

中国人民大学法学院 数字法学教研月报

2025 年第 10 期（总第 22 期）

2025 年 10 月 20 日



本期看点

【数字法治大事件】国家多部门密集出台数字法治领域重磅文件，国家网信办与市场监管总局联合公布《个人信息出境认证办法》，中央网信办与国家发改委印发《政务领域人工智能大模型部署应用指引》，国家发改委等六部门推出《关于加强数字经济创新型企业培育的若干措施》，从数据跨境、AI 治理、企业培育维度完善制度框架；国际层面，意大利修订著作权法，成为世界首部直面人工智能的著作权立法；地方层面，《杭州倡议：“跨境电商平台规则”的指导原则》发布，形成“国家——地方——国际”多层次治理成果，配套专家解读与时评进一步明确明晰规则方向。

【研究动态】本期研究覆盖数字法学核心领域，基础理论聚焦数字法学知识范式转型、数字财产侵害行为处断规则等；个人信息保护探讨民事公益诉讼损害赔偿、爬取行为“过罪化”等问题；数据确权与流通研究健康医疗数据规制、数据资产融资担保等；人工智能涉及生成式 AI 侵权责任、训练数

据合理使用等；平台治理、数字行政与司法、虚拟财产等领域亦有深度研究，为数字法治实践提供理论支撑。

【教研活动】中国人民大学法学院成功举办第六届“未来法治与数字法学”国际论坛，汇聚全球八十多位专家学者共议 AI 法治与全球治理；第六届未来法治学生论坛同步召开，聚焦 AI 时代法学范式转型等前沿议题；中国法学会网络与信息法学研究会 2025 年年会在南京召开，围绕“智能时代的数字法治”展开研讨；多地专题研讨会与培训班同步开展，推动数字法学学术交流与人才培养。

【数字法评】

《论生成式人工智能服务提供者的过错推定责任》，《北方法学》2025 年第 5 期，作者：张新宝、卞龙。

《论大模型训练数据的合理使用》，《法学家》2025 年第 5 期，作者：李铭轩。

本期目录

数字法治大事件.....	3	个人信息保护.....	23
国家互联网信息办公室、国家市场监督管理总局 联合公布《个人信息出境认证办法》.....	3	数据确权与流通.....	25
《个人信息出境认证办法》答记者问.....	3	人工智能.....	27
中央网信办、国家发展改革委印发《政务领域人 工智能大模型部署应用指引》.....	5	平台治理.....	29
国家发展改革委等部门关于印发《关于加强数字 经济创新型企业培育的若干措施》的通知.....	6	数字行政与司法.....	30
世界首部直面人工智能的著作权法---意大利法 做出重要修改.....	8	虚拟财产.....	32
《杭州倡议：“跨境电商平台规则”的指导原则》 发布.....	10	教研活动.....	33
专家解读丨《个人信息出境认证办法》公布 完 成数据跨境制度体系新拼图.....	10	人工智能时代全球法治大会丨第六届“未来法治 与数字法学”国际论坛成功举行.....	33
专家解读丨建构个人信息出境认证制度 完善个 人信息出境合规体系.....	12	第六届未来法治学生论坛成功召开.....	38
专家解读丨推动政务大模型部署应用 赋能电子 政务智能化升级.....	13	中国法学会网络与信息法学研究会 2025 年年会 暨第三届数字法治大会在江苏南京召开.....	40
原创时评丨陈增宝：互联网法院管辖新规的三大 亮点.....	16	“数据要素合规高效流通利用机制研究”专题研 讨会成功举办.....	42
研究动态.....	19	学术综述丨上海市法学会数字法学研究会 2025 年学术年会观点摘编.....	43
基础理论.....	19	2025 年全国数据系统局长培训班成功举办.....	46
		数字法评.....	47
		论生成式人工智能服务提供者的过错推定责任47	
		论大模型训练数据的合理使用.....	58

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：梁因格

数字法治大事件

导言：近期，数字法治领域迎来政策密集发布期，国家多部门联合出台个人信息保护、政务AI应用、数字经济企业培育等领域重磅文件，同时国际层面亦有AI著作权立法突破，配套专家解读与时评进一步明晰规则方向，为数字时代的合规发展与智能化升级筑牢制度基础。具体来看，国家互联网信息办公室、国家市场监督管理总局联合公布《个人信息出境认证办法》，中央网信办与国家发展改革委同步印发《政务领域人工智能大模型部署应用指引》，国家发展改革委等部门则推出数字经济创新型企业培育措施，从数据安全、AI治理、产业扶持三大维度完善数字法治框架。国际上，意大利对著作权法进行修改，出台世界首部直面人工智能的相关法律；国内地方层面，《杭州倡议：“跨境电商平台规则”的指导原则》同步发布，形成“国家-地方-国际”多层次数字治理成果。此外，多篇专家解读深度剖析个人信息出境认证制度、政务大模型应用价值，原创时评亦聚焦互联网法院管辖新亮点，为政策落地与行业实践提供专业参考。

国家互联网信息办公室、国家市场监督管理总局联合公布《个人信息出境认证办法》

原载：“网信中国”微信公众号

近日，国家互联网信息办公室、国家市场监督管理总局联合公布《个人信息出境认证办法》（以下简称《办法》），自2026年1月1日起施行。

国家互联网信息办公室有关负责人表示，《个人信息保护法》对个人信息出境认证制度作了基本规定，按照国家网信部门的规定经专业机构进行个人信息保护认证是向境外提供个人信息的法定途径之一。《办法》对个人信息出境认证的适用情形，个人信息出境认证的申请方式、认证要求及证书有效期，专业认证机构应当履行的义务，监督管理要求等作出细化规定，旨在保护个人信息权益，规范个人信息出境认证活动，促进个人信息高效安全跨境流动。

《办法》明确了个人信息出境认证的适用情

形。个人信息处理者通过个人信息出境认证的方式向境外提供个人信息的，应当同时符合下列情形：

一是非关键信息基础设施运营者；二是自当年1月1日起累计向境外提供10万人以上、不满100万人个人信息（不含敏感个人信息）或者不满1万人敏感个人信息，且向境外提供的个人信息中不包括重要数据。规定个人信息处理者不得采取数量拆分等手段，将依法应当通过出境安全评估的个人信息通过个人信息出境认证的方式向境外提供。

《办法》明确了个人信息出境认证的申请方式、认证要求及证书有效期。规定个人信息处理者应当向专业认证机构申请个人信息出境认证，专业认证机构应当按照认证基本规范、个人信息保护认证规则开展认证活动。明确认证证书的有效期为3年。

《办法》明确了专业认证机构应当履行的义务。规定专业认证机构应当向全国认证认可信息公共服务平台报送个人信息出境认证证书相关信息。发现获证个人信息处理者存在个人信息出境情况与认证范围不一致等情形，不再符合认证要求的，应当暂停其使用直至撤销相关认证证书。发现个人信息出境活动违反法律、行政法规和国家有关规定，应当及时向国家网信部门和有关部门报告。

《办法》明确了监督管理要求。规定专业认证机构自取得认证资质之日起10个工作日内，应当向国家网信部门备案，国家市场监督管理总局和国家网信部门对个人信息出境认证活动进行监督。省级以上网信部门和有关部门发现获证个人信息处理者个人信息出境活动存在较大风险或者发生个人信息安全事件的，可以依法对获证个人信息处理者进行约谈。

《办法》对违反《办法》规定的法律责任、适用效力等作出了规定。

《个人信息出境认证办法》答记者问

原载：“网信中国”微信公众号

近日，国家互联网信息办公室、国家市场监督

管理总局联合公布了《个人信息出境认证办法》（以下简称《办法》）。国家互联网信息办公室有关负责人就《办法》有关问题回答了记者提问。

问 1：请介绍一下《办法》的出台背景。

答：随着全球数字经济飞速发展，数据跨境流动成为推动数据要素全球配置、开展高水平国际合作竞争的关键。《个人信息保护法》规定，按照国家网信部门的规定经专业机构进行个人信息保护认证是向境外提供个人信息的法定途径之一。为落实法律规定要求，2022年11月，国家市场监督管理总局、国家互联网信息办公室公布了规范性文件《关于实施个人信息保护认证的公告》。在此基础上，制定《办法》，完善我国个人信息跨境流动管理制度，有利于保护个人信息权益，促进个人信息合规利用，为数字经济高质量发展提供法治保障。

问 2：《办法》的主要内容是什么？

答：《办法》主要对下列内容进行了规定：一是明确立法目的依据、适用范围。规定个人信息处理者通过个人信息保护认证的方式向中华人民共和国境外提供个人信息，适用本办法。二是明确个人信息出境认证的适用情形。规定个人信息处理者通过个人信息出境认证的方式向境外提供个人信息应当符合的情形，即非关键信息基础设施运营者自当年1月1日起累计向境外提供10万人以上、不满100万人个人信息（不含敏感个人信息）或者不满1万人敏感个人信息，且个人信息不包括重要数据。三是明确个人信息出境认证的申请方式、认证要求及证书有效期。规定个人信息处理者应当向专业认证机构申请个人信息出境认证，中华人民共和国境外的个人信息处理者申请个人信息出境认证的，应当由其在境内设立的专门机构或者指定代表协助进行申请。专业认证机构应当按照认证基本规范、个人信息保护认证规则开展认证活动。明确认证证书的有效期为3年，证书到期需继续使用的，个人信息处理者应当在有效期届满前6个月提出认证申请。四是明确专业认证机构应当履行的义务。规定专业认证机构应当向全国认证认可信息公共服务平台报送个人信息出境认证证书相关信息，发

现个人信息出境活动违反法律、行政法规和国家有关规定的，应当及时向国家网信部门和有关部门报告。五是明确监督管理要求。规定专业认证机构自取得认证资质之日起10个工作日内，应当向国家网信部门备案，国家市场监督管理总局和国家网信部门对个人信息出境认证活动进行监督。

问 3：通过个人信息出境认证的方式向境外提供个人信息的适用情形如何？

答：《办法》明确个人信息处理者通过个人信息出境认证的方式向境外提供个人信息的，应当同时符合下列情形：一是非关键信息基础设施运营者；二是自当年1月1日起累计向境外提供10万人以上、不满100万人个人信息（不含敏感个人信息）或者不满1万人敏感个人信息；三是向境外提供的个人信息，不包括重要数据。同时，规定个人信息处理者不得采取数量拆分等手段，将依法应当通过出境安全评估的个人信息通过个人信息出境认证的方式向境外提供。

问 4：个人信息处理者在申请认证向境外提供个人信息前应当履行什么义务？

答：落实《个人信息保护法》《网络数据安全条例》规定要求，《办法》对个人信息处理者在申请认证向境外提供个人信息前应当履行的义务作了细化。规定应当按照法律、行政法规的规定履行告知、取得个人单独同意、进行个人信息保护影响评估等义务。个人信息保护影响评估重点评估以下内容：一是个人信息处理者和境外接收方处理个人信息的目的、范围、方式等的合法性、正当性、必要性；二是出境个人信息的规模、范围、种类、敏感程度，个人信息出境可能对国家安全、公共利益、个人信息权益带来的风险；三是境外接收方承诺承担的义务，以及履行义务的管理和技术措施、能力等能否保障出境个人信息的安全；四是个人信息出境后遭到篡改、破坏、泄露、丢失、非法利用等的风险，个人信息权益维护的渠道是否通畅等；五是境外接收方所在国家或者地区的个人信息保护政策和法规对出境个人信息安全和个人信息权益的影响；六是其他可能影响个人信息出境安全的

事项。

问 5：《办法》对专业认证机构提出了哪些要求？

答：《办法》对专业认证机构提出了如下要求：一是应当按照认证基本规范、个人信息保护认证规则开展个人信息出境认证活动。符合认证要求的，应当及时出具认证证书。二是应当在出具认证证书或者认证证书状态发生变化后 5 个工作日内，向全国认证认可信息公共服务平台报送个人信息出境认证证书相关信息，包括认证证书编号、获证个人信息处理者名称、认证范围以及证书状态变化信息等。三是发现获证个人信息处理者存在个人信息出境情况与认证范围不一致等情形，不再符合认证要求的，应当暂停其使用直至撤销相关认证证书。四是发现个人信息出境活动违反法律、行政法规和国家有关规定的，应当及时向国家网信部门和有关部门报告。五是应当自国家市场监督管理总局批准取得个人信息保护认证资质之日起 10 个工作日内向国家网信部门办理备案手续，并对所备案材料的真实性负责。六是应当对在履行职责中知悉的个人隐私、个人信息、商业秘密、保密商务信息等依法予以保密。

问 6：《办法》如何对专业认证机构与获证个人信息处理者进行监督管理？

答：《办法》对专业认证机构与获证个人信息处理者的监督管理作了如下规定：一是国家市场监督管理总局和国家网信部门对个人信息出境认证活动进行监督，开展定期或者不定期的检查，对认证过程和认证结果进行抽查，对专业认证机构进行抽查和评价。二是省级以上网信部门和有关部门发现获证个人信息处理者个人信息出境活动存在较大风险或者发生个人信息安全事件的，可以依法对获证个人信息处理者进行约谈。获证个人信息处理者应当按照要求整改，消除隐患。三是任何组织和个人发现获证个人信息处理者违反本办法规定向境外提供个人信息的，可以向专业认证机构、网信部门和有关部门投诉、举报。四是违反《办法》规定的，依据《个人信息保护法》《网络数据安全管

理条例》《认证认可条例》等法律法规处理；构成犯罪的，依法追究刑事责任。

问 7：我国数据跨境流动的制度设计情况如何？

答：《网络安全法》《数据安全法》《个人信息保护法》《网络数据安全条例》对数据跨境流动作出基本规定，个人信息处理者因业务等需要，确需向中华人民共和国境外提供个人信息的，应当具备下列条件之一：（一）通过国家网信部门组织的安全评估；（二）按照国家网信部门的规定经专业机构进行个人信息保护认证；（三）按照国家网信部门制定的标准合同与境外接收方订立合同，约定双方的权利和义务；（四）法律、行政法规或者国家网信部门规定的其他条件。国家互联网信息办公室先后出台《数据出境安全评估办法》《个人信息出境标准合同办法》《促进和规范数据跨境流动规定》，明确了数据出境安全评估、个人信息出境标准合同等管理制度的实施路径，同时建立了自贸试验区数据出境负面清单制度。《办法》的出台，明确了通过认证方式向境外提供个人信息的具体实施路径，标志着《个人信息保护法》明确的数据出境安全评估、个人信息保护认证、个人信息出境标准合同等出境制度设计的全面落地，也标志着我国数据跨境流动制度体系的全面建立。

中央网信办、国家发展改革委印发《政务领域人工智能大模型部署应用指引》

原载：“网信中国”微信公众号

贯彻党中央、国务院决策部署，落实《关于深入实施“人工智能+”行动的意见》要求，为安全稳妥有序推进政务领域人工智能大模型部署应用，中央网信办、国家发展改革委近日联合印发《政务领域人工智能大模型部署应用指引》（以下简称《指引》），为各级政务部门提供人工智能大模型部署应用的工作导向和基本参照。

《指引》坚持以习近平新时代中国特色社会主义思想为指导，深入贯彻落实党的二十大和二十届

二中、三中全会精神，全面贯彻习近平总书记关于网络强国的重要思想，坚持系统谋划、集约发展，以人为本、规范应用，共建共享、高效协同，安全稳定、务求实效，有序推进人工智能大模型技术、产品和服务在政务领域的部署、应用和持续优化。

《指引》强调场景牵引。政务部门可围绕政务服务、社会治理、机关办公和辅助决策等工作中的共性、高频需求，因地制宜、结合实际，选择典型场景进行人工智能大模型探索应用。

《指引》强调规范部署。政务部门应根据不同政务场景需求与现有技术基础，审慎选择人工智能大模型实施路径。应以统筹集约的方式开展政务领域人工智能大模型部署，地市应在省（自治区、直辖市）统一要求下开展部署应用，县级及以下原则上应复用上级的智能算力和模型资源开展应用和服务。应探索构建“一地建设、多地多部门复用”的集约化部署模式，统筹推进政务大模型部署应用，防止形成“模型孤岛”。应加强政务数据治理，持续提升数据质量，支撑政务大模型的优化训练。

《指引》强调运行管理。政务部门应统筹减负和赋能，避免盲目追求技术领先、概念创新，切实防范“数字形式主义”。应建立健全全周期管理体系，明确应用方式和边界，落实人工智能大模型“辅助型”定位，防范模型“幻觉”等风险。应将持续迭代优化作为人工智能大模型部署应用的关键环节，建立常态化更新机制。应扎实做好安全管理，建立安全责任制度，明确安全职责和任务，提升人工智能安全风险应对能力。应严格落实保密要求，防止国家秘密、工作秘密和敏感信息等输入非涉密人工智能大模型，防范敏感数据汇聚、关联引发的泄密风险。

《指引》指出，要加强组织实施，加快推进政务领域人工智能大模型国家标准体系建设和重点标准研制，及时总结推广典型场景和创新应用。开展监测评估，构建政务领域人工智能大模型部署应用全流程监测评估体系，持续迭代优化。做好培训宣传，增强工作人员应用能力和水平，提升全民数字素养。

国家发展改革委等部门关于印发《关于加强数字经济创新型企业培育的若干措施》的通知

原载：“国家数据局”微信公众号

国家发展改革委等部门关于印发《关于加强数字经济创新型企业培育的若干措施》的通知

发改数据〔2025〕1154号

各省、自治区、直辖市及计划单列市、新疆生产建设兵团发展改革委、数据管理部门、财政厅（局），中国人民银行上海总部，各省、自治区、直辖市及计划单列市分行，各金融监管局，中国证监会各派出机构：

为加快培育数字经济创新型企业，让更多企业在数字经济新领域新赛道跑出加速度，推动涌现更多瞪羚企业、独角兽企业，国家发展改革委、国家数据局、财政部、中国人民银行、金融监管总局、中国证监会组织制定了《关于加强数字经济创新型企业培育的若干措施》。现印发给你们，请遵照执行。

国家发展改革委

国家数据局

财政部

中国人民银行

金融监管总局

中国证监会

2025年9月4日

关于加强数字经济创新型企业培育的若干措施

数字经济创新型企业（以下简称“数创企业”）是以数据为关键生产要素，以数字技术创新、应用场景创新、数据价值创新为核心驱动力，具备高敏捷性和高成长性的企业，是发展新质生产力的重要实践主体。为加强前瞻性布局，优化鼓励探索、开放包容的创新生态，强化对数创企业的发现和培育，推动在数字经济领域涌现出更多的瞪羚企业、独角兽企业，提出如下举措。

一、健全数创企业源头发现机制。国家发展改革委、国家数据局牵头搭建数创企业培育库，每

年遴选发现一批创新能力强、发展潜力大的数创企业分级分类入库，实行动态调整。各级数据管理部门加强对入库企业的服务对接和监测分析，有针对性提供各类培育政策、开展精准服务。组建涵盖相关部门、地方、国有企业、民营企业、科研院所、高等院校、投资机构、孵化机构等各类专业人员的数创企业专业化遴选培育组，构建“政府+企业+创新+投资”四合一的专业化遴选培育机制。

二、强化多维用数保障。鼓励地方加快建立公共数据授权运营机制和应用创新生态，促进公共数据可持续供给及开发利用。在保障数据安全合规前提下，支持数创企业公平参与公共数据资源开发利用，探索以成本共担、收益共享等方式，保障数创企业开展公共数据资源开发利用创新实践早期用数需求。鼓励有条件地区探索发放“数据券”“算法券”，降低治数用数成本。鼓励国有企业、行业龙头企业、平台企业等构建产业链上下游、平台生态圈数据服务平台，开发提供普惠性数据产品和技术工具，对面向数创企业发展需求，提供普惠便利数据服务的企业和行业可信数据空间予以重点支持。鼓励数创企业依托公共数据资源登记平台开展数据资源登记。

三、强化算力资源供给支撑。深入实施“东数西算”工程，落实有关政策文件要求，坚持国家枢纽节点算力规模部署，持续优化热点应用区域需求保障。加快构建全国一体化算力网，支持地方协同参与、共同建设，在国家统一标准指导下，推动全国算力资源有序池化，并网运行，打造集算力统筹监测、统一调度、弹性供给、安全保障于一体的新型算力网基础设施。引导各类算力资源与数创企业需求高效精准对接，鼓励国家枢纽节点面向数创企业提供低成本、广覆盖、可靠安全的算力服务，降低算力使用门槛。

四、提升原始创新能力。鼓励国有企业、行业龙头企业、平台企业等带头推进联合创新发展，整合推进各类创新资源与服务向数创企业开放共享，加快推进重点产业领域专利池建设，强化产业链上下游、平台生态圈的融通创新。鼓励地方因地制宜，

在知识产权、研发投入等方面对数创企业提供支持，在全国范围内开展知识产权公共服务惠企行动，助力数创企业创新发展。依托数字产业集群建设，发挥集群主导产业优势，吸引各类创新资源集聚，打造有利于数创企业专业化特色化产业化发展的创新生态。

五、完善成果转化机制。鼓励科研院所、高等院校等建立以企业为主导、需求为牵引、产学研深度融合的“有组织科研+有组织成果转化”科技创新成果转化机制。鼓励地方因地制宜，建设一批专业化、市场化的成果转化服务机构，面向数创企业提供项目遴选识别、验证评估、二次开发、中试熟化、创业孵化等公益性成果转化服务。探索建立数创企业新技术、新产品、新方案、新服务“首购首用”专项政策，帮助和推动数创企业成果更快开拓“首市场”，在市场验证中加速成熟。鼓励地方通过推荐目录、路演推介、应用大赛等促进成果供需对接。

六、强化场景和机会供给。支持地方结合城市全域数字化转型、数字产业集群建设等重点工作，将城市发展战略、产业发展布局、公共资源开发等与企业市场机遇融合对接、精准匹配，深入挖掘技术应用新场景、服务支撑新场景、数据赋能新场景，不断加大应用场景和机会清单的开放供给，形成一批可复制可推广的高质量场景。鼓励跨区域跨领域的场景共建、成本共担、收益共享。引导国有企业、行业龙头企业、平台企业等推进场景开放，为数创企业新技术、新产品、新方案、新服务提供测试、展示、应用机会。以场景和行业痛点为牵引，强化算力、数据等要素协同，加强对人工智能、垂类大模型等数创企业培育。

七、强化企业出海服务。加大数字经济领域国际合作力度，加强对国际数字经济标准、数据政策等方面的研究与对接交流，为数创企业提供更好的政策和制度保障。加大支持数创企业参加国际展览活动和资源技术对接力度，依托中国国际大数据产业博览会、上合组织数字经济论坛、全球数字经济大会等国际合作平台，集中展示数创企业创新成

果。鼓励国有企业、行业龙头企业、平台企业及相关行业协会商会等强化产业链生态圈协同出海，带动数创企业积极发挥自身优势拓展海外市场。支持服务机构、行业协会商会发挥资源能力优势，在市场分析、标准规范、风险提示、合规指南建设等方面加强对数创企业的指导和服务。

八、优化投融资服务。鼓励金融机构结合数创企业投融资需求，按照市场化原则提供金融服务。强化创投资金引导，优化完善国有创业投资考核评价机制，探索将培育数字经济瞪羚企业、独角兽企业情况纳入考核机制。鼓励地方完善经营主体信用评价服务体系，鼓励有条件的金融机构构建符合数创企业特点的信用评价模型，完善风险评价机制，优化金融产品和服务模式，切实加大数创企业金融支持力度。在依法合规、风险可控前提下，规范银行与投资机构的合作，为数创企业提供多元化金融服务。促进银企对接，向金融机构推荐优质数创企业项目。加大力度支持符合条件的优质数创企业上市融资。

九、建立开放包容审慎的创新环境。结合企业行业特点，稳慎探索推行“沙盒监管”模式，分级分类制定“沙盒监管”规则，鼓励在风险可控的前提下开展先行先试。规范涉企检查，推进精准检查，防止重复检查、多头检查，探索推行非现场监管，最大限度减少对企业生产经营的不必要干扰。引导数创企业守法自律经营，探索柔性执法机制，对首次轻微违规问题依法优先采取引导协商、行政指导、信用承诺等方式处理。支持地方优化涉企政务服务，推进涉企问题高效解决，通过智能技术实现惠企政策自动适配、免申即享、直达秒兑。

十、强化队伍建设。支持高等院校、职业院校面向实际需求优化调整数字经济相关学科专业设置，推动数字人才梯队建设。探索建立企业数字人才认定评价体系，通过目录认定、授权认定、专才认定等方式，结合行业评判、市场评价、社会评议等评价标准，多维度开展数字人才评价。健全数字人才流动机制，鼓励高端人才在产学研间有序流动，推动校企人才互聘兼职、项目合作。持续加

强数字人才综合服务保障，支持地方搭建上下贯通、系统联动、部门协作的人才综合服务体系。

各地区发展改革部门、数据管理部门会同财政、人民银行、金融监管等部门结合本地数字经济发展实际，加强统筹协调，打造有助于数创企业起跑加速的政策与服务体系，加强对数创企业全生命周期支持，更好促进数创企业发展成长为瞪羚企业、独角兽企业。国家发展改革委、国家数据局、财政部、中国人民银行、金融监管总局、中国证监会加强跟踪评估、总结反馈，适时对地方典型案例和经验做法予以宣传推广。

世界首部直面人工智能的著作权法---意大利法做出重要修改

原载：“知识产权创新与竞争研究中心 CUPL”微信公众号

在意大利，著作权法的全称为：第 633 号法律《保护作者权利以及与行使作者权有关的权利》（1941 年 4 月 22 日）。

经过多轮修订，最新修订日期为 2025 年 9 月 25 日。本次修订，使得意大利著作权法成为世界上首部在立法中就人工智能辅助创作进行直接回应的立法，修改之二是明确规定人工智能模型或系统以文本与数据挖掘目的使用作品构成合理使用。

意大利著作权法的修改详情：

重要修改之一：作品的定义

本法第 1 条：在第 1 条第 1 款“智力作品”前添加“人类”二字，并在条文中明确添加“包括通过人工智能工具创作的作品，只要其构成作者智力劳动成果”。

修订之后，本法第 1 条对作品的定义为：

本法保护的作品是具有创作性特征的人类智力成果，包括文学、音乐、视觉艺术、建筑、戏剧和电影，无论其表达方式或形式如何，即使是借助人工智能工具创作的，只要是作者智力劳动的成果即可。

根据 1978 年 6 月 20 日第 399 号法律批准并实施的《保护文学和艺术作品的伯尔尼公约》，计算

机程序也作为文学作品受到保护。在材料的选择或编排上构成作者智力创作的数据库，也同样受到本法保护。

重要修改之二：人工智能模型或系统使用作品

本法第70(6)条后，新增加第70(7)条。

第70(7)条：在遵循《保护文学艺术作品的伯尔尼公约》（根据1978年6月20日第399号法律批准实施）的前提下，通过人工智能模型或系统，包括生成式人工智能，以文本与数据挖掘为目的，对于其合法访问的网络或数据库中的作品或其他材料进行复制和摘录，根据本法第70(3)条和第70(4)条的规定，是被允许的。

编者注：

1. 意大利法的每一次修改，不会改变原法的条文编号，如果有新增内容，仅在在某个条文中增加本条的条目。比如，第70条；第70(1)条；第70(2)条……。每一条内会按需分成段落，成为第一款、第二款。

2. 本法第70(3)条规定的情形可以理解为：研究机构和文化遗产机构出于科学研究目的而实施的文本与数据挖掘，是受法律允许的。

3. 本法第70(4)条规定的情形可以理解为：在没有权利人明确声明不允许的情况下，任何机构实施的文本与数据挖掘，是受法律允许的。

4. 意大利著作权法的本次修订，与2025年10月10日生效的第132/2025号人工智能法案第25条“人工智能辅助生成的作品的作者权保护”中的规定相对应。

附：

第70(3)条的内容是：

1. 允许研究机构和文化遗产机构出于科学研究，以文本和数据挖掘为目的，复制其合法访问的网络或数据库中的作品或其他材料，并允许公开披露以新的原创作品形式呈现的研究成果。

2. 就本法而言，文本和数据挖掘是指任何旨在分析大量文本、声音、图像、数据或数字格式元数据以生成信息（包括模式、趋势和相关性）的自

动化技术。

3. 就本法而言，文化遗产机构是指向公众开放或可供公众使用的图书馆、博物馆和档案馆，包括属于教育机构、研究机构和公共广播组织的机构，以及保护电影和音频遗产的机构和公共广播组织。

4. 就本法而言，根据法律规定，研究机构是指大学包括其图书馆、研究机构或任何其他以开展科学研究或开展研究为主要目的的实体。包含科学研究的教育活动，其类型如下：

a) 非营利性运作，或其章程规定将利润再投资于科学研究，包括以公私合作的形式；

b) 追求欧盟成员国认可的公共利益目标。

5. 那些能够对优先获取科研成果施以决定性影响的商业机构，不被视为是研究机构。

6. 根据第1款制作的作品或其他材料的副本应以适当的安全级别保存，并且仅可将其保留和用于科学研究目的，包括验证研究结果。

7. 权利人有权在不超过目的所需的范围内采取适当措施，以保障托管作品或其他材料的网络和数据库的安全性和完整性。

8. 第6款和第7款中提到的措施也可根据权利人组织、文化遗产机构和研究机构之间的协议来制定。

9. 与本条第1款、第6款和第7款相冲突的协议无效。

第70(4)条的内容是：

1. 在不影响第70(3)条规定的前提下，对于其合法访问的网络或数据库中的作品或其他材料进行复制和摘录是被允许的。

当著作权人及相关权利人以及数据库所有者未明确保留对作品和其他材料的使用权时，进行文本和数据挖掘是被允许的。

2. 根据第一款进行的复制和摘录，仅可在文本和数据挖掘所需的时间内保存副本。

3. 对于本条所述活动，在任何情况下均应保证其安全级别不低于第70(3)条所述活动所规定的安全级别。

编者注：意大利语中，“文本与数据挖掘”中的“挖掘”一词，与上述条文中，与复制行为并列的“摘录”一词，使用的原词，均为 estrazione

《杭州倡议：“跨境电商平台规则”的指导原则》发布

原载：“中国跨境电商综合试验区”微信公众号

9月26日，第四届数字贸易法治论坛在杭州正式拉开帷幕。本次论坛由浙江省委全面依法治省委员会办公室指导，浙江省商务厅、浙江省司法厅、中国国际私法研究会、浙江工商大学共同主办，联合国贸法会等国际组织及国内外法学界、相关部门代表齐聚，围绕“全球数字贸易规则的制订与中国贡献”深入交流，现场还签署多份合作备忘、发布《全球数字经贸规则年度观察报告(2025)》，成果丰硕。

论坛期间，中国国际私法研究会、浙江省法学会数字法治研究会、杭州市跨境电商协会等单位领衔发布《杭州倡议：“跨境电商平台规则”的指导原则》（以下简称《杭州倡议》），为全球跨境电商平台规则制定提供关键指引。

《杭州倡议》的核心亮点的在于，它直面当前跨境电商平台治理痛点——传统法律框架适配难、平台与用户权利义务失衡等问题，明确了8条指导原则，为各方提供清晰行为准则：

坚守底线原则：以诚实信用为基础，平台规则修改需公示征求意见，不可随意变更；恪守比例原则，避免滥用权限，保障经营者合法经营权。

保障公平透明：要求平台公平对待用户，杜绝算法歧视与差别待遇；规则需提前披露，自动化决策逻辑要解释说明，确保公开透明。

明确权责与救济：强调平台权责一致，需遵守经营地法律与公共秩序；建立便利救济机制，给予用户异议申辩机会，保障外部救济可及性。

推动共治与公益：倡导协商共治，重要规则修改需征求用户意见；同时维护公共利益，尊重各法域伦理与交易习惯，兼顾可持续发展。

《杭州倡议》不仅为跨境电商平台治理提供规

范依据，更助力全球数字贸易规则协调，为浙江打造法治中国示范区、建设全球数字贸易中心提供坚实支撑。

专家解读 | 《个人信息出境认证办法》公布 完成数据跨境制度体系新拼图

原载：“网信中国”微信公众号

作者：申卫星 清华大学法学院教授

个人信息跨境流动关系个人权益保护、国际经贸活动秩序和国家司法管辖权，“按照国家网信部门的规定经专业机构进行个人信息保护认证”是《个人信息保护法》第三十八条所规定的数据出境合法途径之一。近日，国家互联网信息办公室、国家市场监督管理总局联合公布《个人信息出境认证办法》（以下简称《办法》），进一步完善了数据跨境流动规则，打通了个人信息依法出境的新通道，便利了个人信息的有序跨境流动。

一、《办法》确立了个人信息出境的新方向

认证是由认证机构证明产品、服务、管理体系符合相关技术规范、相关技术规范的强制性要求或者标准的合格评定活动。“认证”之所以可以成为个人信息出境的合法途径之一，是因为它把跨境传输所需的适当保障具体化为经认可第三方审核+持续监督+对数据主体可执行的承诺与救济的组合。

（一）认证是国际经贸领域监管互认的主流工具

个人信息出境的主要需求是国际经济贸易往来，“认证”则是国际经贸领域常见的第三方合格评定方式。WTO框架下的《技术性贸易壁垒协定》开篇提到，认识到国际标准和合格评定体系可以通过提高生产效率和便利国际贸易的进行而在这方面作出重要贡献，因此期望鼓励制定此类国际标准和合格评定体系。为此，我们在各种进出口产品中都能看到很多认证标志，以此作为跨国监管互认的通行凭证。可以说，认证是跨国组织常见且容易被接受的安全管理方式，将认证机制引入数据跨境工作能够增强跨国组织的认可度。

作为个人信息出境监管的先行者,欧盟《通用数据保护条例》(GDPR)在第46条要求数据处理者或控制者在向第三国或国际组织进行个人信息传输时必须提供适当的安全保障措施,其中就包括GDPR第42条规定的认证机制。为此,欧洲个人信息保护监督机构在2022年发布了《关于认证作为跨境传输工具的07/2022号指南》。我国《个人信息保护法》和《网络数据安全管理条例》均将“认证”作为个人信息出境的有效途径之一,《办法》的出台将促进我国个人信息出境认证机制的完善。尽管中国和欧洲国家尚未达成关于数据充分性保护的协定,“认证”作为国际上具有互信基础的监管工具,将有助于促进各国认证结果互认。

（二）认证是持续完善公私合作治理的价值体现

认证能够通过发挥市场化的专业力量支持政府实现公共治理目标。借助市场化的第三方机构开展认证工作,不仅可以提升治理的专业水平、降低政府成本,同时也能够提升认证申请人的自由选择度和获得感。在这个过程中,监管部门通过认定、备案和管理认证机构来扩展监管能力,认证机构则依靠专业服务的信誉来获得社会各方认可。

认证是企业主动合规体系建设的体现。认证制度遵循自愿性、市场化、社会化服务原则,故而个人信息保护认证不会额外增加企业负担,而是为有志于积极从事合规体系建设的企业提供了检验合规工作效果的测试场。随着个人信息保护法律法规的贯彻落实,各类数据处理者主动深化合规体系建设的意识不断增强。符合相关条件的企业可以主动申请开展认证,从而对自身内部合规决策进行检查复核,提升企业的数据安全治理能力。

（三）专业认证机构切实履行跟踪调查职责是其功能实现的保障

《办法》建立了全链条管理机制来压实专业认证机构的责任。专业认证机构向国家网信部门办理备案手续时,应当提交对获证个人信息处理者进行的个人信息出境活动符合认证标准情况的持续监督机制。在出具认证证书之后,专业认证机构发现

获证个人信息处理者存在个人信息出境情况与认证范围不一致等情形,不再符合认证要求的,应当暂停其使用直至撤销相关认证证书,并予以公布。

《认证认可条例》通过设定严格法律责任监督认证机构勤勉尽责。如果认证机构未对其认证的产品、服务、管理体系实施有效的跟踪调查,或者发现其认证的产品、服务、管理体系不能持续符合认证要求,不及时暂停其使用或者撤销认证证书并予以公布的,可能面临5万元以上20万元以下的罚款、没收违法所得;情节严重的,可能被要求停业整顿,直至撤销批准文件。同时《办法》规定,国家市场监督管理总局和国家网信部门对个人信息出境认证活动进行监督,开展定期或者不定期的检查,对认证过程和认证结果进行抽查,对专业认证机构进行抽查和评价,从而压实专业认证机构责任。

二、个人信息出境认证制度与相关制度共同构筑我国数据跨境监管的完整体系

《办法》构建的个人信息出境认证是通用认证体系和数据出境监管中的一块新拼图,需要对《办法》和相关制度的适用关系予以说明。

（一）个人信息出境认证和通用个人信息保护认证

《个人信息保护法》所规定的“认证制度”在第三十八条和第六十二条均有体现,本次的《办法》为第三十八条落地提供了支撑。《个人信息保护法》第六十二条所规定的“支持有关机构开展个人信息保护认证服务”则在2022年就已经启动。国家互联网信息办公室、国家市场监督管理总局在2022年11月联合发布了《关于实施个人信息保护认证的公告》,决定实施个人信息保护认证,鼓励个人信息处理者通过认证方式提升个人信息保护能力,并明确了《个人信息保护认证实施规则》。该规则对个人信息保护认证的认证依据作出规定,提出“个人信息处理者应当符合GB/T 35273《信息安全技术 个人信息安全规范》的要求,对于开展跨境处理活动的个人信息处理者,还应当符合TC260-PG-20222A《个人信息跨境处理活动安全认

证规范》的要求”，并明确了不含跨境处理活动的个人信息保护认证标志、包含跨境处理活动的个人信息保护认证标志2种认证标志。

根据以上规定，个人信息出境认证是特殊类型的个人信息保护认证，除符合通用个人信息保护认证的标准要求外，还应符合关于个人信息跨境的特殊要求，并具有专门的认证标志。《办法》公布后，将在部门规章这一更高层面上明确个人信息出境认证的特殊要求，与《关于实施个人信息保护认证的公告》相衔接，共同构建个人信息出境认证的制度规范体系。同时，《办法》也规定，本办法施行前制定的关于个人信息出境认证的相关规定与本办法不一致的，按照本办法执行，明确了制度间的适用效力。

（二）个人信息出境认证和标准合同、安全评估的关系

《个人信息保护法》明确了安全评估、个人信息保护认证、标准合同等个人信息出境制度。《促进和规范数据跨境流动规定》具体划定了三种途径的适用范围，纳入“安全评估”范围内则不能适用“标准合同”“认证”，而“标准合同”和“认证”的适用范围均为：非关键信息基础设施运营者自当年1月1日起累计向境外提供10万人以上、不满100万人个人信息（不含敏感个人信息）或者不满1万人敏感个人信息，且个人信息不包括重要数据。

值得注意的是，出境认证还可以适用于境外直接处理境内自然人个人信息的企业。《办法》对此专门规定，中华人民共和国境外的个人信息处理者申请个人信息出境认证的，应当由其在境内设立的专门机构或者指定代表协助进行申请，为境外个人信息处理者提供明确的操作指引。

与一次性签约的“标准合同”不同，认证是针对特定处理活动的体系化“打包合规”。认证由经认可的认证机构持续监督，且认证结果可撤销，形成外部“合规张力”，帮助对外提供方在尽职调查与持续监控上节省人力、事务成本，而不必反复做同质化检查。此外，通过“标准合同”途径对外提供个人信息的，需要个人信息处理者在标准合同生

效之日起10个工作日内通过数据出境申报系统备案，这将促使频繁向不同接收方提供个人信息的境内主体选择“认证”方式以减轻事务压力。

党的二十届三中全会《决定》要求我们“建立高效便利安全的数据跨境流动机制”。《办法》的出台，标志着中国个人信息出境管理体系的全面建成，这将进一步优化我国法治化营商环境，促进数字经济健康可持续发展。

专家解读 | 建构个人信息出境认证制度 完善个人信息出境合规体系

原载：“网信中国”微信公众号

作者：丁晓东 中国人民大学法学院副院长、教授

《个人信息保护法》规定个人信息处理者因业务等需要，确需向境外提供个人信息，需通过安全评估、订立标准合同或进行个人信息保护认证。

此前，通过《网络数据安全管理条例》《数据出境安全评估办法》《个人信息出境标准合同办法》以及《促进和规范数据跨境流动规定》，个人信息出境制度已经具有较为完整的规定。《个人信息出境认证办法》（以下简称《办法》）则在法律法规规定基础上，衔接《关于实施个人信息保护认证的公告》的规定，细化了个人信息出境认证的具体实施规则，标志着我国三种个人信息出境制度的全面落地。

从国际视角看，认证机制在个人信息出境中发挥重要作用。早在2011年，亚太经济合作组织（APEC）推出了跨境隐私规则（CBPR）认证机制。目前，日本、新加坡等国采用了全球CBPR认证作为跨境数据传输的手段，获得CBPR认证的全球公司可以直接在这些国家之间跨境传输个人信息。获得认证的企业可在CBPR成员之间直接进行个人信息跨境传输。2016年欧盟《通用数据保护条例》则是首次将认证机制规定为数据出境的机制之一，并发布了对应指南。韩国等其他国家也引入认证机制作为个人信息出境的一种途径。但总体上，现有的个人信息出境认证主要停留在规定层面，实践中由

于具备认证资质的认证机构比较缺乏、认证细则不明确等原因，企业较少采用此种机制进行个人信息跨境传输。《办法》立足中国个人信息保护规定和监管经验，在以下方面给出了中国方案，切实推动个人信息出境认证制度落地。

其一，明确适用个人信息出境认证的条件，衔接相关规定内容。《办法》延续了《促进和规范数据跨境流动规定》相关内容，从三方面规定了适用认证机制的条件：一是主体方面，明确不适用于关键信息基础设施运营者。二是信息类型方面，明确不适用于重要数据。三是处理信息数量方面，规定自当年1月1日起累计向境外提供的个人信息需满足在10万人至100万人之间，或如果含敏感个人信息，则敏感个人信息需不满1万人。

其二，彰显个人信息保护认证的自愿性和市场化，明确个人信息保护认证的定位。《办法》体现了个人信息出境认证的四个特性。一是自愿性，认证申请系由个人信息处理者自愿向专业认证机构提出，认证不是强制性的出境手续，不会给企业带来合规负担。二是市场化，认证主体是市场化的专业认证机构而非国家监管部门，这在很大程度上有助于激发认证活动的活力，增强认证制度的适用性。三是社会化，认证活动以需求为导向，由政府、专业认证机构、申请人以及社会公众共同合作实现。四是统一性，主管部门制定个人信息出境认证相关标准、技术规范、认证规则，统一认证证书和标志，专业认证机构按照认证基本规范、个人信息保护认证规则开展个人信息出境认证活动，能够保证认证门槛的一致性，避免申请人为了提升认证通过率，选择采用标准更低的认证机构。

其三，明确专业认证机构认证重点要求，规范个人信息出境认证程序。首先，明确规定认证机构评定依据，规定专业认证机构应当按照认证基本规范、个人信息保护认证规则开展个人信息出境认证活动。其次，规范认证证书管理，明确认证证书出具、有效期及延期申请、信息报送等内容。最后，规定专业认证机构对认证证书的暂停与撤销，以及重大情况报告制度。《办法》同时规定投诉举报机

制，明确任何组织和个人发现获证个人信息处理者违反本办法规定向境外提供个人信息的，都可向专业认证机构、网信部门和有关部门投诉、举报，强调社会各主体在监督方面的作用，以期最大程度保障个人信息跨境安全。此外，《办法》还规定中国境外的个人信息处理者通过其在境内设立的专门机构或指定代表协助，同样可以申请个人信息出境认证，为个人信息保护认证申请提出了具有针对性、操作性的指引。

《办法》的出台是我国个人信息出境领域立法和监管体系的重要完善，标志着我国个人信息出境制度体系已构建完成。《办法》结合我国认证制度与个人信息保护的相关法律规定，形成了一套契合中国国情、具有较强针对性和可操作性的个人信息出境认证制度。未来，随着《办法》的实施，我国个人信息跨境治理将迈向更高水平，在保障安全的基础上促进个人信息高效跨境流动，进一步激发我国数据要素活力，为我国数字经济发展提供坚实基础。

专家解读 | 推动政务大模型部署应用 赋能电子政务智能化升级

原载：“网信中国”微信公众号

作者：王钦敏 第十二届全国政协副主席、国家电子政务专家委员会主任

人工智能大模型是新一轮科技革命和产业变革的重要驱动力量，是当前世界科技竞争的制高点。以习近平同志为核心的党中央高度重视我国新一代人工智能发展，习近平总书记在二十届中共中央政治局第二十次集体学习时强调，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。国务院印发《关于深入实施“人工智能+”行动的意见》（国发〔2025〕11号），明确要求推动人工智能与经济社会各行业各领域广泛深度融合。政务领域在我国信息化发展进程中一直发挥着先导和引领作用，在人工智能时代也要抢抓历史机遇，推动人工智能技术在政务领域深度应用，

加快赋能国家治理体系和治理能力现代化。近日，中央网信办、国家发展改革委联合印发《政务领域人工智能大模型部署应用指引》（以下简称《指引》），为人工智能技术推进电子政务创新发展、提升国家治理能力、带动全社会人工智能创新应用提供了方向与路径。

一、深刻认识政务领域部署应用人工智能大模型的重大意义

系统谋划部署应用政务大模型是顺应时代潮流、把握发展主动的战略抉择，具有多重深远意义。

（一）部署应用人工智能大模型是顺应全球科技革命大势、引领经济社会智能化发展的必然要求。人工智能大模型作为当前科技创新的前沿阵地和竞争焦点，正以前所未有的深度和广度推进经济社会运行方式的变革式发展。全球主要国家纷纷将发展人工智能上升为国家战略，加速在经济、社会等关键领域的布局应用。政务领域作为国家治理体系的核心环节，是技术赋能的重要“试验场”和“示范田”。部署应用人工智能大模型，是贯彻落实网络强国战略的关键举措，是主动拥抱智能化浪潮、赢得未来发展主动权的战略支点。通过率先在政务领域构建自主可控的大模型技术体系和应用生态，不仅能够提升政府治理能力，更能为我国在全球科技竞争中占据有利地位、引领人工智能发展方向奠定坚实基础。

（二）部署应用人工智能大模型是提升政府履职效能、推进国家治理体系和治理能力现代化的重要途径。传统的电子政务模式在应对日益复杂的社会治理需求和提供精准高效的公共服务方面面临诸多挑战。人工智能大模型凭借其在语义理解、逻辑推理、知识生成和复杂任务处理等方面的强大能力，为创新政府运行机制、破解治理难题提供了全新路径。通过构建集约化、智能化的政务大模型应用开发支撑体系，能够系统优化电子政务的业务流程，增强决策的科学性和前瞻性；能够显著提升行政执行效率，实现跨部门、跨层级的业务高效协同；能够推动公共服务模式向个性化、精准化、普惠化方向跃升，有效破解“数字鸿沟”难题。其核

心在于推动政府履职方式实现从经验判断向智能驱动、从分散管理向协同治理、从事后处置向事前预防的深刻转型，为推进国家治理体系和治理能力现代化注入强劲的智能新动能。

（三）部署应用人工智能大模型是培育新质生产力、赋能经济社会高质量发展的有力支撑。人工智能是发展新质生产力的核心引擎之一。政务大模型的深度开发应用，不仅能直接提升政府运行效率和服务水平，更能通过释放数据要素价值、优化营商环境、赋能产业创新，服务经济社会高质量发展。一方面，政务大模型作为强大的“数字治理中枢”，能够高效整合和分析海量政务数据、经济运行数据、社会民生信息和生态环境数据等，为宏观经济调控、产业政策制定、社会环境治理和市场风险预警等提供精准决策支持。另一方面，政府率先垂范应用以大模型为代表的生成式人工智能技术，能够有效带动全社会人工智能创新应用的落地推广、普及深化，激发市场活力和社会创造力，为培育壮大人工智能产业、构建现代化产业体系提供广阔的应用场景和强大的需求牵引。

二、准确把握政务领域人工智能大模型部署应用的核心原则

《指引》为政务大模型的科学部署和有效应用确立了清晰的方向和路径，其核心原则集中体现在四个方面。

一是坚持系统谋划、集约发展，统筹推进资源高效配置。政务大模型部署应用是一项系统工程，涉及算力、数据、技术、场景等多方面要素。要坚持“系统谋划”，加强顶层设计和整体规划，做好跨部门协调协作，避免“一哄而上”和重复建设。要坚持“集约发展”，依托“东数西算”和全国一体化算力网，统筹规划、协同推进智能算力基础设施布局，整合分散资源，构建集约高效的政务智能算力支撑体系，实现算力资源的弹性调配与动态扩展，显著降低整体建设和运维成本，提高资源利用效率。

二是坚持以人为本、规范应用，聚焦需求提升履职效能。部署应用政务大模型的根本目的在于提

升政府履职能力和服务水平。要坚持“以人为本”，因地制宜、结合实际需求，优先在共性、高频需求场景取得突破，构建生成式人工智能数据层、模型层、应用层等，综合考虑安全、伦理、评估优化，建立行业标准、完善监管机制，形成可复用、可推广的典型案例，并以目标导向、场景驱动方式赋能政府履职方式从经验判断向智能驱动升级。同时，要注意弥合可能出现的“智能鸿沟”，让智能技术朝着更加普惠包容的方向发展。要坚持“规范应用”，明确大模型的“辅助型”定位，编制面向工作人员的提示指南和应用规范，防范技术滥用风险和“数字形式主义”，真正提升使用效能。最终目标是通过智能化手段提升政策制定的科学性、行政执行的精准性、公共服务的普惠性，提升治理能力。

三是坚持共建共享、高效协同，构建开放协同创新生态。政务大模型技术复杂、投入大，需要汇聚多方数据资源和技术力量。最近，国务院公布《政务数据共享条例》，作为第一部促进政务数据共享流通的行政法规，对政务数据共享工作提出总体要求，并就健全管理体制、优化目录管理、细化共享使用规则、统一平台支撑、强化保障措施等提出了具体要求。为了系统构建政务领域人工智能大模型，要坚持“共建共享”，在集约建设的基础上，依托权威高效的政务数据共享协调机制，系统化推进高质量政务数据集的共建共享，推动垂直模型、知识库等资源的“一地建设、多地复用”，避免资源浪费和技术壁垒。要坚持“高效协同”，加强政企协作，探索通过购买服务、按需付费等方式引入企业技术能力，充分发挥市场主体的技术和快速迭代优势，形成“政产学研用”协同发展的良好生态。

四是坚持安全稳妥、务求实效，筑牢安全健康发展底线。政务领域涉及国家安全、公共利益和个人隐私，安全是底线。要坚持“安全稳妥”，建立政务大模型安全责任制度，完善安全管理流程和应急处置预案。特别要加强对数据处理、模型训练和应用各阶段的安全防护，防范数据泄露、模型“幻觉”、虚假信息、深度伪造、算法偏见等风险，确保技术应用在安全可控的范围内运行。要坚持“务

求实效”，通过建立监测评估体系，对应用效能进行评价，重点对应用前后进行对比分析，确保政务大模型真正赋能治理现代化。

三、务实推动政务领域人工智能大模型落地见效的关键举措

习近平总书记指出，“人工智能带来前所未有的发展机遇，也带来前所未有风险挑战。要把握人工智能发展趋势和规律，加紧制定完善相关法律法规、政策制度、应用规范、伦理准则，构建技术监测、风险预警、应急响应体系，确保人工智能安全、可靠、可控。”各级政务部门要以《指引》发布为契机，切实增强责任意识，加强多方协作，保障各项要求真正落地并取得实效。

一是构建协同治理架构。政务领域人工智能大模型的部署应用绝非单一部门或地区能够独立完成。这就要求加强党的集中统一领导，建立健全跨部门、跨层级的协调推进机制，明确主导部门与参与主体的权责划分，以此凝聚共识、系统推进。各级网信、发改、财政、工信、科技、数据管理等部门应加强沟通协作，共同研究解决政务大模型在规划设计、建设实施、场景应用、安全防护、效果评估等全生命周期中出现的重大问题。省级层面要发挥承上启下作用，结合本地区的经济社会发展实际与政务工作特点，确定实施路径与方法，同时加强对下级地区和部门的指导与规范，确保全国范围内形成上下联动、协同推进的良好格局。

二是健全制度规范体系。科学完备的制度规则是保障政务大模型健康有序发展的基石，因此亟需加快配套制度与标准规范的研制工作。一方面，应聚焦政务大模型的运行体系、数据治理、模型训练、应用部署、业务协同、安全防护、运维管理、权责界定和授权机制等关键环节，研究出台相应的管理办法、技术指南与操作规范，为政务大模型的实践应用提供清晰的行为导则与操作框架。另一方面，为保障政务大模型部署应用取得实效，有必要研究构建科学、客观、全面的动态评估评价体系。该体系应涵盖模型能力、应用效能、技术评估、成果评价等多个维度，以此对政务大模型的部署应用成效

进行系统衡量,进而及时发现存在的问题,定期分析咨询堵点难点的解决方案,总结可复制推广的经验,做好政务大模型的持续优化,最终形成“部署-应用-评估-优化”的闭环发展机制,促进先行先试应用成果的落细落地、应用推广。

三是强化基础支撑能力。政务大模型的深度应用与效能发挥,离不开坚实的基础支撑体系与强大的能力保障机制。首先,要夯实数据基础,持续推进全国一体化政务大数据体系的建设,统筹协调政务数据资源,健全统一规范的标准体系、技术规范、安全防护、运维监管等电子政务服务体系,深化政务数据资源的共建共享与开放利用,着力构建权威、准确、动态更新的高质量政务数据集和知识库,为政务大模型提供优质的数据“原料”。其次,要优化算力网络,统筹布局智能算力基础设施,提升算力资源的集约化利用水平与动态调度效率,为政务大模型的训练与推理提供强有力的算力支撑。同时,要加强人才队伍建设,一方面强化对领导干部和工作人员的数字化能力和人工智能素养培训,提升对人工智能技术的认知理解能力与实际应用水平;另一方面,注重培养既具备深厚政务业务知识又掌握人工智能技术的复合型人才,为政务大模型的长期可持续发展提供坚实的智力支持与人才储备。

四是夯实安全保障要求。大模型的深度应用正重塑国家治理的技术生态,其安全问题已超越技术议题,上升为关乎数字主权与治理现代化的战略命题。必须坚持“发展与安全并重、创新与规范协同”的辩证逻辑,将安全基因深度植入政务大模型部署应用全生命周期。在架构层面,需加快构建自主创新技术体系;在制度层面,要前瞻性设计覆盖数据权责、算法透明、应用合规的规则框架;在能力层面,亟需提升对新型风险的识别、研判与驾驭能力。唯有以系统思维平衡效率与安全、创新与稳健,方能使人工智能大模型应用真正成为提升治理效能的“加速器”而非衍生风险的“变量”,为构建安全可信的电子政务提供坚实底座。

《指引》的出台,标志着我国电子政务发展进

入智能化跃升的新阶段。各地区、各部门需深刻把握人工智能发展的历史规律和战略价值,以体系化布局、集约化建设、规范化应用为导向,统筹推进技术赋能与制度创新,为构建泛在可及、科学高效、智慧便捷、公平普惠的政务治理体系提供坚实支撑,奋力谱写中国式现代化建设的电子政务新篇章。

原创时评 | 陈增宝:互联网法院管辖新规的三大亮点

原载:“数字法治杂志”微信公众号

作者:陈增宝 杭州互联网法院院长、二级高级法官

编者按

互联网法院管辖迎来新规!最高人民法院于2025年10月11日发布《关于互联网法院案件管辖的规定》(法释〔2025〕14号),该规定将于2025年11月1日起施行。新规直面数字时代司法新需求,以4个条文对互联网法院的案件管辖范围、协议管辖规则、上诉审理机制等关键问题作出系统性完善,意义深远。

为帮助理论与实务界同仁快速把握新规精髓,本刊特别邀请全国第一家集中审理涉网案件的试点法院——杭州互联网法院的陈增宝院长,围绕新规的三大亮点进行解读,以飨读者。

互联网法院管辖新规的三大亮点

互联网法院的设立,是司法主动适应互联网发展大趋势的一项重大制度创新,被誉为“司法领域里程碑式的事件”,已成为展示我国司法改革成果的重要窗口。2025年10月11日,最高人民法院发布了《关于互联网法院案件管辖的规定》,这是自2018年出台《关于互联网法院审理案件若干问题的规定》后,时隔七年首次对案件管辖范围作出的调整完善,是全面充分落实党中央关于互联网法院功能定位和改革部署的重大举措。该管辖新规亮点纷呈,一经发布就受到社会广泛关注和一致好评,对于进一步加强互联网法院建设、促进网络空间法治化治理、以司法之力护航数字经济高质量发展具有

重要的现实意义。在此，笔者结合工作实际，围绕新规的三大亮点谈几点学习体会。

一、着力推动互联网法院的核心功能迭代升级，彰显了从“机制创新”迈向“规则输出”的发展新方位

此次互联网法院管辖新规的最大亮点在于从全面充分落实党中央关于互联网法院三方面功能定位的高度，在互联网法院前期改革已有效实现探索在线审理机制、构建完善在线诉讼规则等工作目标的基础上，推动互联网法院的核心功能从“审理机制创新”向“治网裁判规则输出”迭代升级。下一阶段的工作重心，将朝着改革部署的第三个目标——通过审理各类新型网络案件推动网络空间治理法治化这一方向迈进，彰显了互联网法院从“机制创新”迈向“规则输出”的发展新方位。

从持续深化互联网法院建设的改革路径看，此次管辖调整完善呈现一体化、体系化、动态演进的迭代发展特征。2017年8月18日，全国首家互联网法院在杭州挂牌成立。2018年7月6日，中央深改委审议通过《关于增设北京互联网法院、广州互联网法院的方案》，决定推广杭州互联网法院试点经验，对互联网法院作出了总结推广“网上纠纷网上审理”的新型审理机制、探索构建适应互联网时代需求的新型诉讼规则、通过审理各类新型网络案件推动网络空间治理法治化三方面功能定位。随着互联网法院在线审理机制等前期改革创新和探索成果先后被在线司法“三大规则”吸收推广以及2021年修正的《民事诉讼法》所确认，互联网法院的工作重心应当与经济社会发展相适应，更加契合互联网发展规律，呈现动态迭代性。在前两项功能定位目标已取得实质性成效情况下，此次管辖调整立足党中央关于互联网法院功能定位，推动互联网法院工作重心及时转向审理各类新型网络案件，从而提炼形成具有中国互联网司法特色的裁判规则体系，以“规则之治”助推网络空间治理法治化，具有重大而深远的意义。

二、全面优化互联网案件管辖格局，服务保障数字经济高质量发展的着力点更加精准、契合

此次互联网法院案件管辖新规的第二大亮点在于紧跟互联网、大数据、人工智能等前沿技术发展步伐，通过将网络数据、网络虚拟财产、网络个人信息保护和隐私、网络不正当竞争与数字经济发展关系更加紧密的新型、复杂、疑难纠纷纳入互联网法院集中管辖，为数字新业态司法保障、人民群众数字权益保护划出新“重点”，为互联网司法更精准、契合地回应新型网络社会问题，保障数字经济高质量发展提供更加有力的制度支撑。

具体来说，此次管辖调整格局优化展现以下三种形态：一是“新增”，如首次以司法解释方式确立网络数据纠纷的“专属性纠纷形态”，回应经济社会乘“数”而上的时代之变。在互联网时代，平台、数据和算法作为互联网经济社会的核心要素，成为驱动数字经济增长和网络文明进步的关键引擎。面对时代之变，此次管辖调整通过将涉数据权属认定、市场交易、权益分配、利益保护等新型纠纷整合为网络数据纠纷这一“纠纷束”，确立由互联网法院专属管辖。这在相关法律制度尚未完全形成相应配套规范体系的情况下，有利于在以后审判中统一数据权益归属、数据行为合法性判断、具体责任认定与赔偿判定等裁判思路 and 治理标准，为数据治理提供有力支撑。二是“聚合”，如将分散在各地法院的网络虚拟财产、网络个人信息保护、网络隐私权和网络不正当竞争纠纷等新型网络纠纷统一归互联网法院集中管辖，推动其持续当好依法治网规则“试验田”。随着互联网发展与社会进步，与数字经济关系密切、互联网技术特性明显的新型网络纠纷大量涌现，通过此次管辖调整，新规确定了由互联网法院聚焦审理前沿性、具有规则示范意义的网络案件，率先对最前沿、最典型的问题作出回应，持续在新兴产业领域提供明确清晰和具有权威性的规则指引，在技术与法律多维互动关系中始终保持与网络发展同频共振，为规范引领、促进保障新质生产力发展提供更精准的司法服务。三是“分流”，如在对网络购物中产生的产品责任纠纷、涉网人格权纠纷等已形成较为成熟的审判经验基础上，互联网法院审判的内容与对象也应当相应进

行更新与变化,加之全国法院“一张网”建设全面上线、在线诉讼平台功能的持续完善和在线诉讼经验的全面推广等,从制度、技术和能力层面为普通法院以“网上纠纷网上审”方式办理传统性、一般性网络案件奠定了良好的基础。因此,此次管辖调整将上述一般性、传统性、规则明确的网络案件调整至普通法院,有利于进一步扩大互联网法院前期探索的在线司法模式和实体裁判规则的功能辐射效应。

从互联网法院肩负的职责使命看,此次管辖调整是互联网法院深入践行习近平法治思想,贯彻落实全面依法治国、依法治网战略的必然要求和现实需要。当前,我国网民规模已突破11亿人,面对网络技术与数字经济迭代发展浪潮中伴随出现的诸如网络暴力、电信诈骗、利用AI技术侵害公众利益等网络社会新矛盾、新问题与新风险,作为网络空间治理的“压舱石”,互联网法院如何发挥“司法治理中枢”功能,以互联网司法之力加强网络生态治理,及时回应数字时代人民群众对司法的新需求与新期待,助力营造清朗有序、正能量充沛的网络空间,是当前迫切需要研究解决的重大课题。在审理机制创新实现实质性突破、新形态网络法律纠纷竞相涌现的双重态势下,此次调整更加注重发挥互联网法院在新时代互联网规则之治中的突破带动引领作用,进一步落细落实习近平法治思想,为进一步发挥互联网法院司法便民利民功能、公正高效便捷解纷、促进网络空间依法治理、服务保障数字经济健康发展提供了遵循、指明了方向。

三、强化新型网络案件管辖的系统集成,互联网司法审判专业化、国际化的特性更加凸显

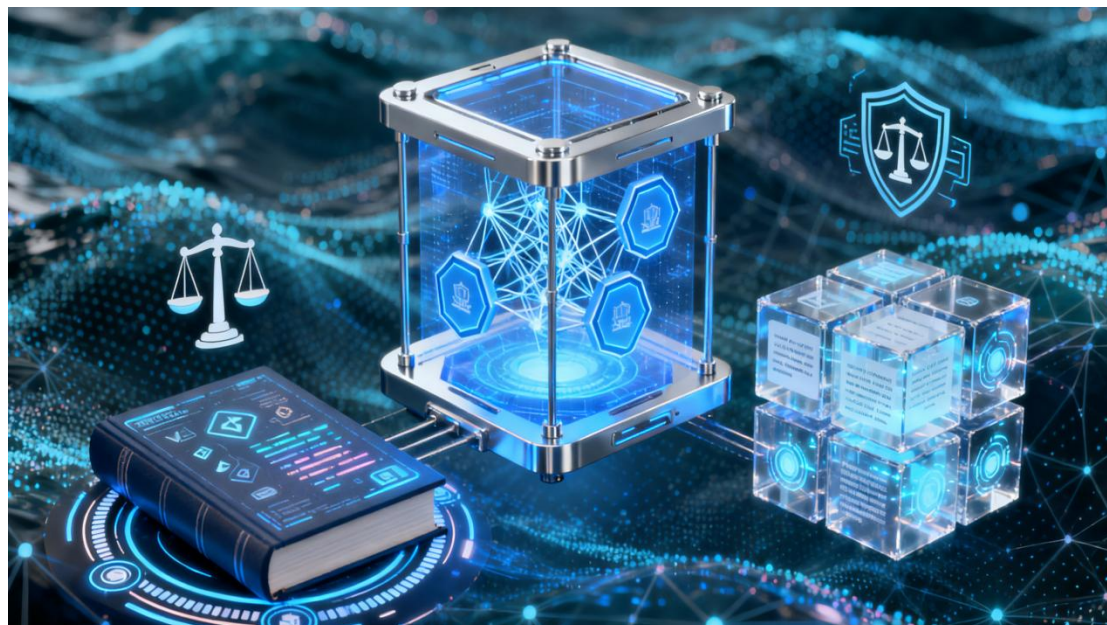
此次互联网法院管辖新规的第三大亮点在于突出司法之“治”的系统集成导向,在调整完善一些新型、前沿性民事案件管辖范围的同时,进一步对应地将相应行政案件、涉外涉港澳台案件统筹纳入互联网法院管辖范围,使互联网司法审判专业化、国际化的特性更加彰显,为复合型司法审判人才培养、司法系统化治理划出“新重点”,有利于形成更加体系化的专业治理框架。

从司法治理的特性需求来看,技术创新推进了数字经济发展走向专业领域细分与全球化,与之相嵌套的互联网司法也必须以专业化、国际化为目标建构体系化治理模式。此次涉网行政案件管辖调整彰显了专业化的“管网”“治网”功能,将涉及网络数据监管、网络个人信息保护监管、网络不正当竞争监管以及网络交易管理、网络信息服务管理等行政案件交由互联网法院集中管辖,有利于支持监督网络监管行政执法。同时,在当前数字化、全球化加剧的背景下,此次管辖调整将互联网法院管辖范围拓展至涉外涉港澳台网络数据纠纷、不正当竞争纠纷等,让互联网法院可以通过集中审理具有首案效应的案件,进一步探索数据权益、平台治理、网络市场竞争秩序、我国公民和法人海外利益保护等领域具有先导示范意义的裁判规则,为深度参与国际网络空间治理、推动加强涉外法治建设、助力加快构建网络空间命运共同体贡献更多司法力量。

从司法人才队伍培养来看,此次管辖调整对互联网司法审判人员的能力素养提出了更高要求。通过将前沿、重点网络领域的新型案件成体系地纳入互联网法院管辖范围,并将传统、一般性网络案件剥离交付普通法院审理,改变了司法资源的配置与供给格局。凭借案例资源、审判资源进一步聚集的新优势,互联网法院在前期已形成一批“既懂法律又懂网络”的人才储备基础上,将持续在以裁判树规则、定标尺、促治理的过程中加强数字法治人才培养,通过人案互促和深化司法智库战略合作,着力打造一批政治素质过硬、业务能力精湛、团队效应突出、标杆作用显著的互联网司法“专家型”领军人才和“实务型”业务人才,重点打造一支既有大局观念又有国际视野,既通晓国内法律又熟悉国际法规则,善于处理涉外法律事务,能够走向世界、适应战略性数字发展需要的涉外法官队伍,努力推动互联网法院成为数字法治人才“蓄水池”和数字法治研究实践基地。

(技术编辑:何芮)

研究动态



基础理论

1. 不确定性下数字法学的知识范式转型与路径展开（卫子豪）

来源：《华东政法大学学报》2025年第5期

不确定性与确定性是具有哲学方法论的范畴，人类行为机制的最深层依据，是从事物的确定性认知不确定性的最大可能。数字时代加剧了不确定性，确定性的认知桥梁难以建构，隐含于法学研究的认知和数字世界的认知难以沟通，形成对数字法学研究的限制，引发范式危机。这源于语言本身的模糊性、数字事实的多元性和价值判断的主观性。据此，应转变在不确定性中求确定性的知识范式，以不确定性为背景，以法学和技术在哲学观上的互变共进为基础，充分拓补法学的推理能力和价值判断的指引作用，在不确定性中尝试建构对事实的理解和对知识的获取。同时在数字法学实证研究中应遵循前述不确定性知识范式，在数字技术运用与数字法学知识库构建过程中，不坚持不变性和唯一正确性，而是以持续演化、开放性的知识模块，实现法律知识的系统化整理，以理论关怀和技术支撑反哺数字法学基础理论研究，实现数字法学基础理论

与法律智能化的协同发展。

2. 从占有到控制：数字财产侵害行为的竞合处断规则（陆一敏）

来源：《华东政法大学学报》2025年第5期

网络账号作为数字化身份，因缺乏独立经济价值，应被排除于财物范畴。以虚拟币、游戏装备等为典型的虚拟物品和以比特币、NFT等为代表的数字资产，因其排他的控制性和可流通性，应当被纳入刑法“财物”的语义射程。数字财产侵害行为在传统占有转移理论模式下难以适配盗窃罪中“打破占有—建立占有”的物理性行为构造。借由“控制”这一在解释学上具有更广泛应用场景的概念，推动由“占有”到“控制”的范式转化，突破有体物对象的桎梏，将“原控制管理权限失效—新控制管理权限建立”支配力关系的转移，实质性地嵌入取得型财产犯罪的构成要件框架。数据代码与数字财产呈“一体两面”的二元结构，单一行为可能同时侵害数据安全与财产法益，应以想象竞合处断。

3. 数字金融领域非法吸收公众存款罪的适用限缩（马永强）

来源：《环球法律评论》2025年第5期

当下非法集资犯罪的规制场域已从传统领域延伸至以加密货币和区块链金融为代表的数字金融领域，这加剧了非法吸收公众存款罪固有的扩张风险，并对罪刑法定原则构成严峻挑战。为避免刑法不当介入数字金融领域导致入罪范围泛化，亟需对非法吸收公众存款罪进行精准的适用限缩。为此，应重新审视本罪的保护法益，立足于本罪适用场景的结构性变迁，充分认识传统单一法益说的局限，明确本罪旨在保护宏观金融稳定和微观公众权益的双重法益内涵，并在此基础上结合具体教义学争议对本罪构成要件进行限缩解释。对于非法集币行为，应坚守“存款”的文义内涵与法律关系本质，审慎认定其入罪范围；对于“非法性”要件，应厘清行政违法与刑事违法的界限，以是否实质侵害双重法益作为出入罪的关键，严格界定其边界；判断“扰乱金融秩序”要件时，则应考虑被害人自我答责的因素，对高风险金融投机活动中的损害后果进行合理归责。

4. 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。

5. 消费者安宁权：数字时代私人生活安宁权的新表达（边琪）

来源：《中国法律评论》2025年第5期

《消费者权益保护法》第29条第3款禁止经营者未经同意向消费者发送商业广告。该条规定的是消费者安宁权，而非个人信息权益。它是私人生活安宁权在消费者保护领域的新表达。注意力隐私权理论揭示了数字时代的商业广告通过注意力侵入、控制和剥削的方式侵害消费者安宁利益。对此，消费者安宁权是重要的抵御工具，它保护虚拟空间中消费者的生活安宁和精神安宁。借鉴相邻关系理论，消费者安宁权侵害行为的判定分为行为合理性及行为损害实质性效果两方面。行为合理性的判定围绕“消费者同意”这一核心要素而展开。行为损害实质性效果体现为是否超过消费者的容忍义务。欠缺消费者同意而发送商业广告的行为是不合理行为。但对于性质轻微且基本未带来不利益的不合理行为，消费者有容忍义务，不构成消费者安宁权的侵害行为。消费者安宁权的保护不能仅依赖纯粹的私法救济路径，而需要增加事前预防与国家保护路径。

6. 从限制到调控：人工智能背景下未成年人保护模式的转型与重塑（刘晓春）

来源：《财经法学》2025年第5期

未成年人使用人工智能应用的行为具有高度交互性、私密性、工具性和枢纽性等特点，对未成年人保护提出了新挑战。我国未成年人模式已经取得的实践成果，在实施效果上依然面临质疑。我国未成年人模式的运行机理包括身份识别、权限设置和家长赋能三个方面，具有明显的限制型特征，面临输入、处理、输出三个层次的信息困境，导致保护要求难以落地。人工智能背景下，需要将未成年人模式进行改造，转向调控型保护模式，即基于人工智能技术带来的信息处理能力和成本降低，破除信息困境，从身份识别、风险防范、发展促进、家长赋能四个层次展开系统性重塑，并建构相应的制度保障，实现个性化、精准化、赋能化的未成年人保护。

7. 数据来源者权利及其实现——基于数据共生的视角（王年）

来源：《财经法学》2025年第5期

数据来源者权作为“数据二十条”规定的一项权利，其是否以及如何转化为一项实定法上的权利，必须具有充分的理由。隐私权、个人信息权益、知识产权、商业秘密等法定在先权益虽在一定程度上解决了数据处理者财产权来源合法的问题，但也忽略了数据本质上是来源者与处理者共同生成的事实。基于数据来源者与数据处理者的共生关系，应在承认数据处理者对汇聚多个来源的数据集合享有一般性财产权的基础上，承认和保障数据来源者对数据的公平使用权。数据来源者是对数据生成具有较大贡献而不实际持有数据的主体。数据来源者权包括知情权、访问权、转移权、更正权和删除权等权能。应采取“相关利益”标准来确定数据来源者权的客体范围，并通过为数据来源者、数据处理者和第三方数据使用者配置相应的权利义务规则来实现数据来源者利益的法律保护。

8. 网络用户参与数据要素收益分配的路径改进：基于用户角色定位的理论构想与现实抉择（于楚涵）

来源：《财经法学》2025年第5期

当前互联网行业通行的“网络用户作为消费者”的传统模式，借助“免费数据换免费服务”的交易，将网络用户创造的数据转化为平台独占的资本，剥夺了用户参与数据要素收益分配的机会，导致收益分配不公，有碍于数字经济的高质量发展。针对这一问题，理论研究层面发展出“网络用户作为劳动者”与“网络用户作为投资者”两种改进路径。前者将数据视作网络用户的劳动成果，根据用户的劳动贡献向其支付报酬，对分配不公进行了有限调整；后者将数据视作网络用户的资本，根据用户的资本贡献向其分享收益，对分配不公进行了充分调整。二者在实现数据要素收益公平分配、促进数字经济高质量发展以及实施的难易程度方面各有利弊，应在实践中综合运用以避短扬长。综合运用二者的基本思路是：网络用户普遍有权作为劳

动者，以获得劳动所得的方式参与数据要素收益分配；只有作出显著数据贡献的网络用户有权作为投资者，以分享投资回报的方式参与分配。

9. 算法个性化定价的反垄断法立场与分析路径（焦海涛）

来源：《财经法学》2025年第5期

作为数字市场中的一种特殊定价方式，算法个性化定价常被具体化为大数据“杀熟”而受到质疑，并被倾向于认定为反垄断法中的价格歧视。算法个性化定价其实具有非常模糊的经济效果：商家可能因压榨了更多的消费者剩余而获得更高利润，但支付意愿低的消费者也可能因此以可承受的价格购买到想要的产品。算法个性化定价由此在商家和消费者之间以及消费者内部产生了分配效应，并同时可能有助于产出扩张。如果反垄断法以社会总福利(或消费者福利)为最终目标，就不宜一揽子禁止算法个性化定价，而应结合市场结构、整体经济效果逐案分析。算法个性化定价在两种情况下可能构成垄断行为：如果对支付意愿高的消费者产生了明显的剥削效应，可能构成价格歧视或超高定价；如果对上下游市场产生了明显的排他效果，可能构成排他性滥用。反垄断法只能解决部分算法个性化定价问题，全面保护消费者利益还需要消费者权益保护规则的适用，尤其是其中的公平交易权和知情权，并加强对商家算法使用行为的监管，确保算法个性化定价的透明度和消费者的退出权。

10. 人工智能时代应对虚假信息的欧洲宪法方式（Giovanni De Gregorio, Oreste Pollicino）

来源：German Law Journal, Vol.26, Issue 3 (2025)

人工智能系统的扩展加剧了虚假和捏造内容等虚假信息的传播，引起了全球范围内政策制定者的关注。然而，解决虚假信息会导致宪政民主国家对作为民主社会活生生的核心的言论自由的范围产生质疑。如果一方面，这一宪法权利被认为是公共当局干预限制虚假信息传播的障碍，但另一方面，捏造内容和纵技术（包括深度伪造）的传播越

来越受到自由主义观点的质疑。在线平台的作用进一步丰富了这一宪法挑战，这些平台通过调解其在线空间中的言论，是马赛克的重要图块，描绘了潜在的监管策略和公共执法对付虚假信息的局限性。在这个框架内，这项工作认为，欧洲应对虚假信息的宪法方法在全球范围内定义了一种独特的模式。欧盟制定了一项结合程序保障、风险监管和共同监管的战略，正如《数字服务法》、《加强虚假信息行为守则》和《人工智能法》等举措所证明的那样。欧洲方法位于自由主义和非自由主义模式之间，提出了一种基于风险缓解以及公共和私人行为者之间合作来解决虚假信息的替代宪法愿景。

11. 数字市场的烟幕弹——选择操控与同意的错觉

(Naimy Paul)

来源：Journal of European Competition Law & Practice, Vol.16, Issue 3 (2025)

“免费且永远免费”曾是 Facebook.com 近二十年的标语。Meta 决定在其运营期间推出不止一项而是两项订阅服务（并随之更改标语），这表明多年来它并未免于竞争或反垄断监管的影响。然而从本质上说，当用户每一次点击、滚动和触摸都在用数据买单时，最初的承诺便早已名不副实。对这种广告支持商业模式的认知，以及“科技巨头”日益膨胀的权力，凸显了反垄断立法在遏制有害影响中的重要性。多年来针对这些企业的反垄断诉讼层出不穷，成效却参差不齐。问题部分在于：新商业模式带来新危害，却用旧工具和理论应对。近期立法有所助益，但传统竞争法仍需更深入理解数字市场运作机制，方能实现全面有效的分析——首要之务便是重视数据选择权的同意机制。

12. 通过引入独家折扣来遵守《数字市场法案》？

(Niels Frank, Mitja Kleczka)

来源：Journal of European Competition Law & Practice, Vol.16, Issue 5 (2025)

2024年3月7日，《数字市场法案》（DMA）的守门人义务正式生效。该法规旨在确保数字领域

市场的公平竞争。自此，被欧盟委员会指定为守门人的企业必须确保其核心平台服务活动符合 DMA 规定。作为回应，苹果公司发布报告阐述其允许第三方应用商店进入 iOS 环境的方案，以此符合《数字市场法案》第6条第4款规定。针对选择通过第三方应用商店分发应用的开发者，苹果推出了全新的费用与佣金体系。经审查这些变更后，欧盟委员会以苹果新收费结构“可能违背其履行《数字市场法案》第6条第4款义务的初衷”为由，启动了进一步调查。2024年6月24日，欧盟正式启动违规程序，指出苹果针对第三方应用开发者及应用商店的新合同费用与配套规则“未能确保有效遵守《数字市场法案》”。

13. 《数字市场法案》时代下的并购补救措施：该法案对欧盟并购管制条例（EUMR）设计承诺条款的影响

(Klein Lilian)

来源：European Competition Journal, Vol.21, Issue 2 (2025)

《数字市场法案》（DMA）与《欧盟并购控制条例》（EUMR）是互补性工具，可同时适用于涉及守门人的并购交易。然而，DMA 与 EUMR 之间潜在的冲突尚未得到充分探讨。本文旨在阐明 EUMR 与 DMA 在设计守门人收购承诺方案时的相互作用。研究认为，DMA 可能影响 EUMR 的补救措施设计，因为这两项工具在应对守门人行为的有害影响时可能采取相似策略。据此，新的《数字市场法案》义务将通过两种途径影响《欧盟并购条例》未来的承诺设计：首先在损害理论阶段，因《数字市场法案》的威慑效应；其次在补救措施设计阶段，基于比例原则。因此在守门人收购场景下，《数字市场法案》可能限制《欧盟并购条例》设计并购承诺的裁量权。

14. 分离还是整合？竞争评估中的数据保护：系统性文献综述

(Vandendriessche Robin, Buts Caroline)

来源：European Competition Journal, Vol.21, Issue 1 (2025)

个人数据在数据保护法与竞争法中的双重角色催生了复杂的互动关系，引发了是否应将数据保护考量纳入竞争评估的争论。本文运用 PRISMA 方法论，系统梳理了 2014 至 2023 年间的相关文献，探讨了竞争评估中此类考量方式的近期演变。我们认为，随着竞争法执法实践的持续演进，相关讨论已逐渐聚焦于两种更具整合性的方法。尽管法律研究（通常回避最激进的整合方法）仍主导着现有文献，但跨学科研究自 2014 年起呈现增长态势。最终，我们归纳出九项可能影响竞争评估的数据保护考量因素，并将其划分为五大类别。我们最终主张采取务实方法：在以数据保护法为规范基准的指导下，将数据保护考量纳入评估范畴。

15. 基础模型引发的竞争担忧：科技巨头的新盛宴？（Mitra Shourya）

来源：European Competition Journal, Vol.21, Issue 2 (2025)

本文探讨生成式人工智能与竞争法的交汇点，重点关注基础模型(FMs)和大型语言模型(LLMs)。研究分析了行业动态，并识别出关键竞争问题，如市场准入壁垒、捆绑销售、市场支配地位滥用及并购行为。文章强调供应链的重要性，考察了基础模型如何融入搜索软件、聊天机器人和生产力工具，特别指出计算能力与数据收集等准入壁垒。研究指出基础模型可能需要新的市场界定方法，或需为数据建立独立的相关市场。论文还讨论了捆绑销售与优势地位滥用相关案例，强调因传统搜索引擎与 AI 聊天机器人界限模糊，捆绑行为的举证难度较大。文章阐释了收购案竞争评估可能面临的变革——数据作为行业高度灵活的商品正重塑评估标准。结论部分呼吁加强对该行业的监管审查力度。

个人信息保护

1. 个人信息保护民事公益诉讼损害赔偿疑难问题研究（肖梵）

来源：《国家检察官学院学报》2025 年第 4 期

个人信息保护民事公益诉讼损害赔偿在理论和实践中，存在赔偿范围及性质不明确、损害赔偿认定缺乏量化标准、惩罚性赔偿适用的合法性争议等问题。数字时代风险社会中个人信息公益诉讼损害赔偿的性质，既非简单的公法问题，也非单纯的私法问题，而是处于公法和私法的边界之上。个人信息公益诉讼损害赔偿的价值追求，是修复受损公共利益的状况、预防未来侵害公共利益风险的发生，其正当性是将赔偿用于受损公共秩序、环境的修复、治理，某种意义上如何修复、预防，决定了如何赔偿。未来应探索建立以修复、预防为核心的“双层”个人信息公益诉讼损害赔偿认定架构，突出强调治理的理念，以不断完善损害赔偿的适用规则；同步探索构建数字生态环境损害鉴定评估机制，培育专业的鉴定机构对数字领域发生的个人信息、虚拟财产、数据和信息安全等的私益和公益损害，进行修复方案制定和费用评估。

2. 信息流通背景下个人信息爬取行为“过罪化”的实务现状与理论重塑（徐光华）

来源：《政治与法律》2025 年第 7 期

个人信息爬取行为作为新兴的自动化个人信息收集方式，在实务中存在“过罪化”倾向：经授权的个人信息爬取、仅违反企业规定的爬取行为均被认定为侵犯公民个人信息罪，对公开信息的爬取行为也丧失出罪空间。其根源在于忽视信息流通中的主体意愿、社会价值及刑法保障流通的功能。需结合信息流通的时代背景，对个人信息爬取行为进行分类评价，适度限制入罪范围。对已授权个人信息的爬取，因主体的知情同意不应认定为犯罪。对公开的个人信息的爬取，应当综合考量个人利益和社会利益来判断其正当性。对于涉私密性的个人信息的爬取，应当判断是否严重侵害个人利益，不能仅以违反企业内部规定为由认定爬取行为成立侵犯公民个人信息罪。应适度限制对个人信息爬取行为的入罪范围，在信息流通与个人信息保护之间寻求平衡，防止刑法过度打击导致信息产业萎缩，妨

害我国数字经济发展。

3. 论民事诉讼中的个人信息处理活动（程啸）

来源：《政治与法律》2025年第9期

民事诉讼中的个人信息处理活动广泛存在，涉及起诉、答辩、举证、质证、执行、裁判文书公开等多个环节。民事诉讼中的个人信息处理活动应当适用《个人信息保护法》，自然人之间的个人信息处理活动除外。当事人为完成举证责任提供个人信息的法律依据包括为履行合同所必需、实施人力资源管理所必需、处理合法公开的个人信息等情形。当事人因履行举证责任而处理个人信息不属于《个人信息保护法》第13条第1款第3项规定的为履行法定义务所必需的情形。法院为调查收集证据而处理个人信息，应当遵循合法、必要和目的限制等原则，并采取对个人权益影响最小的方式。外国法院在民事诉讼证据开示程序中要求当事人跨境提供个人信息的，不属于我国《个人信息保护法》第13条第1款第3项中的法定义务。外国法院强制性取得我国境内存储的个人信息，原则上必须通过国际司法协助程序，除非当事人在符合我国法律关于个人信息出境规定的情形下自愿提供。非法处理个人信息的行为严重侵害个人合法权益、违反法律禁止性规定或者严重违背公序良俗的，由此形成或者获取的证据不能作为民事诉讼中认定事实的基础。

4. 再论我国个人信息保护的立法模式选择（姚岳绒）

来源：《政治与法律》2025年第9期

在形式上，《中华人民共和国个人信息保护法》采用了将国家机关和私人组织一体化规范的立法模式。由于国家机关的个人信息处理条款具有宣示性、象征性，该法在实质上仍是一部私法规范，属于单一立法而不是统一立法。个人信息保护立法应以宪法上的基本权利为基础，而不是止步于民事权利。将个人信息保护立法的权利基础确定为基本权利的核心在于，通过基本权利功能和国家保护义务，构建出一幅个人信息保护立法的全景图。民法

所保护的是权利免受来自私权主体的侵害，而宪法重在保护权利免受来自公权力侵权的风险。民事权利与基本权利不是独立、并行的关系，不是非此即彼，而是由基本权利决定民事权利。个人信息保护立法的最佳状态是权利与权力实现动态平衡，公民权利有足以制约国家权力的能力，而国家权力又足以防止公民权利的滥用。个人信息保护宜采用分散立法模式，而不是强行合体。在《中华人民共和国个人信息保护法》实施后，应当重点调整公权力领域的个人信息处理行为，国家机关处理个人信息有单独立法的必要性与紧迫性，建议优先构建行政领域内的个人信息保护法制。

5. 个人信息竞争法保护的理据探析与条件限定（吴太轩）

来源：《政治与法律》2025年第9期

当前，对于“竞争法应否保护个人信息”这一议题存在“否定论”与“肯定论”两派观点。“否定论”中的“致害原因否认说”“立法目的相悖说”“保护效果有限说”均存在片面性，无法做到逻辑自洽。“肯定论”虽从“竞争法保护具有必要性”“竞争法保护具有可行性”两方面进行了分析，但其论证的宽度与深度有待加强。个人信息的竞争法保护除了具有规范依据（即保护的必要性与可行性），还具有扎实的理论支撑——法益综合化系列学说，亦契合侵害个人信息与竞争违法行为日益竞合的现实场景，且域内外既有的竞争法保护个人信息的努力实践为其积累了参考经验。竞争法理应介入个人信息的保护进程。为了预防竞争法的过度干预，应对竞争法介入个人信息保护的条件进行限定。基于竞争法的法本位与法目的，可将其限定为“侵害个人信息权益+扰乱市场竞争秩序”双重条件。

6. 数字经济时代个人信息犯罪的穿透式治理（赵炳昊）

来源：《中国刑事法杂志》2025年第4期

数字经济时代个人信息在注重保护的基础上

更要强调流通利用。然而,传统个人信息犯罪治理不当偏重强保护而过度限制个人信息处理,未能合理评价技术因素的复杂影响而导致罪与非罪定性难题,混淆预备与实行行为的界限而令个人信息犯罪处罚泛化。通过引入“实质重于形式”的穿透式治理,可实现个人信息犯罪形式治理向实质治理的转型,但穿透式治理的运用应仅限于出罪。个人信息犯罪的穿透式治理,应穿透权利表象,通过“基于风险的方法”将未对信息主体的身体或财产等实体法益产生抽象危险的信息处理行为予以出罪;应穿透技术表象,通过妥当性审查重点排除利用技术微调即可实现治理目标的行为入罪,通过必要性审查重点排除不具有显著外溢性和公共性危险的技术处理行为入罪,通过均衡性审查重点排除部分轻微违法处理或泄露个人信息等技术处理行为入罪;应穿透行为表象,在侵犯公民个人信息罪构成要件中增设不成文的主观违法要素,以行为人是否存在后续犯罪的意图限缩个人信息犯罪的处罚范围。

数据确权与流通

1. 健康医疗数据应用的法律规制（武翠丹）

来源：《华东政法大学学报》2025年第5期

健康医疗数据作为数字时代的关键生产要素，其合规应用已成为推进健康中国战略实施的核心路径。当前关于医疗数据收集与应用的规范机制在权属界定、流通规则及责任分配等维度呈现结构性缺失，成为制约医疗数据要素价值释放的制度瓶颈。健康医疗数据的法律规制应遵循三维路径：其一，在规范基础层面，确立“人格权益—财产权益—公共利益”的权益衡平框架，明确医疗机构基于数据加工行为的有限财产权主体地位；其二，在权利配置层面，以三权分置的监管范式为治理基础，通过分级授权机制实现数据要素的差异化配置；其三，在实施机制层面，创新动态合规治理工具，包括建立医联体数据共享机制、设立非个人数据收益

25 / 69

分配阈值制度、引入数据调取与使用的比例原则审查机制，以期推动医疗数据治理范式从防御型合规向赋能型治理转型，为健康医疗数据要素市场化提供制度供给。

2. 个人分享数据红利的理论基础与行权机制（林涸民）

来源：《政治与法律》2025 年第 8 期

个人直接参与数据交易并分享数据红利，是打造可信数据空间，实现共同富裕的重要路径。个人数据具有财产属性，且是个人“线上劳动”的结果，数据财产权理论与数字劳动理论共同构成个人分享数据红利的理论基础。当前数字经济以数据运营者为中心，数字经济治理应改变这一商业结构，建立“个人数据账户+数据交易所”的可信交易模式。《个人信息保护法》规定的个人数据可携权是建立个人数据账户的行权基础。个人数据账户制度既可使个人直接从数据交易中获益，又可促进可信数据流通，增加市场上的优质数据供给。个人数据账户制度的顺利实施有赖新型数据基础设施的支持。数据交易所应成为个人数据账户支持机构，注重强化对个人数据账户的安全管理，整合互联网与地方性数据资源，并提升个人数据账号质量。当发生数据泄露、数据传输故障等纠纷时，数据交易所应根据过错程度对个人数据账户持有者承担违约或侵权责任。

3. 企业数据资产融资担保的法律实现（罗亚文）

来源：《政治与法律》2025年第9期

企业数据资产融资担保应以规范意义上的担保物权构造为建构路径，以区别于企业信用融资、应收账款质押融资、数据存储设备担保等实践形式。数据资产为担保标的，数据持有权及数据使用权为适格担保客体，数据加工使用权、数据产品经营权应统合为数据财产权规范构造中的“数据使用权”。与抵押担保相比，权利质权形式更符合现行担保体例、数据自身特征及实践形式要求。数据质押担保不宜采“准占有”或权利凭证交付的设

权路径,应将登记作为数据权利质权的公示方式。在数据资产移转、数据许可使用等数据交易采登记对抗主义的情况下,数据质押担保采登记生效规则更为适宜。基于对数据资产之特殊性以及担保企业经营模式等因素的考量,质押担保期间出质人可转让、对外许可使用数据资产,但应设置衡平质权人交易风险的适用规则,以兼顾企业的融资效果及质权人的担保风险。可通过强制管理、“预处置”的缔约方式保障质权人的债权实现。数据资产入表的认定标准可作为质押标的识别的借鉴参考。应关注数据自身特征对估值定价的影响,基于不同参考因素及动态考量思路选择适宜的估值定价方法。

4. 论数据财产权与在先权利的关系(钱子瑜)

来源:《中外法学》2025年第4期

数据财产权与在先权利是两个独立的权利。数据财产权来源于合法的数据收集行为,而非来源于在先权利人的授权,相关在先权利会对数据财产权的权能造成限制,但是对数据不具有直接的支配关系。出于权利的价值位阶与保护次序等因素,在先权利应当获得优先保护,在先权利系数据财产权的限制而非无效事由,数据权利人对于在先权利的实现负有容忍义务。数据权利人违反容忍义务,将导致数据财产权的过度扩张,在先权利无法实现。在先权利亦有边界,应当根据信息的公开程度、数据的匿名化程度以及在先权利的行使是否正当等因素综合判断。由于人格权较财产权有更高的价值位阶,在先人格权利人可以任意撤回“同意”,尽管数据权利人的信赖利益会因允诺的违背而受损,但是在先人格权利人无需支付补偿。前进损失可通过市场机制调节,未来数据产业商业模式将由“免费”转向“付费”,数据财产权与在先权利将实现再平衡。

5. 专家知识数据化与文科智能(邱泽奇,李昊林,张平文,乔天宇)

来源:《中国社会科学》2025年第7期

数智技术正在触发知识生产范式变革,文科智

能是人文社会科学发展的未来方向。高质量数据是发展文科智能的前提。把人类积累的知识数据化是获取高质量数据的路径之一。其中,专家知识是一类不断迭代、浓缩、有着专门结构的知识,也是更具创新价值的知识。法学领域的专家知识在文科中具有典型性。以法学领域为例,专家知识的数据化可以采取“三步走”路径,即框定专家知识范围、重构专家知识结构以及设计专家知识转化技术路径等。“三步走”的专家知识数据化路径具有可推广性,既有助于夯实文科智能的数据基础,也可为推动文科知识生产范式变革迈出坚实的一步。

6. 从本权到衍生权:企业数据产品的刑法保护(马寅翔)

来源:《中国刑事法杂志》2025年第4期

关于企业数据产品的刑法保护,学界虽普遍认同可以提供针对数据管理秩序的基础保护,但在能否进一步提供财产性保护的问题上,则存在以刑法谦抑论与法益保护论为核心的观点对峙。基于维护市场公平竞争、保障企业经济权益的刑法功能定位,提供财产性保护殊为必要。在现行刑法框架下,较为务实的选择是在坚持违法一元论的基础上,采用间接财产保护模式。该模式不以前置法上的数据确权为前提,亦不以行为定型替代法益论证,而是回归企业数据产品的价值生成与交易机制,将关注重心从数据本权转向为前置法所承认的数据产品衍生权利,主要包括知识产权、合同债权等财产性权利。通过将此类衍生权利作为保护法益,为企业数据产品提供间接财产保护成为可能。根据保护法益的不同类型,可将侵害企业数据产品的犯罪划分为两类:一类是针对数据产品衍生权利的侵犯财产权益犯罪,具体包括侵犯知识产权罪与侵犯财产罪;另一类是针对数据产品本身所涉管理秩序的扰乱公共秩序罪。根据犯罪竞合理论,应坚持以侵犯知识产权罪为优先、以侵犯财产罪为重点、以扰乱公共秩序罪为兜底的罪名适用顺序,为企业数据产品提供分级分类的刑法保护。

7. 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。

8. 数据作为知识产权客体的思辨与模式选择（曹新明，范晔）

来源：《知识产权》2025年第6期

我国正积极探索构建数据知识产权保护规则,但对其所涉若干基础性问题尚未达成共识,亟待回应。数据具备知识产权客体的智力成果属性,将其纳入知识产权法体系具有可行性。数据知识产权登记能够为数据知识产权保护规则提供实践基础,但在数据财产赋权立法缺位的情况下,应明确其制度功能为“证权”而非“确权”,进而澄清登记客体、模式和效力等存有争议的问题。数据之上存在多元利益格局,类型化思维是厘清数据知识产权保护问题的关键。由于数据与既有知识产权客体之间存在重叠,可将数据划分为知识产权类数据和其他数据,依托现有知识产权法和竞争法制度,为类型化的数据分别匹配相应的、更适宜的产权制度。

9. 公开数据适用商业秘密保护的可能及实现（陈兵，林逸玲）

来源：《知识产权》2025年第6期

当前,法律层面数据权益归属制度的缺失,加之《反不正当竞争法》“一般条款”自由裁量空间大,实践中仍面临较大的不明确定性,导致数据持有者公开数据的意愿低,很大程度上阻碍了数据开放流通。有必要寻求以商业秘密保护公开数据,规制和督促“非法获取”行为向“合法使用”行为转变,提升数据持有者公开数据的意愿。实践中公开数据的“公开性”表征使其作为商业秘密的“秘密性”要求之认定面临挑战,公开数据的商业秘密权利主体也有待商榷,将算法作为商业秘密保护的理论与具体制度仍不充足,使得公开数据的商业秘密保护存在困境。基于此,有必要适当扩大公开数据秘密性的认定基准,尝试建立有限共享秘密性,鼓励数据持有人依法公开其原始数据、数据集、数据产品以换取数据使用赋能与增值的效用,并明确其作为商业秘密权利主体享有或(和)共享其数据权益。

10. 知识产权视域下的数据财产权建构（赵军）

来源：《知识产权》2025年第6期

我国数据财产权在理论上和实践上依然存在分歧。鉴于数据的非竞争性、可复制性和非排他性,知识产权的理论洞见可运用于数据财产权之中。借鉴界定专利权的方法,“数据采集或加工方法”成为划定数据财产权客体的依据;汲取著作权的经验,建立以自动取得为主,数据登记为辅的赋权原则,形成数据提取权、再利用权以及禁止规避技术措施权三种基本权利。知识产权理论还为数据财产权人私人利益与社会公共利益的平衡提供了丰富工具,避免权利冲突、强制许可、合理使用和保护期限等构成数据财产权限制的组成部分。

人工智能

1. 生成式与中国式的互塑：人工智能赋能法治现代化的路径探析（张善根）

来源：《政治与法律》2025年第8期

生成式人工智能兼具技术属性和社会属性。技术属性是所有科学技术的一般属性，本质上是一种人造物。社会属性则是生成式人工智能的特有属性，其本质是在人类之外建构出新型的知识生产者。DeepSeek 作为我国现象级的生成式人工智能产品，意味着通过生成式内生中国式成为可能。生成式人工智能可以在两个层面赋能中国式法治现代化：基于技术的赋能和基于知识生产者的赋能。技术赋能是人工智能赋能法治的基本模式，主要应用于人民法院、人民检察院等政法职能部门，建构的是“AI+法治”赋能机制。其可以在世界范围内普遍适用，一般不受法律文化、观念、价值的影响。基于知识生产者的赋能革新了赋能模式，建构的是“AI 即法治”的赋能机制。其不仅可以推动 AI 赋能法治的路径从专业化走向大众化，而且可以通过 AI 内生法治知识体系，直接参与并助推中国式法治现代化的建设。若要生成式人工智能发挥双层赋能功能，前提是在法治规训和法治对齐两个向度对生成式人工智能进行双重塑造。

2. 认知的流形模式与跨学科统一（马迎辉）

来源：《中国社会科学》2025 年第 7 期

流形是一种引发现代知识革命的观念构想，以其形式性、内蕴性与拓扑性开启了一种新的思维范式，对现代科学和哲学都产生了深远影响。流形学习作为揭开人工智能算法黑箱的突破口，有望在哲学—认知科学的基础上实现对深度学习更高效、更具可解释性的建模。神经流形作为解释大脑认知行为的方案，有望推进人工智能的神经网络设计。纯粹意识展现为一种多维流形，其内时间性与先天相关性为建构人工智能和脑神经活动的意义空间提供了可能。流形不仅展现了人工智能、神经科学与意识哲学之间的贯通结构，而且展现了自身认识从意向流形的奠基性存在，到神经流形的质料存在，再到流形学习的形式存在的发生构造序列，这为认知研究的跨学科统一带来了可能。

3. 人工智能生成的内容是作品吗？——以学术规

范和著作权法的关系为视角（王迁）

来源：《中国法律评论》2025 年第 5 期

实验表明，人工智能根据 745 字初始提示词生成了近万字的法学论文，再根据两次增加的提示词形成的终稿，在经过简单的格式调整和替换明显伪造的注释后，通过了三位外审专家的评审。依“人工智能文生图第一案”判决的结论和逻辑，该论文也应视为该生独立完成的作品。然而，认定学生凭借其使用人工智能生成的上述学位论文可获得学位，将明显违背学术常识，而做出相反结论则抵触“人工智能文生图第一案”的判决，从而形成悖论。人工智能生成的内容源于其自身的算法和数据训练，提示词相对于该生成内容是思想，因此，学生设计提示词由人工智能生成的内容并非其创作的作品和独立完成的学术成果，不能被认定为作者和取得著作权，且其以本人为作者呈现该内容的行为违反学术规范。

4. 从合理使用到合理学习：人工智能训练数据的著作权困境与规范重构（孙济民）

来源：《中国法律评论》2025 年第 5 期

针对人工智能系统使用受著作权法保护的数据进行训练的行为及其所造成的困境，当前多数研究主张应通过扩张合理使用制度的适用范围加以应对。尽管这一做法能够部分回应现实需求，且具备经济正当性，但是以交易成本为主轴的分析框架并不能全面评价数据使用的正外部性。因此，有必要遵循人工智能的技术逻辑，围绕“合理学习”原则形成综合分析框架。在此基础上，相关法律规范应当参照人类学习活动进行重构，立法应区分输入端与输出端以形成事前与事后两层次规制方案。其中，输入端应准许开发者为训练人工智能系统而居间使用受保护作品；输出端应摒弃内容相似性判断，通过行政监管和司法程序确保开发者遵守注意义务，维护公私利益平衡。

5. 人工智能预训练中大规模抓取个人信息的合法性困境与出路（王苑）

来源：《中国法律评论》2025年第5期

人工智能（AI）预训练中大规模抓取个人信息的法律根据存疑。目前欧盟已明确数据控制者正当利益系该场景下的合法性基础，但我国实在法上并无类似条款。而大规模抓取的无差别性导致抓取时已公开个人信息与元数据或敏感个人信息难以界分，因此已公开个人信息的合理使用规则在此情形下无法适用，同时敏感个人信息单独同意规则亦欠缺实践可行性。《个人信息保护法》第13条第1款第5项公共利益条款适用范围狭窄，可通过扩张解释具有普惠价值的AI预训练有公共利益，寻求现行法的合法性支撑。但鉴于AI发展下的多重利益格局，应考虑在未来AI相关立法中明确“预训练场景下个人信息的原则可抓取”，辅以配套制度确保抓取的合法、正当、必要及安全。

6. 人工智能素养越低，对人工智能的接受度越高

（Tully Stephanie M., Longoni Chiara & Appel Gil）

来源：Journal of Marketing, Vol.89, Issue 5 (2025)

随着人工智能（AI）重塑社会，理解影响AI接受度的因素变得日益重要。本研究旨在探究哪些类型的消费者对AI更具接受度。与四项调查揭示的预期相反，跨国数据及六项补充研究发现，人工智能素养较低的人群通常对AI更具接受度。这种“素养越低接受度越高”的关联性，并非源于人们对AI能力、伦理性或对人类潜在影响的认知差异。其根源在于：人工智能素养较低者更倾向于将人工智能视为神奇存在，当目睹人工智能完成看似需要独特人类特质的任务时，他们会产生敬畏之情。与该理论相符的是，这种关联性通过“将人工智能视为神奇存在”的认知实现中介作用，且在无需独特人类特质的任务中表现出调节效应。研究结果表明，企业若将营销策略和产品开发重心转向人工智能素养较低的消费者群体，可能获得显著收益。此外，试图消除人工智能神秘感的努力反而可能削弱其吸引力。

7. 人工智能发明权：承认真正的发明者需要变革

（SCAMBIA NICHOLAS）

来源：St. John's Law Review, Vol.98, Issue 6 (2025)

本文探讨了人工智能（AI）在专利法领域不断演变角色，尤其聚焦于“Thaler诉Vidal案”的深远影响——联邦巡回上诉法院在该案中裁定，根据《专利法》规定，仅自然人可被认定为发明人。作者主张通过立法变革允许人工智能被指定为发明人，认为现行法律阻碍创新且未能体现DABUS等先进人工智能系统的能力——这些系统能够自主生成具备专利性的发明。文章提出了对《专利法》的必要修订方案，包括重新定义“发明人”概念、调整法律用语选择、更新“构思要求”以适应人工智能的独特能力，最终倡导建立承认人工智能对创新贡献的法律框架。

平台治理

1. 论数字时代多边平台的相关市场界定（王晓晔）

来源：《中外法学》2025年第4期

相关市场界定是竞争分析的核心环节之一，也是反垄断法中技术性最强的领域。在数字时代，由于多边平台存在巨大的间接网络效应，平台一边往往提供“零价格”服务，多边平台的相关市场界定要比传统经济下单边市场的相关市场界定复杂得多。考虑到多边平台的一边往往以“零价格”提供服务，界定市场需要重视非价格参数，SSNIP测试在实践中就需要变通为SSNDQ。数字时代的多边平台可简单区分为交易性与非交易性，交易性平台如电子商务平台的两边可以包括在一个产品市场之内。但是美国Amex案说明，将交易性平台机械地视为单一市场不能充分体现平台两边用户群各自不同的需求替代，不能适应现实中复杂的市场环境。数字生态系统的相关市场现在也成为人们关注的问题。鉴于数字生态系统主要考虑系统内部各企业之间以及各种产品和服务之间的互补性，这本质上削弱了界定相关市场的必要性和实用性，国际上迄今尚未出现界定数字生态系统相关市场的案件。

2. 移动应用市场的需求和排行榜排名的作用 (Benedetta Gianola)

来源：Journal of European Competition Law & Practice, Vol.16, Issue 3 (2025)

在移动应用市场中，排行榜是关键的推荐系统，它使排名靠前的应用更具可见性，并降低了消费者的发现成本。我运用工具变量方法，探究应用排名变化如何影响消费者下载应用的意愿。随后通过模拟应用排名与实际观察值存在差异的情景，进一步分析应用可见性对消费者福利的影响。理论预测：若应用商店排名完全由与消费者效用高度相关的因素决定，任何排名调整都将导致福利损失。但本研究发现，在谷歌应用商店中，特定排名变化反而带来正向福利增益，表明谷歌现行排名机制可能基于与消费者偏好不完全一致的指标。这些发现为反垄断及监管领域关于平台通过推荐系统影响消费者行为的讨论提供了新视角。

3. 在线平台上的黑暗模式及其与《数字服务法》和《不公平商业行为指令》的相互作用 (Benjamin Raue)

来源：Journal of European Consumer and Market Law, Vol.14, Issue 3 (2025)

“黑暗模式”这个时髦新词，其实指的就是人人都熟悉的套路：你被置于某种精心设计的境地，最终做出原本可能不会做出的决定。超市收银台旁摆放糖果的经典案例便是明证——当孩子在队伍里哭闹时，多数家长宁可买块巧克力，也不愿忍受孩子大声发泄的不满。这催生了德语借词“Quengelwaren”。数字世界的“Quengelwaren”则表现为：那些反复弹出、令人困惑且通常令人恼火的提示——询问我们是否允许 Cookie 等追踪技术跨平台追踪我们的行为和个人生活；又如班贝格高等地区法院近期审理的案例所示，那些纠缠不休且可能误导消费者的问题：询问音乐会门票购买者是否需要购票保险，而大多数人既不需要也不想要这种保险。黑暗模式可定义为利用人类决策启发式思维与行为弱点的操作手法。如首例所示，黑暗模式在

模拟世界中早已存在。但数字环境为黑暗模式的运用提供了助推器：相较于设计超市收银区等实体场景，数字流程的测试、设计与重构更为便捷。通过在线方式测试用户对界面变更或决策流程设计的反应也更具可行性。在线平台甚至可能根据群体乃至个别客户的特定优势与弱点定制设计方案。欧盟通过多种途径应对这一现象：欧盟委员会在其《不公平商业行为指令》(UCPD)的解释与应用指南中提及该概念。在监管层面，《数字服务法案》(DSA)第25条及《数据法案》(DA)第4条(4)数据法案(DA)直接禁止使用黑暗模式——且可认为人工智能法案第5(1)(a,b)条亦属此列。

数字行政与司法

1. 数字行政与行政职权配置的变革逻辑 (王锡铨)

来源：《环球法律评论》2025年第5期

在数字政府建设中，数字技术对行政活动的全方位渗透，不仅对行政的手段、方式及程序带来革新，也对行政权的配置结构及组织法秩序带来巨大冲击。传统的行政组织法通过权责法定划分出静态的行政权配置结构，但在政务数据汇集、大模型、算法等技术深度融入行政活动的场景中，行政权与数字技术的结合催生出事实上的权力配置新样态。数字技术的广泛应用使行政组织内部的权力运行呈现出跨部门、跨层级、跨地域等特征，催生出“数据控制效应”。具体表现为：在内外维度上，通过数字基础设施的公私合作出现“权力溢出效应”，导致公共决策权向私营主体转移，引发问责盲区；在纵向维度上，数字平台与算法节点引发权力在行政系统垂直方向上的重新配置，产生“权力对流效应”，冲击层级行政结构；在横向维度上，跨部门的数据汇集模式造成“权力交叉效应”，使职权分工趋于模糊化。“数据控制效应”冲击了传统的权责法定原则，可能带来行政权集中化、去边界化的危险，这进一步带来权责失衡和问责失灵的风险。面对这些挑战，数字行政背景下的行政组织法必须

进行相应的制度变革,应以权力来源的法定性为根基,以权力运作的适配性为枢纽,并以程序理性与责任可追溯性为双重保障,实现“数治一法治”的协同演化。

2. 数字赋能政府监管的法治风险及其因应(张青波)

来源:《环球法律评论》2025年第5期

数字技术赋能政府监管,产生了若干法治风险:政府收集和共享信息对个人信息权益构成威胁,程序缺省和沟通障碍漠视当事人程序权益,数字化行政依托的算法不当侵犯相对人权益。为了因应这些法治风险,应当坚守公众参与、正当程序、比例原则,引入设计型法治的理念,重视组织法和程序法的规制。政府收集信息应遵循比例原则,还应符合设计型法治以及组织、程序上的要求;政府共享信息应恪守目的特定原则和比例原则。自动系统应告知当事人程序权利,预留听取当事人意见的程序节点,确保对当事人受告知权、陈述权、申辩权和听证权等程序权利的尊重。对数字化行政依托的算法,应限制算法自动决策范围、赋予相对人对算法自动决策的拒绝权并确保算法决策必要,以审计验证、检测和人工核实等措施确保准确,以预警、评估、说明处理规则、后检测等措施避免和消除歧视,以公开、解释和审计跟踪等措施打破黑箱。

3. 自动化行政中环境算法决策的司法审查研究(任洪涛)

来源:《政法论坛》2025年第5期

随着自动化行政的发展,环境算法决策的应用日渐广泛。然而,由于环境算法决策运行机理的独特性以及环境问题本身的复杂性,法院在对环境算法决策进行审查时遭遇了“消极审查”“审查不足”和“审查不能”的困境。为有效应对行政法治工具变革所带来的司法审查挑战,创新和优化司法审查制度显得尤为重要。在法律规范框架内,必须重申司法审查作为监督和制约行政权核心机制的功能定位,坚持司法谦抑与司法能动相结合的审查

原则,深入考量环境算法决策在主体、内容以及程序等审查要件上的特殊性,确保在协调权力行使与自由裁量空间的过程中寻求适当的界限。在具体制度创新上,应当聚焦于优化审查模式、革新审查方式以及完善结果处理措施,以确保审查过程的公正性与有效性。

4. 信息科技企业参与侦查的角色定位与风险防范(周玥)

来源:《中国刑事法杂志》2025年第4期

信息科技企业参与侦查具有必要性和正当性,这是数字时代技术治理在司法领域的典型表现。信息科技企业并非侦查主体,但不宜将其角色简单定位为见证型第三方,在实践中实际承担技术协助者、数据提供者和特殊的案外人三种角色。信息科技企业参与侦查,对刑事诉讼的结构带来深刻影响,面临权力异化悖离诉讼目标、主体多元化导致司法归责难题和企业自身发展困境三方面的风险。为满足数字时代的侦查需求,同时防范风险,应从理论上明确信息科技参与侦查的界限,并以第四次刑事诉讼法修改为契机,明确信息科技企业的诉讼地位,完善相关侦查措施规定和构建有效的监督体系。

5. 数字时代刑事辩护制度面临的挑战及其应对(郑曦)

来源:《中国法律评论》2025年第5期

在数字时代,刑事诉讼发生数字化转型,新型技术在刑事诉讼领域得到广泛应用,给刑事辩护制度带来了诸多挑战,包括技术门槛前的辩护能力缺失、刑事辩护相关规则滞后、辩护权利保障的难度增大等。面对这些挑战,应构建技术、规则与权利协同互动的框架,实现规则与技术的协调、技术对权利的赋能、规则对权利的保障,从而为维护刑事辩护制度提供思路。具体而言,可以从提升刑事辩护能力、完善刑事辩护规则、强化辩护权利保护等方面着手,积极应对上述挑战,以实现刑事辩护制度对刑事诉讼数字化转型的适应。

6. 读脑技术与免于自证其罪的权利：对证言证据与实证证据区分的挑战（Javier Escobar Veas）

来源：German Law Journal, Vol.26, Issue 4 (2025)

脑部读取技术的发展令人期待，人们终于有望实现谎言检测。然而，这些新技术的存在也引发担忧：当局可能未经同意就读取人们的思想，获取可在刑事诉讼中用作指控证据的信息，这种情境引发了对可能侵犯自证其罪权利的质疑。本文旨在分析：通过未经同意使用读脑技术获取的指控性信息，是否会违反传统解释下的自证其罪权。根据该传统解释，该权利的适用范围仅限于“证言证据”，而排除“实物或物理证据”。

7. 为数据科学在反垄断执法和平台监管中的应用规划发展路径（Graef Inge, Laitenberger Ulrich & Prüfer Jens）

来源：European Competition Journal, Vol.21, Issue 2 (2025)

竞争监管机构及其他监管部门已开始运用数据科学监测市场并核查规则遵守情况。《数字市场法案》和《数字服务法案》的出台，使得监管机构运用数据科学变得更为重要——鉴于监管工作所需规模之大、复杂性之高。现有数据科学应用案例表明，监管机构能够调整组织架构与工作流程，充分释放技术工具在监管活动中的潜力。与此同时，数据科学应用在数据可靠性与个人隐私方面也带来挑战。除参与具体调查外，数据科学更有潜力推动监管工作转型，向更主动的执法模式转变。监管机构间应进一步促进数据科学专业知识与工具的共享，以加强协作并共享资源，共同监管平台经济。

虚拟财产

1. 数字资产、MiCA 和欧盟投资基金法（Dirk A. Zetsche, Filippo Annunziata & Julia Sinnig）

来源：European Business Organization Law Review, Vol.26, Issue 3 (2025)

本文基于近期通过的《加密资产市场条例》（MiCA），分析了欧盟数字资产投资基金（数字资产基金）的法律环境。在融合既有与 MiCA 引入的新型（甚至未经检验）金融监管的背景下，我们厘清了《可转让证券集体投资计划指令》与《另类投资基金管理人指令》的适用范围，并阐明了数字资产基金领域针对加密资产服务提供商的新规要求。在 MiCA 出台前，多数投资基金对数字资产持审慎态度；据我们所知，迄今整体投资敞口仍处于较低水平。本文旨在探讨 MiCA 生效后是否会增加投资基金对数字资产的敞口，以及可能的增加方式。我们发现，MiCA 通过提升法律确定性，削弱了基金经理回避数字资产投资的首要顾虑。然而，第二大顾虑——运营风险依然存在，仍可能限制传统服务商对数字资产的投资意愿。对于托管机构而言，这种情况尤为突出：无论 MiCA 是否生效，运营风险都将转化为责任风险。在缺乏稳定、稳健且具有韧性的托管机构前提下，数字资产对基金投资者的吸引力难以提升。然而 2024 年至 2025 年上半年美国市场的最新动向表明，全球加密货币投资基金生态正悄然形成。

（技术编辑：李佳丽、麻卓妍）

教研活动

人工智能时代全球法治大会 | 第六届“未来法治与数字法学”国际论坛成功举行

2025年10月3日至4日，第六届“未来法治与数字法学”国际论坛在中国人民大学逸夫会堂成功召开。本届论坛是中国人民大学举办的“第四届21世纪世界百所著名大学法学院院长论坛暨人工智能时代全球法治大会”的主题论坛之一。论坛由中国人民大学法学院与京东集团联合主办，由中国人民大学未来法治研究院、京东集团法律研究院和吉林大学全面依法治国研究中心承办，支持单位包括人民法院出版社《数字法治》编辑部、中国人民大学法学院数字法学教研中心、教育部哲学社会科学创新团队“新科技革命与未来法治创新团队”和国际数字法学协会。

在一天半的论坛中，来自中国、美国、英国、德国、法国、意大利、瑞士、波兰、丹麦、日本和新加坡等国家的八十多位专家学者围绕人工智能的法治建设与全球治理、智能社会的知识产权与数据制度、网络平台的智慧治理与数字劳动、数字文明的未来法治与制度形态等主题，进行了发言和研讨，吸引了国内外两百余位相关专业和行业的嘉宾与会，共议未来法治与数字法学面临的机遇和挑战。



论坛开幕式现场

中国人民大学校长**林尚立**发表致辞。他表示，作为法学教学与研究领域重镇，中国人民大学坚持以习近平法治思想为指导，积极投身数字法治建设，

贡献了“独树一帜”的人大力量。以此次论坛为契机，全球专家聚焦“未来法治与数字法学”等主题展开前沿探讨。希望与会学者聚焦治理需求，以法治之策解决实践问题，紧跟数字技术发展进程，推动研究从“书本之法”走向“治理之策”；站稳价值立场，以法治之规引领技术向善，以前瞻性法学理论为技术发展立规引航，确保技术进步始终服务于人民的美好生活；深化交流合作，以法治之桥凝聚全球智慧，促进全球平台共建、规则共研、标准共商，推动数字文明发展成果更好惠及各国人民，携手迈向普惠共赢的数字文明新图景。



中国人民大学校长 林尚立

中国法学会党组书记、常务副会长，第十四届全国人大常委会委员、宪法和法律委员会副主任委员**王洪祥**在致辞中表示，习近平总书记高度重视数字中国建设。法治作为国家治理体系和治理能力现代化的重要依托，必须积极回应数字时代需求，为数字经济发展筑牢法治基石，注入法治动能，提供法治保障。王洪祥指出，要坚持理念引领，构建以人为本、智能向善的数字法治价值体系；强化制度供给，完善人工智能等新兴领域立法，构建科学完备、统一权威的法律体系；深化人才培养，推动法学与数字科技交叉融合，夯实数字法治的治理基础；加强实践应用，以数字科技赋能法治建设，促进法治与科技良性互动；扩大国际合作，共商共建全球数字治理体系。中国法学界将积极分享数字法治实践经验，与各国专家携手努力，共同推动数字文明成果更好地造福世界人民，为构建普惠共赢的数字文明新图景贡献智慧力量。



中国法学会党组书记、常务副会长，
第十四届全国人大常委会委员、
宪法和法律委员会副主任委员 王洪祥

中国法学会网络与信息法学研究会会长、最高人民法院咨询委员会副主任委员姜伟发表致辞。他表示人工智能治理已成为国际社会共同面临的紧迫议题。人工智能技术兼具巨大发展潜力与全球性风险，其风险具有隐蔽性、跨域性和长期性特征，单一国家难以独立应对，必须依靠国际合作协同治理。当前全球人工智能治理规则仍呈现碎片化态势，构建包容、公平的全球治理框架尤为关键。因此应致力于构建协同共治的全球人工智能治理体系，尊重各国的多样性，寻求治理规则的互操作性，从共识度高的领域切入，逐步扩大合作范围。联合国是开展这一全球治理进程的主渠道，应坚持共商共建共享原则，确保所有国家特别是发展中国家能够平等参与、共同受益。姜伟会长强调，中国积极参与全球人工智能治理合作，先后参与或发起了多项国际倡议宣言，并持续通过对话合作构建一个公平、包容、有效的全球人工智能治理秩序，让技术发展成果更好惠及全人类。



中国法学会网络与信息法学研究会会长、
最高人民法院咨询委员会副主任委员 姜伟

中国人民公安大学校长、中国法学会副会长王轶发表致辞。他回顾了中国人民大学未来法治研究院的设立历程与发展成就。作为中国法学院校中首个回应新一轮科技革命和产业变革的研究机构，人大未来法治研究院自创立之初就确立了“始终奋进在时代前列”的使命担当。王轶表示，人大未来法治研究院在研讨中深刻认识到，生产力的提升、生产工具的改善、生产要素的丰富推动文明形态的转型，而文明转型必然带来法治的变迁。人类走过了以耕种土地为主要生产生活场景的“黄色文明”阶段，经历了以开发利用海洋为特征的“蓝色文明”阶段，如今正迈向以太空探索为标志的“玄色文明”新纪元。在未来法治研究中应当及早布局太空法治研究，及早着手推进太空治理，着力提出该领域的中国方案。最后，王轶校长表示中国人民公安大学愿与中国人民大学密切协作，共同贯彻落实习近平总书记的四大全球倡议，为推动构建人类命运共同体、创造人类更加美好的未来贡献智慧与力量。



中国人民公安大学校长、教授，
中国法学会副会长 王轶

论坛协办方代表、吉林大学法学院院长何志鹏致辞表示，全球参会同仁可利用此次会议的契机，以法治思维深入探讨人口安全与教育变革、人工智能技术发展带来的多重挑战；以法治规则平衡地缘技术差异与社会公平的关系，关注人口逆增长背景下法学人才培养的可持续性，思考人工智能普及对劳动力结构的影响；以比较法视角和国际视野回应制度差异和现实国情，正视发达国家与发展中国家在人工智能应用上的差距，防止数字鸿沟加剧财富、

科技、能力鸿沟，确保技术始终服务于人类共同福祉；以法治道路应对挑战，开辟普惠共赢的数字法治新路径，运用国际法与国内法的现有基础，推动全球法学界共研数字时代治理规则、共商法治建设标准，在应对科技挑战中为未来法治找准方向，给出人类智慧前沿的良好答案。



吉林大学法学院院长、教授 何志鹏

论坛联合主办方代表、京东集团副总裁**胡焕刚**在致辞中介绍了京东的业务创新、技术创新和国际业务增长。外卖业务中，京东已与15万全职骑手签订劳动合同，为其缴纳五险一金，保障外卖员劳动权益。在技术创新上，京东已经形成了AI全景图，实现了AI技术全链条覆盖，发布了三大全新AI产品、四大场景深度AI应用和多个技术平台，未来三年将持续投入AI的研发和运用，带动各个产业形成万亿规模的人工智能生态。胡焕刚总裁呼吁，完善新就业形态劳动者权益保障立法，建立适应数字时代的劳动保障体系；推进法治与科技深度融合，为数字创新提供制度保障。京东集团将持续发挥平台作用，积极参与互联网法治进程，履行社会责任，加强行业与社会共治，为数字经济高质量发展和法治中国建设贡献力量。



京东集团副总裁 胡焕刚

开幕式

第六届“未来法治与数字法学”国际论坛开幕式由中国人民大学法学院党委书记**沃晓静**主持。

在成果发布环节，中国人民大学未来法治研究院助理研究员，交叉科学研究院/高瓴人工智能学院博士后**李铭轩**，介绍了中国人民大学涉外法治大模型2.0版本。

中国人民大学法学院副院长、未来法治研究院副院长**丁晓东**教授和吉林大学法学院**李拥军**教授共同主持了特邀演讲环节。

在特邀演讲环节中，吉林大学哲学社会科学资深教授，中国法学会法学教育研究会会长，教育部社会科学委员会法学学部召集人**张文显**以《数字法治体系的内部关系》为题进行了演讲；北京市人民检察院检察长**朱雅频**以《把握数智时代检察履职的鲜明特质》为题进行了演讲；中国人民大学国家一级教授，中国法学会副会长、学术委员会副主任，中国法学会民法学研究会会长**王利明**以《人工智能侵权及应对》为题进行了演讲；中国政法大学数据法治研究院院长、教授**时建中**以《政务数据的共享、开放、授权运营相关制度的检思》为题进行了演讲；对外经济贸易大学副校长、教授**梅夏英**以《人工智能立法的形式和内容》为题进行了演讲；德国马克斯·普朗克比较法与国际私法研究所高级研究员、教授**雷纳·库姆斯**以《人工智能事件——人工智能出错时的探讨》为题进行了演讲；中国工程院外籍院士，新加坡工程院院士，清华大学讲席教授、智能可穿戴设备中心主任**西拉姆·拉马克瑞斯纳**以《医疗设备、可穿戴设备与人工智能——网络安全的法律前沿问题》为题进行了演讲；瑞士日内瓦大学法学院教授**雅克·德·维拉**以《AI与版权法：欧洲视角》为题进行了演讲。

第二单元：主题演讲

中国人民大学法学院教授、未来法治研究院执行院长**张吉豫**和人民法院出版社编审、《数字法治》编辑部主任**兰丽专**主持了第二单元主旨演讲环节。

在主旨演讲环节，英国杜伦大学法学院教授**威**

廉·露西以《法哲学的核心之问会因社会变迁而改变吗?》为题作了演讲;中国人民大学法学院吴玉章高级讲席教授,中国法学会网络与信息法学研究会副会长、学术委员会主任**张新宝**以《作为知识产权的数据与作为数字经济要素的数据之“并联包容”关系探讨》为题进行了演讲;前马克斯·普朗克外国与国际刑法研究所所长,弗赖堡大学教授**汉斯约格·阿尔布莱希特**以《对人工智能与刑法的思考》为题进行了演讲;浙江大学数字法治研究院院长、教授**孙笑侠**以《论机制型数字权利》为题进行了演讲;西北政法大学“长安学者”特聘教授,《法律科学》主编**杨建军**以《人工智能时代的就业替代与法律应对:中国方案建构》为题进行了演讲;浙江大学光华法学院院长、教授**胡铭**以《智能体在浙大法学教育中的运用》为题进行了演讲;北京师范大学法学院教授,数字法学研究中心主任**汪庆华**以《平台权力的来源、形塑与法律规制》为题进行了演讲。

第三单元:“人工智能的法治建设与全球治理”

第三单元“人工智能的法治建设与全球治理”由《中国法律评论》常务副主编、编辑部主任**袁方**和中国人民大学法学院教授、未来法治研究院副院长**王莹**主持。

在第三单元中,四川大学杰出教授、实证法学与智慧法治四川省重点实验室主任**左卫民**,美国康奈尔大学法学院教授**弗兰克·帕斯奎尔**,上海交通大学凯原法学院院长、教授,数据法律研究中心主任**彭诚信**,英国贝尔法斯特女王大学全球创新、知识产权与科技法中心主任、教授**吉安卡洛·弗罗西奥**,上海交通大学凯原法学院教授**郑戈**,中国社会科学院法学研究所研究员、《环球法律评论》副主编**支振锋**,中国政法大学人工智能法治研究院院长、教授,联合国高级别人工智能咨询机构专家**张凌寒**,中央财经大学法学院副院长、“龙马学者”特聘教授**刘权**,吉林大学法学院教授**朱振**,日本早稻田大学法学院教授,知识产权法律制度研究中心联合主

任**克里斯托夫·拉德马赫**尔,美国南卡罗来纳大学法学院副教授**布莱恩特·沃克·史密斯**进行了演讲,北京航空航天大学法学院副教授**赵精武**,英国埃塞克斯大学法学院助理教授、剑桥大学商业研究中心研究员**左振斌**进行了与谈。



第三单元“人工智能的法治建设与全球治理”

嘉宾合影

第四单元:“智能社会的知识产权与数据制度”

第四单元“智能社会的知识产权与数据制度”由《中国法学》杂志社副编审**任彦**,中国人民大学法学院教授、未来法治研究院执行院长**张吉豫**主持。

在第四单元中,中国政法大学民商经济法学院教授、中国法学会知识产权法学研究会副会长**冯晓青**,德国马克斯·普朗克创新与竞争研究所资深研究员、萨格勒布大学教授**西尔克·冯·莱温斯基**,库亚维-波美拉尼亚大学副教授**马里乌什·克日什托费克**,德国马克斯·普朗克创新与竞争研究所高级研究员**达利亚·金**,香港大学法律学院教授**孙皓琛**,吉林大学法学院教授**侯学宾**,吉林大学法学院副教授、司法数据应用研究中心研究员**齐英程**,澳门大学助理教授**塞莉娅·菲利帕·费雷拉·马蒂亚斯**进行了演讲。中国科学院科技发展战略研究所所长、科技战略咨询研究院研究员**肖允丹**,同济大学法学院副教授、知识产权与竞争法研究中心主任**张韬略**进行了与谈。



第四单元“智能社会的知识产权与数据制度”
嘉宾合影

第五单元：“网络平台的智慧治理与数字劳动”

第五单元“网络平台的智慧治理与数字劳动”由山东大学法学院大数据法治研究中心副主任、教授**郑智航**和中国人民大学未来法治研究院助理研究员、交叉科学研究院/高瓴人工智能学院博士后**李铭轩**主持。

在第五单元中，中国人民大学法学院教授、劳动法和劳动保障法研究所所长、中国社会法学会常务副会长**林嘉**，意大利费拉拉大学法学院教授、欧洲科学院人文与艺术学院（社会科学、法律与经济学部）候任院长**阿尔贝托·德·弗朗切斯基**，中国社会科学院大学法学院副教授、互联网法治研究中心主任**刘晓春**，中央财经大学法学院教授、中国社会法学会常务理事**沈建峰**，中国社会科学院法学研究所社会法研究室副主任、研究员**王天玉**，中国人民大学法学院助理教授、未来法治研究院研究员**阮神裕**，巴黎西岱大学法学院教授**皮埃尔·贝里奥兹**，英国牛津大学法学院人工智能法律与监管领域教授**伊格纳西奥·科福内**进行了演讲。北京邮电大学人工智能法律研究中心主任、教授**崔聪聪**，上海师范大学哲学与法政学院副教授**王奇才**进行了与谈。



第五单元“网络平台的智慧治理与数字劳动”
嘉宾合影

第六单元：“数字文明的未来法治与制度形态”

第六单元“数字文明的未来法治与制度形态”

由北京交通大学法学院副教授、数据法学研究中心主任**付新华**和中南财经政法大学《法商研究》编辑部副编审**王虹霞**共同主持。

在第六单元中，日内瓦大学法学院副院长、副教授**格洛丽亚·加焦利**，东京大学法学政治学研究科教授**后藤元**，上智大学法学院国际关系法系教授**森下哲朗**，吉林大学法学院教授**李拥军**，中国政法大学比较法学研究院教授、中美法学研究所所长**冯恺**，丹麦哥本哈根大学法学院副教授**向文**，厦门大学知识产权研究院教授**董慧娟**，吉林大学法学院副教授**吴梓源**，香港大学法律学院助理教授**关韬睿**进行了演讲。吉林大学法学院副院长、教授、理论法学研究中心、司法文明协同创新中心主任**杜宴林**，北京理工大学智能科技风险法律防控工信部重点实验室副主任、办公室主任、副教授**陈姿含**进行了与谈。

在闭幕式中，中国人民大学法学院教授、未来法治研究院副院长**王莹**，吉林大学法学院教授、吉林大学全面依法治国研究中心副主任**朱振**，京东法律研究院执行院长**吴国基**依次进行了致辞。

最后，中国人民大学法学院院长**杨东**教授总结致辞，祝贺本届论坛圆满落幕，并期待未来能够继续与国内外社会各界进一步深化交流合作，共同推动未来法治与数字法学的持续发展。



第六单元嘉宾及闭幕式合影

本届“未来法治与数字法学”国际论坛中还发布了《中国自主数字法学知识体系原创性理论与

标识性概念》以及12卷《人大未来法治与数字法学研究系列成果》，展示了《隐私与数据的法律研究》《数字财产权利的法律构造》《数字时代的权利理论》《人工智能的法律研究》《人工智能的伦理和治理》《平台责任：网络平台的治理机制研究》《数字时代在线解纷机制：理论重塑与实践创新》《计算法学方法初阶》等人大未来法治研究丛书及《个人信息保护法教程》等数字法学教材，呈现了人大未来法治研究院和法学院数字法学教研中心在未来法治与数字法学领域的阶段性成果，并进一步加强了“国际数字法学协会”与社会各界的深度交流与合作。

第六届未来法治学生论坛成功召开

2025年10月4日，由中国人民大学未来法治研究院、教育部哲学社会科学创新团队“新科技革命与未来法治创新团队”、人民法院出版社《数字法治》编辑部、《人大法律评论》编辑部联合主办的第六届未来法治学生论坛成功举办。本次论坛包括开幕式、专题发言、会议总结三个环节。

开幕式

开幕式由《人大法律评论》主编**谭耀淇**主持。

在开幕式上，中国人民大学法学院副院长**万勇**、中国人民大学未来法治研究院执行院长**张吉豫**、人民法院出版社副编审、《数字法治》编辑部主任助理**黄晓云**先后致辞。



中国人民大学法学院副院长 万勇

万勇教授对莅临论坛的各位专家和同学表示

热烈欢迎和衷心感谢，并鼓励《人大法律评论》向世界上有影响力的法律评论刊物学习。万勇教授还提出了三点意见。第一，法学研究要重视自主知识体系建构。第二，法学研究者要面向世界，具有国际视野。第三，法学论文需要处理好理论深度和实践导向的关系。

张吉豫教授对第六届未来法治学生论坛的举办表示祝贺和感谢。首先，张吉豫教授指出，法治要适应快速发展的社会事实。举办论坛的初衷是鼓励青年学者把自己的想象力和法学知识结合，将社会关切转变为优质论文。其次，张吉豫教授表示，学生论坛为老师和学生都提供了一个思想交流的平台，为学术界作出一定贡献。张吉豫教授建议学生论坛可以与“未来法治与数字法学”国际论坛形成一定的协作。最后，张吉豫教授再次对支持论坛的各位老师表示感谢，也特别感谢《数字法治》编辑部对论坛一如既往的支持。



中国人民大学未来法治研究院执行院长 张吉豫

人民法院出版社副编审、《数字法治》编辑部主任助理**黄晓云**首先代表《数字法治》编辑部对各位老师、同学的到来表示热烈欢迎，向为论坛举办付出辛勤努力的师生致以诚挚的谢意，对未来法治研究院的鼎力支持表示感谢。面对被技术深切重塑的现代社会，黄晓云老师提出三点感想。第一，要做贯通古今的思考者，把古老智慧和新兴科技结合起来找到解决未来问题的“金钥匙”。第二，要做脚踏实地的实践者，以正确的法治理论为引领，以务实精神推动规则优化，绘制更精确有效的法治蓝图。第三，要做心怀温暖的担当者，既要守护正义，也要鼓励创新。



人民法院出版社副编审、
《数字法治》编辑部主任助理 黄晓云

专题一：人工智能时代的法学范式转型

本环节由《人大法律评论》编辑符大卿主持。

第一位报告的同学是清华大学法学院博士研究生董佳乐，她报告论文的题目是《自动化行政中的事实认定》。报告人立足于“如何应对自动化行政中事实认定的人机协同困境”这一核心命题，强调应在准确把握人类思维与机器思维差异的基础上，以比较优势为底层逻辑，构建类型化、场景化的规制路径。

中央财经大学法学院“龙马学者”特聘教授刘权，华东师范大学法学院讲师、晨晖学者陈靓雯分别对董佳乐同学的论文进行了评议。

第二位报告的同学是东南大学法学院博士研究生朱有辰，他报告论文的题目是《从“法益识别”到“行为分型”：企业数据犯罪刑法规制的范式转型》。

中国人民大学交叉科学研究院副院长龚新奇，北京航空航天大学法学院副教授赵精武，中国社会科学院法学研究所助理研究员蔡宇姬分别对朱有辰同学的论文进行了评议。

第三位报告的同学是中国人民大学法学院博士研究生李亚兰，她报告论文的题目是《人工智能生成内容的思想——表达二分法分析》。

人民法院出版社副编审、《数字法治》编辑部主任助理黄晓云，北京大学智能学院助理研究员、《网络法律评论》编辑部主任、《科技与法律（中英文）》编辑部副主任辜凌云，中国人民大学法学院讲师孙靖洲对李亚兰同学的论文进行了评议。

专题二：新兴技术风险的法律挑战

本环节由《人大法律评论》编辑柯晨亮主持。

第四位报告的同学是中国人民大学法学院博士研究生咸春亭，他报告论文的题目是《论“深度伪造”作为新型伪证的法律规制》。

中国人民大学纪检监察学院副教授王燃，清华大学法学院助理研究员刘静分别对咸春亭同学的论文进行了评议。

第五位报告的同学是西安交通大学法学院本科生龚运博，他报告论文的题目是《从具体损害到统计风险：大语言模型幻觉对侵权可预见性标准的挑战》。报告人着眼于“大语言模型‘幻觉’现象的法律规制路径”这一核心问题，指出该类幻觉作为基于统计概率的系统性生成缺陷，已成为制约人工智能应用与发展的显著风险。

吉林大学法学院副教授齐英程，上海对外经贸大学法学院讲师吴涛分别对龚运博同学的论文进行了评议。

会议总结

闭幕式首先由山西省委党校报刊社副社长、《理论探索》副主编、编辑部主任杨在平对论坛进行总评议。杨在平老师认为本次论坛论文整体展现出青年学者对前沿法治问题的敏锐洞察，选题兼具创新性与理论价值，但在学术规范与论证深度上仍有提升空间。杨在平老师认为这些论文展现了年轻学者良好的学术潜质，若能在学术规范与理论建构上进一步锤炼，将在未来法治研究中取得更大突破。

中国人民大学法学院副教授黄尹旭对会议进行总结。黄尹旭老师对参与本次论坛的嘉宾表示衷心的感谢，并对各位汇报人提出了恳切的建议。



论坛成员合照

中国法学会网络与信息法学研究会 2025 年年会暨第三届数字法治大会在江苏南京召开

10月12日，由中国法学会网络与信息法学研究会主办，东南大学承办，《数字法治》编辑部、南京通达海科技股份有限公司协办的“中国法学会网络与信息法学研究会 2025 年年会暨第三届数字法治大会”在江苏南京召开。本次年会的主题为“智能时代的数字法治”。

大会现场

大会开幕式于上午8时30分在南京国际会议大酒店举行，由中国法学会网络与信息法学研究会会长，最高人民法院咨询委员会副主任姜伟主持。中国法学会党组成员、副会长杨万明，东南大学校长、中国工程院院士孙友宏为本次年会致开幕辞。



中国法学会网络与信息法学研究会会长，
最高人民法院咨询委员会副主任 姜伟

杨万明副会长指出，网络与信息法学研究会团结凝聚广大法学法律工作者在推动网络与信息法学研究繁荣发展、参与咨政建言、促进国际交流合作等方面取得了一系列重要成果。新征程上，要强化政治引领，深入学习贯彻习近平法治思想和习近平总书记关于网络强国的重要思想，坚持用党的创新理论武装头脑、指导研究，自觉将数字法治建设置于党和国家事业发展全局中思考谋划，确保研究会工作始终沿着正确政治方向前进。要持之以恒聚焦重大理论与实践问题，深入研究中国式现代化进程中网络和信息化的重大问题，构建具有中国特色、体现时代精神、融通中外的数字法学理论体系和话语体系，为数字法治实践提供坚实的理论支撑。要

守正创新，加强人才队伍建设，探索建立稳定支持、长周期评价机制，培养一批政治素质好、理论水平高、实践能力强的数字法治领军人才和中青年骨干，为我国数字法治建设提供源源不断的人才保障和智力支持。



中国法学会党组成员、副会长 杨万明

孙友宏校长首先代表本次年会的承办方东南大学向莅临大会的各位领导、专家学者和嘉宾朋友们致以诚挚欢迎。他指出，东南大学依托深厚的工科发展优势，着力打造“工科+法学”的交叉融合体系并取得显著成果。展望未来，数字法治必须精准锚定国家重大战略需求，推动建构中国自主的知识与学科体系，培养复合型数字法治领军人才。



东南大学校长、中国工程院院士 孙友宏

上午会议第二阶段为第三届理事会第三次会议环节，由中国法学会网络与信息法学研究会副会长、中央网信办网络法治局副局长、一级巡视员尤雪云主持。

会议进行了理事会成员选举后，通过了中国法学会网络与信息法学研究会常务副会长兼秘书长、同济大学法学院特聘教授、中国社会科学院法学研究所研究员周汉华所作的理事会工作报告。



中国法学会网络与信息法学研究会常务副会长兼秘书长、同济大学法学院特聘教授、中国社会科学院法学研究所研究员 周汉华

上午会议第三阶段为主旨演讲环节，由东南大学法学院院长熊樟林教授主持。

东南大学首席教授、研究生常务副院长，新一代人工智能技术与交叉应用教育部重点实验室主任耿新以《本能、进化、创意——“学习基因”带来的新质 AI 能力》为主题，从技术本源出发，围绕机器学习基因的研究背景、技术路线，阐述了“学习基因”在降低大模型部署成本，以及增强模型本能、进化、创意等新质能力方面的重要作用。



东南大学首席教授、研究生常务副院长，新一代人工智能技术与交叉应用教育部重点实验室主任 耿新

人民法院出版社党委书记、社长，《数字法治》主编余茂玉以《融合众长推动人工智能与司法工作稳健融合》为题，介绍了“法信法律基座大模型”研发和《数字法治》发展情况，表示与各方深入研讨交流如何通过法治和技术“双轮驱动”推动人工智能规范运行，依托《数字法治》期刊助力构建政产学研用协同创新发展体系，让人工智能这台“超级引擎”更好为建设数字中国、法治中国贡献

力量。



人民法院出版社党委书记、社长，
《数字法治》主编 余茂玉

华东政法大学教授王迁以《人工智能生成内容与作品认定》为题，围绕 AI 生成内容是否构成著作权法意义上的作品展开论述，结合法律定义与实验案例，指出生成式 AI 所产生的内容受自身算法、训练数据及硬件配置等因素影响，缺乏人对表达要素的直接控制；用户输入的提示词更接近创作思路，因此 AI 基于提示词生成的内容不构成著作权法所定义的作品。



华东政法大学教授 王迁

上午会议第四阶段为“数据确权”高端对话环节，由中国人民大学法学院院长、教授杨东主持。中国法学会网络与信息法学研究会常务副会长兼秘书长、同济大学法学院特聘教授、中国社会科学院法学研究所研究员周汉华，中国法学会网络与信息法学研究会副会长、学术委员会主任、中国人民大学法学院教授张新宝，华东政法大学法律学院教授、互联网法治研究院院长高富平担任对话人。

周汉华副会长从当前社会处于剧烈转型期的背景出发，强调就“数据确权”问题进行学术讨论的重要性，指出网络技术发展日新月异，按照传统

的民法物权分析方法研究数据确权问题值得商榷，并进一步提出从“网络效应”角度出发，突出数据使用权，以此促进数据流通、实现网络数据价值的最大化。



中国法学会网络与信息法学研究会常务副会长兼秘书长、同济大学法学院特聘教授、中国社会科学院法学研究所研究员 周汉华

张新宝副会长首先指出，不应该将“数据确权”与“数据利用”对立起来，数据确权能够保障数据生产主体的经济利益，激发有关企业和个人的积极性，最终促进数字技术发展并惠及数据使用者；数据的流动性并不意味着其不可量化或定性，现有立法和司法实践已开展关于“数据确权”的诸多尝试和探索。



中国法学会网络与信息法学研究会副会长、学术委员会主任、

中国人民大学法学院教授 张新宝

高富平教授以《重新定义数据保护》为题进行分享，主张数据作为事实不受保护，使用数据的权利和结果也即认知成果受保护，并以数据集为例指出，鉴于数据的流动性和形态不确定性，数据集的整理加工者赋权既不可行，也不利于实现数据社会效用最大化的目标，进而得出“数据确权是一个伪

命题”的结论。



华东政法大学法律学院教授、
互联网法治研究院院长 高富平

本次年会在12日下午设置十二个分论坛，主题分别为“习近平法治思想引领下的网络强国建设”“网络与信息法学自主法学知识体系构建”“人工智能的司法应用与保障”“人工智能立法”“人工智能与知识产权”“人工智能多元治理”“数据权益的法律保护”“公共数据利用的法治保障”“新型网络犯罪治理”“网络犯罪治理的前沿议题”“网络平台治理的理论框架”“数字社会治理的策略与实践”。

“数据要素合规高效流通利用机制研究”专题研讨会成功举办

为深入贯彻习近平法治思想和网络强国战略，落实国家数据要素市场化配置改革部署，2025年9月28日，由人民法院出版社《数字法治》期刊编辑部与中国政法大学人工智能法研究院联合主办，中国人工智能产业发展联盟政策法规工作组、AI善治学术工作组、人民法院出版社电子音像出版社协办的“数据要素合规高效流通利用机制研究”专题研讨会在北京沈家本故居召开。本次会议旨在搭建“政产学研用”深度融合的交流、对话平台，聚焦数据要素流通利用中的制度设计、合规实践与司法保障等核心议题，推动数据要素在法治轨道上释放价值。

人民法院出版社副总编辑、《数字法治》副主编袁登明作开幕致辞。他介绍了《数字法治》期刊作为全国性数字法治领域唯一学术期刊“聚焦数字法治研究 助力法治中国建设”的定位与使命，

并指出，法院出版社依托自身资源优势，已逐步构建起支撑数字法治建设的多层次体系：包括承担全国数字法院人工智能研发总集任务、建设并运营国内最大的法律知识与案例大数据平台（法信），以及打造“理论平台+实践基地”融通的科研交流矩阵。

本次会议邀请了来自政府相关部门、法院系统、高校科研机构及产业界的嘉宾参会。与会代表围绕数据要素流通的基础制度、合规审查与风险防控、激励与保障机制等关键内容展开深入交流，凝聚了多项共识。

法院系统代表结合最高人民法院首次发布的数据权益专题指导案例，阐述了司法实践中如何提炼裁判规则、发挥司法智慧，在立法尚未完善的背景下为数据流通提供裁判思路与制度参考；学术界专家呼吁，当前应避免过度纠缠于行为法与财产法的理论之争，而应聚焦于在现有法律框架下破解AI训练数据版权、个人信息“单独同意”边界、数据来源合法性认定等现实的数据流通障碍；产业界代表则从实务出发，围绕数据交易所法律责任边界、AI训练数据授权链不完整、爬取行为合规标准、数据匿名化技术规范等痛点，呼吁加快构建可操作、可预期的制度与标准体系。

在会议总结环节，人民法院出版社《数字法治》编辑部主任**兰丽专**、中国政法大学教授**张凌寒**指出，本次会议体现出各界关注的问题高度一致，普遍认为当前数据流通领域的制度供给仍显不足，亟需通过跨领域协作推动机制完善，促进数据要素在法治轨道上高效流通与利用。

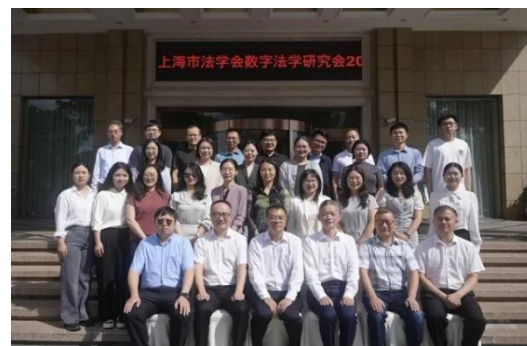
此次专题研讨活动为相关政策完善、司法实践与产业创新提供了有益参考。会议期间，与会嘉宾参观了中国现代法律制度奠基者沈家本故居展览，并在人民法院出版社《法本——中华优秀传统文化法律文化图说》画册等法律文创产品的展陈中，深切感受到了数字法治的未来图景与传统法律文化的时空交错与精神传承。

未来，人民法院出版社《数字法治》期刊将继续打造常态化、品牌化的高端系列研讨会——“数字法治圆桌对话”，搭建“政产学研用”深度融合、

交流对话的智识融通平台，推动真问题调研、共识凝聚与制度构建。

学术综述 | 上海市法学会数字法学研究会 2025 年学术年会观点摘编

2025 年 9 月 26 日上午，上海市法学会数字法学研究会 2025 年学术年会在上海市松江区新桥镇维也纳国际酒店成功召开，会议主题为“数字法理与数据治理”。会议由上海市法学会数字法学研究会主办，华东政法大学数字法治研究院、《数字法学评论》编辑部承办。



部分与会成员合照

会议开幕阶段，由上海市法学会数字法学研究会会长、华东政法大学数字法治研究院院长、教授**马长山**主持，上海市法学会党组书记、专职副会长、《东方法学》主编**施伟东**，华东政法大学党委副书记、纪委书记、教授**陈玉刚**致辞。

第一单元 浦东数据立法的前沿问题

会议第一单元的主题为“浦东数据立法的前沿问题”，由华东政法大学数字法治研究院副院长、副教授**韩旭至**主持。



华东政法大学数字法治研究院副院长、
副教授 韩旭至

围绕着浦东新区数据立法的项目，课题组前期撰写的两份专报，上海对外经贸大学副教授**高阳**首先介绍了两份专报——《浦东数据产品知识产权立法的可能性空间》《上海探索公共数据知识产权登记的制度创新可能性》的撰写背景、基本内容与实践意义。其次，阐述了以区块链赋能数据产品登记为立法项目研究对象的难点与堵点，但是，囿于汇集型数据产品所承载的信息可能属于公共领域或知识产权保护的排除客体，具有一定经济价值但未达到数据知识产权登记要求的数据产品在实践中亦存在保护的必要性，依托于区块链的技术特征，课题组建议以《区块链赋能数据产品登记若干规定》为立法项目的备选方案，具有研究和实践价值。

华东政法大学数字法治研究院助理研究员**时诚**在分析比较浦东数据产品立法的三种方案的基础上，认为颁布《上海市促进浦东新区可信数据空间建设若干规定》是更现实、可行的选择，并提出了可信数据空间的基础架构。

华东政法大学数字法治研究院特聘副研究员**徐则林**指出，第二章可信数据空间培育推广围绕“培育”和“推广”两大主线展开，分类推进企业、行业、城市、个人、跨境可信数据空间建设和应用，从支持企业可信数据空间建设、推动行业可信数据空间发展、鼓励城市可信数据空间建设、稳慎探索个人可信数据空间试点、稳慎推进跨境可信数据空间试点五个维度系统构建可信数据空间的发展框架，形成多层次、全覆盖、稳慎推进的立法结构。

上海师范大学哲学与政法学院副教授**杨帆**针对《上海市促进浦东新区可信数据空间建设若干规定》第三章“空间筑基”的立法设计问题提出了个人看法。

上海政法学院人工智能法学院讲师**余圣琪**围绕《上海市促进浦东新区可信数据空间建设若干规定》中保障措施的内容，从制度协同、资金支持、治理生态、开放对接四个维度出发，提出建议。

在与谈环节，同济大学法学院副教授**陈吉栋**，华东师范大学法学院副教授**陈肇新**，中建八局总承包公司总法律顾问**杨有艳**，华东政法大学法律学院博士研究生**赵国彦**分别发表了与谈意见。

最后，上海市数据局数据安全与合规处副处长**杨立哲**对本单元进行了总结。

第二单元 数字司法与数字政府

会议第二单元的主题为“数字司法与数字政府”，由上海市长宁区人民法院副院长**沈建坤**主持。



上海市长宁区人民法院副院长 沈建坤

上海市高级人民法院、三级高级法官**戴曙**以《司法如何捍卫数字正义》为题发表了演讲。

华东政法大学数字法治研究院副教授**于柏华**指出，数据集合缺少确定性或区分性，并非特定化客体，因此不可能在之上设立任何类型排他权。

华东政法大学纪检监察学院特聘副研究员**薛小涵**认为，政务服务跨境通办机制是数字政府建设的重要内容，体现了数字政府服务群众的智慧化。

上海政法学院法律学院副教授**李泽**以《数字政府平台主体间的权益平衡》为题发表了演讲。

在与谈环节，华东政法大学法律学院博士研究生**雷菲萍**、上海青浦银知金融纠纷调解中心理事长**柴丽华**分别发表了与谈意见。

第三单元 平台治理与算法治理

会议第三单元的主题为“平台治理与算法治理”，由华东政法大学中国法治战略研究院教授**杨显滨**主持。



华东政法大学中国法治战略研究院教授 杨显滨

上海市锦天城律师事务所高级合伙人、全国律师网络与高新技术委员会副主任**吴卫明**指出，在算法与大模型治理中，训练数据已成为核心干预切入点。当前治理体系从早期关注生成内容，逐步转向重视训练数据的合规性，这类似于以数据“教育”模型成长。然而，过度严格的数据清洗可能限制模型能力发展，而仅依赖输出端“栅栏”式管控又显被动。因此，需在数据合规与模型能力之间寻求平衡，通过科学的边界设定，既确保伦理与安全，又促进技术健康发展。这一平衡点的探索，需要学术界与立法界持续深入研究。

上海中联律师事务所合伙人、华东政法大学数字法治研究院特聘研究员**方懿**博士以《区块链网络的去中心化治理》为主题发表了演讲。

华东政法大学数字法治研究院特聘副研究员**童云峰**指出，法律既要合理规制大模型技术风险，还要为之塑造容错机制，避免刑法过度适用。

河南财经政法大学法学院讲师**郭新政**以《跨境平台“算法殖民”的生成逻辑与反垄断规制路径——以全球数字治理为视角》为题发表了演讲。

在与谈环节，上海市人民检察院第一分院三级高级检察官**张宏图**、华东政法大学法治战略研究院博士研究生**陈禹竹**分别发表了与谈意见。

第四单元 上海优化营商环境立法

第四单元的主题为“上海优化营商环境立法”，由上海市法学会数字法学研究会副会长、华东政法大学科研处处长、教授**陆宇峰**主持。



上海市法学会数字法学研究会副会长、
华东政法大学科研处处长、教授 陆宇峰

华东政法大学数字法治研究院副院长、副教授**韩旭至**围绕《上海市优化营商环境若干规定（建议稿）》第一章“营商网络基础设施环境”提出了措施建议。

华东政法大学数字法治研究院特聘副研究员**陈超**提出，当前网络“黑嘴”伤企、不实信息扩散等问题突出，严重破坏营商网络环境。为此，相关条文围绕优化营商环境政务服务环境，明确网络平台、商会协会、政府等多方责任，鼓励平台完善言论管理与不实信息处置机制，支持协会建立侵权预警体系，同时建立多方协同培训机制、规范“一网通办”数据管理等机制，精准破解企业维权难、涉企网络暴力治理效率低等痛点，为优化营商环境提供制度依据。

华东政法大学中国法治战略研究院**杨显滨**教授聚焦营商网络人格权保护，提出了针对性建议。

华东政法大学数字法治研究院特聘副研究员**童云峰**指出，为了保护企业主体的网络的人格权，需要赋予市场主体投诉权，市场主体的商业信誉受法律保护。与此相对，需要设置网络服务提供者的处置义务，网络服务提供者应当建立市场主体商业信誉保护机制。

华东政法大学中国法治战略研究院助理研究员**时诚**认为，优化营商环境财产权保护环境，应当重点加强数据权益保护、数据流通交易和网络知识产权保护。

华东政法大学数字法治研究院特聘副研究员**徐则林**指出，第五章聚焦营商环境中的公平竞争问题，旨在提升营商网络公平竞争环境。

在与谈环节，上海市委网信办政策法规处干部**祁东楠**，抖音集团法律研究与合作负责人**丁道勤**，上海大学法学院副院长、副教授**袁曾**，上海政法学院经济法学院副教授**商建刚**分别发表了与谈意见。

在学术总结环节，上海市委网络安全和信息化委员会办公室法规处处长**俞厚未**对优化营商环境立法阶段进行了总结。

最后，上海市法学会数字法学研究会会长、华东政法大学数字法治研究院院长、教授**马长山**对本

次年会进行了学术总结。本次会议虽从小切口入手，却是构建数字法治大厦的重要基石，体现了学者的使命与贡献。

2025年全国数据系统局长培训班成功举办

9月15日至17日，2025年全国数据系统局长培训班在中共湖北省委党校（湖北省行政学院）成功举办。

本次培训班以深入学习贯彻习近平总书记关于人工智能和数据发展的重要论述，加快全国一体化数据市场建设，推动数据要素赋能人工智能创新发展为主题，围绕数据要素市场化配置改革这条主线，从理论和实践层面深入分析阐述，旨在切实提升全国数据系统主要负责同志的战略思维、履职能力和协同水平，形成推动数据事业高质量发展的工作合力。

国家数据局党组书记、局长**刘烈宏**围绕“深入推进数据要素市场化配置改革 扎实做好数据工作 纵深推进人工智能发展”讲授开班第一课，出席全国数据系统投资工作深入研讨，对投资工作做了部署。国家数据局党组成员、副局长**陈荣辉**主持全国数据系统投资工作深入研讨。专题授课环节，中国社科院大学教授**江小涓**、清华大学法学院教授**申卫星**、北京交通大学教授**张向宏**、中国科学院计

算机网络信息中心副主任**周园春**、上海市信息投资股份有限公司副总裁**山栋明**、上海数据交易所总经理**汤奇峰**、银河通用机器人创始人兼CTO**王鹤**、国家数据局数据资源司副司长**宋宪荣**及国内有关人形机器人企业创始人，分别围绕一体化数据市场与分布式应用创新、数据产权制度建设、数据产业培育、AI就绪科学数据赋能人工智能发展、高质量数据集建设、数据市场新趋势新变化、统筹深化数字经济国际合作重点领域和地方试点工作以及智能体机器人的当下和未来等进行专题授课。现场教学环节，学员们实地观摩了湖北人形机器人创新中心、北京大学武汉人工智能研究院、湖北九峰山实验室等在数据要素赋能人工智能创新发展方面的创新应用成果。培训期间，学员围绕培训主题进行了深入研讨。学员们一致认为，本次培训主题鲜明、内容丰富、务实高效，课程具有很强的理论性、专业性和指导性，对于提升数据系统领导干部理论水平和业务能力、推动数据要素市场化配置改革具有重要意义。各省（自治区、直辖市）、计划单列市、新疆生产建设兵团和各省会城市数据管理部门主要负责同志，国家数据局各司负责同志、国家数据发展研究院负责同志共70余人参加培训，湖北省地市数据局主要负责同志列席。

（技术编辑：林诗敏）

数字法评

论生成式人工智能服务提供者的过错推定责任

此处删除了原文脚注，全文请参见《北方法学》2025年第5期，转载或引用请注明出处。

作者：张新宝，卞龙

摘要：生成式人工智能不符合产品与服务区分视角下的产品概念，其功能主要是输出特定内容和信息，服务提供者不应承担产品责任。生成式人工智能不符合适用无过错责任的基本前提，虽然服务提供者可以更有效地控制技术风险，但是生成式人工智能通常不存在较高程度的危险性，不宜出于保护公众权益的考虑要求服务提供者内化全部的权益侵害风险。鉴于破除受害人举证困境、适度强化服务提供者义务和责任等需要，应当引入过错推定责任。服务提供者可以通过证明其在训练端和输出端都尽到了与避免侵权输出相关的注意义务来推翻过错。注意义务的标准应适当严格，以激励服务提供者采取充分的风险预防措施。

生成式人工智能治理不仅需要事前的监管，而且需要明确的事后责任规则。责任规则通过塑造相应的激励机制影响新技术的发展方向，预防风险并提供救济机制。^[1]其中，归责原则的确定是依法准确认定服务提供者侵权责任的关键，决定着民事责任的构成与配置，体现着民法的公平与正义价值。^[2]合理确定服务提供者适用的归责原则，应当以生成式人工智能的技术原理为分析基础，不仅需要避免在侵权法理论内部出现逻辑矛盾，而且需要回应侵权风险预防和受害人救济的现实需求。《欧盟产品责任指令》将软件纳入产品的范畴，可能会对生成式人工智能提供者适用严格责任。我国的《生成式人工智能服务管理暂行办法》等文件则将生成式人工智能提供者视为服务提供者，认为生成式人工智能属于服务而非产品。理论界对于生成式人工智能侵权应适用的归责原则也存在较大的分歧，需要更充分地讨论以形成共识。本文将明确指出，为何不应将生成式人工智能纳入产品的范畴继而适用

产品责任，或者要求服务提供者承担直接的无过错责任，而是应当要求其承担过错推定责任。本文将明确服务提供者如何证明其尽到了相应的注意义务，以推翻法律对其过错的推定。

一、生成式人工智能侵权不应适用产品责任

欧盟委员会于2022年通过制定《人工智能责任指令》和修订《产品责任指令》的提案，以解决人工智能的责任问题。^[3]修订后的《产品责任指令》将人工智能在内的软件纳入产品责任的适用范围，虽然没有明确指出是否适用于生成式人工智能，但是至少没有将其明确排除。对此，欧洲议会研究服务处（EPRS）的补充评估指出，《产品责任指令》可能无法解决生成式人工智能的责任问题，需要通过扩展《人工智能责任指令》的适用范围来解决。^[4]理论界对于该问题的认识也不一致。主张适用产品责任的观点认为，区分人工智能产品和服务不可行也不合理，包括生成式人工智能在内的人工智能提供者应当承担产品责任。^[5]另有观点认为，用于医疗、导航等专业领域的生成式人工智能可以适用产品责任，用于生成诗歌、日常问答等则不适用产品责任。^[6]本文认为，生成式人工智能侵权不应适用产品责任。

（一）生成式人工智能不符合法律上的产品概念

传统产品责任制度中产品的概念被包含于物的概念，^[7]信息技术的发展使得有体物之外的客体也可以成为产品，不过应符合“用于销售”的概念要素。我国《产品质量法》第2条规定：“本法所称产品是指经过加工、制作，用于销售的产品。”

“销售”通常被理解为以经济目的而将产品以出卖的方式交付他人的行为，^[8]其概念的核心是“交付”，构成了销售产品与提供服务的本质区别——交易的主要目的是转移所有权。产品的销售过程表现为出卖人向买受人交付特定产品、买受人支付对应的价款，发生产品所有权的转移。但是，提供服务的过程通常不存在价款之外的交付行为，不会发生所有权的转移。如果服务的内容是物的定制，虽然最终目的也是物的交付，但是整个过程仍属于服务；

如果服务过程附带了物的交付,例如商家在提供服务过程中附赠了产品,应当认为该产品独立于服务过程,可以对其适用产品责任。产品的概念保持着开放状态,信息技术的发展早已使产品摆脱了有形的界限,无形与有形并非区分产品与服务的有效标准。数据、虚拟财产等的价值已经得到理论和实践的认可,并且成为可以被批量交易的对象。其所有权可以在交易过程中发生转移,满足产品概念中“用于销售”的要求,只要同时满足“经过加工、制作”的要件,便可以构成产品。

下载和使用生成式人工智能软件时不存在交付过程,使用者支付了相应费用之后并未取得完整软件副本的所有权,仅取得了通过互联网接入并使用软件的资格,所以不应当将其纳入产品的范畴。软件作为智力成果存在适用产品责任的空间,^[9]但是并非在所有情况下都可以构成产品。软件不仅可能作为产品被用于销售,而且可能只是被用作提供服务的工具,应当结合具体情况进行判断。^[10]软件通常不会造成大范围的损害,不会危及消费者之外的主体,无法适用产品责任的情形可以通过违约责任解决赔偿问题,受害人的权益也能够得到保障。软件的类型主要包括下载到计算机本地并且完全在本地运行的软件、下载到本地但是需要通过接入互联网使用的软件,以及不需要下载直接通过互联网接入并使用的软件。对于不需要接入互联网使用的软件,用户将该软件完整的程序和文档下载到本地服务器,取得了该软件的完整副本。对于需要接入互联网使用的软件,用户仅下载了运行该软件的必要程序和文档,并非该软件的完整副本,不同用户接入的实际上是同一个软件,所以并不符合产品的概念。过去以光盘作为载体并且可以直接在本地运行的软件,比较符合产品的概念,如今日常使用的各种应用软件则更符合服务的特征。另外,若本身不符合产品定义的软件内置于有形产品,内置的软件可以决定该产品的安全性,可以考虑对“有形产品+软件”的整体适用产品责任。基于上述原因,生成式人工智能在当前的产业模式下应当被定性为服务而非产品。一方面,生成式人工智能软件本

身不符合产品的定义。生成式人工智能软件可以通过下载客户端使用,或者直接在网页端使用,但是使用者并没有取得软件副本的所有权,仅是通过互联网接入提供者的服务器,不同使用者使用的仍是同一个软件。生成式人工智能服务提供者往往是按照字符数量或者使用时长收取费用,不符合产品销售的特征。使用者支付的是使用软件的费用,而非购买软件副本的对价。另一方面,内置的生成式人工智能软件无法决定有形产品的安全性。生成式人工智能无法驱动有形产品,无法对他人的生命、健康、财产等权益造成损害,不应有形产品与内置软件形成的整体适用产品责任。

(二) 信息错误造成他人损害不应承担产品责任

错误信息不具有直接的危险性,适用产品责任有失妥当。过去,法院几乎不会通过产品责任制度来要求出版商等主体对错误信息造成的损害承担责任。^[11]主要有两种路径可以对此作出解释:第一,信息可以构成产品,但是通常不可能存在缺陷,所以相关主体不需要对损害进行赔偿;^[12]第二,信息本身无法构成产品,所以不适用产品责任。若采第一种路径,可以对事实性信息、应用型信息等错误造成的损害适用产品责任;若采第二种路径,应当根据过错责任原则解决此类错误信息造成损害的赔偿问题。对于思想、观点等信息,无论是否将其视作产品,可能都难以通过产品责任追究生产者和销售者的责任。因为涉及言论自由的保护,出版商等无须对思想和观点的准确性负责,难以认为其存在缺陷。但是事实性、应用性等信息存在准确性要求,发生错误可能会被视为缺陷产品。相应的观点认为:信息产品是大批量生产、大规模销售的产品,出版者、发行者有能力分散损失,而且信息提供者相较于消费者处于信息的优势地位,更有能力防范风险。^[13]美国即存在将航空地图视作产品并判决出版商承担严格责任的案例。^[14]本文认为,信息不构成产品责任制度中的产品,因为信息错误不具备直接的危险性。相关主体应当对错误信息承担过错责任,而非产品责任。^[15]纯粹的信息不会直接造成损

害,需要以自然人的行为作为介质才能对外界产生作用。^[16]信息接收者需要甄别信息的真实性,并对误信特定信息真实性造成的损害负责,除非信息提供者负有保证该信息真实性的义务。信息可以被无限复制,具有解释上的诸多可能性,要求信息的提供者对信息错误间接引发的损害承担无过错责任,无疑会极大抑制信息的生产和传播。即使是主张构建信息产品责任制度的观点也认为,不对生产者与销售者适用无过错责任。^[17]比较法上,出版商等对于因瑕疵信息而给消费者造成的损害,通常也是以其违反了注意义务为由承担过失侵权责任。^[18]

生成式人工智能的功能是输出特定内容,服务提供者不对信息错误引发的损害承担产品责任。生成式人工智能输出的内容是对用户提问的创造性回答,应当由使用者对信息的可信度进行判断。

“生成式人工智能的运行全过程处于黑箱状态,存在超越、背离和独立于人的行为的倾向性,由于运行和变换的不客观及不被理解,算法黑箱的存在也让生成式人工智能中出现了逻辑隐层,即人工智能算法系统逻辑处于不可知的隐蔽状态。”^[19]基础大模型通常包含数千亿参数,映射规则相当复杂,输入与输出之间的逻辑关系难以被清晰解读。服务提供者当前技术发展阶段不可能确保输出的准确性,尤其是对专业问题的回答。即使生成式人工智能被应用于正式场合,生成的内容也只是作为人类决策的参考,无法替代人类的专业判断。应当看到,安全和可靠不过是生成式人工智能的发展目标,服务提供者并没有法律上的义务确保信息准确,除非事先对此作出了明确承诺。如果要求服务提供者承担产品责任,意味着需要对所有误信生成内容而产生的损害负责,无疑会使其面临过重且难以预测的潜在风险。

(三) 适用产品责任难以达到理想的规制效果

产品责任制度的核心目的在于控制产品批量生产和销售造成的大范围风险,最大限度地保障公众的人身和财产安全,同时给予受害人充分的救济。^[20]站在功能主义的立场,即使生成式人工智能不满足产品的定义和特征,仍然有可能将其纳入产品责

任的适用范围,以达到强化受害人权益保护的目标。生成式人工智能意外输出侵权信息是现阶段无法完全避免的问题,并且会随着技术方案的不断优化逐渐得到解决,目前不能将输出侵权信息的生成式人工智能一概视为有缺陷。当然,若生成式人工智能的安全性没有达到行业的通常水平,即技术发展的当前阶段所普遍应当具有的安全性,则可能会被认定存在缺陷。可是在生成式人工智能侵权的场景,产品责任的制度设计或许并不会像其被应用于有形产品时一样有助于受害人的救济,并且无法高效激励提供者采取措施提高服务的安全性,难以发挥出应有的制度功能。

虽然无过错责任的适用提高了受害人获得赔偿的可能,但是生成式人工智能缺陷的证明会给受害人造成较大的困难。制造缺陷是最容易被证明的产品缺陷类型,适用无过错责任会大幅减轻受害人的证明负担并提高受害人胜诉的概率。然而,人工智能系统的缺陷往往是由设计选择不当造成,包括训练数据中的偏差、不完善的算法等,所以人工智能系统的缺陷更有可能归因于有缺陷的设计而非制造。^[21]于是,适用产品责任制度的实际效果就会有所折扣。相较于制造缺陷,设计缺陷和警示缺陷比较难以证明,因为其中包含了过错证明的成分。美国法律研究院和多数法院对设计缺陷和警示缺陷的证明采用的是“风险—效益”标准,^[22]即通过衡量产品风险的大小与降低风险的成本确定产品是否存在缺陷。^[23]“风险—效益”标准鼓励生产者在产品的效用和可能产生的风险之间、严厉的警告和产品的可欲性之间进行衡量,找到最为合理的解决方案,而不是为了避免事后承担严苛的责任而过分保守。^[24]因为公众没有理由指望生产者为了产品的安全付出过分高昂的成本而牺牲生产效益,生产者有理由在安全与效益之间保持平衡。^[25]对于设计缺陷的判断而言,需要通过衡量更安全设计所需的成本与此种设计可以避免的损失,判断是否存在更合理的替代设计方案。若生产者改进设计所需投入的成本大于其能够降低的风险,则生产者没有义务去改进设计;若生产者能够以合理成本降低产品的

风险,则说明存在合理的替代设计,即产品存在设计缺陷。对于警示缺陷而言,生产者并不需要对产品危险进行过于严厉的警告,因为强烈的警告虽然能避免较多的危险,但是会导致销量的降低。生产者仅须对应当预见到的与产品相关的风险作出适当、合理的警告,其程度应与产品的潜在危险性成正比。可见,“风险—效益”标准实质上就是要证明生产者存在过错。^[26]生成式人工智能技术对于普通公众而言难以理解,恐怕难以证明缺陷的存在。尤其是设计缺陷的证明,带给受害人的挑战最为严峻。^[27]所以,适用产品责任或许并不会使受害人更容易获得赔偿。作为配套制度,可能还需要通过引入表见证明规则、^[28]改变对设计缺陷和警示缺陷的认定标准,^[29]或者建立证据披露与缺陷推定制度等方式,减轻受害人在缺陷证明上的负担。

二、服务提供者不应承担无过错责任

目前有观点认为,应当以无过错责任原则追究生成式人工智能服务提供者的侵权责任,^[30]以抵消人工智能固有的危险,激励服务提供者减少有害的人工智能活动。^[31]另有观点认为,作为一种新型社会性风险的制造者,生成式人工智能服务提供者应适用无过错责任承担著作权侵权损害赔偿赔偿责任。^[32]如上所述,生成式人工智能服务提供者不应为信息错误造成的损害承担侵权责任。分析是否需要通过无过错责任追究服务提供者的侵权责任,主要是考虑输出侵害个人信息、著作权、商业秘密等权益的信息造成损害的场景。提供生成式人工智能服务不属于《民法典》所规定的适用无过错责任的具体责任类型。而且,提供生成式人工智能服务难以被理解为“高度危险作业”,无法适用《民法典》第1236条关于高度危险责任的一般条款。^[33]可见,生成式人工智能服务提供者的无过错责任无法从既有制度中找到合适依据,需要考虑是否需要将其规定为新型的无过错责任类型。本文认为,无论是基于无过错责任的一般原理还是基于特殊的法政策考量,生成式人工智能服务提供者都不应当承担无过错责任。

(一) 生成式人工智能不具有较高级别的危险

性

主张生成式人工智能服务提供者承担无过错责任的观点基本是从受害人保护、危险控制、分散损失,以及从行为中获益等角度进行阐述,^[34]但是忽略了对生成式人工智能应用本身危险程度的分析。适用无过错责任的归责事由包括危险与控制力。^[35]其中,危险责任的根源是可能给不特定主体造成损害的特殊危险。^[36]特定活动或者物的危险性只有达到了比较高的程度,难以通过尽到合理注意有效控制,行为人或者物的所有者等才具备承担无过错责任的基本前提。^[37]危险责任的理论基础在于,为了自己的利益开启和控制较高程度危险的主体,应当对因此产生的损害负责。^[38]危险活动和危险物具有较高致人损害的可能性,如果相关主体不存在过错但造成了损害后果,不予赔偿显然有失公平,因此需要引入无过错责任来解决不幸损害的分配问题。^[39]侵权责任法正是从分配正义的理念入手,从损害发生的危险源上来确定应由何方承担损害。^[40]作为过错归责的例外情形,无过错责任原则必须有极其充分的理由才可以适用——行为的危险性达到难以有效控制的程度。^[41]否则,社会生活中危险无处不在,动辄以开启了危险为由要求相关主体承担无过错责任,必然会动摇过错责任的根基。

服务提供者通常不应应对侵权信息承担无过错责任,最核心的理由是生成式人工智能通常不具备较高级别的危险性。生成式人工智能不具备物理实体也无法驱动物理实体,不会如自动驾驶等人工智能应用直接威胁到自然人的生命、健康等人身权益,^[42]而且生成式人工智能软件不会如杀毒软件等通过运行对计算机等有形财产造成损害,可见其侵权风险仅在于输出的内容和信息本身。目前实践中主要存在两类生成式人工智能应用:第一类是直接回答用户的提问或者操作,根据用户的要求直接输出内容的生成式人工智能;第二类是将用户提供的内容作为素材进行转化和再创作的应用,例如将用户提供的图片转换成其他风格,根据用户提供的人物形象和文字生成视频等。转化型生成式人工智能使用的素材只要合法合规,输出的内容通常不会直接

构成侵权。虽然某些深度伪造的内容可能被不法分子用于实施诈骗等侵权和犯罪行为,但是此种风险主要来源于实施下游侵权和犯罪的行为人,不应将具有深度伪造功能的生成式人工智能应用视为需要纳入无过错责任适用范围的危险源。创造型生成式人工智能存在直接的侵权风险,因为其会自主输出新的信息,但是通常也不会达到高度危险的程度。在没有介入使用者行为的情况下,服务提供者只要在训练端依法收集数据、构建算法时充分考虑如何避免侵权信息的输出,生成式人工智能自主输出侵权信息的概率相对有限。如果介入了使用者的侵权行为,例如刻意引导模型输出、实施外部攻击、传播侵权内容等,输出侵权信息的危险性此时主要来源于使用者,服务提供者通过规范使用者行为可予以有效控制。

一个基础模型可能会在一千个生成式人工智能应用中被使用,但是可能只有一个是高风险应用程序。^[43]目前,尚无充分证据表明国内的生成式人工智能应用存在比较严重的侵权问题,实践中也已经出现不使用个人信息或者受著作权保护作品进行训练的应用,生成式人工智能的侵权风险可以得到控制。而且,我国采用大模型备案审查制度,可以有效避免存在特殊风险的生成式人工智能应用进入市场。过去,网络中介服务平台上存在大量侵权信息,而且存在借助网络服务实施其他侵权甚至犯罪行为的情况。由于互联网实名制没有完全落实,网络行为的匿名性等特征容易造成受害人找不到实际侵权人而得不到救济的困境。网络中介服务提供者无疑制造了危险源,甚至其危险程度可能高于生成式人工智能服务。但是,长期以来理论和实务都没有将网络中介服务视作难以控制的危险,进而要求服务提供者承担无过错责任。生成式人工智能服务提供者和网络中介服务提供者虽然在性质上存在一定差异,但是生成式人工智能服务所开启的侵权风险程度并不会明显高于网络中介服务。所以,应当认为生成式人工智能的危险程度较为有限,不宜通过无过错责任原则追究服务提供者的侵权责任。

(二) 具备危险控制能力等不属于归责事由

具备危险控制和损害分散能力、损益相一致原则、降低预防侵权的社会总成本等不属于归责事由,其作用是在相关主体制造了较高程度危险的情况下作为解释适用无过错责任合理性的实质理由。无过错责任原则对受害人的保护是以减损加害人的行动自由为代价,倘若没有充分的理由,必然会造成对加害人不公平的结果。^[44]开启了较高程度危险的行为本身具有较强的可归责性,此时将损害分配给危险的开启者承担,符合公平原则的要求。如果潜在责任主体没有制造较高程度的危险,仅以具备危险控制和损失分散能力等作为依据,无法证成服务提供者应承担无过错责任。应当注意,引入无过错责任并非侵权责任法应对技术风险的唯一途径。安全保障义务的设置同样是出于相关主体具备危险控制和损失分散能力、从行为中获得经济利益、降低预防侵权的社会总成本等方面的考虑。^[45]若潜在责任主体只是开启了一般程度的危险,同时具备危险控制能力,应当考虑要求其承担一定程度的安全保障义务,而非要求其承担无过错责任。同样以网络中介服务提供者作为参考,网络中介服务提供者开启了网络空间的权益侵害风险,具备技术和经济能力进行内容审核来控制危险,可以通过价格机制和保险机制等分散损害,网络中介服务提供者和生成式人工智能服务提供者在控制危险和分散损害的能力上不存在显著差别,但是我国法律没有要求网络中介服务提供者承担无过错责任。同样,不能仅因为生成式人工智能服务提供者具备较强的危险控制和损失分散能力就要求其承担无过错责任。

无过错责任原则表面上能够简化生成式人工智能致害的损害赔偿问题,但是不得不警惕,

任何技术在早期阶段可能都会带来不可预测的危险,研发者、经营者等主体相较于普通公众都具有比较强的危险控制和损失分散能力,而且其目的通常也包括从研发和经营活动中获得经济利益。按照此种以危险控制等为核心的论证模式,未来商业性的技术研发者和经营者等可能都要为技术应

用而造成的他人损失承担无过错责任。届时,归责的主要考虑因素不再是过错。只要开启了危险且具备较强的技术和经济能力,便可以出于降低技术风险等考虑要求相关主体承担无过错责任,不仅不符合基本法理,而且会极大制约技术的发展。

(三)可能造成技术与权益保护的失衡

通过上述分析可以认为,生成式人工智能服务不具备适用无过错责任的前提条件。但是有无必要出于受害人保护等方面的考量,例外地要求服务提供者承担无过错责任?主张无过错责任的观点还关注服务提供者和使用者皆无过错,但是对于输出侵权信息造成他人损害情形下的损害赔偿问题,认为受害人的损害在此种情形下得不到救济,显然有失公平。^[46]生成式人工智能目前存在意外输出的问题,可能会因过拟合等原因输出训练语料中的未公开个人信息、作品内容,暂时难以完全通过技术手段避免此种意外记忆。如果坚持过错追责,受害人在服务提供者无过错但是生成式人工智能输出了侵权信息的情形无法获得救济。本文认为不应以此为由要求服务提供者承担无过错责任。首先,适用更严格的归责原则固然会更有利于受害人的保护和损害的预防,但是会相应地限制行为人的行为自由和加重加害人的责任。保护受害人和预防损害并不能作为适用无过错责任原则的理由,而只能作为其目的。其次,意外输出问题仍然是低概率事件。目前国内外涉及侵权输出的案件和事件,几乎都是因为服务提供者未经同意使用他人未公开的个人信息或者作品训练模型,并且没有尽可能在输出端避免生成侵权内容,服务提供者在这些情形都存在过错。再次,可以要求服务提供者尽到较高标准的注意义务,采取相应的技术措施来减少意外输出。目前不能判断在服务提供者尽到较高程度的注意义务之后是否还会存在意外输出的情况。随着语料和算法的优化、服务提供者的义务得到明确和落实,可以预见意外输出风险将会不断降低。最后,行为人无过错、受害人得不到救济的情况在其他场景下同样会存在,并不会造成不公平的后果。不应刻意强调和放大生成式人工智能服务提供者无过

错却造成损害的例外情形,继而以救济受害人为理由引入无过错责任。

面对生成式人工智能等技术带来的权益侵害风险,侵权责任法不能过度强调受害人的保护和救济,否则可能会造成技术与潜在受害人保护之间的失衡。生成式人工智能服务具有极大的正外部性,不宜要求服务提供者完全内化负外部性。信息技术和应用通常都会带来权益侵害风险,但是技术研发毕竟不同于其他纯粹营利的活动,具有较高的公共价值,应当在技术发展初期允许合理水平的风险存在,不宜在较早阶段就通过加重服务提供者责任的方式降低权益侵害风险。政府和社会公众在生成式人工智能发展过程中都是受益者,应当共享技术发展的成果,共担技术发展的风险,不应将技术风险造成的所有不利后果都分配给服务提供者。无过错责任虽然有利于受害人获得赔偿,但是可能会导致“寒蝉效应”,导致某些生成式人工智能模型无法落地应用,减缓甚至阻碍技术的创新和发展。^[47]基于技术发展的考虑,对于生成式人工智能服务提供者仍应坚持过错追责。^[48]此外,行政监管和行政责任的承担在技术风险的化解中发挥着重要作用,尤其是对于信息技术所造成的大规模风险,行政手段相较于侵权责任法可能更为有效,侵权责任法不需要也没有能力完全化解生成式人工智能带来的权益侵害风险。

三、服务提供者应承担过错推定责任

生成式人工智能应用仅具有一般程度的危险性,过错责任原则足以应对服务提供者的侵权责任问题,不过需要考虑应当适用一般过错责任还是过错推定责任。考虑到潜在受害人客观存在的举证困境、训练端的数据制度构建生成式人工智能的技术特征、已经存在的制度基础等因素,引入过错推定责任来追究服务提供者的侵权责任是比较适当的选择。

(一)破除举证困境的客观需要

公众的权益能否得到保障需要考察两个方面:第一,相关制度是否能够激励服务提供者采取充分的措施防范侵权风险;第二,发生侵权事件时受害

人能否顺利获得救济。于前者而言,虽然服务提供者在训练端对个人信息、作品等内容的使用会在输出端引发一定的侵权风险,但是只要明确服务提供者应当在输出端尽到较高标准的注意义务,服务提供者便会尽可能通过优化算法和语料库等方式来提高模型的安全性,以防范侵权风险的发生。于后者而言,受害人可能会在过错证明上面临较大困难,导致一般过错责任原则无法保障受害人顺利获得救济。同时,受害人难以依法获得赔偿会使服务提供者采取预防措施的激励减弱。

服务提供者的过错在部分情况下比较容易证明,能够根据服务的具体内容和细节进行判断。以陈某与上海易某网络科技有限公司侵害作品信息网络传播权纠纷案为例,^[49]被告上海易某网络科技有限公司提供的服务内容是使用他人拍摄的视频作品为用户进行换脸,然后输出换脸后的视频。被告显然在明知是他人作品的情况下进行使用,能够从被告的行为判断出其存在侵害他人著作权的故意。服务提供者的过错在部分情况下比较难以证明,因为其并未实施能够体现主观过错的积极行为。创造型生成式人工智能的输出过程具有概率性和随机性,输出侵权信息并不能说明服务提供者存在过错,受害人还需要另行证明服务提供者未尽到注意义务。然而,缺乏专业知识和能力的普通公众可能难以顺利完成举证:首先,服务提供者在训练端的过错表现为没有依法使用训练数据,输出端的过错表现为对注意义务的违反,不具备专业知识的受害人可能难以知晓服务提供者需要尽到的义务内容。其次,用以证明加害人过错的证据往往处于服务提供者自己的控制之下,相关的证据内容可能存在较高的专业知识壁垒。若要求受害人对服务提供者的过错进行证明,举证可能性相当有限,继而会使受害人面临较高败诉风险。所以,需要通过举证责任倒置保障受害人能够顺利获得救济。《欧盟人工智能责任指令(提案)》规定了证据披露制度,^[50]以缓解受害人的证明压力。考虑到相关技术术语对普通公众而言难以理解,即使披露了相关技术事实,受害人是否有能力从中提炼出能够证明加害人过

错的有效证据存在一定疑问,采用过错推定更能解决受害人的举证难题。

(二) 符合公平原则的内在要求

举证责任倒置是对“正置”举证责任的局部修正,有利于案件实体真实地再现,符合民事诉讼实体公正与程序公正的价值目标。^[51]过错推定责任的目的并不一定是加重加害人的责任,而可能主要在于避免出现对受害人不公平的局面。首先,生成式人工智能服务提供者更接近案件相关证据,具备更高的专业能力和水平,^[52]过错推定只不过是矫正了双方的“武器不平等”状态。现代民法以社会为本位,强调具体正义,过错推定便是此种社会保护思想的体现。^[53]表面上看,过错推定责任将举证不能的风险转移给加害人承担,可能会不当加重加害人的责任。但是,网络环境下作为证据的事实对生成式人工智能服务提供者而言比较容易收集和保存,举证不能的概率比较低。其次,服务提供者作为侵权风险的开启者、控制者和经济上的受益者,适当地增加其义务和责任以保护公众的安全,符合公平原则的要求。再次,目前普遍认为应当在训练端的制度安排上适当侧重语料的获取和使用,相应地应当要求服务提供者在输出端妥当保护公众的合法权益。^[54]基于整个训练和服务过程的统筹考量,要求服务提供者在输出端承担相对更重的义务和责任,以降低甚至消除在训练端造成的权益侵害风险,能够实现整体法律制度设计的公平与合理。

普通公众在传统侵权法律关系中对于证据的掌握情况和举证能力比较接近,“谁主张谁举证”的一般过错责任原则具有天然的合理性。然而,信息业者作为侵权人、普通公众作为被侵权人的新型侵权法律关系具有一个突出特征,即双方技术能力明显不对等导致彼此证据掌握情况和举证能力差距悬殊,需要考虑此时以一般过错责任为原则是否妥当。尤其是在公众权益保护已经面临较大挑战的时代背景之下,信息业者的过错推定责任是更公平的选择。

(三) 符合生成式人工智能的技术特征

创造型的生成式人工智能输出侵权信息主要

存在两种可能性：第一，模型学习了个人信息、作品、商业秘密等，并且在过拟合等情况下意外输出了与语料相同或相似的内容。第二，模型没有学习相应的个人信息等内容，但是偶然“创造”出了侵权信息。如果是第一种情形，服务提供者有义务采取措施避免侵权输出，否则应当认为其存在过错；如果是第二种情形，服务提供者对于侵权内容的输出通常不存在过错。根据生成式人工智能的技术原理，模型通过学习确定各个参数的映射规则，然后基于其学习后形成的功能回答使用者的提问。假设生成式人工智能未学习相应的作品等而输出了侵权信息，意味着其完全是基于概率而创造出了与语料相似或者完全相同的内容，此种可能性相对较低。尤其是侵害著作权、商业秘密等情况，因为其内容构成相对比较复杂，模型在未学习相应内容的情况下偶然“创造”出侵权信息，应当属于极小概率事件。而且，生成式人工智能训练需要尽可能丰富的语料，客观的行业需求决定了服务提供者会尽量广泛地收集高质量数据作为训练语料。所以，侵权信息的输出更有可能是因为服务提供者使用了相应内容作为语料，而非纯粹的偶然。此外，服务提供者只要在输出端充分采取了过滤等技术措施，生成式人工智能输出侵权信息的概率应当比较低，侵权信息的生成更有可能是因为服务提供者未妥当尽到注意义务。根据以上技术特征，推定服务提供者在生成式人工智能输出侵权信息时存在过错具有合理性。

（四）已经存在相应的制度基础

生成式人工智能侵害的主要是非物质人格权益和知识产权等，过错推定责任已经在部分法律规范或者司法实践中得到认可。首先，我国《个人信息保护法》第69条已经明确规定了侵害个人信息权益的过错推定责任，受害人无须对加害人过错进行证明。其次，司法实践中对于著作权侵权的认定已经存在推定过错的做法，具备直接适用过错推定责任的合理性基础。实践中判断侵害著作权等普遍采用的是“实质性相似 + 接触”规则，^[55]不再单独考虑行为人的过错。原告只要证明特定内容与其在

先作品构成实质性相似并且被告具备接触其在先作品的条件和事实，便可以认定被告侵害了其著作权。具体而言，构成实质性相似通常代表着存在权利损害事实，并且存在侵害著作权的行为。行为人接触在先作品，意味着在后的作品可能并非作者独创。^[56]由于接触与否的事实主要由行为人掌控，司法实践中常以作品已经公开发表为由，推定接触的存在，并由行为人承担不存在接触的证明义务。^[57]于是，司法实践中对行为人过错的认定往往以接触事实的认定作为代替，甚至简化为作品已经发表的客观事实认定。虽然此种过错推定是司法推定，不完全等同于法律直接作出的推定，但是都会产生无须证明加害人过错的效果。作为生成式人工智能侵权的重点问题，侵害著作权和个人信息权益在立法或司法中已经采纳了过错推定，规定生成式人工智能服务提供者的过错推定责任具有一定的制度基础，可以实现类似问题处理上的一致性和连贯性。

四、服务提供者对无过错的证明

生成式人工智能服务提供者的过错需要从其客观的行为进行判断，难以直接考察其主观状态。服务提供者对生成式人工智能可能会输出侵权内容具有概括的预见能力，应当尽到合理的注意义务。此种以注意义务为重心的侵权追责模式可以为生成式人工智能服务提供者提供稳定的预期，^[58]兼顾受害人的保护与产业的发展。过错推定原则意味着若生成式人工智能输出侵权信息造成损害，直接推定服务提供者未尽到注意义务。^[59]服务提供者可以通过证明在训练端使用语料符合相关规定并且在输出端尽到了注意义务来证明其不存在过错。

（一）服务提供者在训练端的行为规范

服务提供者获取语料时需要遵守《著作权法》和《个人信息保护法》等法律法规，但是如果严格适用相关规定，服务提供者将会面临语料获取的困境，不符合促进数据要素利用的政策导向。鉴于生成式人工智能的发展需求，应当在训练端慎重认定相关主体的过错，以保障服务提供者能够获得充分的训练数据。对于作品使用行为而言，应当将其纳入合理使用的适用范围，但是需要确保不会输出影

响原作品阅读市场的内容。^[60]如果服务提供者未经许可使用了他人作品而模型未输出影响原作品阅读市场的内容,应当认为服务提供者不存在过错。对于个人信息使用行为而言,依据促进数据要素利用的政策导向,应通过宽松解释《个人信息保护法》的相关规定或者作出例外规定,缓解对个人信息语料获取行为的限制。^[61]如今尚未建立生成式人工智能的专门制度,法院应当在适用相关规定时灵活处理。^[62]

虽然需要在训练端适当侧重于技术的创新发展,但是服务提供者仍然需要遵守相关的行为规范。个人信息保护方面,服务提供者需要注意以下问题:其一,需要遵守“告知—同意”规则,不得在完全未经个人同意的情况下处理未公开的个人信息。虽然无须在每次进行个人信息数据流转时都取得个人同意,但是至少应当在个人信息收集时告知个人可能会将其个人信息用于不特定场景的生成式人工智能训练并取得同意。其二,需要受到必要性原则的限制。服务提供者需要尽可能丰富的数据作为语料,必要性原则似乎难以适用。若能够在实现目标功能的前提下尽可能使用更少的个人信息数据,则仍然需要受到必要性原则的限制。其三,只有在具备充分必要性的情况下才可处理敏感个人信息,避免意外输出给自然人的尊严、人身和财产等造成比较重大的损害。著作权保护方面,机器学习行为通常不会影响作品的正常使用、不会不合理地损害著作权人的合法权益,服务提供者通常可以基于合理使用制度来使用作品,但是应避免输出影响原作品传统阅读市场的内容。于特殊情形,机器学习对作品的使用会不合理地损害著作权人的合法权益,服务提供者需要依法取得著作权人的授权许可。^[63]

(二) 服务提供者在输出端的注意义务

生成式人工智能服务提供者在整个服务周期中需要尽到各种不同内容和目的义务,违反义务与侵权输出之间并非都存在因果关系。例如,服务提供者负有保障服务透明度的义务,但是违反该义务不会造成侵权输出。服务提供者欲证明其不存在过

错,仅须证明其尽到了与避免侵权输出相关的注意义务,主要包括:其一,采取优化算法、关键词过滤等有效的技术手段以避免侵权输出;其二,提醒和规范使用者的行为,避免引导模型输出侵权信息;其三,及时检测和修补安全漏洞,提高模型安全性;其四,收到权利人通知时采取必要措施,避免侵权信息再次生成,建立便捷的通知渠道,及时处理受害人的请求。^[64]问题的关键在于注意义务标准的明确,决定服务提供者证明其尽到了何种程度的注意义务才能推翻对过错的推定。

较高标准的注意义务可以促使服务提供者提供更安全的模型,采取更充分的技术措施减少侵权信息的生成,但是相应地会增加服务提供者的负担,包括技术研发投入的成本和可能承担的责任。本文认为,要求服务提供者在输出端尽到较高标准的注意义务具有合理性。^[65]原因主要有三个方面:第一,服务提供者的数据使用行为开启了非相互性风险,应当承担较高标准的注意义务,以尽可能内化风险、保护公众的合法权益。公众无力预防生成式人工智能造成的侵权风险,或者需要付出不成比例的成本,所以承担较高标准的注意义务是降低预防侵权社会总成本的需要。第二,尽到较高标准的注意义务是兼顾产业发展与权益保护的必然选择。训练端的制度构建需要重点考虑语料获取需求,侧重于技术的创新发展,相应地需要在输出端要求服务提供者尽到较高标准的注意义务,尽可能消除训练活动造成的权益侵害风险,避免造成对权益保护的失衡。第三,服务提供者尽到较高标准的注意义务是保障语料获取需求的需要。如果服务提供者不能在输出端妥当保护公众的合法权益,其在训练端享受“优待”的合理性就不再充分,可能会影响训练端语料获取相关制度的构建和运行。出于上述考虑,不能仅从风险预防的效率来判断服务提供者是否尽到了充分的注意义务,即使风险预防成本高于可能收益也不能当然认为超出了合理的注意义务水平。

注意义务的标准也不应过高,应以激励服务提供者采取充分的预防措施为目标。标准过于严格会导致服务提供者几乎不能反证无过错,以致于达到

无过错责任的效果。“通过设定较重义务来降低风险固然可以达成实效,但也应考虑到在技术发展的现实需要下,适度容忍合理风险的必要性。”^[66]目前生成式人工智能研发仍然面临技术上的难点,包括训练数据的筛选、算法的建构,不可能做到绝对的安全。所以注意义务的设置是为了在现有技术能力下尽可能降低风险,而非完全消除风险。服务提供者是否采取了充分的注意义务应当结合以下因素判断:第一,服务提供者的规模大小和技术水平;第二,生成式人工智能服务所处的具体领域;第三,当前相关技术措施所处的发展阶段。分别来看,企业的规模越大,采取的技术措施就应当越充分;提供服务的类型存在越高的侵权风险、涉及的权益越重要,承担的注意义务标准就应越高;服务提供者应当采取与技术发展水平相符合的技术措施,并且随着技术的发展不断更新和优化。综合来看,不同规模的服务提供者的技术能力也会存在差异,不能要求较小规模的服务提供者与较大规模的服务提供者投入相同的成本来提高其服务的安全性。但如果提供的是关键领域或特殊用途的生成式人工智能服务,公众对服务安全有更高的要求,应当对不同规模的服务提供者等量齐观,以尽可能确保服务的安全性。

(三) 服务提供者证明无过错的具体方法

服务提供者通常需要证明其在训练端和输出端都尽到了与避免输出侵权信息相关的注意义务,才能推翻对过错的推定。此外,服务提供者也可以通过证明没有使用相关作品、个人信息等内容作为语料来证明无过错,因为此时模型是基于其创作功能偶然输出了涉嫌侵权的内容。以著作权和个人信息侵权为例进行说明。对于著作权而言,服务提供者可以通过证明其没有使用原作品或者已经在输出端尽到了相应的注意义务来反证过错不存在。训练语料的范围就是生成式人工智能模型所接触到的信息范围,若服务提供者能够证明训练语料中不包括该原作品,则能够推翻过错。但是由于服务提供者可能通过删除相关数据、修改记录等方式“伪造”证据,所以更可靠的做法是要求服务提供者证

明不可能接触原作品来证明其没有使用原作品作为语料。对于个人信息而言,服务提供者需要证明使用个人信息符合相关法律规定,并且在输出端尽到了注意义务来反证不存在过错,或者通过证明未使用相关个人信息来证明无过错。

还需要注意,服务提供者在训练端的行为会影响其在输出端的义务。前文对注意义务的讨论建立在给予服务提供者数据使用行为一定豁免或例外规定的基础之上,由于其在训练端的数据使用行为未经充分同意或者授权许可,服务提供者应负有避免侵权输出的义务,消除其制造的权益侵害风险。但是若服务提供者在训练端告知权利人将会使用其作品或未公开个人信息等训练生成式人工智能并且模型会输出相关内容,并且已经取得了权利人授权许可或同意,则服务提供者不再负有避免输出相关内容的义务。例如,服务提供者在训练端使用作品时取得了著作权人的授权许可,而未基于合理使用,并且对模型可能会输出相关内容进行了充分告知,则其在输出端可以在授权许可的范围内输出与原作品构成实质性相似的内容等。所以,服务提供者证明其取得了著作权人的授权即可推翻对过错的推定,无须再证明其在输出端尽到了注意义务。同样,如果服务提供者可以证明已经在训练端告知个人模型会输出其未公开的个人信息,并且取得了其单独同意,则服务提供者不负有避免输出个人信息的注意义务,仅需证明其在训练端的个人信息处理行为满足其他规定,便可以推翻对过错的推定。

五、结论

生成式人工智能民事法律制度的建构,应当在输出端强调对公众权益的保护,明确服务提供者的义务和侵权责任,尽可能降低权益侵害风险,保障受害人顺利获得救济。生成式人工智能本质上属于服务而非产品,产品责任的制度设计无法有效应对生成式人工智能带来的权益侵害风险,所以将其纳入产品责任的适用范围并非有效的解决方案。无过错责任通常以行为的高度危险性作为归责事由,应当谨慎适用,不应将其视为纯粹的风险控制工具。适当严格的归责原则和注意义务设置可以促使生成式人工智能服务提供者研究并主动采取有效的措施来避免侵权内容的生成,不仅可以在较高程度

上提高技术的安全性，而且可以避免对服务提供者“活动程度”的不当限制。过错推定责任的适用以解决受害人的举证困境为主要目标，同时也适当加强了服务提供者的责任。注意义务内容和标准的确

定可以为服务提供者保障公众合法权益提供行为指引，同时为其反证无过错提供明确依据。总之，“过错推定责任 + 较高标准注意义务”的制度组合是应对生成式人工智能侵权风险的理想选择。

论大模型训练数据的合理使用

此处删除了原文脚注，全文请参见《法学家》2025年第5期，转载或引用请注明出处。

作者：李铭轩

摘要：大模型训练数据的主要来源是网络上的公开数据，开发者一般通过爬取公开网页和收集开源数据来大规模获取训练数据。随着数据财产权益保护的强化，获取海量训练数据的主要方式面临着合法性挑战。数据财产权益人众多、数据使用行为难追溯导致交易成本升高，大模型开发者无法通过市场机制获得数据财产权益人的许可来确保训练数据的合法性。在市场失灵的情形下，允许开发者合理使用数据进行大模型训练，可以增进社会福利，且一般不会损害数据财产权益人的市场利益。采取集体管理或法定许可等替代方案给数据财产权益人带来的收益非常有限，却会产生更高的制度成本，并给我国大模型的发展造成不利影响。因此，我国应当建立大模型训练数据的合理使用制度，为技术发展提供合法性预期。在规则设计上，大模型训练数据合理使用的对象应限于公开数据；目的应限于预训练；方式应包括训练涉及的数据处理行为；应允许数据财产权益人以技术措施选择退出合理使用。

引言

近年来，ChatGPT、DeepSeek等基于大语言模型（以下简称“大模型”）的人工智能应用快速发展，深刻地改变着人们的生活。大模型是“一种由包含数百亿个及以上参数的深度神经网络构建的语言模型”。^[1]与以往的人工智能模型不同，大模型展现出极为强大的语言能力，并具备处理不同任务的通用能力，为实现通用人工智能（即强人工智能）创造可能。^[2]因此，大模型已成为人工智能领域最重要的技术之一。

数据是大模型发展的关键要素。大模型对数据规模和多样性都有更高的需求。大模型的性能存在着“规模定律”（Scaling Law），随着训练数据量的增长，模型的性能也会提升。^[3]例如，美国开放

人工智能（OpenAI）公司的GPT系列模型，其训练数据量从GPT-1的5GB增至GPT-2的40GB，再迅速扩张到GPT-3的45TB。^[4]此外，大模型训练所需的数据类型也日趋多样，包括网页、书籍、百科全书等不同类型的数据库。^[5]训练数据的多样化不仅有利于增强大模型的通用能力，也有助于保障大模型的公平性与包容性。^[6]因此，大规模、多样化的训练数据已成为大模型发展的重要基础。

然而，大模型训练数据的获取和使用面临着法律上的限制。数据不仅是人工智能的要素，也是法律权益的载体。数据之上承载着著作权、个人信息权益、数据财产权益等多种权益。大模型训练数据的获取和使用行为会落入这些权益的控制范围，这意味着大模型开发者未经权益人许可不得擅自获取和使用数据。这会给训练数据的获取和使用带来较大的限制，进而可能影响大模型技术的发展。因此，许多观点提出，针对大模型训练的情形，应当加强对著作权、个人信息权益的限制，消除法律上的可能障碍。^[4]

不过，现有研究较多关注著作权和个人信息权益保护对大模型训练的影响，却鲜有讨论数据财产权益保护带来的问题。但是，数据财产权益保护对训练数据获取和使用的影响不容忽视。即使大模型开发者可以在著作权或者个人信息权益领域享有侵权豁免，但如果数据持有者能够主张数据财产权益，仍有可能限制大模型开发者对数据的获取和使用。因此，从体系的视角来看，如果要消除大模型训练数据获取和使用的法律障碍，亟须全面地考察数据“权利束”中各种权益的潜在影响，并在制度设计上保持协调与一致。^[5]而且，在实践中，涉及数据财产权益的相关纠纷也已发生。2025年6月4日，Reddit起诉Anthropic，指控其未经许可利用爬虫获取其平台内容数据并使用在大模型训练中，构成违约、不当得利和侵权（包括不正当竞争）。^[6]在该案中，作为原告的Reddit本身并非著作权人或个人信息主体，因此其争议焦点并非大模型开发者是否侵害著作权或个人信息权益，而是其是否侵害数据财产权益。

因此,本文将讨论数据财产权益保护在大模型训练场景下引发的问题,并提出合理的解决方案。尽管我国在数据财产权益的制度设计上尚缺乏共识,^[7]但这并不影响本文讨论的前提:大模型训练数据的获取和使用可能会落入数据财产权益的控制范围,而随着数据财产权益保护的强化,训练数据的获取和使用将受到更多的限制。本文认为,面对数据财产权益保护给大模型训练造成的障碍,引入大模型训练数据的合理使用制度,是一种可行的方案。鉴于此,本文拟围绕大模型训练数据的合理使用这一主题,讨论引入这一制度的前提背景和理论基础,并提出规则建构的具体方案。

一、大模型训练数据合理使用的引入背景

在这一部分,本文旨在论证数据财产权益保护可能会导致大模型训练数据市场失灵。这是引入大模型训练数据合理使用的重要背景。本文将首先考察大模型训练数据的主要来源,然后讨论数据财产权益保护对训练数据获取和使用的影响,最后分析由此可能引发的训练数据市场失灵。

(一) 大模型训练数据的主要来源

大模型本质上是语言模型,其核心目标是对自然语言的概率分布进行建模。^[1]以GPT模型为例,其任务是基于已有的词语序列预测下一个词,即构建条件概率分布。模型训练便是从数据中建模的过程,主要是指模型通过接受输入训练数据,对模型参数进行迭代优化,从而最小化模型预测输出与真实数据标签之间的误差。^[2]因此,GPT等大模型的训练数据实质上是词序列(特征)与下一个词(标签)的配对,而网络上海量的公开文本数据能为训练数据的构建提供丰富的资源。^[3]

从来源上看,大模型训练数据主要来自网络上的公开数据。^[4]大模型开发者通过两种方式大规模地获取网络公开数据。第一,爬取公开网页。大模型开发者可以利用网络爬虫等技术大量获取公开网页的数据,并用来构建训练数据集。例如,开放人工智能公司曾通过爬取Reddit上获得至少3个赞的外链网页,构建了一个高质量网页文本数据集WebText,并使用在GPT系列模型的训练中。^[5]目前,

许多大模型开发者都部署了专门的网络爬虫,如开放人工智能公司的GPTBot等,用来大规模地收集公开网页的数据,为构建训练数据集提供丰富的原材料。第二,收集开源数据。开源数据是指基于开源许可证(Open Source License)发布的数据资源,这些数据的发布者一般允许他人自由且免费使用其收集的数据。大模型开发者也经常收集大量的开源数据,并以此为基础构建训练数据集。例如,大模型开发者通常会使用Common Crawl或其衍生的开源数据集。Common Crawl是一个非营利组织,其通过网络爬虫定期抓取整个互联网的网页,积累了PB级别规模的数据,存储在数据库中供公开访问和下载。^[6]由于数据量过于庞大,大模型开发者在构建训练数据集时,一般会从Common Crawl中提取合适的子集,或使用其他基于Common Crawl构建的开源数据集,如Colossal Clean Crawled Corpus(C4)等。^[7]

到目前为止,绝大多数开发者之所以能顺利地构建适于大模型训练的数据集,得益于上述两种数据获取方式事实上“自由且免费”的特征。首先,大模型开发者可以通过这两种方式自由地获取大量的数据,没有为获取许可付出很高的搜索与谈判成本。其次,大模型开发者也没有为这些数据支付对价,节省了大量的费用。受益于此,许多开发者才能够获取和使用数量足够且种类广泛的数据展开大模型训练,推动这一领域持续创新和发展。

(二) 数据财产权益保护及其影响

本文所讨论的数据财产权益,是指数据持有者对数据所享有的对世性财产权益。这一概念不包括著作权和个人信息权益,也不包含合同约定的有关数据的权利。在我国现行法下,数据财产权益主要是指数据持有者受侵权法和反不正当竞争法所保护的对世性财产利益。除对非公开数据的商业秘密保护外,我国法院也主要依据反不正当竞争法为公开数据提供保护。在实践中,数据持有者经常以《反不正当竞争法》第2条(即一般条款)主张他人获取和使用数据的行为构成不正当竞争。在这类案件中,许多法院都认定他人获取和使用数据的行为构

成不正当竞争，侵害了数据持有者的合法权益。^[1]通过一般条款的适用，我国法院事实上创设了一种针对公开数据的“竞争性财产权益”，使数据持有者能够阻止具有竞争关系的其他主体对其数据的不正当获取和使用行为。^[2]2025年6月，新修改的《反不正当竞争法》在第13条新增一款，确认了司法对数据财产权益的保护实践，将侵害数据财产权益的情形具体化为不正当竞争行为的法定类型。^[3]

我国还一直存在着强化数据财产权益保护的呼声。许多观点认为，现有制度对数据财产利益的保护尚存不足，应进一步采权利保护模式，即通过立法事先确认数据持有者对数据享有财产权。^[4]当然，目前这些方案仍然停留于理论探讨阶段，尚未转化为具体法律制度。但从近年来的政策趋势观察，我国正在积极探索数据财产权制度的构建，并可能在未来付诸实践。2022年12月，中共中央、国务院发布《关于构建数据基础制度更好发挥数据要素作用的意见》（以下简称《数据二十条》），明确提出要“建立保障权益、合规使用的数据产权制度”，在国家政策层面释放了支持数据确权的信号，并为未来立法提供了思路和框架。上述迹象表明，未来我国有可能逐步建立数据财产权制度，进一步强化数据财产权益的保护。

随着数据财产权益保护的强化，当前大模型开发者获取海量数据的两种主要方式将面临严峻的合法性挑战。第一，爬取公开网页的合法性空间不断收紧。爬取公开网页的行为很可能会落入数据财产权益的控制范围，并随着数据财产权益保护的强化而受到进一步限制。在现有的利益保护模式下，数据爬取行为构成不正当竞争行为的认定标准非常宽松。特别是在行为不正当性的认定上，若数据爬取者在爬取数据时违反机器人协议或服务协议，大概率会被认定具有不正当性。^[5]而在实践中，网站通过机器人协议或服务协议限制大模型开发者爬取数据的现象愈发普遍。有研究调查了三个大模型训练最常用的开源数据集C4、RefinedWeb和Dolma，发现在大模型兴起的一年里，许多网站在通过机器人协议和服务协议不断加强对数据爬取

的限制。例如，在最关键的网站中，2024年4月时大约有20%-33%的内容受到机器人协议的限制，而一年前这一比例不到3%；而网站服务协议中有关爬取的限制也增加了26-53%。^[6]较低的认定门槛加上不断扩张的协议限制，使得大模型开发者未经许可爬取公开网页的行为越来越容易落入不正当竞争的范围。而且，如果未来对数据采取权利保护模式，那么即便网站没有通过机器人协议或服务协议明确限制大模型开发者的爬取行为，大模型开发者在爬取网页数据之前也需要获得网站的许可，否则会构成对数据财产权的侵害。这会导致大模型开发者越来越难以通过爬取公开网页来获得海量的训练数据。

第二，使用开源数据的合法性存在不确定性。通过收集开源数据来构建大规模训练数据集也受到极大的限制。首先，许多开源数据集的构建本身就依赖于爬取公开网页。由于爬取公开网页的合法性空间收紧，构建开源数据集的可用资源也进一步收缩。如果在开源数据集的构建过程中存在侵害数据财产权益的行为，那么其本身的合法性就存在问题。此时，大模型开发者使用该数据集就可能面临法律风险。其次，即使开源数据集本身不存在合法性问题，也并不意味着大模型开发者将其用于大模型训练是合法的。其一，大模型开发者在使用时可能会因为使用范围超出发布者许可范围而违法。例如，开源数据集仅允许用于科研目的，而大模型开发者将其用于商业性模型的训练。其二，即便没有超出开源数据集的许可范围，仍不能确保大模型开发者使用该数据集训练大模型就是合法的。这是因为，开源数据集的发布者所规定的许可范围有时也可能会超出其自身所享有的权限。例如，有的网站可能会允许Common Crawl等非营利组织爬取其网页数据，但对大模型开发者专门收集数据的网络爬虫进行了限制。然而，Common Crawl在爬取这些网页数据后，有可能会被其他开发者用来训练大模型，而Common Crawl的使用条款中并未明确限制开发者的这一行为。如果该网站对其数据享有对世性的数据财产权益，开发者使用其数据训练大模型

的行为虽未违反Common Crawl的使用条款，却有可能构成对数据财产权益的侵害。因此，随着数据财产权益保护的强化，大模型开发者使用开源数据构建训练数据集也面临着更多的限制与不确定性。

（三）训练数据市场失灵及其缘由

在数据财产权益保护日益强化的背景下，收集海量训练数据的主要方式面临着合法性障碍，导致使用数据训练大模型的自由受限。在很多时候，大模型开发者必须通过市场机制获得数据财产权益人的许可，才能确保训练数据的合法性。然而，基于大模型训练的现有实践，通过市场机制来获得许可存在较大的难度，主要原因在于高昂的交易成本。交易成本通常包括搜索和信息成本、谈判和决策成本以及监管和执法成本。^[1]在大模型训练数据市场中，这些成本因为种种原因极其高昂，阻碍了交易的发生。

第一，大模型训练数据涉及的数据财产权益人众多，这导致开发者所要付出的搜索和信息成本、谈判和决策成本变得极高。过去典型的数据利用场景往往涉及专门的目的、功能或领域，其所需的数据经常集中在有限的权益人或中介手中。但是在在大模型训练的场景下，训练数据可能涉及的财产权益人规模激增。特别是占比最高的网页数据，其涉及的网站成千上万。例如，在大模型训练常用的C4数据集中，就包含大约1500万个网站，词元（token）数量高达1560亿。如今，许多大模型训练使用的网页数据规模早已大大超出C4数据集的规模，达到了万亿词元的级别，其涉及的网站数量很可能也早已超过C4数据集中的千万级别。^[2]理论上，每一个网站都有可能是一个独立的数据财产权益人，那么大模型训练涉及的权益主体数量将可能是千万乃至上亿级。对大模型开发者来说，要去寻找如此众多且分散的权利人，将付出极大的搜索和信息成本；还要与这些网站进行一一协商或付费，将会付出极高的谈判和决策成本。在这种情形下，如果继续强调对数据财产权益的保护，会给大模型开发者带来沉重的负担，其几乎不可能通过市场机制获取和使用这些数据。

第二，大模型训练数据的使用行为在事实上难以追溯，这导致监管和执法成本也急剧增加。在过去典型的数据纠纷中，许多侵害数据财产权益的行为是对他人数据的提供、展示等具有公开性的行为，很容易被追溯。但是在在大模型训练的场景下，数据使用行为变得更为隐蔽和难以证明。首先，大模型的训练，包括其中涉及的数据使用，一般不在公开的环境下进行，因此外界无法直接观察和监测大模型训练中的数据使用行为。其次，对商业性的大模型开发者而言，有关训练方法的信息至关重要，包括训练数据的来源、内容和处理方法等，开发者会对这些信息采取保密措施以维持竞争上的优势。再者，训练得到的结果是模型的参数权重，并不会保存任何训练数据，虽然开发者最终需要对外提供大模型服务，但是外界很难从模型本身或其生成的结果来反推其使用的训练数据。由于追溯大模型训练数据使用行为的难度很高，要在这种情况下严格地执行数据财产权益保护制度，就必须付出大量监管和执法成本。

可见，高昂的交易成本使得大模型开发者无法通过市场机制取得数据财产权益人的许可。在现实中，伴随这一现象发生的是大模型的“非法兴起”。^[3]由于无法通过市场机制合法获得数据权益人的许可，许多大模型的训练事实上未经许可就获取和使用可能受他人权益保护的数据，存在着极大的违法风险。正因如此，大模型开发者正面临着许多由数据持有者提起的诉讼。^[4]这种“非法兴起”在一定程度上表明，在大模型训练的场景下，市场难以发挥作用来促成人工智能公司与数据持有者之间的交易。在这种情形下，如果严格地执行数据财产权益保护制度，不仅会产生极高的监管和执法成本，还会导致大模型开发者无法获取足够的数据进行训练。换言之，过强的数据财产权益保护可能会导致市场失灵的发生，即大模型开发者无法通过市场机制与数据权益人达成对社会而言有益的交易，产生了无效率的结果。

二、大模型训练数据合理使用的制度证成

大模型训练数据市场可能发生的失灵现象为

合理使用制度的引入提供了必要的前提背景。合理使用被视为解决市场失灵的主要手段之一。然而，市场失灵只是引入合理使用的必要而非充分条件，要证成大模型训练数据合理使用制度，仍需结合其他证成条件进行论证。此外，合理使用也并非解决市场失灵的唯一手段，集体管理、法定许可等制度也是解决市场失灵的可能方案。要证成大模型训练数据合理使用制度，需要比较这些替代方案，证明合理使用是更优的选择。因此，在这一部分，本文将首先回顾证成合理使用制度的理论，并分析引入大模型训练数据合理使用制度的正当性，然后讨论解决市场失灵的替代方案，指出其无法完全取代合理使用。

（一）制度证成的理论

合理使用起源于著作权法，是指在特定条件下著作权人以外的主体可以不经著作权人许可无偿使用作品的制度。在著作权法领域，合理使用的理论得到了最为充分的讨论。著作权法学者早已围绕市场失灵的概念为合理使用制度构建了极具解释力的理论。在其开创性的论文中，温迪·J.戈登（Wendy J. Gordon）教授认为，著作权合理使用应被解释为应对市场失灵的一种制度方案。^[1]著作权法赋予作者对作品的排他性权利，以激励作者的文学艺术创作行为，在此之后，将作品转移给效用最大的使用者的任务主要由市场来完成。然而，由于交易成本等原因，市场并不总能促成对社会而言有益的转让行为，于是出现了低效率或无效率的市场失灵现象。在发生市场失灵的情况下，需要通过市场之外的途径来解决这一问题。合理使用正是通过允许潜在的使用者自由无偿使用作品，使更多的人能够享受作品带来的效用，从而解决了市场失灵问题，提升了社会整体效率。

不过，仅存在市场失灵并不能完全证明合理使用制度的正当性。合理使用是最为宽泛的权利限制，不仅赋予了使用者从事特定使用行为的自由，还免除了其支付对价费用的义务。因此，合理使用的证成需要适用更加严格的条件。戈登教授也意识到这一点，并尝试构建证成合理使用制度的完整理论。

他指出，除了市场失灵的存在之外，某一情形要构成合理使用还需满足两项条件：第一，允许使用者使用作品能够增进社会福利。第二，对著作权人的激励不会因允许使用者继续使用作品而受到实质性损害。^[1]

著作权合理使用的市场失灵理论及其分析框架可以扩展到数据财产权益领域，用于证成数据财产权益的合理使用制度。这主要源于数据财产权益与著作权在原理上的相似性。与著作权类似，数据财产权益不仅是对数据持有者投入的合理回报，也是对其数据生产和供给行为的有效激励。^[2]在确立数据财产权益之后，市场会在大多数时候发挥作用，使数据流转至最能实现其价值的使用者手中。但当市场不能有效地实现这一目标时，便有了合理使用等法律干预手段介入的可能。类比著作权合理使用的证成条件，要证成某一情形构成数据财产权益的合理使用，需满足三项条件：第一，存在市场失灵，即使用者无法通过市场机制支付适当的费用获得使用数据的许可；第二，允许使用者使用数据能够增进社会福利；第三，对数据财产权益人的激励不会因允许使用者继续使用数据而受到实质性损害。

（二）合理使用的证立

在第一部分，本文已经证明了大模型训练数据市场可能发生失灵现象。接下来，要证明大模型训练数据合理使用的正当性，仍需继续证立后两个条件，即允许合理使用能够增进社会福利且不会对数据财产权益人的激励造成实质性损害。

1. 合理使用增进社会福利

评估合理使用是否增进社会福利，最直接的证据来自当下的现实。实际上，大模型“非法兴起”的现实及其带来的社会效益，已经在很大程度上证明了合理使用可以极大地增进社会福利。在大模型“非法兴起”的情形下，开发者获取和使用数据进行大模型训练的状态与合理使用下的状态完全相同——大模型开发者没有获取数据权益人的许可，也没有向其支付费用。然而，在这种情形下，大模型技术迅猛发展，并产生巨大的社会价值。目前，大模型主要被应用于生成式人工智能，极大地

推动了文学艺术等领域的发展。大模型提高了生成式人工智能的内容生成能力,为文学艺术创作提供新的工具。一方面,这一工具降低了创作的门槛,使更多普通人有机会参与文学艺术创作;另一方面,它也为专业艺术家提供新的契机,拓展了文学艺术创作的模式和空间。^[1]除了文学艺术领域的价值之外,大模型展现出的通用性特征使其具备在其他领域产生价值的潜力。通用大模型能够在不同的领域和任务中替代专用模型,降低开发的难度和成本,使更多人能够利用其强大的能力,开发出解决各种实际问题的人工智能应用。这种广泛的应用前景预示着大模型的发展将极大地提升个体的效率以及社会整体的福利。综上,引入大模型训练数据合理使用制度,可以在不增加开发者任何负担的情况下,将当前大模型“非法兴起”的状态合法化,从而保证大模型技术的快速发展和社会福利的持续增长。

2. 合理使用不会造成损害

与著作权合理使用的证成相比,数据合理使用在后一个条件的证成上面临更大的挑战。在著作权法意义上,大模型开发者对作品数据的使用通常构成非作品性使用或非表达性使用,这意味着,其使用行为基本上不会对著作权人享有的市场利益造成实质性损害。然而,从数据财产权益的角度出发,大模型开发者对数据的使用,确实有可能会损害数据财产权益人的市场利益,进而对数据财产权益人的激励造成影响。第一种情形涉及对数据服务或产品市场利益的损害。如果大模型开发者所提供的数据服务或产品可以实质性替代数据财产权益人所提供的服务或产品,那么很可能会给数据财产权益人在现有数据服务或产品市场中的利益造成损害。例如,数据财产权益人付出实质性投入积累了较大规模的法律数据,并基于这些数据向公众提供公开的法律信息内容服务;而开发者未经许可获取上述公开数据来训练大模型,并基于大模型向公众提供相同或相似的法律信息内容服务。^[2]如果这类情形也被纳入大模型训练数据合理使用的范围,很可能导致数据财产权益人在现有数据服务或产品市场中的利益受到较大的损害。第二种情形涉及对数据

许可市场利益的损害。如果在大模型开发者与数据财产权益人之间已经形成了大模型训练数据的许可市场,或存在建立相关数据许可市场的可能性,那么允许开发者合理使用数据会给数据财产权益人在现有或潜在数据许可市场中的利益造成损害。例如,像Reddit等网站的数据虽然是公开的,但其已着手建立公开数据许可的商业模式,要求大模型开发者在使用数据前须获得许可并支付相应费用。^[3]在此背景下,若法律允许大模型开发者无偿获取并使用上述平台的数据,势必会对正在形成中的数据许可市场造成冲击,损害数据财产权益人可能获得的利益。

尽管如此,本文认为,上述情形并不足以完全否定大模型训练数据合理使用的正当性。第一,在大多数情形下,允许开发者合理使用数据训练大模型,不会损害数据财产权益人的市场利益。首先,在大多数情形下,大模型开发者所提供的服务或产品并没有与数据财产权益人所提供的服务或产品存在直接竞争关系,不会构成实质性替代。大模型开发者使用数据训练大模型,主要是为了让大模型学习数据中的知识以提升大模型的语言能力和通用能力,而非为了提供与数据财产权益人相同或相似的服务与产品。只要基于大模型向公众提供的服务或产品与数据权益人提供的服务或产品之间存在着较大差异,引入合理使用一般不会对数据财产权益人的既有市场利益造成损害。实证研究也表明,ChatGPT等基于大模型的产品的主要实际用途与其训练数据主要来源所处的市场领域存在较大的区别。例如,ChatGPT中超过30%的对话被用于创造性写作,但事实上ChatGPT训练数据中创造性写作的数据占比并不高;而新闻类数据在ChatGPT训练数据中的占比相对较高,但只有不到1%的ChatGPT使用与新闻相关。^[4]第二,虽然在大模型开发者和数据财产权益人之间确实存在一些数据许可交易,但是相较于大模型使用的数据规模,这些交易所涉及的数据规模只占很小的部分,主要发生在少数占据优势地位的大模型开发者与个别平台型数据财产权益人之间。换言之,这些交易很可能只是个案,

并不能证明大模型训练数据市场的失灵现象是可以治愈的。事实上,在多数情形下,大模型训练数据市场过高的交易成本使得许多交易根本难以发生,数据财产权益人也很难主张可得利益的损失。因此,允许大模型开发者合理使用数据进行大模型训练,在多数情形下也不会损害数据财产权益人在数据许可市场中的利益。

第二,虽然可能存在少数损害数据财产权益人利益的情形,但是通过准确地界定合理使用的适用范围,可以将这些少数情形排除在合理使用之外,从而保障数据财产权益人的利益不会遭受损害。具体而言,可以通过明确合理使用的对象、目的与方式,将可能损害数据财产权益人利益的情形排除在合理使用范围之外;亦可允许数据财产权益人在满足特定条件下选择退出合理使用,从而使其能在交易成本较低的情况下,与大模型开发者达成许可协议,为形成有效的训练数据许可市场留下可能性。

综上,合理使用可能损害数据财产权益人的情形只是少数例外的情形。只要适当地限定合理使用的范围,在大多数情形下,引入大模型训练数据合理使用并不会对数据财产权益人的市场利益造成损害。

(三) 替代方案的比较

有观点认为,我国可以借鉴著作权领域的经验,通过引入数据财产权益的集体管理或法定许可来降低交易成本,从而解决市场失灵问题。集体管理是指权益人通过集体管理组织对权益对象的使用予以许可、收取相应报酬的制度。^[1]法定许可则是指根据法律的规定可以不经权益人许可而以特定方式使用权益对象,但应当向其支付报酬的制度。^[2]这两者在著作权领域均有着成熟的经验,通过集中行使权益或直接规定许可费的方式,解决了一对一授权情形下可能存在的交易成本过高问题。在此背景下,有观点主张,可将集体管理或法定许可的方案推广至大模型训练场景,以解决大模型训练涉及大量作品数据使用的问题。^[3]同理,当大模型训练数据涉及大量数据财产权益时,也可以采取类似的集体管理或法定许可制度。

无论是集体管理还是法定许可的方案,其对比合理使用的共同特点是仍然要求大模型开发者向

数据财产权益人支付费用。其支持者可能会认为,相比合理使用,这些方案的优势在于兼顾了数据财产权益人的利益:通过给予权益人经济补偿,可以更好地实现利益平衡。然而,本文认为,与合理使用相比,集体管理或法定许可的优势并不明显,而其可能带来的制度成本和代价却不容忽视。

第一,采取集体管理或法定许可方案给数据财产权益人带来的收益十分有限。首先,由于大模型训练数据所涉及的权益人规模非常庞大,即使开发者向权益人整体支付一笔不菲的费用,大多数权益人能够获得的收益也极其微薄。假设某个大模型的训练数据可能涉及1000万个网站,即使开发者愿意支付1亿元的许可费,分摊到每个网站的平均收益也只有10元。而且,考虑到数据上还承载着著作权等其他权益,数据财产权益人可能还需要将这些收益进一步分配给其他权益人,实际上所能获得的收益甚至更少。其次,大模型训练对数据的使用频率也更为有限。著作权集体管理或法定许可所应对的典型情形很多都是对作品的持续、高频利用,这种持续、高频利用使著作权人能够持续累积并最终获得较为可观的收益。然而,大模型训练数据却并不具备这一特征。大模型训练的成本非常高昂,绝大多数大模型开发者不会频繁地使用数据进行训练,这可能会导致权益人无法通过累积使用次数来获得足够高的收益。^[4]

第二,采取集体管理或法定许可方案会付出更高的制度成本。由于这些方案都涉及集体管理组织的建立以及许可费的确定、收取和分配等活动,其制度的建立和运行相比合理使用都需要付出更多的成本。而且,相比著作权的集体管理或法定许可,大模型训练数据的集体管理或法定许可可能会面临更高的制度成本。以集体管理的建立和运行为例。首先,相比著作权集体管理,建立数据集体管理组织需要付出更高的成本。在著作权领域,已经存在着成熟的集体管理组织,并且这些组织已经获得了许多作品的授权,如果要将著作权集体管理扩展到

大模型训练的场景,有相对坚实的基础。但在数据财产权益领域,并不存在类似的基础,如果要建立集体管理组织并获得足够规模的授权,需要投入更多的成本。其次,大模型训练的情形与著作权集体管理所应对的典型情形存在较大差异。一方面,大模型训练所涉及的数据类型更加多样。传统的著作权集体管理组织往往专注于单一类型作品的著作权管理,但大模型训练却会涉及多种类型数据。数据类型的多样化给集体管理带来了更高的运行成本,著作权集体管理组织管理单一类型作品的经验可能很难适用于大模型多样化训练数据的管理。^[1]另一方面,数据集体管理所需管理的对象规模更加庞大。在大模型训练的场景下,涉及的数据和财产权益人规模可能远超传统著作权集体管理组织所管理的规模。我国规模最大的著作权集体管理组织是中国音乐著作权协会,截止至2024年底,其管理的会员规模为14064人,音乐作品规模约为2300万首。如前所述,单个大模型训练数据所涉及的数据财产权益人规模在千万级别以上,网页数量更是高达亿级以上,这远远超过传统著作权集体管理所管理的规模。由于管理对象规模更为庞大,数据集体管理制度的运行成本也将显著增加。

第三,采取集体管理或法定许可方案会给我国大模型的发展造成不利影响。首先,采取集体管

理或法定许可方案会巩固大企业的竞争优势,给初创企业制造极高的进入壁垒。^[2]大规模数据的使用会产生极高的许可费用。从目前实践中达成的训练数据交易案例看,许多大模型开发者与单个平台型数据财产权益人之间的许可费都已经高达每年千万美元,如果要向所有的数据财产权益人付费,数额很可能远超每年千万美元这一级别。如此庞大的成本很少有企业能负担得起,这将大大减少大模型领域的竞争,给该领域的创新带来不利影响。其次,采取集体管理或法定许可方案可能会影响我国在该领域的国际竞争力。有学者指出,如果有的国家选择不要求大模型开发者付费,那么开发者和开发活动可能会向这些监管较为宽松的国家转移,发生“创新套利”现象。^[3]在大模型技术仍不断发

展的阶段,如果我国率先要求开发者承担这一费用负担,可能会影响其在我国研发技术和落地应用的积极性,削弱我国在人工智能领域的国际竞争力。

三、大模型训练数据合理使用的规则构建

在证成大模型训练数据合理使用的正当性后,有必要进一步探讨其具体规则的构建。考虑到与著作权合理使用的差异,数据合理使用更有可能对数据财产权益人的利益造成影响,其适用范围必须严格限定:一方面要以解决市场失灵问题为必要,另一方面也要避免对数据财产权益人的激励造成不利影响。在这一部分,本文将分别从大模型训练数据合理使用的对象、目的、方式及退出等方面,探讨其规则具体内容的建构,并就立法的完善提出建议。

(一) 合理使用的对象:公开数据

根据不特定主体是否可以事实上自由访问并获知内容,数据可以分为公开数据和非公开数据。本文认为,大模型训练数据合理使用的对象应限于公开数据,原因如下:

第一,从必要性的角度分析,仅对公开数据适用合理使用足以解决大模型训练场景下的市场失灵问题。首先,大模型训练数据的主要来源是公开数据,保障大模型训练数据的获取和使用首要是要确保公开数据的获取和使用。其次,大模型训练数据市场失灵的问题也主要发生在公开数据的获取和使用上。权益人众多和使用难追溯的问题更为明显地体现在公开数据的获取和使用上,这些问题是造成交易成本上升的主要原因。但在非公开数据的获取和使用上,交易成本偏高的问题并不明显。一方面,只要公开数据的获取和使用得到保障,大模型开发者一般已经能够获得足够规模的训练数据,虽然获取更多非公开的高质量数据对大模型训练也有帮助,但其规模和涉及权益人数量相对较小,采取一对一的直接授权模式也具有可行性。另一方面,未经许可对非公开数据的获取行为往往会留下更多的痕迹,而且许多非公开数据通常是权益人的独有数据,在证明大模型开发者的不当获取行为时,权益人所面临的举证难度也相对较小。因此,将合

理使用的对象扩张到非公开数据的必要性不大。

第二,从合理性的角度分析,对公开数据与非公开数据的权益保护和限制应有所区别。虽然公开数据和非公开数据均受法律保护,但公开数据所受的财产权益保护程度应当弱于非公开数据。相应地,公开数据所受到的权利限制也应当强于非公开数据。之所以要区别对待公开数据和非公开数据,首先是因为公开数据被认为具有开放性和公共性。一种流行的观点认为,互联网具有固有的开放性:“它允许世界上的任何人发布其他任何人都可以访问的信息,而无需身份验证。当计算机所有者决定托管网络服务器,以使文件可以通过网络访问时,默认设置是允许公众访问这些文件。”^[1]换言之,任何公开在互联网上的信息或数据应默认具有开放性和公共性,允许他人访问和获取这些信息和数据。基于这种观点,对具有公共性的公开数据理应适用合理使用,保障数据的共享流通。其次,未经许可获取和使用公开数据和非公开数据所造成的利益损害程度不同。这与数据持有者对公开数据和非公开数据的利益期待有关。总体而言,数据持有者对非公开数据的利益期待更强。如果要获得非公开数据,经常需要破坏数据持有者采取的保密措施,这对其预期利益的损害和对社会秩序的破坏都更为严重。但相比而言,公开数据更容易被其他主体所获取和使用,在某些情形下甚至可以推定数据持有者默示同意其他主体对数据的获取和使用。

(二)合理使用的目的:用于预训练

大模型的训练过程一般分为预训练(pretraining)和微调(finetuning)两个阶段。预训练是指“使用与下游任务无关的大规模数据进行模型参数的初始训练”。^[2]微调则是指在预训练模型的基础上,针对特定的任务或数据进行额外的训练,通常包括指令微调(instruction fine tuning)和对齐(alignment)。两个阶段的训练目的有所不同。通俗地讲,预训练是让大模型学习广泛的知识,使之具备通用的语言理解和生成能力。微调则是让大模型在特定领域进行专门学习,以便更好地完成特定任务以及与人类的价值观保持一致。本文认为,大模型训练数据合

理使用的目的应限于用于预训练而非微调,原因如下:

第一,市场失灵主要发生在预训练数据的获取和使用上。由于目的不同,大模型开发者在预训练和微调阶段所使用的数据存在较大的差别。在规模方面,预训练会涉及大规模数据的使用,而微调使用的数据量则相对较小。在类型方面,预训练数据不局限于某一领域,微调数据则更多会针对特定的领域或任务。这些差别所导致的直接后果是,预训练数据涉及的数据财产权益人数量和类型更为庞大,其获取和使用可能会产生更高的交易成本,更容易发生市场失灵。而微调数据由于规模较小且往往针对特定的领域或任务,其涉及的数据财产权益人数量和类型都比较有限,一般不会造成特别高的交易成本。因此,大模型训练数据合理使用应主要适用于更容易发生市场失灵的预训练阶段。

第二,对微调数据的获取和使用适用合理使用更有可能造成损害。从数据使用的目的来看,预训练阶段的数据使用比微调阶段的数据使用更具转换性。转换性使用(transformative use)的概念源自著作权法,原意是指以不同方式或基于不同目的对作品进行使用。^[3]数据的转换性使用是指在后数据使用者以不同于数据财产权益人的方式或目的对数据进行使用。在转换性使用的情形下,在后使用者的使用方式或目的与在先权益人的使用方式或目的存在较大差别,因而在后使用者的行为对在先权益人市场利益的影响较小,造成损害的可能性较低。^[4]预训练的目的是使大模型具有通用的语言能力,这与绝大多数数据财产权益人使用数据的目的存在明显差别,具有显著的转换性,一般不会直接影响数据权益人的市场利益。微调的目的主要是使大模型具备特定领域的知识或处理特定任务的能力,从而可以直接应用于特定的服务或产品,其使用的数据往往与这些服务或产品紧密相关。例如,为了开发提供法律咨询服务的模型,开发者可以在通用大模型的基础上使用法律类问答数据进行微调,从而使大模型具备更丰富的法律知识以及更符合法律咨询要求的生成能力。这些微调数据经常来

自提供相同或相似服务或产品的数据财产权益人，包括法律数据库商、法律问答网站等。而基于微调获得模型所提供的服务或产品，很可能与数据财产权益人提供的服务或产品非常接近，甚至有可能构成实质性替代。可见，微调数据的使用转换性程度较低，很可能会严重影响数据财产权益人的市场利益，不宜广泛地纳入合理使用的范围。

（三）合理使用的方式：训练涉及的数据处理行为

大模型训练数据合理使用的方式应限于训练涉及的数据处理行为。数据处理包括数据的收集、存储、使用、加工、传输、提供、公开等。而大模型训练主要涉及收集、存储、使用、加工、传输等数据处理行为。例如，在开始训练之前，开发者需要收集训练所需的数据，将收集到的原始数据加工成机器可读的格式化数据，并存储在一定的介质中；在训练的过程中，则需要将数据传输到训练的服务器上，并使用该数据展开模型训练。这些处理行为可能会落入数据财产权益的控制范围，应当通过合理使用给予侵权豁免，从而方便大模型训练的合法进行。

至于提供和公开行为是否涵盖在大模型训练数据合理使用之中，则需认真考虑和讨论。首先，大模型训练一般不涉及数据的提供和公开。大模型训练数据合理使用所涵盖的范围应当以训练正常进行所需的数据处理行为为限，将超出此范围外的数据处理行为纳入合理使用应当更加慎重。其次，未经许可提供和公开数据有时可能会实质性损害数据财产权益人的利益。从相关案例看，许多数据财产权益人试图阻止的行为主要是对数据的提供和公开行为，而非其他数据使用行为。^[3]这是因为在许多情形下，提供和公开数据会对数据持有者的产品或服务构成实质性替代，严重损害数据持有者的利益。在大模型的场景下，也可能会发生类似的情形。如果大模型开发者在使用其他网络平台的公开数据训练大模型后，又通过大模型直接向用户提供相同或相似的信息内容服务，那么就有可能构成对其他网络平台信息内容服务的实质性替代。显然，

这种行为应当被明确排除在合理使用的范围之外。

但在有些情形下，允许提供和公开大模型训练数据可以增进社会福利，且一般不会实质性损害数据财产权益人的利益。例如，大模型开发者基于法律对于人工智能透明度的要求而公开部分训练数据，有助于增强模型的透明度，提升公众对大模型技术的监督和信任。又如，有企业专门提供大模型训练数据服务，负责数据的采集、清洗与标注，并将其提供给多个大模型开发者使用，允许此类企业向开发者提供训练数据，有助于降低重复采集与处理的成本，提高社会效率。因此，对提供和公开大模型训练数据是否构成合理使用，不宜采取一刀切的方式，而应采取场景化认定的方法，综合该行为对数据财产权益人利益的影响等因素进行评估。

（四）合理使用的退出：以技术措施选择退出

如上所述，大模型训练数据市场失灵并非完全的市场失灵。在特定情形下，大模型开发者可能会与部分数据财产权益人达成数据许可交易。如果构建过于宽泛的合理使用规则，可能会削弱大模型开发者与部分数据财产权益人达成交易的动力，从而损害数据财产权益人的可得利益。通过立法限定合理使用的适用条件，明确将这些情形排除在合理使用之外，是一种可行的方案。然而，立法者往往囿于信息成本，也不能完全预见到所有可能的情形。很多时候，交易双方的当事人拥有更多的信息，可以根据具体的情形做出既符合自身利益同时也最具社会效率的决策。因此，赋予数据财产权益人选择退出合理使用的权利，有助于在交易成本较低的情形下，促成数据财产权益人与大模型开发者达成交易，从而更有效地维护数据财产权益人的利益。

不过，允许选择退出合理使用也可能引发新的问题。反对选择退出合理使用的观点认为，如果选择退出合理使用的成本过低，权益人可能滥用这一选择权架空合理使用制度。关于选择退出合理使用的机制是否合理，学界早有讨论。最典型的例子便是关于合同排除合理使用条款法律效力的探讨。多数观点认为，应当否定此类合同条款的法律效力。其主要理由是，合理使用制度是著作权法中维护著

作权人与公众之间利益平衡的重要机制,如果著作权人可以通过合同排除合理使用的适用,无异于由著作权人单方面重新界定著作权保护的内容和边界,构成“私立知识产权”,进而破坏立法者原本设定的利益平衡。^[1]尤其是在点击合同广泛应用的背景下,网络平台非常容易通过合同手段破坏这一制度性平衡,从而加剧对公共利益的侵蚀。此外,欧盟在文本与数据挖掘合理使用规则中引入了选择退出机制,也引发了争议。根据《数字单一市场版权指令》第4条的规定,商业性主体进行文本与数据挖掘可以构成合理使用,但著作权人也享有选择退出这一合理使用的权利。有观点认为,该条款设定的选择退出手段成本极低,不仅导致无法实现促进文本与数据挖掘发展的制度目的,还可能产生显著的社会负外部性。^[2]

本文认为,为了防止数据财产权益人滥用选择退出机制,应当在引入这一机制的同时,提高权益人选择退出合理使用的门槛。正如许多观点所担心的,如果选择退出机制的门槛过低,在交易成本很高的情形下,即使数据财产权益人只有很小的概率能够获得许可费用,由于只需付出极低的成本就可以规避合理使用的限制,那么数据财产权益人有可能会抱着投机的心态选择退出合理使用,来保留向使用者主张数据财产权益的可能。这会导致越来越多的情形被排除在大模型训练数据合理使用之外,合理使用的制度目的落空。例如,如果通过机器人协议或服务协议的规定就可以选择退出合理使用,数据财产权益人几乎不需要付出任何成本就可以规避合理使用的限制,那么绝大多数权益人很可能会选择修改机器人协议或服务协议,来保留向大模型开发者主张权益的机会。正如实证研究所表明的,实践中通过机器人协议或服务协议限制大模型开发者爬取数据的比例正在不断提高。^[3]相对合理的方案是,允许数据财产权益人通过付费墙、软件锁等技术措施选择退出大模型训练数据的合理使用。这类措施一般需要数据财产权益人付出较高成本,因此权益人在选择是否退出合理使用时,会权衡退出所能获得的收益与成本。^[4]一般情况下,只

有与大模型开发者达成交易的可能性较高且所获收益较高时,数据财产权益人才会选择退出大模型训练数据的合理使用。这时,数据财产权益人所采取的行为策略一般会与最具社会效率的决策相契合,既能更好地维护其自身的利益,也会促进社会整体福利的提高。

(五) 合理使用立法的完善

在现行法下,通过合理地解释和适用反不正当竞争法相关条款,有可能达到与引入合理使用接近的效果。首先,法院通过对一般条款适用要件的解释,可以将一些合理的数据获取和使用行为排除在规制范围之外。一般条款保护数据财产权益需满足多个要件,包括存在竞争关系、行为具有不正当性以及对数据权益人造成实际损害等。在大模型训练的场景下,法院可以通过对这些要件的解释,给大模型开发者获取和使用数据留下适当的合法性空间。例如,如果大模型开发者获取和使用数据的目的主要是为了科学研究,那么一般情形下可以认定其与数据权益人不构成竞争关系。即便大模型开发者获取和使用数据是为了商业性目的,如果它和数据权益人所提供的服务或产品属于两个不太相关的市场,也可以认定两者不存在竞争关系;或只要大模型训练的行为并未对数据权益人造成实质性损害,也完全可以通过对损害要件的解释,使大模型开发者获取和使用数据的行为免责。其次,《反不正当竞争法》新增的数据条款也存在着解释空间。例如,该条款要求经营者不得以“不正当方式”获取和使用其他经营者合法持有的数据,法院未来可以通过对“不正当方式”的解释,将合理使用的情形排除在数据财产权益的控制范围之外。

但是,这一司法解决方案不能替代立法的完善。第一,上述条款的解释和适用有较大的弹性,无法提供立法所具备的确定性。一般条款本身无法为合理使用的认定提供具体指引,法官在解释和适用该条款时有较多的裁量空间,因此会导致合理使用的认定具有极大的不确定性。“不正当方式”作为一个不确定概念,也无法事前清晰地界定其内涵和外延。^[5]这难以为大模型开发者提供稳定的合法性预

期,无法为实现合理使用立法提供确定性的作用。第二,随着数据财产权益的权利化,从法理的角度分析,合理使用作为权利限制,原则上应由立法明确规定。数据财产权益与著作权、人格权一样是私权,从私权神圣、意思自治的私法原理出发,私权的限制属于例外情形,原则上应当交由立法以列举方式进行限定,不宜由司法随意解释创设。^[4]正因如此,无论著作权合理使用还是人格权(包括个人信息权益)合理使用,我国都采取立法列举的方式加以规定。数据财产权益的合理使用也应如此。因此,有必要在立法层面引入大模型训练数据的合理使用规则。

目前,在人工智能和数据领域,相关立法活动正在加快推进当中。2023年和2024年,国务院连续两年将“人工智能法草案”列入预备提请全国人大常委会审议的法律案。全国人大常委会也将“人工智能健康发展等方面的立法项目”纳入2024年和2025年立法工作计划中的预备审议项目。应当利用这些立法机会,认真考虑建立大模型训练数据的合理使用制度,为大模型训练提供明确的合法性基础,

保障人工智能技术和产业的发展。可以考虑的方案包括:第一,在《人工智能法》中单设有关大模型训练数据的合理使用条款。第二,在数据财产权进行立法时设置范围更广泛的数据合理使用条款,并将大模型训练数据的合理使用明确列为其中的一类情形。

结语

数据财产权益保护的强化给大模型训练数据的获取和使用带来了合法性挑战。基于市场失灵理论的分析可知,在多数情形下,允许开发者合理使用数据进行大模型训练,可以增进社会福利,且不会损害数据财产权益人的市场利益。相比集体管理或法定许可等替代方案,合理使用亦是更优的选择。只要对大模型训练数据合理使用的规则进行适当的设计,确保其适用范围是必要且合理的,就能在技术发展与权益保护之间取得有效的平衡。未来,我国有必要在人工智能和数据立法中认真考虑引入这一制度,为数据要素市场与人工智能产业的发展提供更好的法治保障。

(技术编辑:艾薇)