程中工具理性与价值理性的结合与统一。数字安全不仅强调数字静态安全,还强调数字动态安全。安全价值本身就体现了人们对于绝对安全的主观追求和相对安全客观接受的统一。要想使人们对数字技术形成安全信任,就必须让数字技术在经济学上满足生产要素投入、在哲学上符合主客体相互作用原理、在社会学上符合信息传播与个体差异三个基本条件。数字安全价值要想从一项基本价值上升为法律价值还应当具有相应的法律规范基础。我们应当遵循总体国家安全观的基本要求,强调技术与数字安全并重的基本理念,建立多元多层的数字安全治理机制,构建硬法与软法相结合的数字安全法律规则体系。

5. 算法服务提供者版权过滤义务的理论证成与规范构造

(任安麒)

来源:《华东政法大学学报》2025年第4期

在算法技术冲击下,网络服务提供者著作权间接侵权认定规则在司法适用中的难题凸显,算法服务提供者版权过滤义务成为传统“避风港”规则的重要补充。基于算法服务侵权风险的类型化解析、算法过滤技术的实践考察、注意义务的边界厘定,

赋予算法服务提供者版权过滤义务具有必要性、可行性与合理性。在规范构造上,欧盟《数字化单一市场版权指令》中“过滤器”条款的冲突调和范式,以及各成员国实施进程中的本土化改造与制度创新具有借鉴意义。应严格遵循比例原则,构建版权过滤义务主体、适用对象、过滤标准等基本制度;引入最低限度条款、双重事前标记机制与三重事后申诉救济机制,调和矛盾冲突;做好与著作权配套制度、算法备案、新兴技术的衔接工作,实现算法过滤机制的有序运行。

6. 数字时代刑事诉讼理论研究的转型

(郑曦)

来源:《环球法律评论》2025年第4期

数字时代下技术的广泛应用推动了刑事诉讼制度的变革发展,由此衍生出一系列新的法律问题,使得刑事诉讼理论研究的对象发生变化,突出体现在人与工具关系的改变、法律概念的内容更新、刑事诉讼结构的变化三个方面。研究对象变化的背后,是刑事诉讼的价值冲突出现了新的内容,不但传统价值之间的竞争继续,新兴价值也在影响刑事诉讼。为此有必要革新研究方法,除了解必要的数字技术、充分关注法律动态,还需适当运用量化评估、实验研究等科学方法,以应对数字时代刑事诉讼理论研究的的新问题。数字时代下刑事诉讼理论研究的转型势在必行,如何构建理论与实践深度融合的研究范式,为刑事诉讼制度在数字化进程中坚守公平正义与司法为民的核心价值提供理论支撑,是未来刑事诉讼理论研究的重要课题。

7. 数字时代设计型规制的理念及其展开

(高秦伟)

来源:《现代法学》2025年第3期

“通过设计保护隐私”进入法律规范后可以取得良好的效果,面对新技术的发展,数字时代下的政府规制亦要求基于设计而展开。此种设计型规制或通过设计实现规制的理念意味着既需要将设计视为技术架构的一部分,亦需要将设计视为法律的构成,探讨设计的内容及其适法性。这在目前正逐渐成为一些国家对数字技术展开规制的主要理念。设计

型规制是一个持续、迭代的法律与社会过程,它在充分发挥私营部门作用的同时,能够使规制机关全流程、渗透性地规制如算法、人工智能等新技术,尽早预警,并快速、灵活地处置风险。设计型规制可使规制机关在需要不断应对颠覆性技术带来挑战的背景下,既关注公众隐私、安全、生命与健康等议题,又能使规制在步调上与技术发展相协调,从而成为数字行政法规制经济社会的重要方法。在规制展开时,应支持技术人员、企业管理者、用户、外部社会力量和政府等多方主体参与,实现技术与法律的有效互动、合作规制及风险规制的有机衔接。

8. “数字人权”的法理再解构

(刘志强)

来源:《现代法学》2025年第3期

运用黑格尔辩证法的“正题—反题—合题”模式分析“数字人权”,可以揭示其理论困境与实践挑战,并提出创新性综合命题。从正题来看,“数字人权”的人性基础、宪法权威、社会功能存在争议,因此需要反题质疑。“数字人权”的人性基础面临双重悖论。一是正当性悖论,“数字人性”论试图重构人性本质,背离人性,未能充分证明“数字人权”的道德正当性。二是必要性悖论,“人性价值”论未能充分论证“数字人权”作为独立法律概念的必要性。从反题来看,“数字人权”无论是通过规范推导还是解释推演,均难以确立其合宪性基础。一方面,“数字人权”概念的实践效果难以有效促进社会沟通;另一方面,“数字人权”会引发功能异化与焦点偏移的问题。从合题来看,经过正反命题的“辩证”交锋,“数字人权”应实现理论超越,形成综合性命题,即人权概念的“数字转型”。应在价值层面坚持“人性保留”原则,在规范层面维护宪法共识,在实践层面构建公私协作的保障机制,最终形成适应数字时代的人权领域法体系。

9. 论法律领域对法律大模型的能动塑造

(王禄生)

来源:《中国法学》2025年第3期

法律大模型是兼具技术属性与社会属性的大

型技术系统。法律领域通过对价值追求、知识基础、组织偏好、制度准备、认知结构的多重引导,调控法律领域数据、知识、组织、制度、符号等资源的生产、供给与配置,从而能动塑造法律大模型的功能表征、应用场景与扩散路径。中国法律领域的资源禀赋与资源困境在技术动力与技术约束两个层面塑造中国法律大模型的技术风格。因此,不应将法律领域单纯视作法律大模型的应用场景和改造对象,而应赋予其破除法律领域资源困境的能动地位。这就要求以推动中国法律大模型自主创新为总体目标,协同领域内多层次意图形成一致性能动塑造行动。在此基础上,以组织资源为中心化约各资源之间的复杂互动关系,构建法律领域资源的优化方案,打造符合中国式法治现代化核心要求的法律大模型技术风格。

10. 数智时代的过劳问题及其法律因应

(田野)

来源:《中国法学》2025年第3期

数字智能技术广泛应用于劳动管理,在提升效率的同时,也带来了过度劳动加剧的风险。数智时代的过劳具有普遍性、隐蔽性、精神性、过剩与过劳并存等特点。工作时间与私人生活的边界模糊化、数智化驱动的薪酬制度、算法管理的规训与控制是数智时代过劳形成的重要原因,资本的逐利本性是数智时代过劳的政治经济学根源。破解这一困境应当采取系统化的法律治理路径,通过革新劳动基准实现底线控制,通过对劳动者赋权(以离线权为中心)提供过劳自救机制,通过对用人单位课以算法善治义务消除过劳的技术根源,通过将过劳纳入工伤保险和对算法致劳追究侵权责任达成救济目标,由此建构起由基准、权利、义务、救济构成的完整过劳治理体系。

11. 算法时代法律适应性困境与动态构建

(王春业)

来源：《中外法学》2025年第3期

在智能算法通过自主学习、深度学习并显著增强其社会适应性的趋势下，法律规范的灵活性不足而难以适应社会动态变迁的局限性日益显现，导致出现法律规范的适用性困境，甚至在某些领域法律面临被边缘化的危机。算法时代法律适应性困境源于规范结构的静态化特征，立法者完全可以通过运用动态体系论方法，对法律规范进行结构性改造和优化。立法者有必要运用动态体系论方法在法律条款中设置弹性要素，使法律规范获得“学习”的能力；法律解释机关应适时加强对法律规范中要素的解释，实现要素间的动态协调，使法律规范适应不断变化的现实；法律适用者将动态体系化的法律规范转换为算法代码，实现智能化的法律实施机制，可以从根本上破解算法时代法律的适应性困境。

12. 数字产品租赁的规则构建

(陈丽婧)

来源：《财经法学》2025年第4期

向消费者一时供予数字产品用益的交易模式，属于租赁合同的调整范围。数字产品的无形性不妨碍其成为租赁合同的客体。数字租赁应当确立与其他合同类型的适用边界。当消费者使用数字产品时，不受狭义知识产权许可或技术许可合同的调整。数字租赁不涉及权利的最终归属，不同于买卖合同。数字租赁无须再介入个性化的劳务，有别于委托或承揽合同。鉴于现行法尚无特别规范，应依数字产品特质修正租赁法规则。数字产品应至少满足客观的瑕疵标准。出租人的保持义务为更新义务所替代，于此适用举证责任倒置规则。针对打包合同或带数字元素的合同，解除的效力原则上只涉及数字产品部分，例外地可扩张至合同的其他部分。针对数字产品的微小瑕疵，以个人信息提供对价的承租人仍可行使解除权。由于经营者更易掌握数字产品的动态，数字承租人不负瑕疵告知义务。

13. 《网络弹性法案》的监管方法解析：以保护基本权利为名？

(Pier Giorgio Chiara)

来源：European Journal of Risk Regulation, Vol.16, Issue 2 (2025)

内部市场中联网设备的迅速普及，使得其薄弱的网络安全标准备受关注，这体现在广泛存在且常常未修复的漏洞以及成功的网络攻击上。针对网络物理系统的攻击不仅对欧盟的经济产生重大影响，还对消费者的健康、安全和基本权利构成威胁。在联网设备网络安全市场失灵的背景下，欧洲议会和理事会于2024年10月23日通过的《关于含数字元素产品横向网络安全要求的法规》（欧盟法规2024/2847，即《网络弹性法案》CRA）于2024年12月10日正式生效。本文首先阐述了支撑该欧盟网络安全法律文件的三大监管基础选择（即横向方法、基于风险的方法、产品安全方法），随后探讨了《网络弹性法案》在多大程度上提升了基本权利的保护，正如委员会提案的说明备忘录中所声称的那样。

14. 《网络团结法》：欧盟法律体系下新的欧盟网络安全团结机制的框架和展望

(Susanna Villani)

来源：European Journal of Risk Regulation, Vol.16, Issue 2 (2025)

网络安全不仅是个别国家，也是整个欧盟都需要解决的问题。在最近通过的第2025/38号条例（即所谓的《网络团结法》）的基础上，该研究旨在通过回答以下问题来分析建立预防和应对网络事件的超国家能力：相关法案中团结是如何以及在多大程度上具体被拒绝的？该法案规定的机制如何与成员国在更广泛的安全领域的特权具体互动？

15. 网络安全与打击网络犯罪：合作伙伴还是竞争对手？

(Laura Bartoli)

来源：European Journal of Risk Regulation, Vol.16, Issue 2 (2025)

2000年初，网络安全漏洞被归类为“互联网犯罪”，因此通过刑事司法系统的工具进行管理。《布达佩斯网络犯罪公约》制定了新的定罪条款和新的程序指南，更新了数字时代的刑法和刑事诉讼程序类别。不幸的是，这种风格已被证明是不够的。为了应对越来越多的威胁，欧盟已转向更加先发制

人、基于行政法的网络安全方法，以保护关键基础设施和行业免受破坏性攻击。然而，刑事层面并没有被取代：相关的国际文书仍然存在，而且最近已扩大到涵盖更多领域。本文将研究《联合国网络犯罪条约》等新一波网络犯罪立法，试图确定两种不同对比策略与滥用网络行动之间的相互作用和摩擦。

16. 量子计算与法律：应对量子飞跃的法律影响

(Kasim Balarabe)

来源：European Journal of Risk Regulation, Vol.16, Issue 2 (2025)

量子计算的快速发展为法律领域带来了无与伦比的机遇和挑战。本文研究了量子计算的关键法律影响，重点关注知识产权、数据安全、监管和道德考虑。量子算法和硬件的独特特性给现有专利体系带来了重大挑战，需要一个清晰一致的框架来保护量子创新，同时促进合作。量子计算对当前加密方法的威胁凸显了对前瞻性数据保护政策和采用后量子密码学的迫切需求。随着量子技术的不断发展，政策制定者必须与利益相关者密切合作，制定适应性强、基于原则的法规，在促进创新和降低风险之间取得平衡。此外，量子计算的社会和伦理影响也不容忽视；优先考虑能够带来重大社会公益的应用程序并建立强有力的道德准则至关重要。为量子时代做好法律劳动力的准备需要共同努力培养量子素养和专业知识。通过采取积极主动的跨学科方法，法律界可以在塑造量子未来方面发挥至关重要的作用，确保这项变革性技术维护法治、保护个人权利并促进社会的更大利益。

17. “数字时代的增值税”一揽子计划的支柱3——单一增值税登记：其起源、内容和对商业实践的影响

(Pawel Mikula)

来源：ERA Forum, Vol.26, Issue 2 (2025)

欧盟 (EU) 增值税 (VAT) 立法的重大立法修改最近获得通过——“数字时代的增值税”一揽子计划。这些修正案旨在使增值税适应数字经济，减少增值税欺诈，增加税收，创造企业之间的公平竞争环境并减轻官僚负担。这些变化分为3个支柱。

第三个支柱主要旨在限制增值税纳税人需要在多个欧盟国家注册增值税的情况。本文概述了支柱3的起源和目标，介绍了其内容，并通过几个实际示例评估了其功能。

数据确权与流通

1. 数据分层确权的法理构造——基于流通效率与利益平衡的视角 (时建中)

来源：《环球法律评论》2025年第4期

数据产权制度的模糊性导致数据要素市场面临供给侧激励不足、流通端动力匮乏与利用端价值受限的三重困境，其根源在于现行权利配置未能适配数据的双重属性。数据应分层确权，通过主体、客体、内容与权利强度的四维解构，构建兼顾数据要素流通效率与多元利益平衡的框架。主体分层以“内容主体—行为主体”二元界分厘清权益归属，明晰权益边界。客体分层通过“关于”类数据与“控制或处理”类数据的分类实现精准治理，既强化原生信息利益保护，又激励行为主体积极处理数据。权利内容分层依托“三权分置”解构权能模块，适配数据形态演进。权利强度分层通过梯度化设计，在数据双重利益与流通效率之间建立动态平衡。这一体系突破了“重流通轻安全”或“重保护轻利用”的二元对立，系统回应了数据要素市场化的核心矛盾。

2. 论数据来源者与数据处理者的权利义务关系

(申卫星)

来源：《环球法律评论》2025年第4期

数据来源者与数据处理器作为数据生成过程中最主要两类主体，二者之间的权利义务关系构成数据产权配置的逻辑起点。从权利维度观察，其数据权益配置呈现出对数权与对人权的双层结构。数据来源者对原创数据、操作数据等享有持有、使用、收益、处分的对数权，是数据来源者的一阶权利，同时配备向数据处理器主张获取或复制转移数据的对人权作为二阶权利。数据处理器通过持有权、使用权、经营权三权分置实现各类数据利用场景的权能组合配置，达至不同维度内对数据的利用支

配,并以平行持有下的内部关系调和权能冲突。从义务维度观察,数据处理者对数据来源者负有协助义务及保障、克制义务,同时数据来源者对数据处理者、数据处理者之间互相负有不破坏和不恶意竞争的义务。

3. 数据产权登记制度的体系化构建

(张素华)

来源:《环球法律评论》2025年第4期

数据产权登记是数据基础制度的重要组成部分,直接关乎数据要素市场化配置能否实现。这一制度的构建涉及登记的功能定位、客体、效力、登记内容审查以及登记机构的设立与准入等诸多内容和要素。这些要素之间看似独立,实则环环相扣,互为基础。登记的功能定位反映了数据的本质属性与制度内核;登记对象与内容的选择则决定了应采效力对抗主义。在明确区分登记行为和为登记服务的行为的基础上,登记对象和内容的实质审查由登记服务机构承担,登记机构仅负责形式审查,但受“红旗规则”的制约。登记机构的设置宜采数据局与数据交易所的协作并轨制,相应机构的设立在当下以行政审批为宜。与此同时,数据产权登记制度还应与公共数据资源登记在登记客体、效力体系以及治理规则上保持协同,共同助力全国统一数据登记体系的建设。

4. 数据主权的博弈与理论重构

(汤铮)

来源:《环球法律评论》2025年第4期

随着全球数字化进程的加速,数据主权已成为国际社会关注的焦点。数据主权是国家主权权能在数据治理上的具象化,为国家进行有效的全球数据治理提供理论基础。然而,数据主权的理论并不成熟,导致其在具体应用中出现了诸多分歧和博弈。数据主权的实施原则,如数据所在地原则、效果原则、行为发生地原则和数据控制者原则,在应对数据的流动性、分布式存储和跨境传输过程中出现的监管冲突时,均面临不同程度的局限性。同时,数据本地化和数据出境治理这两大实施机制在实践中也引发了国家间的管辖权冲突。应当以国家对无

形物的主权及保护义务为基础,建立以合理联系与国家责任为核心的主权实施机制。这样中国可更有力地主张全球数据治理话语权,并在多边框架下推动达成更有弹性和包容性的规则共识。

5. 论数据发展权

(许恋天)

来源:《现代法学》2025年第4期

数据确权旨在促进数据资源合规高效流通与开发利用,然而,既有确权方案实现了“从所有到使用”的推进,未能完成“从所有到利用”的跨越。数据资源规模化开发利用从根本上受限于个人信息保护、数据安全等用途管控措施。为了充分协调数据资源用途管控与开发利用的复杂关系,有必要借鉴“土地发展权”理论,确立一种规范数据资源开发活动、促进多方主体合理利用的新型数据权利——数据发展权。数据发展权是从数据所有权中分离出的调整数据资源用途变更、开发利用和增值收益分配关系的独有权利,并不以数据资源所有权的配置为逻辑前提。基于我国数据资源社会所有及开发利用主体多元的实际,可以采取数据发展权设立的国家、平台、个人三元分置模式,以此统筹构建公平与效率兼顾的数据发展权运行体系,进而实现国家、平台、个人数据发展利益的公平合理分配。

6. 数据财产权排除强制执行的权益结构

(陈爱飞)

来源:《中国法学》2025年第3期

案外人执行异议及执行异议之诉是案外数据财产权益人请求排除强制执行的程序基石,数据财产权益的实体性质与对抗效力则是其能否成为足以排除强制执行的民事权益的核心要素。基于“数据三权”分置模式,我国可在尊重和保护数据原始主体优先权益的基础上,综合运用新型财产权理论与传统物权理论,从物权本权性数据财产持有权、物债二分的数据财产使用权、以许可使用为代表的数字财产经营权等维度,确立一种三权分置型排除强制执行的权益结构。同时,考虑到三权分置型排除强制执行结构仍然与执行理论及实务中足以排除强制执行民事权益的类型界分存在适配不足,我

国可进一步确立限定物权型排除强制执行结构,以提升数据财产权排除强制执行的权益结构与案外人执行异议及执行异议之诉法理的契合性。

7. 企业数据资产的“财产—经济”分置刑法保护 (赵桐)

来源:《中国刑事法杂志》2025年第3期

企业数据资产与传统财产不同,呈现出“数据资源—数据集合—数据产品”的产权结构,由此为刑法的解释适用带来新挑战。财产犯罪的保护法益是财产上的支配关系。在企业数据资产中,只有体现排他性支配关系的企业数据产品才能够成为财产犯罪的保护对象。行为人通过非法获取、转移或处理企业数据产品,剥夺了企业对数据产品的控制,从而非法获利的,构成转利性财产犯罪。行为人通过技术手段毁灭破坏企业数据产品,进而剥夺企业对数据产品的支配能力的,构成毁坏型财产犯罪。企业数据资源与企业数据集合属于“公共品”,无法被排他性支配,须转向综合的经济犯罪保护方案,即由单一的财产权保护转向企业数据资源、企业数据集合、企业数据产品上的“财产—经济”法益分置保护。对于企业数据资源,应采取秩序性法益刑法保护方案,在行为人破坏技术性保护措施场合,论以计算机犯罪;在同业竞争性经营者非法获取企业数据资源的场合,处以非法经营罪、破坏生产经营罪等。对于企业数据集合,应采取秘密性法益刑法保护方案,以侵犯商业秘密罪进行规制,但须遵循最小干预原则对之严格解释适用。对于企业数据产品,由于其既属于财产又可能构成独创性汇编作品,在符合著作权的独创性认定标准时,可采取著作权法益刑法保护方案。

8. 数据财产权益强制执行的论证与规则构建 (樊创)

来源:《法律适用》2025年第6期

当前数据的资源价值和财产属性已成社会共识,随着以数据企业为被执行人的案件增多,数据财产权益能否执行关乎当事人胜诉权利实现和司法公信力维护。现实中数据财产权属争议、不易查控、个人信息安全等问题,使得数据财产权益能否作为

被执行人责任财产尚需论证。从数据财产属性、市场可流通性、维护当事人权益和促进数据治理方面审视,数据财产权益执行具有正当性和必要性,应当突破传统规则将其纳入可执行财产范围。结合数据“权利束”理论、中国特色土地物权制度实践和“数据二十条”中数据三权分置,数据财产权益可以数据“合法持有”作为执行责任财产认定和权益归属判断基准,其所具有的数据基础性权能、外观可识别性、可对权利合法性评价、减少执行异议等功能,使其具备执行查控、价值评估、处置交付等程序适格性。在具体规则构建上,应聚焦数据利用与安全,以“合法持有”平衡多元主体权利,构建具体执行实施与“数据安全—异议救济”双层规则体系。

9. 人工智能时代数据集合财产性权益的侵权救济 路径 (徐小奔)

来源:《法律适用》2025年第6期

高品质的数据集合是人工智能时代数字经济发展的关键生产资源。在现有财产保护体系内,著作权保护、商业秘密保护、反不正当竞争法保护都为数据集合的财产性权益提供了侵权救济路径,当事人根据实际情况可以选择不同的方案进行维权,也形成了有益的司法经验。这些经验有助于厘清数据集合财产性权益的特点,即数据价值的集合性,明确数据处理者和数据处理行为是数据集合财产性权益保护的激励对象。但是,建立在传统信息产权保护基础上的现有方案也存在保护不周延、不充分之处,难以应对发展越发迅猛的数字经济产业。未来,有必要对数据集合进行赋权立法,在区分单一数据财产与数据集合财产的基础上,重点为数据集合提供更充分的法律保障。

人工智能

1. 人工智能训练侵犯作品复制权吗?

(李春晖)

来源:《华东政法大学学报》2025年第4期

关于使用数据训练人工智能(AI)是否侵犯著作权问题,存在否定著作权法可适用性的“釜底抽薪”进路和承认可能侵犯著作权但主张以合理使用为代表的权利限制的“先进后出”进路。前者有利于新技术、新业态发展。复制权系训练用数据著作权保护问题的核心。就AI训练本身而言:(1)其技术本质决定了训练过程不是对作品的复制;(2)AI的伦理地位决定了其学习过程类似自然人,因而不可能是复制;(3)学习过程中的临时复制不在复制权范围内则基本无疑义。就作为AI训练前导的训练数据准备中的复制而言,一方面,其不具有传播目的和效果,不应被解释为著作权法上的复制;另一方面,其实质仍为临时复制。训练数据复制权问题也是整个知识产权法中实施、使用链条各环节独立权能化的表现之一。后者在历史上具有合理性,但随着科技与市场环境的变化,弊端愈来愈明显,知识产权法逻辑应考虑回归民法规则和民事侵权理论。基于以上理由,著作权法不应适用于AI训练对数据的使用,后者应否规制、如何规制,应重新进行利益衡量和价值取舍。这可在著作权法框架下进行,亦可在数据立法框架下解决。

2. 自动驾驶汽车“电车难题”伦理困境的刑法解答

(魏超)

来源:《政治与法律》2025年第6期

目前的科技无法完全避免自动驾驶汽车“电车难题”的发生,必须在其上路前便预设碰撞程序。个性化伦理设定不但会造成“囚徒困境”,而且会动摇民众对法律的信心,故有必要依照法律规范制定相应决策。乘客作为自动驾驶汽车的受益者,与险情存在因果关系,并非置身事外的中立者,不具备法律上优先保护之理由。保全多数人的功利主义与掺杂多种计算方式的综合考量理论在本质上均属于“为生命定价”之行为,忽略了不同个体间的差异,存在侵犯人性尊严之嫌疑,同样不足取。在

现代社会,原则上应当由个体自行承担风险,故不应为自动驾驶汽车设置对生命的紧急避险程序,仅应当配置减速程序以减轻其遭受的撞击损害。

3. 人工智能在数字政府建设中的应用及其法治化路径 基于DeepSeek在数字行政中应用场景的分析

(王静、张红)

来源:《中外法学》2025年第3期

以DeepSeek为代表的生成式人工智能正成为我国数字政府建设中新的技术变量。各层级政府都将人工智能应用作为数字政府建设的重要切入点。当前政府应用DeepSeek等主要在内部办公、智能问答、政策推送和辅助行政管理等场景。区别于企业和个人,政府应用人工智能技术在目标与任务、数据处理、算法应用的要求、技术应用产生的影响等方面具有特殊性,同时也隐含着合法性、政务数据安全、加深偏见与差别化待遇以及技术避责等风险。人工智能在数字政府建设中的应用及其法治化,应当遵循安全优先原则、合法原则、透明原则和责任原则。应当积极对行政系统的内部治理进行调适,提供以“软法”为先导、兼容“硬法”的法律规则供给、推行监管沙盒策略、采取风险规制的措施并强化对采购和部署的监督。

4. 论用户主张AIGC版权的“最低限度创造性标准”

(蒋炯)

来源:《中外法学》2025年第3期

AIGC的特定表达细节可以被视为用户的贡献,因为它是用户对AI的内容生成潜力加以个性化固定的结果。AI并非纯粹的消极工具,但其非消极性为人所欲、忠人之事,应当被理解为AI作为工具的优势而不是贬低人类贡献的理由。用户利用AI生成文艺内容的行为与委托或者提供被演绎内容的行为缺乏可比性,因而不能机械地套用后两种行为的默认规范来进行评价。针对作品和作者的解释方案应当在不违背法律文本的前提下尽可能实现良好的政策效果。“较低贡献对应薄版权、较高贡献对应厚版权”乃界权成本收益分析在版权领域的体现,在AI时代无需也不应改变。针对创作激

励与公众自由平衡这一目标,核心在于确保知识产权的保护程度与权利人的实际贡献相匹配,防止知识产权的排他性范围过度扩张、保护力度过强。

5. 人工智能立法过程中的刑事法因应

(马国洋)

来源:《中国刑事法杂志》2025年第3期

人工智能立法过程中对于刑事法的关注契合人工智能适应性治理的基本原理,可以基于规则论和工具论两个层面分别从刑法和刑事诉讼法展开。人工智能立法与刑法的协同涉及多场景协作,应强调人工智能立法的前置法和指引性作用,而刑法应秉持谦抑性的理念,就人工智能立法予以场景性的支持。人工智能立法与刑事诉讼法的协同更多关注的是单一场景人工智能应用问题,应要求高位阶的人工智能立法设置专门的场景性规则,刑事诉讼法则应对于人工智能立法的作用进行确认,并进一步通过司法解释等规范细化场景治理方案。在具体制度设计上,一方面,应以发展契合技术性程序正义的人工智能为基准,将可信、安全等原则融入刑事诉讼特色并进行层次化排布;另一方面,应基于人工智能生命周期,分别为人工智能规划和设计、数据选择与处理、算法设计与模型建立、技术审查与检验、技术部署与使用等环节,设计具有刑事诉讼场景适配性的规则。

6. 技术浪潮下法律职业的未来展望——兼论人工智能时代法律职业的变与不变

(江必新)

来源:《法律适用》2025年第6期

以 ChatGPT、DeepSeek 等为代表的新一代生成式人工智能正深刻重塑法律职业。法律职业之“变”体现在多个维度,包括法律职业人工作方式与工具、法律职业角色与分工、法律服务模式与市场格局、法律职业生态、法律制度和机制等方面的深层变革。但是,法律职业的“本质”并未随人工智能技术的发展而改变。法律职业之“不变”主要表现在法律职业的终极目的与核心价值、法律人特有的思维模式与专业知识、法律工作的情感交流与心灵沟通、法律共同体的人文关怀与社会责任等方面。法律职业的“变”与“不变”是相对的,而非绝对对立。

二者在人工智能浪潮推动下进行双向映射与互相渗透。在“变”与“不变”的张力中,要通过构建以人机协同为核心的法律职业新范式,在技术理性与法律价值之间实现平衡,在变革中坚守,在坚守中进化,完成法律职业的深度重构,推进法治中国建设。

7. 论人工智能时代商业方法专利客体的审查标准

(邹龙妹、邓卓)

来源:《知识产权》2025年第4期

人工智能时代,由于商业方法通过算法与技术的深度融合实现模式创新,商业方法更强调技术整合与应用场景的结合,其“技术性”判断显得更为模糊。为针对性地解决我国当前审查实践中存在的技术性要求过于严格、可操作性不强等问题,涉人工智能商业方法专利客体适格性的审查,更需要整体视角下准确把握“技术”的内涵,坚持“无技术肯定不行”的原则,明确“客体审查”与“创造性审查”分离机制,发挥“技术性”标准的“门槛式”作用,合理界定方案“利用自然规律”标准的“不经人为干预”内涵,排除的是人的主观属性,可利用人的自然属性,判断商业需求是否实现向技术问题的转化。

8. DeepSeek 模型蒸馏的著作权法正当性重勘

(林秀芹)

来源:《知识产权》2025年第4期

针对 AI 模型 DeepSeek-R1 蒸馏引发的著作权纠纷展开法理分析。研究基于蒸馏技术特征与法律要件双重维度,解析知识蒸馏的技术本质——通过模型间参数迁移实现知识迁移的学习方法,其 AI 创新型训练的本质与技术中立性可以否定著作权侵权的指控。依据《著作权法》及相关原理,论证模型参数作为功能性技术方案不属于著作权保护范畴。通过中、美、欧、日等立法和司法实践中关于机器学习数据合理使用的“文本与数据挖掘例外条款”,揭示当前法律在技术中立原则与权利人利益保护间的制度张力。根据“三步检验法”和“转换性使用”等分析框架研究表明,DeepSeek-R1 的蒸馏不构成侵犯著作权,此外, AI 模型的蒸馏行为定性应当促进“技术创新—权利保护—公共领域”

的动态平衡，为此，应当肯定蒸馏行为的正当性。

9. 法律与技术的协同演化——以日本文本与数据挖掘权利限制条款为例 (刘影)

来源：《知识产权》2025年第4期

基于技术与制度协同演化理论来应对机器学习对于著作权权利限制条款的挑战，有助于突破技术作用于制度或者制度作用于技术的单向视角局限，填补制度与技术相互作用的动态视角盲区。日本著作权法的立法创新体现在两方面：一是通过权利限制条款的结构重塑，来应对人工智能等新兴技术对现行著作权法的制度挑战；二是通过将文本与数据挖掘特定技术场景纳入权利限制范畴，为人工智能技术进步创设更为充分的制度允纳空间。我国《著作权法》中权利限制条款的结构重塑，应在制度与技术协同演化的总体思路下展开，根据著作权人的不利益程度划分成为三大类，并为之单独设计权利限制条款。机器学习对于作品的使用不会给著作权人带来规范性损害，其权利限制条款的设计思路可分为三步：首先通过抽象性条款来增强法律适用的灵活性，然后通过列举方式将以文本与数据挖掘为代表的特定技术纳入其中，最后通过但书条款的设计来确保著作权人的利益不受显著损害。

10. 重思“表达”概念的规范意义——兼论人工智能用户指令行为的法律性质 (李琛)

来源：《知识产权》2025年第5期

作品即表达，表达的本质性构成是符号。不同的作品类型运用不同的符号媒介，即作品语言。表达能力即运用作品语言的能力，这种能力具有稀缺性，“著作权法只保护表达”因此获得了正当性。由于作品语言之间存在差异，在大多数情况下，不同作品类型之间不构成同一或转换关系。人工智能用户的单纯指令行为只是表示了对生成结果的需求，没有运用相应的作品语言，不属于表达行为。可以用文字“说画”的观点，违反了作品语言对作品类型制约的规律。生成式人工智能的便利性，构成该技术使用的自然动力，无须法律额外激励。无论基于概念分析还是政策分析，都不应把人工智能用户的单纯指令行为解释为创作。

11. 人工智能时代商业秘密保护规则重塑

(冯晓青、李可)

来源：《知识产权》2025年第5期

人工智能凭借其强大的数据处理与自主学习能力，极大提升了信息的生成与利用效率，同时也对商业秘密保护提出了新挑战。激励理论在人工智能环境下并未失效，商业秘密的开发、使用和后续优化仍需被激励。从法政策学的视角看，人工智能生成的信息在符合商业秘密法律要件的情况下，理应受到商业秘密制度的保护。在权益归属方面，将用户确认为人工智能生成商业秘密的权益人，契合商业秘密保护立法宗旨与政策目标。在此基础上，现行法律制度应当结合人工智能的认知水平，优化对商业秘密可知悉性的界定，强化对商业秘密保护措施的要求，并明确人工智能手段在商业竞争中的合法应用边界，以确保商业秘密保护制度在人工智能时代的有效性。

12. 人工智能训练中作品数据来源者利益共享机制研究 (张嘉鑫)

来源：《知识产权》2025年第5期

作品数据凭借其高质量特性，在AI训练中得到了广泛应用，并催生了巨大的经济价值。由于AI训练不应被纳入著作权人的专有权控制范围，著作权并非数据来源者主张其财产利益的权源。但作为数据的初始生成者，数据来源者主动提供或参与贡献作品数据原材料，理应获得相应的财产利益分配，确有必要引入利益共享机制，发挥其权衡、矫正与激励功能。在性质上，利益共享机制赋予了针对作品数据的新型积极利用方式，并基于不当得利制度构造相应的权利义务关系。依据场景对利益共享机制进行合理配置，数据来源者在AI商用场景以及数据持有者与AI训练者合意有偿使用场景中均可主张利益共享。为保障利益共享得以实现，应当通过设置披露义务消除实现利益共享的前端障碍，并以集体治理模式帮助数据来源者获得财产利益。

13. 人工智能法案的强有力治理：人工智能办公室、人工智能委员会、科学小组和国家当局

(Claudio Novelli et al.)

来源：European Journal of Risk Regulation, Vol.16, Issue 2 (2025)

没有执法，监管就什么都不是。这尤其适用于充满活力的新兴技术领域。因此，这篇文章有两个野心。首先，它解释了欧盟新的人工智能法案(AIA)如何由各个机构实施和执行，从而明确了AIA的治理框架。其次，它提出了一种规范性治理模式，为确保AIA的统一和协调执行和立法的履行提供了建议。本文探讨了国家和欧盟机构如何实施AIA，包括欧盟委员会等长期机构以及AIA下新成立的机构（例如人工智能办公室）。它调查了它们在超国家和国家层面的作用，强调欧盟法规如何影响机构结构和运作。这些法规不仅可能直接规定机构的结构设计，而且还间接要求执行AIA所需的行政能力。

14. 侵权人工智能：根据国际、欧盟和英国版权法对人工智能生成的输出承担的责任

(Eleonora Rosati)

来源：European Journal of Risk Regulation, Vol.16, Issue 2 (2025)

与人工智能的开发和使用相关的其他问题相比，对人工智能（“人工智能”）产生的产出在版权及相关权下面临的责任方面的分析被忽视了。本研究通过探讨国际、欧盟和英国法律下的相关问题来填补这一空白。具体来说，该研究涉及可作的复制、责任分配和辩护的可用性。分析最终表明，虽然很明显每个案件都需要根据其自身的是非曲直来决定，但生成式人工智能输出阶段提出了版权法下的几种责任。如果决策者和相关利益攸关方的目标是确保人工智能的平衡和可持续发展，那么与人工智能产生的产生和传播有关的问题就需要得到充分关注，并在辩论中发挥比迄今为止更大的作用，无论是在风险评估和合规方面，许可计划，或在有争议的情况下。

15. 真正基于风险的人工智能监管如何实施欧盟的

人工智能法案

(Martin Ebers)

来源：European Journal of Risk Regulation, Vol.16, Issue 2 (2025)

欧盟(EU)的《人工智能法案》(AI Act)声称基于基于风险的方法，以避免过度监管并尊重立法比例原则。本文认为，基于风险的监管确实是人工智能监管的正确方法。然而，与此同时，该文件表明，《人工智能法案》的重要条款并未遵循真正基于风险的方法。然而，这并不是无法解决的。《人工智能法案》提供了足够的工具来支持面向未来的立法，并按照真正的基于风险的方法实施它。在此背景下，本文分析了《人工智能法案》应该如何按照其基于风险的方法的初衷来应用和实施，以及世界各地的立法者在监管人工智能方面可以从《人工智能法案》中吸取哪些经验教训。

16. 《人工智能法案》过山车：立法过程中基本权利保护的演变和监管的未来

(Francesca Palmiotto)

来源：European Journal of Risk Regulation, Vol.16, Issue 2 (2025)

本文追溯了欧盟人工智能法案（以下简称《人工智能法案》）的立法过程，对其形成过程中所做的选择进行了实证和批判性的描述。它特别关注导致最终文本中增加或减少基本权利保护的动态及其对基本权利的影响。本文采用过程追踪方法，阐明了这项具有里程碑意义的立法背后的制度差异和协议。然后，它分析了政治妥协对基本权利保护的影响。它旨在传达的核心信息是，在阅读《人工智能法案》时要考虑到其制度环境和政治背景。正如本文所示，欧盟三个机构的不同政策目标和任务，加上前所未有的重新起草水平和达成政治协议所需的时间短，影响了《人工智能法案》的制定。展望未来，该文件指出了实施、执行和司法解释在加强人工智能时代基本权利保护方面的作用。

17. 算法透明度、工会信息权与战略性诉讼

(Riccardo Maraga)

来源：Italian Labour Law e-Journal, Vol.17, Issue 1 (2025)

本文旨在分析人工智能的引入为企业内部员工参与所开辟的新空间。具体而言,本文重点探讨了近期由国家及欧洲法律体系引入的工会信息与咨询权,特别关注意大利法律(第1-bis条,1997年5月26日第152号立法令)以及《欧盟法规》2024/1689所规定的集体信息权利,该法规确立了人工智能的统一规则,以及欧洲议会和理事会关于通过数字平台改善工作条件指令。最后,本文对上述监管干预措施进行了批判性反思,并概述了未来员工参与公司决策的可能发展方向。

18. 颠覆性但兼具包容性的人工智能:从劳动法角度看解决方案与边界

(Federica Palmirotta)

来源: Italian Labour Law e-Journal, Vol.17, Issue 1 (2025)

本文聚焦于推动包容性人工智能应用,这种应用不仅要超越人工智能带来的风险,还要积极致力于保障劳动者的福祉和基本权利,例如不受歧视的权利,这得益于所有利益相关方从一开始就参与其中。为此,从劳动法角度出发,本文对“包容性人工智能在工作场所的应用”这一概念进行了阐述,并提供了一些示例,以评估其在近期《欧盟法规》(EU) 2024/1689中的风险金字塔中的法律地位。其次,它分析了可在欧洲法律框架内调动的监管技术,以落实这一包容性原则,因为这些技术要求雇主积极履行预防和消除歧视的义务,从而确保包容性。

平台治理

1. 平台数据互联互通的竞争法治理模式拓展

(陈汇臻)

来源:《现代法学》2025年第4期

高质量互联互通是平台经济快速发展的基础。当前,平台互联互通治理框架更关注竞争法事后规制,与超大型平台事前监管相互配合,形成了“约束”型治理模式。然而,随着平台经济和生成式人工智能的深入发展,互联互通需求逐步向数据层面延伸。单纯基于否定性禁止的治理模式,难以凭借滞后的商

业道德标准评价高度内部化的数据行为,或者制定出差异化的互联互通方案。因此,有必要在竞争法体系下,针对数据增设“引导”型竞争法治理模式。该模式应赋予相关经营者自主决定数据互联互通水平的权利,将数据生产要素中的剩余控制权界定为公有,并将数据互联互通提升为平台竞争规制制度的核心目标之一。建议将数据互联互通确定为直接测度市场力量的因素,以及排除不正当竞争行为违法性的合理事由,并纳入企业合规制度及反垄断法基本救济措施,最终形成治理平台互联互通的“约束—引导”二元框架。

2. 网络平台恢复义务的理论证成与责任制度完善

(于波)

来源:《现代法学》2025年第4期

尽管我国已经引入“通知—必要措施”规则多年,但“恢复”是否属于网络平台的法定义务仍存在较大争议。恢复是一项法定义务,这既是相关法律条文的确切文义和体系解释后的当然结论,也是知识产权利益平衡论和法经济学视角分析的应有之义。在义务性质上,恢复义务是一项程序义务,其核心价值在于判断网络平台的主观过错。恢复义务的违反并不直接引发侵权实体法责任,实体法责任的界定仍应根据侵权责任构成要件来分析。恢复义务的法律责任条款可参照前端环节的“采取必要措施”条款进行完善,通过增加“法律责任”要素,为网络平台处理权利人与用户的纠纷提供明确的法律指引。

3. 网络环境中著作权侵权归责的基础定位

(王艳芳)

来源:《知识产权》2025年第5期

在网络环境中,著作权侵权归责以“避风港”规则为基本抓手,以保护著作权与适当定位网络服务提供者侵权责任为总基调,以兼顾著作权产业与互联网产业协同发展为基础目标。治理网络著作权侵权以“通知—删除”规则为主要路径,并辅之以网络服务提供者“知道”或者“应当知道”(有理由知道)的间接侵权归责规则。巨型互联网平台的兴起以及人工智能、算法、过滤等技术的发展,对现行网络著作权侵权归责模式提出新挑战,但不足

以构成颠覆性力量。裁判者受既有制度设计的约束,应当敬畏立法边界和抑制轻率的创新冲动,对于网络著作权侵权的归责需要谨守中立理念,重点依靠“通知—删除”(通知—采取必要措施)的基本机制,准确把握网络服务提供者间接侵权责任的定位,防止通过降低侵权归责标准、附加积极义务等方式,变相加重网络服务提供者的侵权责任,而破坏著作权制度的利益平衡格局。

4. 平台惩戒性管理行为的救济机制研究——场景化的“权力—权利”平衡 (唐俊麒)

来源:《财经法学》2025年第4期

平台经济的发展离不开平台经营者的技术与组织,但平台严厉的处罚手段也不断侵害着平台用户的权利。平台经营者出于治理需求而具备强大的“管理权力”,依此作出的“处罚”发挥着超出私人事务管理的公共治理效果,也打破了传统的社会治理结构和“权力—权利”制衡体系。当前平台纠纷解决主要依托格式条款限制、赋予平台用户权利、引入政府监管等倾斜保护机制,尚未考虑到平台经营者所拥有的管理权力以及平台自我规制溢出的公共性,难以调和二者间的不平等状态。当前制度既未充分实现对弱势用户权益的倾斜保护,也难以兼顾平台经营者在特定情形下的被动弱势地位。为调整平台惩戒性管理行为中“平台经营者—平台用户”的特殊法律关系,应通过审慎适用比例原则,将法官的利益衡量过程具体化与合理化,从而探寻利益平衡的“新倾斜方案”。

数字行政与司法

1. 大语言模型应用中的司法偏误与认知干预

(李学尧)

来源:《政治与法律》2025年第5期

近期的实验研究和司法实践都表明,单纯依赖简单的人机协作,难以有效遏制法官在自我确认偏误与大语言模型“强化输出”之间的共振效应。为避免大语言模型在司法实践中沦为单向度的技术支配,在挖掘传统法学方法论中关于多元价值论证的基础上,可结合认知科学的防偏误机制,提出“认

知协同司法决策模型(六步法)”。该模型通过在确立争点、规范检索、事实认定、裁判形成、价值审查与最终公开说理的全流程中引入对立论证、逆向思维、强制反驳清单等操作环节,帮助法官保持深度审查与反思能力,抵消大语言模型单一输出的强化效应。案例模拟表明,这一模型在提升信息搜索效率的同时,能够强化裁判对社会多元价值的考量,确保法官主体地位与裁判公正性。

2. 可信时间戳认证电子数据的法律效力与审查认定 (谢登科)

来源:《法律适用》2025年第7期

电子数据已成为信息网络时代诉讼活动中认定事实的核心证据,其“三易”特征使得保障真实性成为关键问题。可信时间戳认证是我国司法实践中重要的技术性鉴真方法,其技术原理包括权威授时、哈希值校验和数字签名等,主要保障电子数据的形式真实性。根据《互联网法院规定》第11条第2款,可信时间戳认证仅能推定数据未被篡改,但无法证明其实质真实性。法院需重点审查当事人是否按标准申请可信时间戳认证、认证书及电子数据是否经验证通过。对于实质真实性,实务中需结合其他证据,通过生活经验和逻辑法则审查,并可引入有专门知识的人辅助当事人质证。对可信时间戳认证规则的正确理解与适用,有助于数字经济时代司法活动中电子数据的审查认定。

3. 在和平抗议中使用面部识别技术:大规模监控行为的风险及其对保护人权的影响

(Giulia Gabrielli)

来源:European Journal of Risk Regulation, Vol.16, Issue 2 (2025)

在执法领域,越来越多地将基于人工智能(AI)的监控技术(例如面部识别)用于国家和公共安全目的,这引起了人们对在缺乏适当保障措施的情况下可能带来的滥用和任意性的潜在风险的严重关切。在国际一级,国际组织、人权监测机制和民间社会一贯强调生物识别系统对保护和促进人权和基本自由的影响,特别是关于大规模监视可能导致侵犯隐私权和集会自由的危险。本稿将评估欧洲人

权法院在格鲁欣诉俄罗斯案（11519/20）中的立场 的背景下。

和最近的监管尝试，评估适用于将这些技术用于国
家安全目的的国际人权和标准，特别是在和平抗议

(技术编辑：李佳丽、麻卓妍)

教研活动

讲座回顾 | 斯坦福大学 Graham Webster & 耶鲁大学 Samm Sacks： 中美人工智能与数据治理

2025年7月22日，受中国人民大学法学院、中国人民大学未来法治研究院、中国人民大学涉外法治研究院邀请，美国斯坦福大学史波格利国际研究所研究学者 **Graham Webster**，耶鲁大学蔡中曾中国中心高级研究员 **Samm Sacks** 出席中国人民大学未来法治研究院“全球荣誉讲席系列讲座”、中国人民大学法学院数字法学教研中心“数字法学系列讲座”，并围绕“中美人工智能与数据治理”展开了精彩的主题讲座。会上，两位研究员就美国近期在人工智能与数据隐私领域的立法和规制动态进行了探讨。

中国人民大学法学院副院长、教授**丁晓东**出席并主持本次会议。中国人民大学吴玉章高级讲席教授**张新宝**，北京大学法学院教授**王锡铤**，清华大学法学院教授**申卫星**，中国人民大学法学院副教授、未来法治研究院执行院长**张吉豫**，中国人民大学法学院教授、未来法治研究院副院长**王莹**等出席本次讲座。

会上，Webster 研究员首先对美国人工智能治理框架进行了分析。拜登政府时期的人工智能治理框架包括两个重点，即人工智能安全性与竞争政策。这一阶段人工智能安全主要通过行政命令推动，内容多为“软法”——不具法律约束力。规制机构会与企业合作制定自愿性治理规范，重视对大模型的测试与监督。同时中美竞争成为了美国人工智能政策发展的主要推动力，尤其在人工智能芯片与军事能力方面。此阶段美国人工智能框架的重点是在“人工智能爆发点”前保持领先。特朗普政府时期对美国人工智能治理框架进行了重大变革，取消了拜登政府时期的诸多行政命令，不再强调“安全”，而转向模糊的“国家安全”话语。



美国斯坦福大学
史波格利国际研究所研究学者
Graham Webster

随后，Samm Sacks 研究员对美国的数据隐私立法动态进行了分析。其对《美国数据隐私与保护法案》（ADPPA）失败的原因进行了探讨。其认为，一方面，美国立法者对州法是否被联邦法取代（preemption）、是否允许个人起诉（private right of action）存在分歧。另一方面，国家安全和地缘政治话语成为推动数据隐私立法的唯一有效论述手段，而这种话语体系未能支撑 ADPPA 走完立法流程。此外，Samm Sacks 研究员还对美国近期的谷歌数据反垄断案进行了分析，并指出强制谷歌向竞争对手开放数据的判决导向可能对数据治理产生深远影响。



耶鲁大学 蔡中曾中国中心高级研究员
Samm Sacks

王锡铤教授对中美人工智能治理框架的差异进行了探讨，认为中美在人工智能领域对安全的理解存在差异。同时，其也指出人工智能立法存在是否统一立法、如何平衡安全与创新、如何构建政府各机构之间的数据共享协调机制等问题。



北京大学法学院教授 王锡梓

申卫星教授探讨了中美人工智能治理逻辑之间的区别以及造成差异的原因。其认为美国更多强调数据隐私、种族歧视、技术滥用等“社会性”安全，中国则更偏重于国家治理能力、社会稳定、技术自主性。同时，其指出人工智能治理不能盲目模仿，而应重视本地化发展。此外，其也就 Sacks 研究员所提出的谷歌案进行了探讨，分析了共享数据的匿名化问题与数据许可制度的应用潜力。



清华大学法学院教授 申卫星

张新宝教授回顾了中国 30 年来对美国技术与制度的模仿与反思。同时，其也指出在人工智能领域的“灯塔式”崇拜或已过时，人工智能治理需警惕对人工智能风险的过度夸大，并防止受到非理性情绪的影响。



中国人民大学吴玉章高级讲席教授 张新宝

张吉豫副教授首先就人工智能领域的法律扩

散进行了探讨。其认为人工智能治理的国际影响力取决于政策的可移植性与规范化程度。同时，虽然企业自律机制有积极效应，但可能仍需法律框架加以制度化。其也分享了开源大模型在中国的发展，并讨论了开源大模型面临的著作权和平台责任等方面的法律问题。



中国人民大学法学院副教授、
未来法治研究院执行院长 张吉豫

王莹教授就美国人工智能治理政策向两位研究员进行了提问。首先，美国政府机构曾发布的人工智能问责政策在拜登政府结束后是否延续仍然存疑，这类代表“软法”或“愿景引导型”的政策是否已随政府更替而中断，或成为“历史性文件”，不再发挥实际政策影响。另一方面，其关注是否存在机制化的渠道使得科技企业可以参与政策制定，抑或者更多依赖于非正式的政治游说与公共舆论博弈。同时，其提出“非安全化”的路径可能有助于促进中美的人工智能合作，例如在中美之间建立某种共同的衡量基准 (baseline) 或话语锚点，以“繁荣” (prosperity) 替代安全/风险作为人工智能治理目标。



中国人民大学法学院教授、
未来法治研究院副院长 王莹

随后，Webster 研究员与 Sacks 研究员分别就与会专家的提问与观点进行了回应，并对学生观众的提问进行了解答。丁晓东教授对两位研究员的精

彩演讲表示感谢。会后，参会学者与专家进行了合影留念，Webster 研究员与 Sacks 研究员与参会学生进行了小型沙龙讨论。至此,本次讲座圆满结束。



讲座圆满结束

讲座回顾 | 哥大法学院李本教授： 人工智能与中美法院的数据使用

2025年7月4日，受中国人民大学法学院、中国人民大学未来法治研究院、中国人民大学涉外法治研究院邀请，美国哥伦比亚大学法学院副院长、中国法研究中心主任、**罗伯特·立夫 (Robert L. Lief)** 讲席教授**本杰明·李本 (Benjamin L. Liebman)** 出席中国人民大学法学院名家法学讲坛，并围绕“人工智能与中美法院的数据使用”展开了精彩的主题讲座。在讲座中，李本教授从美国法院面临的“诉讼过载”问题出发，解释了美国法院对人工智能应用的态度，分析了中美法院在人工智能应用方面的共性问题与阻碍，并探讨了中国法院的对于大数据的利用。



美国哥伦比亚大学法学院副院长、
中国法研究中心主任、
罗伯特·立夫 (Robert L. Lief) 讲席教授
本杰明·李本

中国人民大学法学院副院长、教授、中国人民大学未来法治研究院副院长**丁晓东**出席并主持本次会议。中国人民大学法学院教授**朱景文**，中国人

民大学吴玉章高级讲席教授**张新宝**，中国人民大学法学院教授、中国人民大学未来法治研究院副院长**王莹**，中国人民大学法学院助理教授**刘洋**，中国人民大学法学院助理教授**彭雅丽**，中国人民大学法学院助理教授**陈靖远**，美国加州大学洛杉矶分校 (UCLA) 法学院助理教授**安雨田**以及来自法律科技企业的专家等出席本次讲座。



中国人民大学法学院副院长、教授
未来法治研究院副院长 **丁晓东**

朱景文教授对法院应用人工智能的前景与潜力表示了赞同与认可，同时指出，人工智能裁判的一个局限性可能在于人工智能难以将法官所掌握的“隐性知识”代码化，从而导致人工智能判决结果和法官判决结果之间具有不一致性。



中国人民大学法学院教授 **朱景文**

张新宝教授介绍了实践中企业为保证缔约对方按约履行合同而采取的策略性履行行为，同时就人工智能法官能否认可此类策略性履行行为进行了讨论。



中国人民大学吴玉章高级讲席教授 张新宝

王莹教授回忆了在哥伦比亚大学访问的经历，并指出人工智能的发展已经给自动驾驶法律规制的研究范式带来了新的变革。



中国人民大学法学院教授、
中国人民大学未来法治研究院副院长 王莹

随后，安雨田助理教授、刘洋助理教授、彭雅丽助理教授、陈婧远助理教授也分别发表了与谈意见。来自北京律数科技有限公司与杭州汉资信息科技有限公司的专家分享了中国法律人工智能与大模型的最新发展成果，并向与会学者们就人工智能在法律应用中的适用前景和限度进行了探讨。

李本教授就与会专家所提问题进行了一一回应，并对参会学生所提问题进行了深入解答。在精彩主题演讲和深入与谈交流后，主持人丁晓东教授对李本教授和各位与谈嘉宾再次表示感谢，并热烈欢迎李本教授未来再次访问中国人民大学，至此本次讲座圆满结束。



讲座圆满结束

第四届“中英法治圆桌会议”在英成功举办

2025年7月14日至19日，应西班牙嘉理盖思律师事务所和英国英中协会邀请，中国法学会党组书记、常务副会长王洪祥率代表团访问西班牙和英国，与英中协会共同举办第四届“中英法治圆桌会议”，围绕“网络空间及技术前沿治理”主题进行探讨，并与西班牙相关法学法律组织开展法治交流。

7月17日，第四届“中英法治圆桌会议”开幕式在英国伦敦举办。英中协会主席伊莎贝尔·希尔顿主持开幕式，中国法学会党组书记、常务副会长王洪祥、英国皇家学会外事秘书艾立森·诺布尔、

伦敦大学学院全球政治与网络安全教授玛德琳·卡尔、英格兰及威尔士律师协会技术与法律委员会联席主席阿克伯·达图等出席并致辞。王洪祥在致辞中提出三点看法：一是汇聚法治智慧，共同推动全球人工智能健康安全有序发展。二是深化法治互信，共同维护国际科技合作体系的开放性和包容性。三是加强法治合作，共同应对新兴技术带来的法律挑战。



第四届“中英法治圆桌会议”成员合影

围绕“网络空间及技术前沿治理”主题，福建省高级人民法院分管日常工作的副院长、一级高级法官**欧岩峰**、中国政法大学副校长、数据法治研究院院长**时建中**、北京大学法学院教授、数字法治研究中心主任**王锡铤**、清华大学法学院教授、智能法治研究院院长**申卫星**、广东外语外贸大学法学院教授**王燕**、中国人民大学法学院副教授、未来法治研究院执行院长**张吉豫**、中国社会科学院法学研究所网络与信息法研究室副主任、副研究员**周辉**分别在专题发言，中国法学会培训中心主任王伟主持第一专题。

在西班牙，代表团访问了嘉理盖思律师事务所和马德里律师协会，分别商讨了中西论坛法律委员会后续合作计划，签署新合作备忘录的总体内容，并围绕中欧高质量共建“一带一路”法治保障，就涉外法律服务和法治人才培养合作等开展了深入交流，得到外方高度赞赏和积极响应。访问取得圆满成功。



代表团访问西班牙嘉理盖思律师事务所

中国政法大学与联合国教科文组织联合举办 WAIC2025 “人工智能规则全球对话与协同发展”论坛

7月27日，由司法部指导，中国政法大学和联合国教科文组织联合主办的2025世界人工智能大会“人工智能规则全球对话与协同发展”论坛在上海世博中心成功举办。作为“2025世界人工智能大会暨人工智能全球治理高级别会议”（以下简称“WAIC2025”）的重要主题论坛之一，本次论坛聚焦全球人工智能治理的前沿问题，围绕更常态化的国际对话、更广泛的人工智能应用、更明确的人工智能治理规则、更程度的互利共赢深入交流分

享，凝聚发展共识，推进规则互鉴，深化全球人工智能治理合作。

中华人民共和国司法部副部长**武增**，上海市副市长、市公安局局长**张亚宏**；阿尔及利亚知识经济、初创和微型企业部长**努尔丁·瓦多赫**；联合国秘书长数字与新兴技术特使办公室成员、数字与新兴技术特使高级顾问**梅赫迪·斯内内**；柬埔寨邮政电信部信息通信技术总局局长**萨姆·塞瑟雷**；中国政法大学党委书记**姜泽廷**，联合国教科文组织东亚地区办事处主任兼代表**夏泽翰**出席论坛并致辞。来自国际标准化组织等十余个国家和地区的近200名国际组织负责人、业界领袖、政府部门代表、学术机构代表及专家学者共襄盛举。中国政法大学副校长卢春龙主持论坛开幕式。

武增在致辞中指出，法治是人工智能健康发展的重要保障，要坚持以人为本、智能向善，坚持平等互利、开放包容，坚持立足当下、放眼未来，推动人工智能法治建设和全球治理体系规则完善，让人工智能更好服务经济社会发展，让技术创新造福全人类。

张亚宏在致辞中指出，当前上海正在深化人工智能全产业链布局，将深化规则对话与协同、强化制度创新与引领、拓展治理协同与合作，着力打造“智能向善的治理合作先行地”，为弥合全球智能鸿沟贡献“上海方案”、“中国智慧”。

姜泽廷和夏泽翰代表主办方致辞。姜泽廷表示，中国政法大学将人工智能治理能力建设，作为服务国家发展全局的关键支点，加强人工智能立法研究，积极推动建立常态化、多层次的全球人工智能治理对话机制和平台，推动人工智能领域的法治文明交流互鉴，促进各国在监管经验共享、技术标准互认、风险评估协同等领域的务实合作。

夏泽翰表示，人工智能迅猛发展，重塑经济社会的同时也带来治理挑战，联合国教科文组织号召：以国际共识框架为治理基准、聚焦南方国家建设弥合数字鸿沟、依托WAIC等平台强化包容性治理，让人工智能成为“向善”力量。

论坛重磅推出中国政法大学数据法治研究院张凌寒教授领衔的两项科研成果——发布“全球人工智能治理规则地图2025”并启动“全球人工智能规则合作研究平台”，推进人工智能协同治理迈上新台阶。“全球人工智能治理规则地图2025”以可视化呈现全球各个国家和地区的人工智能治理规则，破解规则碎片化困局。“全球人工智能规则合

作研究平台”汇聚全球智慧，共同探讨治理方案，推动构建更具包容性与共识性的全球人工智能治理框架，形成更具共识性的全球治理方案。两项成果将持续推动全球各方资源整合，搭建更广阔的人工智能国际交流与合作空间。



“全球人工智能规则合作研究平台”启动

中央网信办网络法治局副局长唐磊、中国政法大学副校长刘艳红等6位实践与理论界专家学者做主旨发言，5位行业专家展开深度圆桌对话。在开放、共享、互促的良好氛围中，与会同仁聚焦真问题，凝聚大共识，探索新路径，为推动人工智能法治建设与健康发展建言献策。



圆桌对话现场

从理论研究到实践落地，从规则构建到全球协同，从国别实践到治理共识。此次论坛是我校落实第九次党代会精神的重要实践，是促进全球法治文明交流互鉴的集中体现，也为全球深化人工智能领域标准化合作搭建了开放、包容、前瞻的对话平台。中国政法大学将以此次论坛为契机，统筹发展与安全，持续推动全球人工智能治理规则与标准的协同创新，为形成更加紧密的全球人工智能治理命运共同体贡献更多实践成果，着力推进我国在人工智能全球治理领域发出“中国声音”、打造“中国名片”、贡献“中国智慧”。

与会专家在采访中对此次论坛给予高度评价。联合国教科文组织东亚地区主任Khan教授对论坛的务实性与创新性予以高度评价，他肯定了中国政

法大学在全球人工智能治理议程中做出的贡献，认为学校长期以来在人工智能法律与伦理研究领域扮演着至关重要的角色，不仅汇聚了国际专业知识，也不断推动法律与政策界对人工智能伦理议题的深入理解。欧盟人工智能法主要起草者、麻省理工学院研究员Mazzini教授强调，高校在人工智能治理中扮演着不可替代的“桥梁”角色，中国政法大学作为集科技与法学于一体的高等学府，在构建负责任、可信赖的人工智能治理体系中具有重要的战略意义。国际标准化组织主席曹成焕，日本上智大学副校长、法学教授森下哲朗均对本次论坛推动的全球人工智能政策对话与合作机制予以高度认可，强调了构建全球协作的治理平台对人类未来至关重要，并呼吁各国以开放姿态参与协作，让技术进步真正惠及全人类。

据悉，2025世界人工智能大会暨人工智能全球治理高级别会议，于7月26日至28日在上海举办。作为全球人工智能领域最具影响力的盛会之一，WAIC2025聚焦人工智能发展关键命题，系统刻画智能时代知识版图与时代坐标。大会期间，姜泽廷代表学校出席2025世界人工智能大会开幕式与主论坛。



论坛现场

“人工智能立法前沿与实践研讨会”在上海召开

2025年7月27日下午，由世界人工智能大会组委会、中国政法大学和《数字法治》编辑部主办的“人工智能立法前沿与实践研讨会”在上海世博中心顺利举行。会议聚焦人工智能技术快速发展背景下的立法需求与治理路径，集聚政府代表与国内外高水平专家学者，共同探讨人工智能法治发展的趋势与挑战。

上海市委常委、副市长陈杰、上海市人大常委会副主任宗明、市政府副秘书长、市经济信息化委

主任**张英**出席会议并指导交流，充分体现了上海对人工智能法治建设的高度重视。

“国际经验分享”环节，欧盟《人工智能法案》起草组组长、麻省理工学院研究员 **Gabriele Mazzini** 博士，韩国科学技术所（KAIST）教授、韩国 PIPC 咨询委员会数据处理标准组组长**金炳弼**教授，日本上智大学副校长**森下哲朗**教授，分别介绍了欧盟、韩国和日本在 AI 立法与数据治理方面的最新制度进展与前沿思考，由北京航空航天大学**赵精武**副教授主持分享。

“国内学界观点”环节，西南政法大学校长**林维**教授、中国政法大学副校长**刘艳红**教授、清华大学智能法治研究院院长**申卫星**教授、北京大学人工智能研究院 AI 安全与治理中心主任**张平**教授、中国信通院政治与经济所正高级工程师**辛勇飞**所长和西北政法大学人工智能法治研究中心主任**杨建军**教授，围绕人工智能法治体系构建中的核心范畴、制度张力与治理路径等议题，进行了深入探讨。本环节由华东政法大学韩旭至副教授主持讨论。

会议特别邀请了最高人民法院、中央网信办、全国人大教科文卫委、教育部、国家数据局、科技部、司法部、全国人大法工委、人民法院出版社和公安部的部委领导与会交流，由中国政法大学张凌寒教授主持对话。

中宣部部委代表，上海市人大财经委、上海市司法局、上海市经济信息化委等上海市相关单位代表，以及上海大学和上海师范大学的学者代表也出席了本次会议。

本次研讨会作为 2025 年世界人工智能大会法治专题的重要组成部分，充分展现了中国在人工智能立法领域的制度创新与国际对话意愿，为推动我国 AI 治理体系和治理能力现代化建设提供了重要学术支撑与政策储备。

“人工智能伦理与法治”论坛在上海圆满收官

2025 年 7 月 26 日下午，以“人工智能伦理与法治”为主题的 2025 人工智能法治论坛在上海世博中心成功举办。论坛由世界人工智能大会组委会办公室指导，上海市法学会、华东政法大学、上海政法学院主办。

上海市法学会会长**姜平**、上海市委政法委常务

副书记**张磊**、华东政法大学校长**肖凯**、上海政法学院院长**刘晓红**出席并先后致辞。上海市法学会专职副会长**施伟东**主持论坛开幕式。

主题演讲

5 位法学法律界及人工智能相关领域的资深专家，围绕可信人工智能发展、人工智能适应型法治、人工智能立法反思、人工智能安全治理、人工智能国际法治等话题发表精彩演讲。华东政法大学刑事法学院副教授吴思远、上海交通大学凯原法学院副教授李潇洋担任论坛主题演讲环节的主持人。

西北政法大学人工智能法治研究中心主任**杨建军**教授、上海交通大学凯原法学院**李学尧**教授、韩国科学技术所**金炳弼**教授、北京邮电大学人工智能法律研究中心主任**崔聪聪**教授、中国网络空间研究院副院长、高级工程师**钱贤良**分别发表了主题演讲。

圆桌对话

4 位上海市人工智能伦理与治理专家委员会的专家委员以“追问、应对、共识”为脉络，从法学、哲学、技术与产业视角展开精彩对话。圆桌对话由上海市法学会专职副会长**施伟东**主持。



圆桌对话现场

上海交通大学凯原法学院院长、特聘教授**彭诚信**、复旦大学哲学学院教授、中国科协复旦大学科技伦理与人类未来研究院院长**王国豫**、蚂蚁集团副总裁、首席技术安全官**韦韬**、亚马逊云科技上海人工智能研究院院长**张峥**分别发表了意见。

技术编辑：林诗敏

数字法评

构思决定论视野下生成式人工智能生成物的版权保护

此处删除了原文脚注,全文请参见《法学研究》2025年第4期,转载或引用请注明出处。

作者:张金平

内容提要:我国法院通过类比摄影作品,初步构建了生成式人工智能(GenAI)生成物的版权司法保护模式,但可能引发三重转向:独创性的判断和实质性相似的比对都从最终表达转向构思,而版权直接侵权责任从执行者转向构思者。然而,GenAI与照相机在执行人类构思上存在根本差异,并非机械固定而是“自主”转化,故类推适用摄影作品构思决定论之“可视化构思+机械固定”二元要件难以支撑三重转向。为此,建议将构思决定论修正为三元要件:一是“可视化构思”延续传统理论;二是“受控固定”要求作者通过合理控制保障其独创性构思是GenAI生成最终表达的有效原因,维系思想与表达二分法;三是“创作记录”作为作者证明独创性贡献的必要证据,既划定权利边界又防范滥用GenAI。

一、问题的提出

人工智能生成物的可版权性问题的在1965年就已经提出,^[1]但一直以来人工智能技术的发展水平都未取得革命性进步,即使英国在1988年创设了全新作品类型“计算机生成作品”,人工智能也没有超出人类预设规则进行创作。然而,2022年底生成式人工智能(简称GenAI)ChatGPT的问世,标志着GenAI可以根据大模型训练得出的统计学规律生成内容——对用户输入的提示词以此规律推断下一个最可能的内容,并以符合该规律的通用表达方式,输出在外观上与传统作品难以区分的“作品”^[2]——打破了纯粹执行人类预设规则的技术瓶颈。那么,对于基于人类提示词而输出的GenAI生成物,是以其由GenAI生成最终表达拒绝提供版权保护,还是能够通过某种理论使其嬗变为用户创作的作品?这并非纯粹的理论问题,而是已经进入中美实

务之中。

目前,美国版权局在三起GenAI生成物登记案中全都驳回了相关登记申请,两起驳回被诉至法院。在“泰勒案”中,史蒂芬·泰勒坚持主张图片《离天堂最近的入口》是由GenAI完全自主生成,并由GenAI享有原始版权,美国版权局和美国法院都以不涉及人类的智力贡献为由驳回泰勒的诉求。^[3]在其他两起登记案件中,登记申请人都主张对GenAI生成物享有原始版权,但美国版权局仍认为GenAI生成物因不涉及人类的智力贡献而驳回。^[4]其中,杰森·艾伦提起了诉讼,主张对GenAI生成的图片《太空歌剧院》享有原始版权,目前该案(下称“艾伦案”)仍在审理中。^[5]

相比之下,我国著作权法不要求作者提起诉讼之前进行作品登记,即便登记也不进行实质审查,目前已有多个生效判决涉及GenAI生成物的著作权保护和侵权问题。其中,在“春风案”中,北京互联网法院进行了开创性探索,首次在全球范围内确立了用户基于提示词对GenAI生成物享有版权(简称“GenAI生成物版权”)的司法保护模式,只要用户从“构思”到最终生成物形成的整个过程中体现了独创性贡献。^[6]此后,“伴心案”、^[7]“武汉文生图案”、^[8]“艺术椅案”^[9]都采用该模式判断原告对涉案生成物是否享有著作权。然而,在“春风案”之后的“广州奥特曼案”中,广州互联网法院认定GenAI提供者对用户利用提示词生成内容的著作权侵权承担直接侵权责任,^[10]似乎引发了GenAI生成物著作权的“权责分离”现象,即用户享有GenAI生成物著作权,但无须为其著作权侵权承担责任。

对于实务界的不同做法,中美学界形成了三种观点。一种是肯定论,即认为生成物构成作品,而且用户对该生成物享有完整的版权,因为该作品体现了人类的智力贡献,即用户在提示词的设计和修改上体现了独创性。^[11]第二种是否定论,即认为生成物不构成作品,因为人类输入的提示词相对于该生成物而言是思想而非表达,而且GenAI根据该指引自动生成了内容,用户无法预期和控制,故不能对其主张版权。^[12]第三种是观望论,即一方面认为

GenAI 生成物版权如成立,可能会导致思想与表达不分、颠覆传统版权理论,但另一方面又担心版权法如不作调整,可能在人工智能时代被边缘化。^[13]

这三种观点充分展示了 GenAI 生成物版权保护的争议性,但现有讨论主要停留在可版权性层面的局部分析,涉及版权侵权的讨论则浅尝辄止,因此各方观点陷入僵局,无法体系性地解决人们心中的疑虑,即按照北京互联网法院的裁判思路,会不会引发版权法的“三重转向”:首先,独创性考察的重点从传统的作品最终表达(即生成物)转向创作构思(即提示词),是否会导致思想与表达不分?其次,两个基于提示词生成的作品之间发生版权侵权纠纷,二者实质性相似的比对对象是否也应从最终表达转向构思(提示词),才能准确划定 GenAI 生成物版权的权利边界?最后,如果生成物版权侵权成立,直接侵权承担者是不是也应当从执行者转向构思者,才能实现“权责一致”,既让用户从具体创作的执行中解放出来,又能督促其负责地使用该工具?在“广州奥特曼案”中,这种转向为何未能实现?

本文基于 GenAI 与照相机执行人类构思的技术原理差异,从类比摄影作品构思决定论的角度,对 GenAI 生成物版权司法保护模式展开体系性反思,旨在追问其在解释论层面^[14]能否逻辑自治地应对上述三重转向带来的挑战。具体而言,本文先阐释支撑摄影作品独创性转向的构思决定论,明确摄影作品版权取决于构思的二元要件及适用条件,然后分析能否通过类推适用摄影作品构思决定论的方式解决 GenAI 生成物版权三重转向问题;若不能类推,应如何修正摄影作品构思决定论,使之既能妥当解决三重转向问题、适应人工智能时代的创作需求,又有利于全面构建我国自主的 GenAI 生成物版权保护模式,提升我国在全球人工智能版权治理规则方面的国际影响力。

二、GenAI 生成物独创性判断之构思决定的类推困境

(一) GenAI 生成物的可版权性逻辑:类推摄影作品

在“春风案”中,原告利用“Stable Diffusion”(简称 SD)进行图片创作,通过正向提示词(如“日本偶像”“酷姿势”)设定创作要求,通过反向提示词(如“绘画”“卡通”)排除不希望出现的内容。SD 根据提示词每次仅生成一张图片,原告根据 SD 生成的初步图片又修改三次提示词得到涉案图片。对此,法院认为,“从原告构思涉案图片起,到最终选定涉案图片止,这个过程来看,原告进行了一定的智力投入……涉案图片体现了原告的智力投入……原告对于人物及其呈现方式等画面元素通过提示词进行了设计,对于画面布局构图等通过参数进行了设置,体现了原告的选择和安排”,因而认定原告对涉案图片享有著作权。此外,法院还指出这与著作权法对照相机生成物提供保护的逻辑一致:“只要运用智能手机拍摄的照片体现出了摄影师的独创性智力投入就仍然构成摄影作品……技术越发展,工具越智能,人的投入就越少,但这并不影响我们继续适用著作权制度来鼓励作品的创作”。^[15]从这个角度而言,我国法院采用了类推解释的方法解决 GenAI 生成物的可版权性问题。

美国版权局则显得相对保守,并未认可用户有关类比摄影作品获得 GenAI 生成物版权的主张。在“《黎明的扎里亚》版权登记案”中,克里斯蒂娜·卡什塔诺娃将其提示词类比为摄影作品的构思,进而对 GenAI 生成的《黎明的扎里亚》插图主张版权,但美国版权局认为,卡什塔诺娃并未采取具体措施就其“主观构思”“给出可视形式”,没能达到美国联邦最高法院在 1884 年“沙乐尼案”中提出的要求。^[16]在后续的“艾伦案”中,杰森·艾伦通过 624 次提示词输入,利用“Midjourney”生成图片《太空歌剧院》,并主张对其“主观构思”“给出了可视形式”,但仍未获得美国版权局认可,理由是艾伦对涉案图片并未用颜料绘画、素描和上色,对“Midjourney”生成图片的控制也没有达到类似摄影师对照片的控制程度。^[17]对此,艾伦提起诉讼,坚持主张其已经满足“沙乐尼案”的要求,并对《太空歌剧院》享有完整的版权。^[18]

从中美实践对比角度而言,我国法院通过类比现行法对摄影作品的保护,在个案中回应了 GenAI 生成物版权的司法保护需求,完成了阶段性历史任务。然而,这种做法并非没有争议,而且我国著作权法并未经历摄影作品可版权性争论的历史阶段,因而有必要进一步深入到最初摄影作品版权保护的正当性论争之中,以获得关于 GenAI 生成物独创性判断对象的启示。

(二) 摄影作品构思决定论

1837 年摄影技术出现,其对当时版权法的挑战并不亚于今天的 GenAI。它是一种在底片上生成图像的技术,对图像形成的唯一力量是光在感光材料上的化学反应。因此,当时的人们普遍认为照片是客观世界自然物体或者人物的具体形象在纸上的复制品,并极力反对人为干预图像的生成。不过,随着照相机的普及,人们对照片也提出了版权主张。在美国第一起摄影生成物版权侵权案(下称“伍德案”)中,原告将摄影生成物(底片)类比为当时(即 1831 年)美国版权法规定的“印刷品”或“雕刻品”,进而对照片主张版权。不过,美国法院并未接受原告的类推解释方法,因为原告虽“结合其想象力和机械技巧实现其构思,但并没有形成法律所要保护的可固定成果(即印刷品或雕刻)”,故驳回其版权主张。[19]

1865 年,美国通过修订版权法明确保护摄影作品,但修法时并未论证摄影作品版权保护的正当性,即照相机生成物是否构成美国宪法意义上的“著作”。[20]以拍摄名人写真而闻名的摄影师拿破仑·沙乐尼,发现他拍摄的著名作家奥斯卡·王尔德的艺术照被盗版,于是起诉被告版权侵权。案件上诉到美国联邦最高法院,原因是被告认为照片既不是美国宪法意义上的“著作”,也不是作者的产物。1884 年,美国联邦最高法院作出终审判决。首先,该院强调美国宪法意义上的“著作”内涵丰富,足以囊括摄影作品,只要其“代表了作者具有独创性的智力构思(conception)”。其次,照相机的工作原理是对镜头前人或物的轮廓的机械固定,在“这个过程中并没有创新、发明或原创性的空间”,

这“对普通照片而言可能确实如此”。再者,该案涉及的肖像画是一幅“有特色且优雅的照片”,并且“原告制作了这幅照片……因为它完全源自其原创的‘主观构思’(mental conception),并对该构思通过下列步骤给出了‘可视形式’(visible form):通过设定奥斯卡·王尔德在照相机前的姿势,选择和安排照片中的服装、窗帘和其他各种配件,安排拍摄主题以呈现优美的轮廓,安排和处理光线和阴影,提示和唤起所需的表情;正是通过这些安排,原告完全制作了涉案作品,他生成了涉案作品”。最终,该院认为,涉案摄影作品是“具有独创性的艺术作品”,完全符合美国宪法的要求,应受到版权法保护。[21]

质言之,美国联邦最高法院认为,摄影作品的独创性由拍摄前的构思决定,至于照相机具体执行该构思的技术如何复杂,无关紧要。[22]对此,不妨称为“构思决定论”。为便于后续判断 GenAI 生成物版权保护的正当性是否可以类推适用摄影作品的构思决定论,在此有必要进一步澄清构思决定论的二元要件:一是可视化构思要件。摄影师对其构思应当通过一系列步骤给出“可视形式”,如对镜头前所要呈现的人物表情、服装、姿态、背景等画面表达元素进行选择和安排,从而有效防止构思停留在脑海中。此外,这些选择和安排应体现摄影师的个性。二是机械固定要件。摄影师的可视化构思最终通过照相机得到机械固定,满足了版权法对最终表达的保护要求。值得注意的是,虽然机械固定是对可视化构思的百分之百转化,但这并非摄影师刻意控制才实现的,而是照相机本身的技术原理使然。正因为第二个要件的客观存在——照相机本身并不参与镜头前所摄对象的选择和安排,摄影师只要在照相机执行前对其构思给出“可视形式”即可,无须对机器执行环节进行人为干预——第一个要件在摄影作品可版权性上才具有决定性意义,即摄影作品之独创性判断无须额外考虑机器执行情况。因此,通过构思决定论,美国联邦最高法院不仅完成了摄影作品独创性判断之对象从最终表达至构思的历史性转向,而且论证了摄影作品版权保

护的正当性,为摄影作品构建了一个“具有浪漫主义色彩的作者”, [23]撑起了整个摄影产业并影响至今。[24]

(三) 类推适用的困境: 构思与执行上的非同质性

正如艾伦主张的,他在 GenAI 执行前将“构思”以提示词的形式给出“可视形式”,按照“沙乐尼案”的构思决定论,他对《太空歌剧院》享有版权。毋庸置疑,类推是用于填补法律漏洞的一种重要解释方法,常用于解决新型疑难案件。然而,类推解释的适用有其重要前提,即二者在“规范上有意义的特征为相同”。[25]根据《著作权法实施条例》第3条,在作品创作行为中,具有著作权法意义的法律特征是作品的构思和执行(即“直接产生”作品的行为)。[26]前者是对所要创作作品的主观构思,后者是将该构思转化为可有形复制的表达。在 GenAI 辅助创作和照相机辅助创作中,机器用户确实都从最终执行中解放出来, [27]但从两种机器辅助创作的技术原理而言,这两种场景下的构思和执行并不相同。

1. 二者在构思上不具有同质性

一直以来,著作权法都坚持思想与表达二分法,即允许同一种思想存在多种表达,通过保护具有独创性的表达,实现作为私权的版权与公众表达自由的平衡。因此,中外有关 GenAI 生成物版权的争议中都会提出一个争议点,即基于提示词主张 GenAI 生成物版权,可能会模糊思想与表达的边界。[28]例如,在利用具有文生图功能的“Stable Diffusion”时,提示词是对所要生成内容的文字性描述,在具备独创性时可以按照文字作品获得保护;但相对于生成的图片而言,这些提示词不论在文字描述上如何详细,都属于图片创作的思想,因为提示词和生成内容之间在作品表达要素上是截然不同的,提示词是文字表达,图片是图形表达。

那么,为什么“沙乐尼案”的构思决定论在摄影作品可版权性上并没有引发思想与表达不分的担忧?这是因为美国联邦最高法院将构思决定论的“构思”进行了限定,而这种限定往往容易因“构

思”与“思想”在术语上的相似性而被忽视。该院强调,摄影师通过照相机执行前的一系列安排,已经就其构思给出了“可视形式”,包括对王尔德的表情、服装、姿态和背景等的安排。[29]换言之,在摄影作品构思决定论框架下,该构思并不是与表达直接对立的思想,而已经是摄影师对所拍摄物体或人物的具体形象在图形元素上作出的选择和安排。

相比之下,在生成式人工智能场景下,用户的构思(即提示词)仅是对所要生成内容的文字性描述,与后续生成内容(如图形)通常并不属于同一作品类型,即使同属文字作品,后续生成的文字的选择和安排也要比提示词复杂得多。对此,我国法院在探索 GenAI 生成物版权司法保护模式时承认,提示词“没有百分之百地告知 Stable Diffusion 模型怎样去画出具体的线条和色彩”; [30]美国版权局也强调,涉案申请人并未像“沙乐尼案”中的摄影师那样采取具体安排就其“主观构思”“给出可视形式”。[31]

2. 二者在执行上不具有同质性

在内容生成层面,照相机的技术原理是机械固定,即将镜头前现实存在的事物转化为底片上的可视呈现。所以,在“沙乐尼案”中,被告否定沙乐尼是作者的重要原因是,照相机的技术原理并没有留给摄影师独立创作的空间,“在机械和材料方面,该过程的其余部分仅仅是机械性的,没有创新、发明或原创性的空间”。[32]不过,美国联邦最高法院将摄影师的独创性所在转向了照相机执行前的构思。借助摄影师拍摄前的构思,摄影师不仅没有改变照相机的机械固定功能,而且还无须担心体现其独创性的可视化构思会受到照相机的干扰,因而照相机彻底变成了忠实于摄影师的创作工具。在这个意义上,也就不难理解“摄影作品”为何被定义为“记录客观物体形象”的“艺术”作品。[33]

相比之下,GenAI 在执行时并非对文字性描述的提示词进行机械固定,而是根据大模型训练所得的统计学规律,在用户设定参数的共同作用下,生成另一种类型作品的具体表达——如“春风案”和

“太空歌剧院案”中的文字转化为图片——或者在文字作品意义上更复杂的表达。这些生成内容的具体表达性元素的选择和安排并非由用户完成，相反，“这些作品的‘版权意义上的传统要素’是由技术决定和执行的”。[34]此外，在 GenAI 执行时可能会出现难以有效转化的情形。考虑到 GenAI 和照相机的执行功能差异，而且艾伦对“Midjourney”生成图片的控制没有达到“类似摄影师对照片的控制程度”，美国版权局否定了艾伦的类推适用摄影作品构思决定论的主张。[35]

综上，按照类推解释的要求，在用户的构思（文字性）与摄影师的构思（可视化）、GenAI 的执行（不受控）与照相机的执行（机械固定）均不具有同等特征的情况下，用户难以主张类推适用摄影作品构思决定论获得 GenAI 生成物的版权。

三、GenAI 生成物版权侵权之构思比对的类推困境

（一）GenAI 生成物版权侵权判断：比对构思还是最终表达

一般而言，版权侵权根据二元要件即“接触+实质性相似”进行判断，但该二元要件均是针对作品的最终表达而言，并不适用于创作的构思。其中，接触要件考察被告是否事先看过或读过原告的作品，但只要原告证明其作品在被告作品创作之前已经公开发表，被告具有接触的可能性。实质性相似则考察被告是否非法挪用了原告享有独创性的表达，如果有实质性相似的部分，而且满足接触要件，则构成版权侵权；如果不满足接触要件，则属于被告独立创作的作品，不侵犯原告版权。[36]GenAI 生成物是否侵害他人按照传统方式创作的在先作品，应当进行“接触+实质性相似”的判断，这在中美都没有争议。[37]

然而，在两个都利用提示词生成的作品之间，应当如何判断版权侵权？按照现行法，法院应遵循传统版权侵权判断方法，仅对最终生成的作品进行比对。[38]然而，如果法院基于提示词承认 GenAI 生成物的版权，此时的侵权比对对象可能就需要考虑转向提示词。例如，美国学者莱姆利指出，“在

基于提示词的版权中，我们想知道的是，判断侵权的对象是否会从作品的相似性转向提示词的‘复制’。如果是这样，在提示词具备独创性时，我复制了提示词就可能侵权。然而，如果我输入了另一个提示词生成实质性相似的结果，那么即便生成的作品相似并用于竞争目的，也不存在侵权”。[39]

在“艺术椅案”中，法院在裁判说理中明确承认 GenAI 生成物的可版权性，但因原告未能举证证明与涉案图片相对应的具体提示词和创作过程，最终并未支持原告的版权主张。不过，考虑到被告参考了原告网上公开的大概提示词，法院仍然展开侵权判断，先比较了提示词，再比较原被告的 GenAI 生成物是否构成实质性相似。[40]很显然，该院已意识到提示词对于 GenAI 生成物的重要影响，因而在侵权比对上先考虑了提示词。如果涉案原告恰好保留了创作过程和相应的提示词，而被告也复制了原告提示词的全部或者绝大部分，考虑到 GenAI 不会对同一提示词生成完全一致的生成物的技术特点，法院是否会因被告复制原告具有独创性的提示词而判决被告侵权？这样的裁判是否会颠覆传统版权法？

（二）摄影作品之比对构思的侵权判断

按照“沙乐尼案”的构思决定论，摄影师可以对照相机生成物主张版权。那么，摄影作品侵权比对是否也转向构思的比对？对于这个问题，版权法已经有上百年的处理历史。

在 1914 年的“格罗斯案”中，摄影师拍摄了一幅艺术照，后转让了其版权。两年后，该摄影师又拍了一张艺术照，相同的模特、发型、姿势和表情，唯一不同的是该模特年长了两岁。受让人起诉该摄影师版权侵权。根据“沙乐尼案”的构思决定论，美国第二巡回法院比对了两张照片中有关模特具体形象的构思，“相同的摄影师，以及几乎一致的模特姿势、灯光、阴影，非常直接地表明第一张照片（的构思）被用到了第二张照片。至于该模特在第二张照片的姿势、灯光、阴影是物理性地呈现在摄影师的眼前，还是说摄影师对这些元素精准组合的主观复制（mental reproduction）已经深刻烙印

在脑海里,以至于他不需要在物理上复制这些元素的组合,在本案中都是无关紧要的”。此外,该院还强调,摄影作品实质性相似的判断者并非专业人士,而是其目标受众即普通人,“艺术家或者鉴赏家肯定可以观察出两张照片的差别……但对于普通的照片购买者而言,第二张照片就是第一张照片的复制品”。因此,该院判决摄影师侵犯了自己已经转让出去的版权。〔41〕显然,法院在判断是否存在抄袭时,将焦点集中于摄影师的构思即对所拍摄模特的具体形象安排,而不在于最终形成的照片本身。

值得注意的是,摄影师在主张他人抄袭时,仍需要证明他人抄袭了其构思中具有独创性的

部分,而不是思想。例如,在1992年的“罗杰斯案”中,原告罗杰斯拍摄了一对夫妻和八只狗的照片,被告将其制作成雕像。法院指出,“受保护的并不是‘一对夫妻和八只狗坐在长凳上’的思想,而是罗杰斯对该思想的表达——具体的安排,特别是灯光以及对这对夫妻和八只狗的具体形象的表达。”〔42〕不过,摄影师对人物具体形象的构思不延及一般的人物特征,如种族、性别、头发颜色,也不扩展到人物的具体身体部位。与此同时,如果被告能够证明其复制的元素满足“固有场景法则”或者思想与表达合并理论,也不构成侵权。例如,在表现都市白领陷入绝望的广告照片中,出现西装笔挺的男士白领、高楼以及该白领在楼顶屋檐边站立准备往下跳等元素,都属于展现都市绝望白领的“固有场景”而不受保护。〔43〕

此外,按照“沙乐尼案”的构思决定论可以推论,在摄影师对拍摄的人物或物体具体形象已作出详细安排的情况下,其他人偷偷跟拍,即便偷拍者在拍摄时有很多独创性选择的空间(如照相机的选择、构图、焦距、曝光时间等),偷拍生成的照片也构成版权侵权,因为偷拍者在现场接触了摄影师的构思,同时照搬了摄影师具有独创性的可视化构思。〔44〕

然而,摄影作品侵权场景下“接触+实质性相似”的比对对象转变,并没有颠覆传统版权法。这

是因为,摄影师的构思得到了机器的彻底执行,因而摄影作品目标受众(即普通公众)在进行实质性相似判断时,比对构思与比对照片不会出现结果上的差异,因而不会背离版权法的保护宗旨,即保护作者对其作品受众市场的控制。〔45〕

(三) GenAI 生成物版权侵权类推摄影作品侵权判断的困境

在类推适用摄影作品构思决定论的情况下, GenAI 生成物版权侵权的判断路径也会很自然地考虑类推适用摄影作品侵权的判断方法,即“接触+实质性相似”的判断对象从最终表达(生成物)转向构思(提示词)。然而,一如类推适用构思决定论进行独创性判断一样,这里的类推适用,也需考虑版权侵权判断规则上二者是否具有相同的法律特征,即 GenAI 场景和摄影场景下的接触和实质性相似是否实质等同。

一是构思接触要件的判断,对 GenAI 生成物的接触并不等同于对提示词的接触。一般而言,对已发表作品特别是网络发表的作品而言,接触要件很容易满足。在摄影场景下,考虑到照片是摄影师构思的精确复制品,对照片的接触等同于对摄影师构思的接触,因而在摄影作品侵权场景下将构思作为接触的对象并不会扰乱既有的版权侵权理论。〔46〕相比之下, GenAI 并非对提示词在文字表达层面的机械固定,而是将其转化为小说、美术作品、视听作品等。所以,对这些生成物的接触无法直接等同于对提示词的接触。

二是构思实质性相似的判断,提示词的实质性相似并不能等同于 GenAI 生成物的实质性相似。一方面, GenAI 生成内容时,除受提示词影响外,还受 GenAI 模型更新、其他可调参数以及 GenAI 对用户使用习惯的理解等因素的影响,因此,即使输入相同的提示词,也可能得到不构成实质性相似的生成物。虽然在摄影作品场景下,在构思确定后,摄影师在按下快门时也有很多的独创性选择空间,但所得照片都是机械固定摄影师的可视化构思,很容易得出在后生成的照片在构思层面不仅仅是实质性相似甚至是完全一样的结论。另一方面,提示词

与GenAI生成物各自的实质性相似判断方法可能完全不同。虽然提示词和摄影的构思都可以在进行思想与表达区分之后再比对独创性的部分,但摄影作品场景下可视化构思的实质性相似和照片的实质性相似的判断方法都是针对图形元素的整体观感法,[47]而提示词的实质性相似采用文字作品的判断方法(如“抽象—过滤—比对法”),生成物根据其所属作品类型而可能采取截然不同的判断方法,如文字作品仍采用“抽象—过滤—比对法”,[48]但美术作品采用整体观感法,[49]视听作品则可能采取非常复杂的“抽象—过滤—比对法”。[50]

三是在GenAI场景下类推适用摄影作品构思的“接触+实质性相似”判断,容易偏离版权法的保护宗旨。版权法保护的是作者在法定期间内对作品受众市场的控制,因此版权侵权的判断者是作品的目标受众。将摄影作品的实质性相似等同于摄影构思的实质性相似,不会偏离版权法的保护宗旨,因为照片是对摄影师之图形元素构思的机械固定,图形元素构思的目标受众仍是照片的普通受众,该受众仍能在接触该构思后进行整体观感法的实质性相似判断。相比之下,GenAI生成作品的实质性相似无法直接推定或等同于提示词的实质性相似,根本原因在于提示词和GenAI生成作品的目标受众并不相同,二者所要保护的目标受众市场也完全不同:提示词的目标受众是创作者,生成作品的目标受众是普通消费者。

综上,即便在可版权性方面可以类推适用摄影作品构思决定论,在版权侵权判断方面,使用提示词生成的两个作品之间也难以类推适用摄影作品构思的“接触+实质性相似”判断。

四、GenAI生成物版权保护与版权侵权的分离

(一) 构思者享有权利但不承担责任的权责分离

在GenAI生成物可版权性方面,如果独创性的考察可以忽略GenAI对生成物的贡献,重点考察用户的构思,进而让用户对该生成作品享有版权,那么在GenAI生成物侵害他人版权的责任承担方面,是否同样可以忽略GenAI执行时使用他人作品的行

为,由提供构思的用户承担版权直接侵权责任?否则,在GenAI生成物之上的权利与责任失衡,可能导致用户滥用GenAI,引发新一轮“版权蟑螂”现象。

“版权蟑螂”已引发广泛争议。例如,将文章、照片或者歌曲主动上传到互联网,然后对未经许可使用者发起律师函或诉讼,导致网络版权侵权案件高居不下,浪费司法资源。[51]如果承认GenAI生成物版权,那么用户是否会滥用GenAI的强大内容生成能力,掀起新一轮“版权蟑螂”现象?例如,输入简单提示词便可轻而易举地生成成百上千的“作品”,然后将这些“作品”散播于公开网络,进而对任何未经许可使用者都主张维权。而且,在中美司法实践中,对于GenAI生成物的版权侵权,版权人通常主张GenAI提供者承担改编权直接侵权责任,而不是主张用户侵权。例如,在“纽约时报案”中,原告主张输入提示词后,OpenAI生成与原告作品实质性相似的内容,应承担改编权直接侵权责任。[52]同样地,在“广州奥特曼案”和“杭州奥特曼案”[53]中,版权人都主张GenAI提供者承担直接侵权责任。这就很可能进一步诱使用户滥用GenAI。

我国现行法没有明确承认对GenAI生成物的版权保护,法院的法律适用标准不一致,也可能导致权责分离现象。例如,将“春风案”和“广州奥特曼”的判决结果联系在一起,就容易给公众造成这样一种印象。在基于该技术变革的新规确立之前,这种潜在权责不一致现象会不可避免地出现。此外,我国著作权法并未像美国那样直接规定侵权作品本身不享有版权,[54]否则也不会产生权责不一致问题。按照我国现行法,未经许可改编他人作品只是在创作来源上存在法律瑕疵,但其本身仍然是享有著作权的作品。[55]因此,在我国,利用GenAI辅助创作的内容,可能会因GenAI执行中不当使用他人作品而构成侵权,但如果GenAI辅助创作生成的内容本身具有独创性,用户仍然可以主张版权。相反,在美国,用户无法在侵权的情况下仍主张版权。[56]

(二) GenAI 提供者参与内容生成的结构性风险

在 GenAI 场景下,出现权责分离的根本原因在于 GenAI 提供者直接参与到用户的内容生成过程之中。GenAI 提供者提供的是在线内容实时生成服务,GenAI 提供者根据用户的提示词请求向用户提供内容生成服务,而用户仅仅是提供了提示词。按照现行版权法规则,直接侵权人应当是直接生成最终表达的主体,所以 GenAI 提供者因生成最终表达而承担直接侵权责任,在这个判断过程中,无须考虑用户对该生成内容是否享有版权。相比之下,照相机是一个物理产品,虽然具备了执行摄影师创作构思的机械固定功能,但是照相机生产者将产品出售之后,就不再参与到摄影作品的具体创作过程之中,例如不会远程修改摄影师拍摄的照片。在这种情况下,摄影师无从将其摄影作品的直接侵权责任转嫁给照相机生产者。

此外,用户对 GenAI 执行其构思的控制力,远不如摄影师对照相机的控制力。摄影师完全控制照相机,因此在对摄影作品享有版权的情况下,没有任何理由和机会将直接侵权责任推卸给照相机生产者。相比之下,用户使用 GenAI 生成内容的典型场景,是通过输入提示词的方式请求 GenAI 提供者实时生成内容。GenAI 之所以能够根据提示词进行内容输出,原因在于其在训练阶段大量接触了他人的版权作品,并学习这些“作品中基本元素之间的重复模式”,然后“应用这些模式来生成新的内容”,因此可能生成与在先作品相类似的表达。[57]对此,用户可以主张不知道 GenAI 对提示词的具体转化方式,更不知道其在转化时或者此前训练时是否接触了他人享有版权的作品。

在这种情况下,GenAI 技术的出现可能会产生权责分离风险。一方面,用户利用提示词对其构思给出“可视形式”,进而按照 GenAI 生成物版权司法保护标准主张版权。另一方面,在生成内容发生改编权侵权时,用户又抗辩其仅提供了简单的提示词,而且对 GenAI 的执行原理和具体执行过程既不知情也无控制能力, [58]并辩称是 GenAI 提供者生

成了侵权内容,进而将责任推卸给 GenAI 提供者。

(三) 摄影作品构思决定论无法回应权责分离现象

在照相机时代,构思决定论的提出者无须考虑照相机的使用会导致权责分离这种 GenAI 技术带来的结构性风险,因为照相机的机械固定过程毫无摄影师发挥独创性的空间,机器的执行对侵权责任的承担并无影响。因此,构思决定论的提出者也就无法提前给出解决方案。值得注意的是,录像机这种作品机械固定设备的出现,曾引发了录像机生产商是否须承担间接侵权责任的问题,从而在一定程度上可能使直接侵权人转移一部分责任。为回应这个问题,历史上著名的“索尼案”让技术中立的录像机设备生产商免于承担帮助侵权责任。[59]基于其被动、系统自动响应用户提示词而生成侵权内容,GenAI 提供者或许可以主张适用技术中立原则,进而从帮助侵权中解脱出来,但这仅仅建立在用户直接侵权成立的情况下才能奏效,因而对用户滥用 GenAI 本身并无影响。

五、构建 GenAI 生成物构思决定论的三元要件

考虑到 GenAI 与照相机在构思要求、执行方式、提供者是否参与内容生成三大方面的差异,产生于一百多年前的摄影作品构思决定论已经无法为 GenAI 生成物版权面临的三重转向问题提供直接且有效的理论依据。然而,考虑到版权法对于人类独创性构思必须在作品最终表达上获得体现的同等要求,基于 GenAI 和照相机将人类构思转化为最终表达的程度差异,可以在承继“可视化构思”要件之内在要求(不是形式上必须“可视”)的情况下,重点将摄影作品构思决定论之“机械固定”要件修正为“受控固定”要件,并适应性增加“创作记录”要件,以为体系性解决上述三重转向问题奠定理论基础。

(一) 构思决定论的价值:构建 GenAI 时代的版权市场

首先,构思决定论顺应 GenAI 时代的版权市场保护需求。随着技术的发展和操作的便利化,利用 GenAI 创作会变得与利用照相机创作一样普及和简

单,利用 GenAI 创作的版权市场也将逐渐发展起来。在摄影技术发明 28 年后,美国在 1865 年率先通过立法保护摄影作品。美国参议院事后的解释是因为发现了摄影产业的市场前景:“在缺少任何报告和听证讨论该问题的情况下,我们找到了国会这样做的潜在合宪性假设。在寻找照片和底片保护正当性的过程中,我们应当注意到国会立法时期,因为马修·布雷迪所拍摄美国内战照片的出名,照片展现出了商业价值。”[60]对于照片的商业价值,霍姆斯法官曾指出,只要“满足任何公众的需求,他们就有商业价值——可以大胆地说,它们没有审美和教育价值——他人胆敢忽视作者权利进行盗版,就足以证明这些照片的价值和成功”。[61]如今,中美 GenAI 生成物版权诉讼表明,利用 GenAI 创作作品已经展现出商业价值:艾伦的《太空歌剧院》不仅斩获了数字摄影类大奖,同时也引来了大量盗版;[62]“伴心案”原告的作品直接被地产商盗用,[63]但其也收到了来自知名企业的新式广告设计邀请。[64]

其次,构思决定论契合作品可版权性要求,能够为 GenAI 时代的版权市场发展提供版权基础理论支撑。利用工具进行创作一直是人类的特长,版权法所要保护的恰恰是人类的智力贡献。例如,美国联邦最高法院强调,照片可以作为作品获得保护,“只要照片体现了作者具备独创性的智力构思”。[65]不可否认,摄影技术中生成图像的光化学反应,以及 GenAI 的内容生成规律,都难以被创作者(即摄影师或 GenAI 用户)掌控,但这些工具都不妨碍人类将其智力贡献融入最终生成的内容中,而且,恰恰是人类的智力贡献,让机器生成的作品具有了独创性。在这个意义上,摄影技术和 GenAI 技术都是人们可以加以利用的创作工具,而构思决定论仍然能够契合人类通过 GenAI 表达其独特构思,并支持用户嬗变为生成作品的作者,从而激励更多人利用 GenAI 进行智力创作。

最后,构思决定论可以通过适当修正,体系性地应对 GenAI 场景下独创性考察对象、版权侵权比对象、版权侵权责任承担主体三方面的挑战。虽

然将摄影作品构思决定论直接类推适用于 GenAI 生成物会出现一系列问题,但这些问题对于版权法而言并非颠覆性的全新问题。沿袭传统版权理论,对构思决定论进行适应性修正,可以让版权法在 GenAI 时代重焕生机。

(二)“可视化”构思要件:重在关键性表达要素的描述

1. 构思是 GenAI 辅助生成作品的独创性来源

现阶段不认可 GenAI 生成物可版权性的最大疑虑是,该生成物是 GenAI 系统自动执行的结果,在这个执行过程中没有体现人类独创性的空间。[66]这种疑虑在新技术刚出现时并不罕见。对这种疑虑的回应是,人类的智力贡献成为 GenAI 辅助生成作品的独创性来源。按照版权法的要求,作品是作者创作的智力成果,其中的“创作”必须是作者先有构思再有执行。只是在传统创作中,构思和执行往往融合在一起,难以区分或者没有必要区分,但到了机器辅助创作的场景,人类的构思和机器的执行可以明确加以区分,而且技术的发展也是构思的执行不断外包给机器的过程。对此,美国联邦最高法院从作者身份的角度强调了作品必须起源于作者,“作品任何内容的起源都归功于作者,他是发起者、制造者、完成科学或文学艺术作品的人”,在摄影场景下“是通过将所要拍摄的人安排在拍摄场地的合适位置,对该拍摄场地进行合理安排进而实际完成构图的人”。[67]类似地,在 GenAI 辅助创作场景下,同样是人类先通过输入提示词表现了构思,GenAI 才被动启动创作执行的程序。因此,只要在构思上体现了人类的独创性,人类的构思便成为 GenAI 辅助生成作品的独创性来源。而且,如同摄影作品的形成一样,恰恰是人类构思的注入,让机器生成的“冰冷”事物有了艺术的气息,能够反映构思者的个性,进而丰富人类的精神世界。[68]相反,如果 GenAI 生成物完全源自 GenAI 自己生成的提示词,那么 GenAI 生成物因完全不涉及人类的智力贡献而不受版权法保护。[69]

2. 提示词对关键表达性要素的描述不要求给出“可视形式”

首先,构思者需通过提示词展现其创作构思,对所生成内容在其所属作品类型的表达性元素上作出具体选择和安排。例如,为了生成一幅画,该提示词需对画的风格、主题、颜色以及所要刻画物体或人物的具体形象在图形元素上作出选择和安排。若是为了生成一篇短篇小说,则需对小说的人物、人物关系、人物个性、核心故事情节进行描述。不过,构思者并不需要将所有细节都加以描述和限定。摄影师拍摄人物写真时,也不需要刻意对人物的指甲作出要求。在 GenAI 场景下,对各个细节都作出描述或限定,反而可能无法最大化 GenAI 辅助人类创作的价值,因为这并不符合 GenAI 转化性执行的技术特性。[70]

其次,构思在通过提示词的文字描述之外,并不要求在 GenAI 执行之前就在形式上给出摄影作品那样的“可视形式”,这种在最终表达层面上的“可视形式”要求由后续要件进一步解决。一方面,摄影构思本身就是对所摄对象在图形元素上的具体安排,必然是可视的。相比之下,提示词仅是文字性描述。反之亦然,摄影场景下的构思要得到照相机的机械固定,就不能是文字性描述,而应当是镜头前可以直接固定的可视化构思。另一方面,在摄影场景下,摄影师对其构思给出可视形式只需在机器执行之前完成,无须对机器执行进行控制,这是因为照相机执行时并不修改构思者已经给出的可视形式。相比之下,GenAI 是转化性执行,构思者无法在机器执行之前就其构思给出在最终表达层面的“可视形式”。

最后,不宜将提示词仅理解为文字作品。按照目前 GenAI 生成物的通用启动方式,输入文字性的提示词是一种常规操作。但不宜因此忽略文字或字符所载的内容是对作者构思的具体表达,即对其所要生成内容的关键表达要素进行描述。这就类似于不能将摄影师为唤起独特表情而对所拍摄人物说的话(比如“来个俏皮的眼神”)当成“口述作品”,而应当视为摄影师对人物表情所作的选择和安排。同理,游戏的连续画面是由游戏开发者通过代码完成的,不能仅将游戏代码机械地解读为文字作品,

因为相对于计算机而言,文字和画面都是代码表示的内容,都可用于执行编程者的构思。[71]

3. 独创性判断对象为提示词组合展现的表达性要素选择与安排

在判断 GenAI 辅助生成作品的独创性时,不能机械适用摄影场景下的独创性判断方法,仅依第一次输入的提示词进行判断。一方面,GenAI 并非机械固定人类构思,而是根据其统计规律对提示词进行理解和转化,构思者的提示词如不准确,可能存在转化不充分的情况,这就需要构思者理解和利用 GenAI 的转化规律,在提示词中作出更准确的描述。另一方面,构思者除了第一次输入提示词展示其主要构思外,后续也会根据 GenAI 输出内容的启发,进一步通过提示词调整、丰富和完善其构思。第一次输入的提示词往往是最初的构思,虽然可以说此时已经产生了独创性,[72]但并不完整,而且仅根据第一次提示词得到的 GenAI 输出结果,往往也无法准确判断是源自机器还是源自构思者。因此,需要就最初提示词、后续为了让最初提示词中描述的关键性表达要素得到有效执行而修改的提示词以及相对应的参数调整进行综合判断,以确定构思者在整体上对 GenAI 生成内容关键表达元素作出的选择和安排是否在最终表达中得到对应的呈现,进而才能在多次提示词和生成内容之间来回考察,以确定 GenAI 生成内容是否具备最低限度的创造性。当然,该创造性并无美学或者艺术上的要求,但简单套用提示词模板而一次性输出的内容,则明显不满足创造性的最低限度要求。

(三) 受控固定要件: 确保构思在最终表达中得到有效呈现

1. 构思者需满足现行法的“直接创作”要求

独创性的第一个要件是独立创作,而独立创作的规范要求是“直接产生”最终表达,这是人格印记在作品中的最直接体现。目前,《著作权法实施条例》第3条第1款明确“创作”仅限于“直接产生”作品的智力活动,第2款将“为他人创作……提供咨询意见”排除在“创作”之外。除此之外,现行法和司法解释没有对“直接产生”作出其他规

定。然而,我国著作权法第3条将“摄影作品”纳入作品之列,而摄影作品恰恰是摄影师交给照相机执行其构思产生的摄影生成物。在这种情况下,只有两种可能解释:一是现行法明确将摄影师交给机器执行其构思的创作方式“视为”直接产生,[73]那么其他类似的机器辅助创作也可以“视为”直接产生;二是通过照相机辅助创作是“直接产生”的唯一例外,其他类似的机器辅助创作不能类推,除非立法明确规定。

人类通过 GenAI 辅助创作作品是否构成“直接产生”,也是当下的主要争议。[74]对此,“春风案”的解决方案是类推解释,既然现行法将照相机辅助创作的摄影作品纳入法定保护作品之列,那么 GenAI 辅助创作的作品也满足现行法的要求,“涉案图片是基于原告的智力投入‘直接产生’”。[75]笔者赞同“春风案”的类推解释方法。如果采用“唯一例外”的解释方法,那么 GenAI 转化构思的程度就必须达到照相机的“百分之百”机械固定,有任何减损都无法满足该唯一例外的类推要求。这对 GenAI 而言是“削足适履”。[76]而且机械解释法律要求 GenAI 转化达到百分之百固定,其结果无异于加速版权法在 GenAI 时代的边缘化。有鉴于此,接下来更重要的工作是结合 GenAI 转化构思的特性加以适当限定,使 GenAI 辅助创作可以“视为”直接产生,但也要防止直接产生的泛化而导致思想与表达不分、颠覆传统版权法。

2. 构思者对构思的受控固定作出合理努力

为了契合 GenAI 转化性执行构思的特性,必须修正摄影作品构思决定论的“机械固定”要件。那么,GenAI 对构思的转化到什么程度,才能既符合我国法的“直接产生”要求,也满足作品的独创性要求?对此,有观点认为,“即便人工智能介入到创作过程,作者提供的提示词仍反映了他们独特的创作意图和人格。因此,对 AI 生成的作品提供版权保护,仍然属于作者人格的延伸。”[77]也有观点指出,输入提示词后,用户仍会进行调整,所以只要有这样的调整就足够了。[78]这两种观点仍然过于“浪漫”,因为按照 GenAI 生成物的逻辑——

以大模型训练获得的统计学规律,推测最符合提示词的下一个内容,所以提示词在这种逻辑框架下类似于提供“创作”指导——提示词仍会被传统观点认为属于思想,并没有转化为最终表达。对此,美国版权局也指出,对于 GenAI 辅助生成的最终表达,必须审查是属于 GenAI 的贡献还是源自用户的贡献,而这“取决于 GenAI 的执行方式,以及使用该工具进行创作的具体方式如何影响最终作品的生成”。[79]

为了不破坏思想与表达二分法这一版权法的根基,构思者须将其构思“受控固定”到 GenAI 生成的最终表达中,而不是由 GenAI 不受控地自主形成最终表达。一方面,这是尽最大可能地接近摄影作品之“机械固定”,从而维护前述“视为”直接产生的扩大解释。另一方面,这也契合版权法对最终表达的保护,防止将思想也纳入作者的私有领地。

关于“受控固定”的解释,英国上议院布雷特法官在 1883 年有关摄影作品作者的裁判逻辑提供了启示:“我所能想到的最有可能构成作者的,实际上是尽可能接近于涉案照片产生原因的那个人”,在摄影场景下就是“监督安排,通过安排人员位置并确定人们所在位置来实际完成构图的那个人。这个人就是照片产生的有效原因”。[80]类似地,GenAI 生成的最终表达中,用户的贡献和 GenAI 的贡献并存,用户想要成为 GenAI 生成物的作者,必须证明其是“尽可能接近于”该最终表达产生原因的人,即能够证明通过合理控制 GenAI 的转化执行对其构思给出了“可视形式”——最终在结果意义上将其构思中的关键表达意图贯彻于 GenAI 生成的最终表达中。例如,在“武汉文生图案”中,法院认为,原告提示词“均与画面的元素及效果对应,生成的图片和原告的创作活动之间具有一定的‘映射性’”,进而认定原告对涉案图片享有 GenAI 生成物版权。[81]相反,仅输入简单的提示词,未对所生成内容的关键性表达元素作出选择和安排,无法对生成物主张版权;仅输入一次详细描述了关键性表达元素的提示词,用户也无法令人信服地证明其是产生第一份生成物的有效原因。[82]

其中,“合理”控制的判断,应当以构思中对所属作品类型的关键性表达元素的选择和安排得到有效贯彻和固定为主,以提示词对关键性表达元素描述的精准程度、构思者对 GenAI 转化规律的把握程度等因素为辅。例如,在创作《太空歌剧院》过程中,艾伦迭代了 624 次提示词:“伴心案”原告指出,“对新手来说,使用 Midjourney 就像开盲盒,永远不知道输入的词能够生成什么样的图片”,但是经过无数个夜晚的学习和实践,他“从最开始获得一张满意的图片需要迭代 80—100 轮,到现在只需要 30—50 轮,有时十几轮就能接近自己想要的结果”。[83]至于在这个过程中,GenAI 生成不可预见的内容,或者对非关键性表达元素擅自作出安排,都不影响构思者主张版权。版权法提供的保护范围,仅限于作者的独创性表达,而不延及这些未预见的部分或者未作出安排的表达。

3. 构思受控固定后作品自动产生

作品创作是事实行为,只要作品创作完成,版权就自动产生。因此,作品创作是定性行为,不是定量行为。在摄影场景下,可视化构思获得机器固定后,摄影创作即告完成,摄影作品自动获得版权保护。所以,机械固定要件是结果要件。同理,在 GenAI 辅助创作中,受控固定要件也是结果要件。只有构思者通过合理控制使其构思中的关键性表达元素得到有效固定时,这个作品创作事实行为才完成,此时的生成物就自动成为版权法保护的作品。作者后续再修改、调整而形成的生成物,都可以作为新的、单独的作品得到保护。这与传统作品的创作一样,“一部作品的总体思想或某一构思只要以某一种形式得到完整的表达,不论其是全部构思还是部分构思,该作品即视为创作完成”。[84]

此外,GenAI 本身在现行法下不具有法律主体地位,GenAI 提供者也未具体参与到用户利用 GenAI 创作作品的过程中(如未与用户共同形成创作构思或达成合作创作的意图),所以 GenAI 和 GenAI 提供者都不是 GenAI 辅助生成作品的适格作者。因此,在 GenAI 最终生成作品的过程中,并未有其他版权法意义上的适格主体的智力贡献,从始

至终体现的都是用户的智力贡献。质言之,输入提示词并经合理控制使其构思在 GenAI 生成最终表达中得到有效呈现的那个人,就是作者。

(四) 创作记录要件: 确认作者身份、权利边界、侵权原因

1. GenAI 技术特性要求留存创作记录

一方面,GenAI 事后无法根据同一提示词生成完全相同的内容。目前 GenAI 多采用扩散模型生成内容,其魅力在于根据同一提示词“能够合成表面上与训练集中任何内容都不相同的新图像”。[85]而且,GenAI 也会持续更新大模型。GenAI 的大模型训练并非一次性完成,而是持续进行,包括增加训练的基础数据、微调训练的算法。OpenAI 从 2022 年发布 ChatGPT3.5 以来,目前已经迭代到 GPT4.5。国内的 DeepSeek 从 2025 年初发布以来,也已经进化到 Deep-Seek-R1-0528 版本。所以,事后输入同一提示词,GenAI 可能生成完全不同的结果。另一方面,GenAI 生成物中人类和机器的贡献可能并存,仅从生成物的最终表达本身无法区分其中的内容是源自谁的贡献。

2. 创作记录是构思者证明作者身份的必要证据

创作记录是证明构思者系 GenAI 辅助生成作品作者的必然要求。[86]GenAI 在内容生成时采用转化性执行,构思者无法沿用摄影作品构思决定论直接就 GenAI 生成物主张版权。相反,用户必须尽合理努力控制 GenAI 对其构思的转化,使其构思中明确的关键性表达要素在 GenAI 生成物中得到有效呈现,仅在这个时候其创作事实行为才宣告完成,作品版权才自动产生。因此,用户如果想要主张 GenAI 生成物版权,应当提供证据证明其创作满足前述两个要件,否则,即便该 GenAI 生成物事实上已经构成作品,也无法证明其作者身份。例如,在“武汉文生图案”中,原告提供了涉案提示词及相应参数,并对整个过程提供了证明,法院因此认定原告的提示词“均与画面的元素及效果……具有一定的‘映射性’”,确认其对涉案图片享有版权。[87]相反,在“艺术椅案”中,原告主张其对“Midjourney”

生成的涉案艺术椅图片享有版权,但未能提供其使用该软件进行创作的过程记录,法院无法认定原告在这个过程中“作出了体现其独创性的个性化选择和修改”,继而判定原告对涉案图片不享有版权。[88]

值得强调的是,创作记录要件乃针对 GenAI 辅助创作的特点而专门提出,本身并不违反

《伯尔尼公约》的版权自动产生原则。如前所述,只要创作者满足了可视化构思、受控固定两个要件,GenAI 辅助创作的表达就构成作品,其上的版权便自动产生。这时,即便仍适用署名推定原则,相对方只要抗辩创作者无法证明其作者身份,就无须支付作品使用费,即便进入诉讼程序,法院在缺乏创作过程证据的情况下也无法支持其作者身份主张。因此,在 GenAI 场景下,创作者的版权维权成本相较传统作品要高得多。为了便利这类作品的权利主张和人工智能辅助创作版权产业的发展,除用户先行自救(即自行留存记录)外,GenAI 提供者或其他配套产业应当开发出便利用户记录创作过程的工具。而且,为了保护其提示词和创作过程技巧方面的商业秘密,可考虑设计简化版的作者身份认证系统,从而方便作者对外许可,除非进入诉讼确实需要证明身份,才需要举证全部创作过程。这一点也体现在“艾伦案”中。艾伦在申请过程中仅仅是“声称”其输入了 624 次提示词,但并未提供相关证据,[89]即便美国版权局已经认可 GenAI 生成物版权,也可能不予登记。可以预见,在美国法院后续审理过程中,艾伦能否向法庭证明其创作过程,是其主张能否获得支持的关键。

3. 创作记录是划定版权边界的基础

GenAI 辅助生成作品虽然基于提示词,但二者之间需要作者通过合理控制的努力才能满足“受控固定”要件,因此,只有提供了利用提示词创作的过程记录,作者才能有效证明其版权的权利边界,即其通过提示词描述的关键表达性要素是否得到有效转化,具体对应到 GenAI 辅助生成作品最终表达的哪一部分。

正因如此,当原被告都利用 GenAI 辅助生成作

品并发生侵权纠纷时,按照传统作品的最终表达进行比对已经无法有效确定侵权与否;同时,考虑到 GenAI 与照相机对构思的不同固定方法,直接采用摄影作品侵权的构思比对也难以有效判断侵权。因此,此时的侵权判断方法应当采用比“接触+实质性相似”更上位的侵权二元要件,即“复制+非法挪用”。[90]由于 GenAI 生成作品的最终受众仍然是普通受众,所以在进行侵权比定时,应当先进行 GenAI 生成作品的比对,如果目标受众认为被告的作品来源于原告的作品或者“产生相同、相似的欣赏体验”,则存在事实层面的复制。继而,进行非法挪用判断,确定这种事实复制是否构成版权法意义上的侵权,即被告是否复制了原告具有独创性的表达。

在 GenAI 场景下,按照构思决定论的前两个要件,非法挪用的判断须对提示词(组合)进行实质性相似的比对。在比定时,须先对提示词进行思想与表达的区别,确定原告提示词在关键性表达元素上的选择和安排是否具有独创性。如果 GenAI 生成作品之间完全一样,而且原告提示词所关联的关键性表达元素具备独创性,则可以直接判定被告侵权,因为基于 GenAI 生成作品的逻辑,GenAI 不可能生成完全相同的作品,所以此时被告是原封不动地抄袭原告的 GenAI 生成作品。[91]如果 GenAI 生成作品之间在事实层面存在实质性相似,则还需要将被告的提示词与原告的提示词进行比对。这时的判断者应当是专业人士(即利用提示词的创作者),因为提示词本身比较复杂并伴有专业性术语,普通公众可能无法读懂。这就类似于计算机的编程语言,虽然计算机软件的目标受众是普通公众,但计算机代码并不是设计给普通公众阅读的,而是设计给计算机识别和执行的,所以计算机软件侵权的实质性相似判断由专业人士完成。[92]如果专业创作者认为被告提示词与原告提示词之间仍然构成实质性相似,即被告使用了原告构思中具有独创性的关键表达特征,则可以判定被告侵犯了原告的版权。[93]

4. 创作记录可明确侵权内容的生成原因

不同于照相机生产商与摄影师创作的分离,因

响应用户提示词实时输出内容，GenAI 提供者与用户建立了持续服务关系，用户才有可能一方面对基于提示词生成的内容主张 GenAI 生成物版权，另一方面在发生侵权时主张自己对 GenAI 的执行不知情，抑或抗辩侵权内容是 GenAI 提供者直接生成的。为此，在受控固定要件的基础上增设创作记录要件，也能够有效防止用户滥用 GenAI。一方面，通过自动化系统一键生成批量作品无法满足受控固定要件，即无法证明其提示词和其他安排是生成内容的有效原因，因而无法获得版权，无法以版权人之名义发起批量诉讼。另一方面，用户如果满足受控固定要件，那么其对于 GenAI 生成物中哪些侵权内容的生成起到主要作用——如是否故意提示按照他人已详细描述并享有版权的故事角色进行小说撰写，只需换一个角色名字——也可以通过创作记录非常清楚地作出判断，用户也就无法逃避侵权责任。相反，如果不满足创作记录要件，用户即便在事实上满足了受控固定要件，也无从主张权利，因为仅凭 GenAI 生成的最终表达无法直观地表明，用户为其构思得到 GenAI 贯彻并在最终表达中得到有效呈现确实作出了合理控制；如果 GenAI 提供者通过后台日志能够初步证明是用户的原因导致侵权内容的生成，用户无法提供相反证据的，也要承担相应的侵权责任。[94]毕竟，“这些负责创作可供他人阅读的有形文学形式的人，可以‘作者’的身份为自己主张版权，因为他们要对‘以这种方式让作品在整体上具备独创性’负责”。[95]

结语

GenAI 生成物版权要获得接受，必然会经历漫长的历程，此点从摄影作品获得《伯尔尼公约》接受的历程可见一斑。摄影技术自 1837 年发明，在 1884 年首次以议定声明的方式获得

《伯尔尼公约》的非正式接受，成为可受版权保护的客体；在 1908 年以新型作品的形式单独规定于《伯尔尼公约》正文第 3 条；在 1948 年纳入《伯尔尼公约》第 2 条规定的“艺术作品”之列，获得与文学作品同等的地位；在 1967 年其定义从 1896 年以来一直采用的“通过摄影程序生成的作品”，

修订为“通过摄影方式表达的作品”，“旨在强调摄影作品的决定性因素是作品表达的方式，而不是技术程序本身”。[96]历史经验表明，通过借鉴和修正摄影作品构思决定论，使其契合 GenAI 的转化性“固定”特点，人类利用 GenAI 辅助创作的作品也终将可以表述为“通过 GenAI 方式表达的作品”，甚至简化为可媲美摄影作品的“智能作品”。从目前我国司法实践的开创性进展来看，中国在这个过程中将扮演重要乃至主导性角色。

数智时代的过劳问题及其法律因应

此处删除了原文脚注，全文请参见《中国法学》2025 年第 3 期，转载或引用请注明出处。

作者：田野

内容提要：数字智能技术广泛应用于劳动管理，在提升效率的同时，也带来了过度劳动加剧的风险。数智时代的过劳具有普遍性、隐蔽性、精神性、过剩与过劳并存等特点。工作时间与私人生活的边界模糊化、数智化驱动的薪酬制度、算法管理的规训与控制是数智时代过劳形成的重要原因，资本的逐利本性是数智时代过劳的政治经济学根源。破解这一困境应当采取系统化的法律治理路径，通过革新劳动基准实现底线控制，通过对劳动者赋权（以离线权为中心）提供过劳自救机制，通过对用人单位课以算法善治义务消除过劳的技术根源，通过将过劳纳入工伤保险和对算法致劳追究侵权责任达成救济目标，由此建构起由基准、权利、义务、救济构成的完整过劳治理体系。

一、科技进步与过劳

累，正在成为现代人的普遍感受与共鸣，全社会的过度劳动问题令人忧心。科技对于劳动有着举足轻重的影响，人们原本朴素的美好期待是：借助科学技术进步推动生产力提高之红利，劳动者将不断获得解放，辛苦繁重的低级工作都由机器代劳，工作时间越来越短，劳动者将有越来越多的空闲得以享受美好的私人生活。^[1]早在 1930 年，著名经济

学家凯恩斯就预测,100年后人类将会过上“富裕有闲”的生活,因无所事事而空虚烦恼。^[2]然而近100年过去了,凯恩斯口中的“休闲时代”并没有如期而至,21世纪反倒成为“过劳时代”。^[3]2019年5月,世界卫生组织首次将过劳列入《国际疾病分类》。^[4]人们非但没有因为技术的进步而变得更加悠闲,相反却陷入更严重的新型技术过劳。

根据国家统计局2025年2月公布的数据,我国企业职工每周平均工作时间达到了49.1小时。^[5]这与《国务院关于职工工作时间的规定》第3条规定的40小时周标准工作时间相比,足足多出了9个多小时。且从统计数据的动态更新来看,工作时长呈现逐步递增的趋势。“996工作制”引发了公众极高热度的关注,从一个侧面反映了我国加班常态化和过劳普遍化的严峻现实。^[6]过劳不只是我国独自面临的挑战,而且是一个世界性问题。在美国一项针对脸书员工的调查结果中显示,员工抱怨最多的是工作时间太长、工作和生活难以平衡、工作压力太大。亚马逊程序员的周工作时长达80-100小时,即便在深夜、周末或节假日也被要求随时待命。在谷歌,员工也被要求24小时待命。优步默许专车司机每周最长工作100小时,每天驾驶超14小时。美国劳工统计局的统计数据显示,数字化生产行业超时工作比例十分明显地高于传统部门。^[7]

“过劳”作为人类社会发展中的一个问题由来已久,在大数据与人工智能时代背景下,过劳呈现出独特的新样态。数字智能手段在企业的劳动管理中被日益频繁地采用,借助这些高科技手段,用人单位有更强的能力对劳动者进行持续的精准控制,工作和私人生活的时空界限被打破,劳动者随算法规训而“起舞”,陷于过劳的怪圈中无法自拔。在平台经济领域过劳问题尤为突出,每日十几个小时奔波于马路上的骑手是过劳的现实写照。^[8]过劳引发的交通事故、猝死等事件频繁见诸报端,引发社会高度关注。

数智时代的劳动有两个面向:从劳动者一侧看,是运用数智技术作为手段从事的劳动;从用人者一侧看,则是运用数智技术实施劳动管理。本文侧重

于后一面向,聚焦算法管理导致的过劳。从范围来看,数智时代的过劳是涵盖整个劳动领域的普遍性问题,算法管理不仅应用于新就业形态,也辐射到传统的就业形态,例如算法智能监控同样广泛存在于典型劳动关系下的各种正规就业岗位。从现实看,新就业形态的过劳问题更为突出,引发的社会关注度更高。因此,从问题导向出发,本文将新就业形态领域的过劳作为切入视角和考察重点。同时,笔者并不希望将眼光仅仅局限于此,而是由新就业形态管窥全局,尝试探索更广阔劳动世界的过劳治理之道。

科技越来越先进了,我们却越来越累了。残酷的现实倒逼我们发出时代性之拷问:科学技术进步给劳动者带来的到底是更多的闲暇生活,还是变一种花样的过劳?数字智能技术到底是劳动者获得解放的福音,还是更沉重的高科技枷锁?数字智能时代过劳的特殊性是什么?究竟是什么原因造成的?摆脱数智时代的过劳困境的出路在何方?如果说经济和劳动的数智化是不可逆转的历史大势,那么如何解决数智技术异化导致的过劳加剧问题,就是数智时代法律必须认真面对和积极回应的挑战。

二、四次工业革命时代的过劳样态

科技与劳动有着千丝万缕的联系。迄今为止人类共经历了四次工业革命——蒸汽化、电气化、自动化和智能化,而过劳问题从始至终贯穿了每一个阶段。在不同的发展阶段,过劳体现出差异化的形态,具有鲜明的时代烙印。在工业4.0时代,过劳呈现出独有的数智化样貌。有趣的是,考察过劳迭代是管窥更广阔的劳动世界面貌变迁和劳动法演进的一条独特线索。

(一) 蒸汽时代的过劳

从资本主义诞生之日起,工人阶级过劳的苦难史就随之开始了。在资本的原始积累阶段,资本家通过野蛮残暴的手段使劳动者失去生产资料,迫使其进入工厂进行超长时间的高强度劳动。在第一次工业革命之前的手工作坊时代,由于技术不发达,生产力落后,资本家只能靠尽可能地延长劳动者的

工作时间来榨取更多的剩余价值。从14世纪到18世纪末的数百年时间里,欧洲国家制定了许多所谓“劳工法规”(实为反劳工法规),其突出特点是强迫劳动,规定最低工时和最高工资,与后来保护劳动者的最高工时和最低工资恰背道而驰。超长工作时间是这一时期血汗工厂的典型特征,工人的每日劳动时间普遍在十几个小时以上。“在资本主义原始积累阶段所出现的过度劳动,是低劳动生产率、资本主义的经济制度以及当时的政治制度的产物。”[9]

瓦特改良蒸汽机开启了蒸汽时代,随着大机器工厂取代手工作坊,第一次工业革命轰轰烈烈地展开。至19世纪上半叶,工业革命宣告完成,其与资本原始积累阶段的结束时间大致吻合。机器的发明和广泛采用极大地提升了劳动生产效率,资本主义发展由此进入快车道。然而,工人过劳的情况却没有因此而得到改观,先进的机器不过是资本家剥削工人的新工具。资本家强迫工人在机器上长时间劳动,所产生的过劳问题比手工作坊时期有过之而无不及。日益加剧的过劳和盘剥使劳动者不胜其苦,也因此激起了强烈的反抗。对抗过劳与剥削的工人运动此起彼伏,声势浩大。[10]在巨大的压力下,资产阶级被迫作出一定程度的妥协,这促成了现代意义上劳动法的诞生。1802年,英国颁布《学徒健康与道德法》,规定纺织工厂不得雇佣9岁以下的学徒,童工的每日工作时间不得超过12小时。尽管现在看来12小时的工作时长仍旧不甚合理,但这毕竟是立法开始限制最高工时的良好开端,因此该法被视为现代劳动法肇始之源。由是观之,对抗过劳的斗争促成了劳动法的问世。

(二) 电气时代的过劳

自19世纪60年代开始,借第二次工业革命之力,人类进入电气时代。电力和更先进机器的采用使社会生产力进一步提高,自由资本主义向垄断资本主义过渡。由于技术的进步和工人运动的不断推动,缩短工时成为大的趋势,劳动立法逐渐完善。1886年,芝加哥工人举行大罢工,要求实行8小时工作制,这一事件后来成为国际劳动节的起源。缩

短工时也是1919年召开的国际劳工组织第一届会议的主题,大会通过了《工业工作时间每日限为8小时及每周限为48小时公约》(国际劳工组织第1号公约)。此后,8小时工作制在全球普遍推行。与《学徒健康与道德法》时期相比,工时有较大幅度的缩短,这显然有利于缓解工人的过劳。这一时期是劳动法的发展完善时期,缩短工时的立法成为重中之重。总体观之,劳动者的生存状况在该阶段有所改善。当然,这并不意味着过劳问题被彻底解决,而是仍在一定程度和范围内存在。电气时代的过劳体现出新的特征,在工作时间的长度被法律严格限定的情况下,雇主对工人在单位时间内劳动强度的控制不断加强。电力支持下的福特制流水线作业加快了工作节奏,劳动者不断重复单一任务,导致身心疲惫。

(三) 自动信息化时代的过劳

20世纪40年代,第二次世界大战结束后,人类开始进入自动信息化时代。这一时期是资本主义历史上的黄金发展周期,由于战后难得的和平、更先进的自动化技术的推广等诸多因素,世界主要资本主义国家在该阶段取得飞速发展。这一时期同时也是劳资关系的甜蜜期以及劳动法的成熟期。工人阶级的状况得到进一步改善,与资产阶级达成历史性的短暂妥协。然而,这一黄金时期在经过30年左右的时间后戛然而止。进入20世纪70年代,受石油危机因素影响,世界经济陷入低迷,工作时间延长的苗头卷土重来。[11]特别是在日本,劳动者的工作强度和压力不断加大,过劳的社会问题日益突出。正是在这样的背景下,过劳作为一个独立的概念被提出,关于过劳的研究成为一个专门的研究领域,而过劳立法也逐步取得进展。过劳是一个全球性的问题,但“过劳”概念最早是在日本提出,并且在相关研究和立法方面获得进展的。[12]总体上看,计算机、机器人等自动信息化工具的采用在一定程度上解放了劳动者的双手,但也带来了新的过劳问题,其对劳动者技能提出了更高的要求,并导致更多的脑力消耗,体力性过劳向脑力性过劳转变。

（四）数智时代过劳的新样貌

以人工智能科技的兴起为标志，人类进入工业4.0时代。大数据和人工智能技术在劳动管理中被广泛应用，给劳动世界带来了翻天覆地的变化。[13]职场智能监控盛行，工作时间的弹性增强，工作与私人生活的时间和空间界限被打破，劳动关系变得愈发模糊化和不稳定，技术势差的扩大导致用人单位和劳动者算法权力的失衡，劳动者容易陷入算法规训的牢笼而形成“自愿”过劳。数智时代的过劳与既往三个时代的过劳相比，呈现出一些新的特征，突出表现在以下几个方面：

第一，过劳的普遍性。从骑手、网约车司机、平台家政工等低端劳动者，到电脑程序员、公司经理等从事脑力劳动的高级技术人员和管理人员等广泛的人群存在过劳问题。[14]其中，依托平台工作的新就业形态劳动者过劳的问题最为突出。与此同时，传统雇佣形态下的典型劳动者也出现过劳加剧的状况。疲惫，似乎成为现代人的通病。

第二，过劳的隐蔽性。早期的过劳往往是由雇主强行命令工人从事超长时间劳动导致的，比较容易识别判断，数智时代的过劳则隐藏得更深。由于工作时间和地点都变得弹性化，工作的自主性看起来增强，是否劳动及其时长似乎是由劳动者自己掌控的，不存在强制工作的命令。即数智时代的过劳常常表现为“自愿”过劳，而其背后的幽灵推手其实是雇主借助大数据分析实施的算法控制和科技狠活“洗脑”。[15]

第三，精神性过劳问题突出。既往时代的过劳往往更多体现为生理性过劳（包括体力性过劳和脑力性过劳），即长时间工作造成身体的生理疲惫累积而不能得到及时有效缓解。在数智时代，紧张、焦虑则日益成为过劳的新变种。数智技术对用人单位产生了科技赋能效应，使其有能力对劳动者实施持续精准的控制与监督。[16]在智能坐垫、智能手环等高科技设备的覆盖下，劳动者随时随地都处于监控之中，承受巨大的精神压力，这使其精神健康受损，乃至导致精神疾病、自残、自杀等严重后果。传统劳动法对劳动者职业安全健康的维护更多地

集中在生理健康层面，对劳动者精神卫生健康的维护则鞭长莫及。

第四，过劳与过剩并存，过剩刺激过劳。数智技术在规训一部分劳动者“自愿”超强度工作的同时，使另一部分劳动者丢掉了饭碗。过劳人与过剩劳动力并存，成为数智时代冰火两重天的独特景象。[17]机器取代人的现象不是现在才出现的，而是在第一次工业革命时期就引起了人们的关注，但是数字智能技术带来的大规模失业的忧虑是前所未有的。数智技术的强大使得许多原本由人力从事的工作可以由算力完成，岗位取代不是杞人忧天而是正在发生的现实。大量过剩劳动力形成的产业后备军，对在职劳动者产生隐形的竞争压力，敦促其更好地表现以保住难得的工作机会，这又反过来使过劳加剧。

简要回顾四次工业革命的历史与过劳的演进史可以发现，科技与过劳有着千丝万缕的联系，始终如影随形地纠缠在一起。科技到底是解放人还是束缚人，恐怕不能简单地一刀切判断。每一种先进科技背后都潜藏着善恶两种可能性，即缓解或者是加剧过劳。技术本来是中立的，但技术背后的资本有善恶。当资本在逐利动机驱使下将先进技术用作压迫劳动者的手段时，技术越先进则过劳越严重——迄今为止的历史不断验证着这一点。若从法律的视角观察，与过劳抗争的历史也就是劳动法产生、发展和完善的历史。由是观之，科技、过劳、立法实为三位一体的关系。关注当下，数智时代的过劳呈现出新的面貌特征，给劳动法提出了新的挑战。这既是劳动法遭遇的困境，也是劳动法淬火重生的进化机遇。

三、数智时代过劳的成因

为何数字智能技术不仅没能解决过劳问题，反而造成新型的过劳？其背后的复杂原因值得深思。

（一）工作时间与私人时间边界的模糊化

从过劳的影响因素及其联系紧密程度来看，工时因素首当其冲。一般而言，工作时间过长是造成过劳的最直接原因。工作时间与私人的休息时间应当保持一定的比例和清晰的边界，工作不应侵入劳

动者的私人生活,以保证劳动者能够得到足够的休息而恢复再生产所需的体力和脑力。早在1866年,在日内瓦召开的第一国际大会上就提出了“8小时工作、8小时休闲、8小时睡眠”的黄金比例分割。以此为导向,现代劳动立法上对工作时间以及加班时间都有十分严格的限制,这些时间是不能逾越的强制性法定基准。数智时代到来之前,在传统的工作模式下,工作时间与私人生活的边界还是相对清楚的,下班之后就是私人时间。然而当数智时代到来时,这种情况发生了改变。在前沿技术的加持下,雇主对劳动者的控制能力增强,甚至将触角延伸到私人领地,使劳动者回家后仍然需要工作的情况愈发普遍。那些原本需要在车间、办公室完成的工作,现在可以通过远程劳动在家完成。在办公场所以外的空间完成工作,带来了工作时间计算的难题。[18]工作时间可能会不断蚕食劳动者的私人时间,从而隐蔽地延长工时,为过劳埋下隐患。

工作的自主性增强,工作时间日益弹性化,工作与私人生活的时间界限变得愈发模糊难辨。[19]其中,平台经济下新就业形态的工作时间模糊化问题尤为突出。由于新就业形态从业者的工作自主性较高,其对于工作时间的长度有较大的弹性控制权。从现实的情况来看,新就业形态从业者超长时间劳动导致过劳的情况十分严重。以骑手为例,每天送单十几个小时的现象很普遍。在骑手等新就业形态从业者工作时间的计算问题上,一个焦点议题是待命时间是否属于工作时间。从业者在登录平台的系统后,除了完成订单的实际劳动时间外,还有部分时间处于等待工作任务的间歇状态。对于这部分时间是否应计入工作时间,理论界和实务界存在不小的争议。[20]再者,由于新就业形态从业者的法律地位不明,是否适用劳动法存疑,适用工时基准对其存在障碍。[21]

(二) 数智驱动薪酬制的负面影响

获得报酬是劳动者工作的基本目标,要获得更多的报酬就需要付出更多的劳动,薪酬制度设计与劳动强度有着密切的联系。以数智技术为驱动的薪酬计算与分配制度呈现出许多新的面向,对过劳的

形成起着推波助澜的作用。

首先,在数智化薪酬机制下,劳动报酬的数量由算法通过大数据分析自动生成。面对自动化薪酬系统,劳动者没有定价权,缺少参与协商的有效程序机制,沦为科技的旁观者。由于算法由企业设计部署,收入分配算法必然更多体现用人单位的意志,对劳动者则不友好。在计件薪酬制下,收入水平是与接单量正相关的。由于缺少合理的劳动定额限制,劳动者要获得更多的收入只有通过不断地增加接单量,这刺激了过度劳动。[22]实践中虽然不少骑手的月收入能达到几千元,但看似不错的收入实际是以超时长、超强度劳动透支身体为代价的。

其次,借助大数据和算法分析,资本能够对劳动报酬的成本支出进行更精准的控制,最大限度地确保付出的每一分钱都不被浪费。根据经典的劳动法理,在传统劳动关系模式下用人单位负担工资续付义务,劳动者在工作时间里只要在岗,即使因暂无工作任务而没有实际从事劳动,用人单位仍需支付工资。假如劳动者每日在岗时间是8小时,其中有4小时实际从事了劳动,那对闲暇的4小时支付报酬在资本方看来就是一种成本浪费。数智技术则能帮助资本避免这种“浪费”,使每一分钱都有回报。劳动者要获得报酬,只有马不停蹄地工作。

最后,实时的绩效评估和动态薪酬制对劳动者构成巨大压力。数智算法的运用使绩效评估不再局限在年终考核,而是每分每秒地实时进行。由于薪酬是动态调整的,劳动者必须时刻提醒自己调整到最佳状态并保持下去,不够良好的表现随时都会反映在绩效的下滑中。相对于对收入水平的影响,其更大的影响恐怕在于对劳动者心理的负面暗示,使之产生巨大的心理压力,乃至发展成过劳。

(三) 算法管理的重压

在工作时间长度特定的情况下,工作的强度是影响过劳的另一重要指标。在数智时代,工作的强度较以往大幅提升,算法管理是导致这一结果的最重要因素。

算法管理是劳动管理的一种新形态,是将人工智能手段运用于劳动力管理的产物。在数智时代到

来之前,对劳动者的管理主要依靠人力实现。前三次工业革命的成果更多地体现为有形生产工具的革新,较少直接作用于劳动力的管理。而数智技术则不仅带来了生产工具的革命,也带来了劳动管理方式的革命。算法管理取代人类HR,涵盖了从招聘到劳动过程管理再到解雇的每一个人力资源管理环节。平台经济是算法管理运用最为普遍的领域,不过算法管理的版图却不限于此,而是扩张至包括典型劳动关系在内的整个劳动世界。[23]最早关注算法管理问题的是一些管理学者,[24]此后法学、经济学、社会学等多学科的加入使算法管理成为跨领域的研究热点。在实践发展和理论研究的助推下,算法管理渐渐进入法律规范的轨道。2021年人力资源社会保障部等8部门联合发布的《关于维护新就业形态劳动者劳动保障权益的指导意见》(人社部发〔2021〕56号)第10条开算法管理规范之先河。此后,国家市场监督管理总局等7部门联合印发的《关于落实网络餐饮平台责任切实维护外卖送餐员权益的指导意见》(国市监网监发〔2021〕38号)、人力资源社会保障部办公厅印发的《新就业形态劳动者劳动规则公示指引》(人社厅发〔2023〕50号)等政策文件将算法管理的规范进一步细化。域外法则更进一步,欧盟2024年底颁布的《改善平台工作条件指令》[25]专设第三章“算法管理”(ALGORITHMICMANAGEMENT)。纵观国内外,算法管理经历了从实践、学理再到立法的进阶之路。

算法管理极大提升了劳动的强度,容易诱发压力型过劳。首先,算法权力失衡导致用人单位和劳动者实力鸿沟扩大,是过劳的根源。在以从属性为根本特征的劳动关系下,用人单位和劳动者的地位本来就是不对等的,在数智技术的赋能效应下,用人单位变得更加强大,劳动者则变得更加弱势。[26]“数字技术不仅强化了雇员对生产资料的依附,还进一步丰富了雇主对雇员的控制手段,雇主对雇员的控制程度空前提升。”[27]数智技术使用用人单位管理劳动者的手段更加高明,控制更加精准和有力,劳雇双方的实力差距进一步拉大。劳动者沦为透明的数字人,被算法管理牢牢掌控,随算法而“起舞”,

这为过劳的形成埋下祸根。

其次,算法管理使用工关系的性质模糊化,就业的不稳定、不安全,权益保障的不足,是过劳滋生的土壤。由于算法管理貌似弹性灵活,使劳动关系的认定变得异常困难。这一问题在平台用工领域表现得尤为突出,大量平台从业者游离在劳动法保护的边缘。[28]再者,算法管理模式下载工就业的碎片化特征使工作非常不稳定,加之权益保障的缺乏,导致从业者产生极大的不安全感,为了摆脱这种不安全感而花费更多时间工作以赚取更多收入,这不可避免地加重了过劳。

最后,算法管理营造的压抑工作环境易导致劳动者紧张焦虑,形成精神性过劳。这突出体现在以下两个方面:一是算法评分机制的规训效应。劳动者如果不能按时按质地完成订单任务,将会受到负面评分的不利影响乃至惩戒。算法评分机制虽貌似不强制,但隐形的规训效应事实上十分强大,使劳动者陷入“自愿”过劳。[29]二是智能监控带来的精神压力。作为目前十分流行的算法管理的手段,智能监控十分普遍。与前几代监控技术不同,智能监控打破了时间和空间的限制,使用人单位有能力持续无死角地监控劳动者的一举一动。[30]这对于无时无刻不被监视的劳动者而言,既是对其人格的不尊重,也是巨大的精神压力源。这种状况长期累积就会导致亚健康与过劳。欧盟《改善平台工作条件指令》第12条特别规定了算法管理下的安全与健康问题,其要求数字劳动平台应当评估自动化监控系统和自动化决策系统对安全与健康风险,特别是与工作相关的事故、社会心理和人体工程学方面可能存在的风险,不得使用对平台工人造成不当压力或以其他方式危及平台工人的安全及身体和精神健康的自动化监控系统或自动化决策系统。需要指出的是,该条特别强调了“压力”和“精神健康”,在立法中印证了数智时代过劳的精神性特征。

(四) 过劳成因的马克思主义解读

在马克思主义看来,资本对剩余价值的追求是过劳产生的根源所在。即使是在数智时代,这一基本规律也仍旧适用和具有解释力。尽管在人类社会

发展的不同阶段劳动世界的面貌是不断变迁的,劳动管理模式与过劳的形态是变化的,但剩余价值规律却是恒久不变的。资本总是以追求尽可能高的剩余价值为终极目标,而增加工人劳动的时间和强度是达成这一目标的必由之路。在《资本论》中,资本提高对劳动者剩余价值剥削程度的方式有二:一是通过延长工作时间获取绝对剩余价值;二是通过采取先进技术提高生产率、加强劳动强度来缩短必要劳动时间,挤出更多时间以获取相对剩余价值。数智时代亦不例外,只不过具体表现不同。

就绝对剩余价值而言,“资本家总想把工资降低到生理上所能容许的最低限度,把工作日延长到生理上所能容许的最高限度”[31]。在数智时代,这种工作时间的延长是通过隐蔽的高科技手段实现的,资本借由数智技术模糊工作与私人生活的界限,使家里家外都成为职场,工作时间与私人时间难以分辨,工时被鬼使神差地延长。再者,貌似弹性的劳动算法管理使既有劳动立法上的工时标准遭遇适用困境,失去工时上限的刚性约束,超长工时遂在实践中频频上演,而资本由此获取了更多剩余价值。就相对剩余价值而言,由于购买劳动力所付出成本的有期限性,资本家“小心翼翼地注视着不让有一分钟不劳动而白白浪费掉”[32]。在避免付酬的劳动时间被浪费方面,数智技术有着独特的优势。资本通过智能监控等手段监视劳动者的一举一动,确保劳动者不在工作时间“摸鱼”。在数字平台经济下,平台按订单计酬,即只对实际提供劳动的时间付酬,没有一分钟的浪费。智能算法助力了资本的精打细算和更多剩余价值的获取,劳动者则被透支了身体。

诚如马克思所言,每一次科技革命带来的生产方式变革,最终“都变成一种加紧吮吸劳动力的手段”[33]。数智技术压迫下的过劳不过是这一故事的新版本,但在表现形式上与其他时代的过劳有所不同。

四、数智时代过劳的法律治理

数智时代的过劳是科技异化的表现,带来劳动者亚健康、焦虑、抑郁、自杀等诸多社会问题。这

与满足人民群众日益增长的对美好生活的向往的国家政策导向背道而驰,也与国际劳工组织倡导的体面劳动背道而驰,因此寻找摆脱过劳困境的出路势在必行。事实上,现行立法中并非没有治理过劳的规定,工时、休假、职业安全健康等制度对于预防过劳皆有重要价值。不过,由于数智化背景下的过劳作为新一轮科技革命带来的新兴问题具有特殊性,那些以前数智时代劳动关系及管理模式为对象而设计的制度架构存在明显的不适应性和滞后性。缘是之故,必须探索新的法律因应之道。在笔者看来,数智时代的过劳问题不是某项单一制度能够解决的,必须建构多个制度相互协同的过劳治理体系。以劳动基准的革新作为底线控制,对劳动者赋权——以离线权为中心,对用人单位课以算法善治义务,对劳动算法导致的过劳损害予以救济,建构由基准、权利、义务、救济组成的治理体系架构,或是法律回应数智科技挑战的一条有效路径。

(一) 底线控制: 数智化背景下劳动基准的革新

劳动基准是劳动者权益保护的基本防线,也是预防过劳的底线。由于数智化使劳动的弹性显著增强,使劳动基准的适用遭遇困境。然而,灵活、弹性不应成为逃逸于劳动基准之外的当然理由。[34]因应劳动数智化的发展趋势,对劳动基准也应当与时俱进地作出调整,以筑牢预防过劳的底线。

由于工作时间的长短与过劳有着最近的直线联系,因此控制工作时间显然是治理过劳的首要切入点。如前所述,在职场数智化背景下,工作时间与生活休闲时间的边界模糊不好计算,是造成过劳的重要原因。为此,必须根据数智化劳动的特殊性对工时基准进行革新,对模糊化的工作时间边界加以再厘定。

第一,工作时间概念的正面界定。工作时间是最核心的劳动基准,是劳动法上的一个基础性概念。然而,到底什么是工作时间,我国劳动法上一直缺乏明确的界定。厘清工作时间的定义,是划清工作与私人生活边界的前提。在本质上,工作时间是劳动者让渡自由而就劳动事项受用人单位约束控制

的时间。工作时间的价值功能有多个面向,包括保护劳动者休息权、作为工资的计量单位以及作为认定工伤的时间尺度等。诚如有学者指出的,工作时间最核心的功能应当是作为劳动基准,树立保护劳动者休息和健康的底线。[35]若在底线控制意义上理解工作时间,则立法的正确表述方式应当是“工作时间不得超过……”。《国务院关于职工工作时间的规定》所采用的“每日工作8小时、每周工作40小时”的表述,并非真正意义的劳动基准,而是标准工作时间(或者说是正常工作时间)。关于工作时间的认定标准,存在学说的争鸣。最严格的标准是只有实际提供劳动的时间才能计入工作时间,更宽松的标准是只要劳动者处于用人单位的控制指挥之下即属工作时间,是否实际从事劳动并非所问。还有学者提出“劳动者解放说”,认为只要劳动者没有获得解放,受用人单位拘束,即可认定为工作时间。[36]在数智时代工作时间日趋模糊化、碎片化的背景下,为了避免工作对私人生活时间的过度侵蚀,宜采相对宽松的认定标准,从工时的本质出发,只要劳动者处于受控制、被占用的职场状态,即使没有实际从事劳动,也应被认定为工作时间。所谓“受控制”,在解释上应当认为也包含算法管理系统的控制。一个颇具争议的问题是:处于工作与私人生活过渡地带的灰色时间是否属于工作时间,如在算法管理系统中的待命时间。由于没有实际的劳动给付,计入工作时间对用人单位似有不公。对此问题,一种解决思路是创设第三类时间。不过这种观点遭到批判,反对者认为不是工作就是休息,二者之间不存在中间范畴。笔者赞同二分说,创设第三类时间容易带来混乱。借用“劳动者解放”理论,劳动者若在算法管理系统中被占用则不得解放,理应纳入工作时间。当然,在具体计算方法上可以通过按比例折算等方式灵活掌握。值得关注的是,人社部办公厅印发的《新就业形态劳动者休息和劳动报酬权益保障指引》(人社厅发〔2023〕50号)提出了“宽放时间”的概念,其可通过集体协商的方式被灵活折算为工作时间。[37]

第二,确立用人单位的工作时间记录义务。从

实践操作层面看,工作时间的证明和计算是一个难题,反映在劳动争议解决程序中则为举证困境。在数智化环境下,由于以消费者需求和市场需求为导向为中心,算法是根据个性化需求分配工作任务的,原本流水线作业模式下的完整工作时间被撕裂成碎片,工时计算之困雪上加霜。为应对这一挑战,一个直接有效的方法是对用人单位课以工时记录义务。[38]由于劳动过程是由用人单位控制的,劳动者对工作时间的举证能力弱,用人单位则有着天然的优势。数智化一方面带来了工时计算的难题,另一方面也提供了用以破解这些难题的技术手段。在数智化管理模式下,劳动者的一举一动都被算法记录,用人单位作为算法的掌控者,在记录和统计工作时间方面不存在技术瓶颈。用人单位应当建立工作时间台账,清晰完整地记录工作时间,而劳动者有查询权,该台账可作为判断劳动者工作时长及是否过劳的重要证据。不少域外国家劳动法上都有工时记录义务的特别规定,但我国立法就此却长期缺位,有必要适时通过立法填补空白。[39]

第三,完善休息规则。如果说设置工作时长限制的根本目的是维护劳动者休息和健康,那么劳动立法直接确立休息规则恐怕是更直接的解决方案。休息规则是劳动基准的另一个维度,甚至是更重要的维度。正如有学者所倡导的,应实现“从工作时间计算为中心到休息时间为中心的立法视角转换”[40]。从休息的视角确立劳动基准在立法上应表述为一定周期内休息的时间不得少于多少日或小时。为避免数智化将休息时间切割成支离破碎的片段,影响实际休息的效果,立法应明确每日连续休息的最短时间。例如,欧盟《工作时间指令》第3条规定,每24小时必须保证11小时的连续休息时间。我国立法也应当采用类似做法。除了休息时间的长度,休息的间隔也很重要。从人的生理角度看,连续工作一定时间就需要获得休息,超长时间连续工作极易诱发过劳。在算法控制的劳动管理系统中,工作沉迷是一个现实问题。作为应对,在技术系统中设置疲劳提示和强制下线、停止工作推送功能,是一个有益的尝试。例如,《关于落实网络餐饮平

台责任切实维护外卖送餐员权益的指导意见》规定，“合理管控在线工作时长，对于连续送单超过4小时的，系统发出疲劳提示，20分钟内不再派单”。在新就业形态之外，正规就业形态劳动者同样也有间隔性休息的需求。

工资是另一类重要的劳动基准，且与工时有着微妙的联动关系。我国现行的工资基准存在诸多问题，是过劳泛滥的隐形根源。最低工资作为基础性的劳动基准，旨在维持劳动者的基本生活所需。实践中有相当大比例的企业实际是压着最低工资线确定劳动报酬的，在这个意义上，最低工资也就成为常态工资甚至是最高工资。从各省的最低工资水平来看，标准总体上是较低的。2024年度上海的最低工资标准只有2690元^[41]，在上海这样的一线城市靠这些收入想维持基本生存都是不容易的，而这已经是全国最高的最低工资标准。在基础工资不能满足生活需要的情况下，劳动者不得不在正常工作时间之外通过加班、兼职等方式谋求额外的收入来源，这就带来劳动时间和强度的增加。加班工资基准对过劳的形成也产生了潜移默化的作用。一方面，用人单位在生产经营中有增加工作时间的弹性需求，而我国的标准工时制比较僵化，只能通过加班实现目的；另一方面，由于基础工资水平低，劳动者有通过加班获得额外收入的动机，且由于加班费的标准比较高，变相刺激了劳动者加班。^[42]尽管我国设置了加班时间上限，但在实践中事实上被架空，普遍加班成为常态，为过劳提供了滋生的土壤。提高基础工资的保障水平，适度降低加班工资的标准，应成为解决问题的方向。^[43]在数智化管理的背景下，还要关注算法定价权的问题。为了确保算法系统确定的工资水平公正合理，应避免用人单位凭借科技壁垒独决，使劳动者能够参与其中，探索适应数智化特征需求的工资集体协商是一条重要的进路。

（二）对劳动者赋权：以离线权为路径的自救

如果说劳动基准是国家从宏观层面对预防过劳的外部底线控制，那么在此之外还有必要在微观层面开辟劳动者自救的机制，赋予劳动者离线权就

是自救的实现路径。事实上，二者具有密切的联系，对于用人单位超出工作长时的指示，劳动者自行离线是贯彻劳动基准的私人实现机制。

顾名思义，离线权即离线的权利，在数智时代特指切断与互联网、各种数字智能控制系统联系的自由权。在广义上，离线权是所有自然人普遍享有的旨在摆脱数字网络过度控制的权利。在狭义上，离线权则特指职场环境下劳动者切断与工作联系以实现工作与私人生活平衡的权利。目前各国的离线权立法主要针对后者。对离线权的理解有两个向度：一是劳动者对工作的拒绝权，侧重权利面向；二是用人单位在非工作时间不得联系、不得打扰，侧重义务面向。事实上这是一个问题的两面，权利的行使离不开义务的履行，就如同硬币的两面。若从权利语境出发，第一种意义的离线权更为直接，立法条文多表述为劳动者有权如何。^[44]由于职场数智化控制与过劳愈演愈烈，劳动者成为被困在工作系统里的透明数字人，因此赋予劳动者离线权显得尤为重要。^[45]具体而言，劳动者的离线表现为在非工作时间拒绝携带电脑等办公设备、从数字网络工作系统退出登录、退出工作群、不接听雇主的电话、不回复雇主的微信和电子邮件、不接受雇主发出的工作指示等。简言之，离线权就是劳动者在非工作时间对用人单位不理睬、不回复、不服从的自救权。学者形象地将之称为“必要的消失”权。^[46]

目前域外已经有不少国家在立法上明确确立了离线权。例如，法国早在2016年就通过第2016-1088号法律对劳动法典进行修订，在国际上首次确立了劳动者离线权。在此之后，意大利、比利时、西班牙、爱尔兰等国家的立法也陆续确立了离线权。2021年欧洲议会发布了《关于欧盟委员会离线权建议的决议》，并附有《欧洲议会和理事会关于“离线权”的指令建议文本》作为附件。可见，在国际上，离线权立法正如火如荼，成为大势所趋。在我国，确立离线权的呼声也日益高涨。^[47]号角已经吹响，只待付诸行动。

离线权究竟如何治理过劳？首先，在根本上，

通过离线权的赋予使劳动者获得了对抗用人单位的砝码,以使职场数字权力的过度失衡得到矫正,从而消除过劳的根源。其次,离线使因劳动管理数智化导致的工作与私人生活边界的模糊化得以厘清。即便用人单位借助数智技术手段将触角延伸至劳动者的私人时间和空间,劳动者也可自行下线,也就切断了工作不当延伸的触角,为工作与私人生活强行划界。离线权是一种单方权利,无需与用人单位协商,是“自己动手,丰衣足食”式的过劳自救法律工具。再次,离线使劳动者切断了针对自己的数据收集和算法控制的渠道。数据和算法是用人单位对劳动者施加控制的两大利器,二者皆以在线为前提条件。劳动者只要在线,就成为源源不断输出信息的数据提供人,进而被困在算法营造的系统里。劳动者一旦离线,数据源和控制通道就被阻断了,劳动者就逃离了系统而得以喘息。最后,离线的直接效果是劳动者得到必要的休息,生理和心理压力得到舒缓,从而有效地预防过劳的形成。事实上,离线权不是一种完全独立的权利,而是一种衍生性权利——即衍生于休息权,离线是为了更好地休息,避免身心过度耗损。[48]离线权是一种富有鲜明时代特色的权利,数智技术既然带来了过劳加剧的风险,相应地就应同时配置用于对治过劳的工具包。当过劳数智化,劳动者的权利配置也应当数智化。被困在系统里的劳动者,通过离线而摆脱系统的控制,也就重获了休息和自由。

离线权的必要性和正当性毋庸置疑,关键在于究竟如何确立离线权?相应的制度规则如何设计?在笔者看来,对此可从解释论和立法论两个维度分别展开。首先,从解释论来看,在工作时间之外,劳动者当然有不工作的权利。因此,即使我国立法尚未明确规定离线权,劳动者也理所当然地自动享有离线的权利。离线权作为子权利根植于其母权利——休息权,而对于休息权我国宪法和劳动法都不乏规定。此外,民法也可为离线权提供一定的解释依据。离线权具有人格权属性,隐私权项下的私人生活安宁权以及个人信息权益,皆可为离线权提供解释空间。[49]虽然劳动者在非工作时间可以

离线是不证自明之理,不过通过制度创设明确确立离线权仍旧十分必要。相比解释论的隐晦不宣,明确赋权模式可以更好地起到权利彰显的作用,刺激劳动者乃至全社会的权利意识觉醒,使更多人明白地认识到自己享有这项权利,并积极行使这项权利。在“996”加班文化盛行的内卷时代,通过立法明确宣示、倡导离线权,对于营造尊重休息权、关爱劳动者身心健康的社会氛围意义重大。

确立劳动者离线权,要想通过高层级立法一蹴而就,在短时间内较难实现。可考虑采取分步走的策略,运用多元化手段逐步确立离线权。第一,通过集体协商方式确立离线权。用人单位与工会或者劳动者代表进行协商,将离线权纳入集体合同,是确立离线权的重要路径。第二,通过行业自律的方式确立离线权。相关行业应当通过行业协会出台行为守则等方式,使尊重劳动者离线权成为一种行业自觉。第三,有关主管机关应及时出台离线权的规则指引,作为软法弥补当下刚性立法缺位的不足。第四,作为终极选择,立法应当确立离线权。具体而言,可在未来制定的劳动基准法或劳动法典中加入离线权的特别条款。法律规范中的离线权规则应当涵盖离线权的适用范围、适用条件、除外规则、行使方式、救济与责任等具体内容。

(三)对用人单位课以义务:以预防过劳为导向的算法善治

在对劳动者赋权的同时,矫正职场算法权力失衡的另一面是对用人单位课以算法善治义务。从技术角度看,算法是人工智能的核心,是导致过劳的罪魁祸首。因此,应有针对性地建构以预防过劳为导向的算法向善义务体系。正如最高人民法院副院长贺小荣在接受《南方周末》采访时指出的,应把劳动管理算法关进法治“笼子”。[50]

令用人单位负担算法治理义务有充分的正当性基础。首先,用人单位通过制定和部署算法,开启了算法致劳的风险,因此理应对由其开启的风险负担防治义务。其次,在算法的运用过程中,用人单位有能力控制风险。尽管算法管理具备一定自动化决策的能力,但用人单位对于算法如何运行不是

只能袖手旁观,而是可以有所作为的,其作为科技系统的掌控者最有能力控制算法运行中的过劳风险,有义务持续监控风险并适时采取相应措施。最后,用人单位从算法管理中受益。算法提升了用人单位劳动管理的效率,增强了对劳动者的技术控制,由此获得了丰厚的经济回报。按照权利义务相一致原则,用人单位不能只获得利益,也应负担算法治理义务。

用人单位的算法治理义务并非空穴来风,而是具有充分的规范依据。国家网信办等9部门联合发布的《关于加强互联网信息服务算法综合治理的指导意见》(国信办发〔2021〕7号)指出:“企业应强化责任意识,对算法应用产生的结果负主体责任。”该条十分明确地昭示了企业作为算法第一责任人的地位。国家网信办制定的《互联网信息服务算法推荐管理规定》第20条规定:“算法推荐服务提供者向劳动者提供工作调度服务的,应当保护劳动者取得劳动报酬、休息休假等合法权益,建立完善平台订单分配、报酬构成及支付、工作时间、奖惩等相关算法。”该条文更加具有针对性地对劳动领域的工作调度算法作出了规范,“应当”的措辞毫无疑问地指向用人单位的算法治理义务。

用人单位负担一系列的算法治理义务,构成一个完整的义务群。第一,算法适恰性义务。用人单位应公平合理地设计算法,尊重生活经验法则,避免过严算法诱发过劳的风险。应当在算法中设计疲劳提示和连续劳动间歇性强制下线规则,保护劳动者的休息权益。《关于落实网络餐饮平台责任切实维护外卖送餐员权益的指导意见》指出:“优化算法规则,不得将‘最严算法’作为考核要求,要通过‘算法取中’等方式,合理确定订单数量、在线率等考核要素,适当放宽配送时限。”该条明确提出抵制“最严算法”,倡导“算法取中”,对于预防过劳意义重大。从程序合法的角度看,由于管理算法的本质是一种劳动规则,[51]其制定应当参照劳动法上的劳动规章践行民主程序——包括征询劳动者意见、公示和告知。

第二,确保算法透明可释的义务。算法黑箱是

算法饱受诟病的原因之一,为揭开算法的黑箱,必须贯彻公开透明原则。劳动者享有算法的知情权和解释权,用人单位则负有相应的算法公开和解释说明的义务。[52]《关于加强互联网信息服务算法综合治理的指导意见》第13条对于算法的公开和解释义务作了明确规定,其对于劳动算法亦可适用。

《新就业形态劳动者劳动规则公示指引》第5条对劳动领域用工算法的公示加以特别规范。从域外法来看,欧盟《改善平台工作条件指令》第9条对算法管理的透明度问题作了大篇幅的规定。算法的公开和解释有助于劳动者了解算法的真相以及自身享有的休息权,从而摆脱算法对超时劳动的引诱与规训。

第三,算法的风险评估义务。鉴于算法在劳动管理中的普及以及其对于劳动者权益的重大影响,有必要对劳动算法进行风险评估。《关于加强互联网信息服务算法综合治理的指导意见》第8条规定了算法的安全评估。由于攸关劳动者的就业与生存,因此在众多算法中劳动算法的风险尤其高。值得特别关注的是,欧盟《人工智能法案》将与劳动就业有关的人工智能系统明确列为高风险等级,是重点评估对象。在我国,《新就业形态劳动者劳动规则公示指引》第7条对平台领域的算法风险评估作了特别规定。对过劳的算法风险评估不应止于算法的制定或修订阶段,而应动态地贯穿于算法制定和运用全流程中的各个环节,周期性地展开。

第四,算法的人力监督义务。算法管理作为一套技术系统,不可避免地存在着一些局限性。算法不能完全脱离人类的监督独立运行,在技术环路之中人力的干预始终应当在场。当用人单位监督发现算法存在导致过劳的较高风险时,应当采取必要措施,包括适当修改和停止使用算法管理系统。

(四) 数智时代过劳救济的两条道路

对已陷入过劳状态的劳动者给予救济,是数智时代过劳治理体系的末端环节。对此,可尝试从两个维度搭建救济机制:一是从劳动法视角将过劳纳入职业伤害保障,二是从私法视角对因不当算法管理导致过劳的用人单位追究侵权责任。

1. 数智时代过劳救济的工伤保险进路

在劳动法视角下,对劳动者职业伤害风险的救济主要是通过工伤保险路径实现的。能否将过劳纳入工伤保险的保障范围,是过劳损害救济的重要思考方向。当然,这也是一个颇具争议的问题。[53]

我国《工伤保险条例》并未将过劳纳入工伤,职业病目录也未将过劳纳入其中。从解释论的视角看,过劳适用工伤保险仅有的希望是过劳性猝死。

《工伤保险条例》第15条第1款第1项规定,“在工作时间和工作岗位,突发疾病死亡或者在48小时之内经抢救无效死亡的”,视同工伤。据此,过劳死如果满足该条设定的条件,则可被视同工伤。当然,条件是比较苛刻的,过劳死并不总是发生在工作时间和工作地点,死亡也并不总是那么急促,视同工伤是小概率事件。相比过劳死,现实中更普遍的状况是过劳而不死、活受罪。过劳死是一个事件,过劳而不死则是一种状态,后者比前者更难被认定成工伤,更加值得关注。

将过劳纳入工伤在解释论上行不通,只有仰赖立法完善。而要真正实现这样的立法,难度极高。从理论上,学者对将过劳纳入工伤保险的呼吁很早就开始了。[54]在富士康员工“十四连跳事件”助推下,过劳自杀工伤认定问题曾引发社会热议。[55]遗憾的是,每一次学术上的激情呼喊都没有产生立法上的涟漪,以至于这一问题被渐渐淡化。造成这一窘境的根源在于过劳认定的难度,特别是过劳与工作之间因果关系的证明,立法上不好确立清晰的认定标准。不同于事故性工伤,也不同于尘肺病等便于医学上评估的典型职业病,过劳是一种难识别的亚健康状态。除了超强度工作,过劳还可能与自身健康、环境影响、行为习惯、婚姻状况、家庭压力、人际关系等个人因素有着千丝万缕的联系。特别是在数智时代背景下,心理性、精神性过劳问题日渐突出,其比生理性过劳更难认定。[56]从过劳救济先进国家的情况来看,过劳的工伤认定有赖于一套测量量表和评估机制的建立。[57]而在我国,这样的量化评估机制尚属空白。

在笔者看来,应当用动态发展的眼光看待过劳

立法问题,在数智时代过劳加剧的新形势下,关于过劳工伤立法的讨论应当被重启并转化为行动。评估一项立法是否必要,一个关键考量因素是其拟解决的问题在现实中的严重程度与迫切程度。注目当下,数智时代过劳问题非但没有消失反倒更加严重了,是一个关系民生人权的真问题、大问题、普遍性问题,过劳工伤立法的需求比以往任何时候都更加迫切。考察可知,日本、韩国等肯定过劳工伤的国家立法也经历了一个转变的历程,从最初的排斥到后来逐渐接受,过劳的范围从心脑血管疾病扩展到焦虑症等精神疾患,过劳因果关系的认定标准也越来越宽松,过劳立法越来越开放。[58]我国作为一个过劳问题同样严重的14亿人口大国,过劳工伤立法攸关人民福祉,不应当一直被留白,现在启动正当时。

过劳工伤立法的瓶颈在于认定标准模糊、因果关系不好证明,但并非不能突破。从实质标准来看,确因工作原因遭受事实上的损害是工伤认定的核心。同理,如果劳动者因长期受到高压的数智化管理而陷入不能缓解的身心疲惫,也符合工伤的实质性标准。更进一步,对过劳的认定可尝试采用分类分层的方案,在立法上逐步推进。首先,死亡是过劳损害的最高级别,相对较易认定。其次,是与长期承受压力相关的身体性疾患,可率先承认心肌梗塞、脑梗塞等心脑血管疾病型过劳构成工伤,之后再逐渐向其他身体疾患扩展。最后,是心理性、精神性过劳疾患,如焦虑症、抑郁症等。此类过劳的认定难度最高,不过压力可导致精神疾患已有医学上的证明,借助评估量表进行个案认定可构成工伤,应逐步被工伤立法涵摄。至于尚不构成疾患的不良情绪和单纯的疲惫状态,则一般不宜被认定为法律意义上的过劳。对于过劳是否因工作原因导致的判断,则应由行政机关和司法机关依据一定的标准进行个案裁量。从日本的经验来看,对过劳因果关系的判断有异常负荷说、过重负荷说、发症促进说等不同标准。[59]从主体的视角,早期多以“同类劳动者”的心理负荷水平作为参照对象,后期又发展出“最脆弱劳动者”和劳动者本人标准。[60]在时

间维度上,要求发病与承受不当工作压力的一段时间不得超过一定间隔,日本1999年颁布的《关于心理负荷(压力)导致的精神疾患工作原因判断指针》最初的规定是6个月。为了消除纯主观判断带来的随机性和不确定性,日本发展出了认定过劳的心理评估量表,近二三十年间经过多次修订不断完善。这些经验都值得我国借鉴,特别是关于过劳评估量表的建设。

我国的过劳工伤立法注定是一项长远艰巨工程,无望在短期内一蹴而就。但路虽远,行则将至,不行则不至。数智时代为过劳工伤立法提供了挑战兼机遇,现在就应当上路。

2. 数智时代过劳救济的侵权法进路

过劳工伤制度的建立和完善乃是远景之展望,但数智时代过劳的救济需求当下就存在。为解决眼下的难题,只能另辟蹊径,通过侵权法对过劳进行救济是一条有希望的路径。相比于工伤的认定,私法上的责任有更高的弹性和涵摄性。当然,通过侵权法对过劳进行救济也存在不小的障碍,但并非无法逾越。

从损害要件来看,数智时代过劳的损害可从人格权中找到权利基础和规范依据。在一定情况下,身心疲惫达到一定程度的,可构成对健康权的损害。《民法典》第1004条规定:“自然人享有健康权。自然人的身心健康受法律保护。任何组织或者个人不得侵害他人的健康权。”值得特别关注的是,该条采取了“身心健康”的表述,即除了生理健康,心理健康也可能在保护范围之内,这对于将过劳纳入健康权损害提供了更大的解释空间。[61]劳动者过劳导致可确诊的生理或心理疾病的,应有条件地认可其构成私法上的损害。此外,休息权也是过劳救济的潜在权利基础。民法典虽未明确规定休息权,不过其或可从一般人格权获得解释,民法学者对于休息可作为一种人格权多有认可。[62]

因果关系认定难是数智时代过劳侵权责任成立的最大障碍。到底是不是用人单位的超负荷工作安排导致了过劳,较难证明。具体到数智化劳动场景,过劳损害的因果关系证明难度更高。与一般的

侵权行为不同,算法侵权本质上是一种规则侵权。

[63]由于用人单位的隐身,很多工作安排都是由算法发出的,而很多算法表面看并非是强制性的,再加上算法的黑箱效应以及劳动者科技认知能力局限导致的举证能力不足,算法与过劳的因果关系证明并非易事。不过,困难并非不能克服。考虑到算法侵权的特殊性,改变因果关系的证明范式以及降低证明标准成为主流导向。在数智时代侵权责任的因果关系已经不体现为自然法则意义上的引起与被引起的规律性联系,而是更多地体现为一种相关关系。[64]在这一路径下,因果关系判断不需要被囚困在是否一个具体的行为实际引起了一个具体损害后果的思维定势之中,只需要此类行为导致此类损害结果达到一定的概率即可。应当看到,一些不合理的算法与过劳有着并不牵强的相关关系。例如,平台通过最严算法不断压缩骑手的送单时间,借订单分配、算法评分及奖惩机制对劳动者进行规训,在这些情形下算法与过劳损害的因果关系并不遥远。可喜的是,在一些政策规范中隐约可见肯定算法致害因果关系成立的立场。例如,《最高人民法院关于为稳定就业提供司法服务和保障的意见》(法发〔2022〕36号)第8条指出:“与用工管理相关的算法规则存在不符合日常生活经验法则、未考虑遵守交通规则等客观因素或者其他违背公序良俗情形,劳动者主张该算法规则对其不具有法律约束力或者请求赔偿因该算法规则不合理造成的损害的,人民法院应当依法支持。”该条所称“因该算法规则不合理造成的损害”,“因”“造成”的表述明显表达了认可算法致害可成立因果关系的肯定立场。特别值得关注的是,在判断因果关系是否成立时,该条意见采用了“生活经验法则”“交通规则”“公序良俗”等客观标准,只要劳动管理算法违背了这些标准,即可推定因果关系成立,不需要劳动者进行其他额外的证明。

过错要件的证明也是数智时代过劳侵权认定的一个难题。由于算法侵害的技术性、隐蔽性,要劳动者证明用人单位的主观过错难度很大。从算法侵权的归责原则来看,我国现行立法没有特别规定,

在理论上则有过错原则[65]、无过错责任原则[66]、过错推定原则[67]和区分说[68]之争。在笔者看来,算法侵权归责原则的抉择不能不考虑算法适用的场景及其风险等级因素。考虑到职场算法的高风险性,算法侵权的归责原则应向有利于劳动者的方向倾斜。我国《个人信息保护法》第69条对个人信息侵权实行过错推定原则,用于劳动管理的算法进行分析时所使用的数据如果包含了可识别到具体劳动者的个人信息,则依据该条适用过错推定原则。如果算法使用的是经过脱敏的数据,就超出了个人信息保护的范围而无法适用该原则。不过,考虑到算法分析比一般的个人信息处理手段有更高的技术门槛,其过错的证明更加困难,按照举轻以明重的法理,理应实行对受害者友好的归责原则。在法无明文窘境下,在解释论上应放宽劳动者对用人单位过错的证明标准。从前述《最高人民法院关于为稳定就业提供司法服务和保障的意见》第8条的规定来看,在表述劳动用工的算法侵权时,特别强调的是算法本身“不合理”,而淡化了用人单位的过错,或者说其具有从算法不合理反推用人单位过错的意思,虽非过错推定原则的正面规定但有异曲同工之妙。而对算法规则是否合理的判断标准,是该规则“不符合日常生活经验法则、未考虑遵守交通规则等客观因素”,特别强调“客观因素”。

数智时代过劳的治理是一个十分复杂的问题,限于篇幅,本文无法穷尽治理体系每一方面的细节。除了上述四种手段之外,过劳治理工具箱中还有其他遗珠未及详述,例如劳动监察。从实践层面看,很多过劳的发生并非由于无法可依,而是执法不严,既有的劳动基准及其他保护制度得不到有效贯彻落实。因此,针对这一问题的劳动监察急需加强。需要特别指出的是,在数智时代劳动监察工作也应当与时俱进,在传统的监察对象之外,增强对用人单位实施算法管理的监督管理,对利用不合理算法压迫劳动者过度劳动的行为予以纠正和惩戒。2024年11月,中央网信办、工信部、公安部、国家市场监管总局四部门联合开展“清朗·网络平台算法典型问题治理”专项行动,其中确立的六大核心任

务之一就是:防范盲目追求利益侵害新就业形态劳动者权益,严防一味压缩配送时间导致配送超时率、交通违章率、事故发生率上升等问题。这一劳动监察任务直指算法导致的过劳问题,是劳动监察因应数智时代发展的积极探索。

五、结语

数智技术运用于职场,改变了劳动世界的面貌。期盼中的劳动者全面自由解放并未如期而至,高科技范儿的新型过劳却不请自来。无处不在的智能监控打破了时空的界限,工作的边界消失,家庭亦是职场。算法管理下的工作虽貌似自由,其实是隐蔽的牢笼。数智时代的过劳不是罕见个例,而是普遍的时代景象。当数量庞大的劳动者群体处于过劳困境之中,体面劳动就成为空谈,社会公平正义受损。整个劳动法的历史都是与过劳的抗争史,面对数智时代过劳的新特征、新挑战,法律也应随之进步和革新。建立和完善数智时代的过劳治理机制,是保护劳动者健康、促进人的全面发展的需要,是中国式现代化背景下社会法体系完善的重要一环。[69]可从劳动基准的底线控制、对劳动者赋权、对用人单位课以算法善治义务、对过劳损害进行救济等多个维度,建构数智时代过劳的法律治理体系,借法律之力矫正数智技术的异化。祈盼理想中科技加持下的“休闲时代”早日来临,愿所有劳动者皆得休息和健康。

(技术编辑:艾薇)