

中国人民大学法学院 数字法学教研月报

2026年第5期（总第29期）

2026年7月28日



本期看点

【数字法治大事件】

人工智能全球治理迈向深化：2026世界人工智能大会暨人工智能全球治理高级别会议召开。围绕人工智能创新发展、安全治理与国际合作发表重要倡议，推动人工智能治理从技术规范进一步延伸至伦理、安全、规则与全球协同层面。

个人信息保护制度持续精细化：国家网信办、公安部联合发布《小型个人信息处理者个人信息保护简化措施规定》，针对小型个人信息处理者合规能力不足问题探索差异化治理路径；同时，数据出境安全管理政策法规问答进一步明确跨境数据流动合规要求，推动个人信息保护与数据有序流通协调发展。

数字基础设施与网络空间治理不断完善：中央网信部门印发IPv6技术创新和融合应用相关实施方案，推动互联网基础资源高质量发展；有关部门围绕互联网基础资源治理部署重点任务，为数字经济发展提供更加稳定、安全、高效的基础支撑。

【研究动态】

人工智能治理与法律回应：学界持续关注生成式人工智能、智能体治理、算法规制与人工智能责任体系等前沿问题，围绕人工智能技术应用中的安全风险、伦理边界与法律责任展开深入研究。

数据治理与数字法治体系建设：研究重点聚焦

数据跨境流动、个人信息保护、数据要素利用与安全治理机制，探讨数字经济发展背景下数据权益配置与监管制度完善路径。

数字时代部门法融合发展：围绕人工智能对民法、刑法、行政法、劳动法等传统部门法的影响，学界进一步探索数字技术背景下法律制度的适应性调整与规范创新。

【教研活动】

人工智能法治前沿交流：多场学术会议围绕人工智能全球治理、智能体安全、平台治理与算法规制等议题展开研讨，推动数字法学与人工智能技术、国际治理等领域的交叉融合。

数字法学自主知识体系建设持续推进：高校与科研机构围绕数字中国建设、网络强国法治保障、数字刑事法治建构等主题开展学术交流，推动数字法学基础理论与实践研究不断深化。

【数字法评】

《人脸识别技术法律规制的国际比较与场景化重构》，作者：王业亮；中国人民大学法学院；载《行政法学研究》2026年第4期。

《人工智能安全管理义务及其承担者的刑事责任》，作者：皮勇；同济大学法学院；载《中国法学》2026年第2期。

本期目录

数字法治大事件	3	数字行政与数字司法	37
习近平出席 2026 世界人工智能大会暨人工智能 全球治理高级别会议开幕式并发表主旨讲话 .. 3		新兴科技产品的法律问题	38
国家网信办、公安部联合公布《小型个人信息处 理者个人信息保护简化措施规定》	5	教研活动	40
《小型个人信息处理者个人信息保护简化措施 规定》答记者问	8	马长山：《数字时代的人权逻辑》讲座	40
专家解读 简化保护措施 兼顾个人信息权益与 小型企业创新发展利益	10	马长山：《数字司法的法治边界》讲座	41
数据出境安全管理政策法规问答（2026 年 7 月）	12	乔仕彤：《互联网平台与中国知识产权保护的崛 起》讲座	43
中央网络安全和信息化委员会印发《深化互联网 协议第六版（IPv6）技术创新和融合应用实施方 案（2026—2030 年）》	13	第十三届“新兴权利与法治中国”学术研讨会47	
《深化互联网协议第六版（IPv6）技术创新和融 合应用实施方案（2026—2030 年）》答记者问17		中国科学技术法学会 2026 年年会	50
明确 16 项任务！工信部等四部门发文推动互联 网基础资源高质量发展	20	北京大学第二届“数字法治的理论视野与实务前 沿”暑期学校	58
研究动态	24	公开个人信息的 AI 使用边界何在？——“AI 善 治·新智新声”人工智能法治在线课堂 暨硕博研 习（春夏季）第六场	62
数字法学基础理论	24	William Hubbard：《人工智能与法律职业、法 学教育》讲座	64
人工智能治理	26	马长山：《数字法学研究的立场和方法》讲座68	
数字时代的部门法研究	30	郑戈：《作为“或然性工具”的生成式 AI：法 学研究 with 论文写作的方法论反思》讲座	69
数据要素的保护与流通	31	劳伦斯·莱斯格：《人工智能治理与治理人工智 能》讲座	72
个人信息保护	34	数字法评	80
网络平台治理	36	人脸识别技术法律规制的国际比较与场景化重 构	80
		人工智能安全管理义务及其承担者的刑事责任89	

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 毕坤阳 朱唯希 吕昊然 邓语鑫 梁因格 王昊

联系方式：RUCdigitalaw@163.com

本期编辑：王昊

数字法治大事件

导言：

在人工智能技术加速演进、数字基础设施持续升级、数据要素价值进一步释放的背景下，我国数字法治建设正围绕人工智能治理、个人信息保护、数据跨境流动和网络基础资源发展等重点领域持续深化制度供给，推动数字技术创新与安全规范治理协同推进。2026年7月，习近平出席2026世界人工智能大会暨人工智能全球治理高级别会议开幕式并发表主旨讲话，系统阐释人工智能发展与全球治理的重要方向，强调推动人工智能朝着有益、安全、公平的方向发展。相关重要讲话不仅体现出我国积极参与全球人工智能治理体系建设责任担当，也表明人工智能治理正在从技术发展议题进一步拓展为涉及安全、伦理、规则和国际合作的综合性数字法治议题。

值得关注的是，人工智能治理理念正在进一步向具体制度规则和应用场景延伸。智能体互信互联互操作全球合作倡议的提出，回应了智能体快速发展背景下技术协同、可信运行和跨主体合作的现实需求。随着人工智能从辅助工具逐步向自主决策、任务执行和多主体协同方向发展，智能体之间的身份认证、数据交互、责任划分和安全保障等问题日益凸显，推动构建开放、可信、可控的智能体治理规则，成为人工智能时代数字法治发展的重要方向。

在个人信息保护领域，《小型个人信息处理者个人信息保护简化措施规定》的发布及答记者问，进一步回应了中小主体开展数字化经营过程中个人信息保护合规能力不足的现实问题。通过在保障个人信息权益与降低合规负担之间寻求平衡，该规定探索建立更加精准、分类化的个人信息保护治理模式，体现出我国个人信息保护制度正在从普遍性规范要求进一步走向主体差异化、场景精细化和监管适配化。同时，数据出境安全管理政策法规问答（2026年7月）的发布，也进一步明确数据跨境流动中的规则适用、合规要求和实践路径，推动数据安全治理与数据有序流通之间形成更加协调的制度安排。

在数字基础设施和网络空间基础资源治理方面，中央网络安全和信息化委员会印发《深化互联网协议第六版（IPv6）技术创新和融合应用实施方

案（2026—2030年）》及相关答记者问，进一步推动互联网基础协议升级与数字基础设施高质量发展。IPv6作为支撑未来网络演进的重要基础设施，其创新应用不仅关系到网络连接能力和技术自主发展，也与数据流通效率、网络安全保障和数字经济发展密切相关。与此同时，工信部等四部门发布相关文件，明确16项任务推动互联网基础资源高质量发展，进一步完善互联网基础资源治理体系，为数字经济持续发展提供更加稳定、安全、高效的基础支撑。

总体来看，本期数字法治大事件集中体现出我国数字治理正在形成以人工智能全球治理为引领、以个人信息保护和数据跨境安全管理为保障、以新型网络基础设施建设为支撑的综合治理格局。从智能技术应用到数据流通利用，从个体权益保护到基础资源治理，数字法治建设正不断突破传统网络治理边界，向覆盖技术创新、产业发展、安全保障和国际合作的系统化治理阶段迈进。

习近平出席2026世界人工智能大会暨人工智能全球治理高级别会议开幕式并发表主旨讲话

原载：“网信中国”微信公众号

7月17日上午，国家主席习近平在上海世界会客厅出席2026世界人工智能大会暨人工智能全球治理高级别会议开幕式并发表主旨讲话。

黄浦江畔，潮启新澜。世界会客厅内嘉宾云集。

习近平同与会的国家元首、政府首脑、国际组织负责人及各国代表团团长集体合影。

在热烈的掌声中，习近平发表题为《携手构建公正合理的全球人工智能治理体系》的主旨讲话。

习近平指出，当前，世界百年变局加速演进，新一轮科技革命和产业变革加速突破，全球人工智能技术创新进入前所未有的活跃期，既蕴含巨大机遇，也面临治理挑战。人类不得不直面时代之问：当机器开始思考，人类如何与之相处？当算法参与决策，安全如何保障？当技术挑战伦理，治理如何跟上？当鸿沟不断拉大，普惠如何实现？中方认为，

各国应当秉持以人为本、向上向善理念，让人工智能成为促进共同繁荣、维护共同安全的一个重要动力源，携手构建公正合理的全球人工智能治理体系。

习近平就此提出4点意见：

第一，坚持开放共赢，驱动创新发展。抓住难得的历史性机遇，鼓励开源开放、合作共享，全面促进人工智能科技创新、产业发展、场景应用，协同推进传统产业改造升级、新兴产业培育壮大、未来产业前瞻布局，让人工智能赋能千行百业。

第二，强化风险意识，确保安全可控。高度重视人工智能引发的各类内生和衍生风险，推动构建法律法规、技术监测、风险预警、应急响应体系，筑牢安全底线，防范滥用恶用，确保人工智能始终处于人类控制之下。共同反对在人工智能领域泛化国家安全概念、把本国安全凌驾于他国安全之上的做法。

第三，鼓励包容并蓄，促进文明互鉴。用全人类共同价值塑造人工智能的价值观，善用人工智能技术增进不同文明的理解和包容，促进不同文明交流互鉴，精心培育各美其美、美美与共的文明百花园。

第四，倡导和衷共济，完善全球治理。践行真正的多边主义，切实发挥联合国的重要作用，加强人工智能发展战略、治理规则、技术标准的对接协调，早日形成具有广泛共识的全球治理框架，让这一前沿技术更好造福人类社会。帮助全球南方国家加强能力建设，弥合数智鸿沟，促进可持续发展，避免在人工智能领域造成新的历史不公。

习近平强调，今年是中国“十五五”开局之年。“十五五”规划为未来5年中国经济社会发展擘画蓝图，也为国际社会提供机遇清单。近年来，中国坚持有效市场和有为政府相结合，加强人工智能科技创新，积极推进“人工智能+”行动，培育各类主体共生共荣的健康生态，智能经济核心产业规模已经超过万亿元人民币，“中国智造”已成为中国式现代化的又一亮丽名片。同时，中国坚持发展和安全并重，不断完善相关法律法规、政策制度、应用规范、伦理准则，确保人工智能安全、可靠、可

控。作为负责任大国，中国在人工智能领域始终致力于做国际公共产品的提供者，源源不断贡献中国方案。

习近平指出，在各方共同努力下，世界人工智能合作组织在上海应运而生。这是中方响应全球南方呼声、团结国际社会积极推动人工智能发展和治理的重大举措，将成为人工智能发展史上的一个重要里程碑。为进一步支持全球人工智能发展、推进全球人工智能能力建设，未来5年，中国将面向发展中国家提供5000个人工智能专题研修培训名额；面向东盟、阿盟、非盟、拉共体、上合组织、金砖国家建设国际人工智能应用合作中心；推动气象智能预警方案“妈祖”在30个国家落地应用。中国愿以更加开放的姿态、更加务实的行动、更加长远的目光，同各方一道把握和应对人工智能发展的机遇和挑战，携手共创人类社会更加美好的未来。

哈萨克斯坦总统托卡耶夫、柬埔寨首相洪玛奈、泰国总理阿努廷、联合国秘书长古特雷斯分别致辞。他们热烈祝贺世界人工智能合作组织协定签署仪式在上海举行，高度评价中国为推动人工智能全球治理作出的重要贡献，积极呼应习近平主席提出的倡议和举措。各方表示，人工智能有望成为人类最大的发展机遇，也潜藏着风险。中国倡导智能向善、普惠包容、安全可控，积极帮助发展中国家加强相关能力建设，推动人工智能国际合作，彰显了负责任的大国担当。人工智能技术应由各国共享共治，使其公平惠及所有国家，服务全人类可持续发展，消弭全球数智鸿沟和南北发展差距。人工智能带来的机遇与挑战跨越国界，需要各方携手合作把握和应对。希望并相信世界人工智能合作组织将为此发挥重要作用，通过开放包容、平等互惠的合作，共同打造安全、繁荣、不让任何国家掉队的美好未来。

大会发表《2026世界人工智能大会暨人工智能全球治理高级别会议主席声明》。

7月16日晚，习近平和夫人彭丽媛在世界会客厅为与会国际贵宾举行欢迎宴会。

蔡奇、丁薛祥、王毅等出席上述活动。

陈吉宁主持开幕式。

国家网信办、公安部联合公布《小型个人信息处理者个人信息保护简化措施规定》

原载：“网信中国”微信公众号

近日，国家网信办、公安部联合公布《小型个人信息处理者个人信息保护简化措施规定》（以下简称《规定》），自2026年9月1日起施行。

国家网信办有关负责人表示，《个人信息保护法》规定，国家网信部门统筹协调有关部门针对小型个人信息处理者，制定专门的个人信息保护规则、标准。制定《规定》，是贯彻落实党中央决策部署，落实《个人信息保护法》有关要求的重要举措，有利于在坚守个人信息安全底线的同时，降低小型个人信息处理者合规成本，提升小型个人信息处理者个人信息保护水平，为中小微企业创新发展提供法治保障。

《规定》明确了适用范围和总体要求。规定在中华人民共和国境内的小型个人信息处理者实施个人信息保护，适用本规定。明确小型个人信息处理者，是指处理不满10万人个人信息的个人信息处理者。规定支持小型个人信息处理者在遵守个人信息保护相关法律、行政法规和国家有关规定基础上，依据本规定采取与小型个人信息处理者规模、能力等相当的简化措施保障个人信息安全，保护个人信息权益。

《规定》明确了小型个人信息处理者个人信息处理规则等的简化措施。规定小型个人信息处理者制定公开个人信息处理规则、履行告知义务、取得个人同意、受理个人行使权利申请、删除个人信息等的简化措施。明确小型个人信息处理者向境外提供个人信息时免于申报数据出境安全评估、订立个人信息标准合同、通过个人信息保护认证的条件。

《规定》明确了小型个人信息处理者履行义务的简化措施。规定小型个人信息处理者开展个人信息保护合规审计、个人信息保护影响评估、建立个人信息保护管理制度、履行个人信息安全事件补救和通知义务等的简化措施。明确通过个人信息保护

认证的小型个人信息处理者，在认证有效期内可以免于开展个人信息保护合规审计。

《规定》明确了监督管理要求。规定对小型个人信息处理者不予处罚、从轻或者减轻处罚的情形。明确支持企业、相关社会组织、专业机构等帮助小型个人信息处理者提升个人信息保护能力水平，鼓励履行个人信息保护职责的部门降低小型个人信息处理者合规成本。规定网信部门、公安机关和其他履行个人信息保护职责的部门可以对小型个人信息处理者履行个人信息保护义务情况进行监督检查。明确违反规定的法律责任以及施行时间。

《规定》以附件形式提供了《小型个人信息处理者个人信息保护合规审计自查表》《小型个人信息处理者个人信息保护影响评估表》，便于小型个人信息处理者自行开展个人信息保护合规审计、个人信息保护影响评估。

国家互联网信息办公室

中华人民共和国公安部

令

第25号

《小型个人信息处理者个人信息保护简化措施规定》已经2026年6月26日国家互联网信息办公室2026年第14次室务会会议审议通过，并经公安部同意，现予公布，自2026年9月1日起施行。

国家互联网信息办公室主任 庄荣文

公安部部长 王小洪

2026年7月22日

小型个人信息处理者个人信息保护简化措施规定

第一条 为了支持中小微企业创新发展，简化小型个人信息处理者履行个人信息保护义务的措施，根据《中华人民共和国个人信息保护法》、《网络安全安全管理条例》等法律、行政法规，制定本规定。

第二条 在中华人民共和国境内的小型个人信息处理者实施个人信息保护，适用本规定。

本规定所称小型个人信息处理者，是指处理不

满10万人个人信息的个人信息处理者。

第三条 支持小型个人信息处理者在遵守个人信息保护相关法律、行政法规和国家有关规定基础上,依据本规定采取与小型个人信息处理者规模、能力等相当的简化措施保障个人信息安全,保护个人信息权益。

第四条 小型个人信息处理者个人信息处理规则至少应当包括下列内容:

(一)小型个人信息处理者名称或者姓名;

(二)个人行使权利的受理部门或者人员及联系方式;

(三)个人信息的处理目的、处理方式,处理的个人信息种类、保存期限等。

小型个人信息处理者线下收集个人信息的,可以通过在经营场所醒目位置张贴公告等简便方式公开个人信息处理规则;线上收集个人信息的,可以通过服务协议、产品服务客户端弹窗和网站公告等方式公开个人信息处理规则。

小型个人信息处理者处理不满十四周岁未成年人个人信息的,应当制定专门的个人信息处理规则。

第五条 支持园区、产业基地、商业物业等服务管理单位,为其服务管理范围内开展相同线下业务的小型个人信息处理者统一制定并在显著位置公开个人信息处理规则。同意遵守统一个人信息处理规则且在该规则中被列明的小型个人信息处理者,可以不再制定个人信息处理规则。

第六条 小型个人信息处理者同时符合下列条件的,可以仅通过公开个人信息处理规则向个人履行告知义务,个人信息处理规则应当通过加粗字体、放大字号、标记不同颜色等显著方式向用户提示,并便于查阅和保存:

(一)处理个人信息(不含敏感个人信息)为提供产品或者服务所必需;

(二)不向其他个人信息处理者提供且不对外公开个人信息,并在个人信息处理规则中明示。

法律、行政法规、部门规章对于小型个人信息处理者处理敏感个人信息另有规定的,从其规定。

第七条 小型个人信息处理者公开个人信息处理规则并履行告知义务后,个人因获取产品或者服务需要,在充分知情的前提下自愿、主动向或者自愿、主动配合小型个人信息处理者提供获取产品或者服务所必需的个人信息的,小型个人信息处理者即可按照公开的个人信息处理规则处理其个人信息;为特定目的处理敏感个人信息的,小型个人信息处理者应当在个人信息处理规则中告知处理敏感个人信息的必要性以及对个人权益的影响,取得个人的单独同意。

法律、行政法规、部门规章对于小型个人信息处理者处理敏感个人信息另有规定的,从其规定。

第八条 同时符合下列条件的小型个人信息处理者,可以不再制定个人信息处理规则、履行告知义务:

(一)小型个人信息处理者仅通过网络平台开展个人信息处理活动,且不向网络平台外的其他个人信息处理者提供;

(二)网络平台已针对小型个人信息处理者个人信息处理活动,制定发布了相应的个人信息处理规则,并与小型个人信息处理者约定了各自的权利和义务;

(三)小型个人信息处理者声明遵守网络平台按照本规定要求制定的个人信息处理规则,且处理个人信息为提供产品或者服务所必需,未超出前项规则载明的处理目的、处理方式和个人信息种类范围。

符合前款规定条件的,网络平台已开展个人信息保护合规审计、个人信息保护影响评估,覆盖小型个人信息处理者依托该网络平台开展的个人信息处理活动的,小型个人信息处理者可以不再重复开展。

小型个人信息处理者处理个人信息的目的、处理方式和种类超出网络平台个人信息处理规则范围的,应当按照本规定要求另行制定个人信息处理规则并履行告知、合规审计和影响评估义务。网络平台调整个人信息处理规则的,应当及时通知相关小型个人信息处理者。

第九条 小型个人信息处理者因合并、分立、解散、被宣告破产等原因需要转移个人信息的，可以通过在经营场所醒目位置张贴公告、短信提醒等简便方式告知接收方的名称或者姓名以及联系方式；自主提供线上产品或者服务的，还应当通过其产品服务客户端弹窗公告等方式，告知接收方的名称或者姓名以及联系方式；依托网络平台提供产品或者服务的，可以通过其在网络平台内的经营者页面公告、小程序公告等方式告知。

小型个人信息处理者应当至少提前30个工作日公开发布前款规定的告知事项，公开时长不少于30个工作日。

第十条 小型个人信息处理者向境外提供个人信息，符合下列条件之一的，免于申报数据出境安全评估、订立个人信息标准合同、通过个人信息保护认证：

（一）为订立、履行个人作为一方当事人的合同，如跨境购物、跨境寄递、跨境汇款、跨境支付、跨境开户、机票酒店预订、签证办理、考试服务等，确需向境外提供个人信息的；

（二）按照依法制定的劳动规章制度和依法签订的集体合同实施跨境人力资源管理，确需向境外提供员工个人信息的；

（三）紧急情况下为保护自然人的生命健康和财产安全，确需向境外提供个人信息的；

（四）为履行法定职责或者法定义务，确需向境外提供个人信息的；

（五）关键信息基础设施运营者以外的个人信息处理者自当年1月1日起累计向境外提供不满10万人个人信息（不含敏感个人信息）的；

（六）法律、行政法规或者国家网信部门规定的其他条件。

前款所称向境外提供的个人信息，不包括重要数据。小型个人信息处理者向境外提供个人信息的，应当按照法律、行政法规的规定履行告知、取得个人单独同意等义务。

小型个人信息处理者确需向中华人民共和国境外提供个人信息，依法向网信部门申请数据出境

安全评估的，可以由所在地省级网信部门形成评估结论建议并报国家网信部门核准。

鼓励履行个人信息保护职责的部门、数据跨境服务中心等，为小型个人信息处理者个人信息出境提供咨询等服务。

第十一条 小型个人信息处理者可以通过公开个人行使权利的受理部门或者人员及联系方式，建立个人行使在个人信息处理活动中权利的的申请受理和处置机制。

第十二条 停止产品或者服务的小型个人信息处理者应当采取必要措施删除个人信息，确无能为力删除个人信息的，可以向所在地有关主管部门报告并请求提供帮助；主管部门不明确的，可以向所在地设区的市级网信部门报告。

第十三条 小型个人信息处理者可以按照本规定附件《小型个人信息处理者个人信息保护合规审计自查表》的简便方式，至少每五年开展一次个人信息保护合规审计，并保存合规审计自查表至少五年。

法律、行政法规对于处理未成年人个人信息的合规审计另有规定的，从其规定。

第十四条 小型个人信息处理者可以按照本规定附件《小型个人信息处理者个人信息保护影响评估表》的简便方式开展个人信息保护影响评估，并保存影响评估表至少三年。

第十五条 小型个人信息处理者可以通过在组织管理文件中明确个人信息保护内部管理要求、个人信息安全事件应急处置要求等简便方式，建立个人信息保护管理制度、个人信息安全事件应急预案。

第十六条 发生或者可能发生个人信息泄露、篡改、丢失的，小型个人信息处理者应当立即采取补救措施，依照法律、行政法规规定通知个人，确因客观条件限制无法通过其他方式通知个人的，可以仅通过在经营场所醒目位置张贴公告、在产品服务客户端弹窗和网站公告等简便方式通知个人，并按照规定通知履行个人信息保护职责的部门；涉嫌犯罪的，应当及时向公安机关报案。

第十七条 支持个人信息保护认证机构针对小型个人信息处理者开展认证工作,提高服务质量。通过个人信息保护认证的小型个人信息处理者,在认证有效期内可以免于开展个人信息保护合规审计。

第十八条 小型个人信息处理者开展个人信息处理活动有下列情形之一的,应当不予处罚:

(一)违法行为轻微并及时改正,没有造成危害后果的;

(二)有证据足以证明没有主观过错的,法律、行政法规另有规定的,从其规定;

(三)其他依法不予处罚的情形。小型个人信息处理者开展个人信息处理活动,初次违法且危害后果轻微并及时改正的,可以不予处罚。依法不予处罚的,履行个人信息保护职责的部门应当视情采取约谈、发送提示函等监管措施。

第十九条 小型个人信息处理者开展个人信息处理活动有下列情形之一的,应当从轻或者减轻处罚:

(一)主动消除或者减轻违法行为危害后果的;

(二)主动供述履行个人信息保护职责的部门尚未掌握的违法行为的;

(三)发生个人信息安全事件时,及时通知个人和采取补救措施,并主动通知有关部门的;

(四)配合履行个人信息保护职责的部门查处违法行为有立功表现的;

(五)其他依法从轻或者减轻处罚的情形。

第二十条 支持企业、相关社会组织、专业机构等通过组织培训、讲座、普法活动、咨询辅导等方式,帮助小型个人信息处理者提升个人信息保护能力水平。鼓励履行个人信息保护职责的部门为小型个人信息处理者提供安全便捷个人信息处理的基础设施、技术工具、咨询服务等,降低小型个人信息处理者合规成本。

第二十一条 网信部门、公安机关和其他履行个人信息保护职责的部门可以采取抽查评估、审计报告等方式对小型个人信息处理者履行个人信息保护义务情况进行监督检查,小型个人信息处理者

应当予以配合。

网信部门、公安机关和其他履行个人信息保护职责的部门发现小型个人信息处理者违法处理个人信息、多次发生个人信息安全事件的,应当依照《中华人民共和国个人信息保护法》、《网络数据安全管理条例》等有关法律、行政法规予以处理,依照有关法律、行政法规的规定记入信用档案,并予以公示。

第二十二条 本规定自2026年9月1日起施行。

《小型个人信息处理者个人信息保护简化措施规定》答记者问

原载:“网信中国”微信公众号

近日,国家网信办、公安部联合公布《小型个人信息处理者个人信息保护简化措施规定》(以下简称《规定》)。国家网信办有关负责同志就《规定》回答了记者提问。

问1:请介绍一下《规定》的出台背景?

答:党的二十大报告提出,加快建设网络强国、数字中国,强化网络、数据等安全保障体系建设,加强个人信息保护,支持中小微企业发展。党的二十届四中全会提出,加强个人信息保护,支持中小企业和个体工商户发展。《个人信息保护法》第六十二条规定,国家网信部门统筹协调有关部门针对小型个人信息处理者,制定专门的个人信息保护规则、标准。

出台《规定》,是贯彻落实党中央关于加强个人信息保护和支持中小微企业发展工作部署,落实《个人信息保护法》有关要求的重要举措,有利于在坚守个人信息安全底线的同时,降低小型个人信息处理者合规成本,提升小型个人信息处理者个人信息保护水平,为中小微企业创新发展提供法治保障。

问2:《规定》的适用范围是什么?

答:在中华人民共和国境内的小型个人信息处理者实施个人信息保护,适用《规定》。《规定》所称小型个人信息处理者,是指处理不满10万人个人信息的个人信息处理者。

问3:《规定》中的“处理不满10万人个人信息”,相关数量如何进行统计?

答:处理不满10万人个人信息,在计数时,不包含10万人本数,按照个人信息处理者当前处理个人信息涉及的自然人数量进行累计统计,不包括已经删除的个人信息。

问4:《规定》的主要内容是什么?

答:《规定》主要对下列内容进行了规定:

一是明确小型个人信息处理者个人信息处理规则等的简化措施,规定小型个人信息处理者制定公开个人信息处理规则、履行告知义务、取得个人同意、受理个人行使权利申请、删除个人信息等的简化措施。

二是明确小型个人信息处理者履行个人信息保护义务的简化措施,规定小型个人信息处理者开展个人信息保护合规审计、个人信息保护影响评估、建立个人信息保护管理制度、履行个人信息安全事件补救和通知义务等的简化措施。

三是明确监督管理要求,规定对小型个人信息处理者不予处罚、从轻或者减轻处罚的情形,鼓励履行个人信息保护职责的部门降低小型个人信息处理者合规成本,明确违反规定的法律责任。同时,《规定》以附件形式提供了《小型个人信息处理者个人信息保护合规审计自查表》《小型个人信息处理者个人信息保护影响评估表》,便于小型个人信息处理者自行开展个人信息保护合规审计、个人信息保护影响评估。

问5:《规定》对小型个人信息处理者制定公开个人信息处理规则,提出了哪些简化措施?

答:《规定》明确小型个人信息处理者个人信息处理规则应当包括的内容,并提供线下、线上公开个人信息处理规则的简便方式。按照《规定》要求,小型个人信息处理者个人信息处理规则至少应

当包括下列内容:

(一)小型个人信息处理者名称或者姓名;

(二)个人行使权利的受理部门或者人员及联系方式;

(三)个人信息的处理目的、处理方式,处理的个人信息种类、保存期限等。小型个人信息处理者线下收集个人信息的,可以通过在经营场所醒目位置张贴公告等简便方式公开个人信息处理规则;线上收集个人信息的,可以通过服务协议、产品服务客户端弹窗和网站公告等方式公开个人信息处理规则。

《规定》明确,支持园区、产业基地、商业物业等服务管理单位,为其服务管理范围内开展相同线下业务的小型个人信息处理者统一制定并在显著位置公开个人信息处理规则。同意遵守统一个人信息处理规则且在该规则中被列明的小型个人信息处理者,可以不再制定个人信息处理规则。

问6:《规定》对小型个人信息处理者履行告知、取得个人同意义务,提出了哪些简化措施?

答:小型个人信息处理者同时符合下列条件的,可以仅通过公开个人信息处理规则向个人履行告知义务:(一)处理个人信息(不含敏感个人信息)为提供产品或者服务所必需;(二)不向其他个人信息处理者提供且不对外公开个人信息,并在个人信息处理规则中明示。

对于公开个人信息处理规则并履行告知义务的小型个人信息处理者,《规定》明确,个人因获取产品或者服务需要,在充分知情的前提下自愿、主动向或者自愿、主动配合小型个人信息处理者提供获取产品或者服务所必需的个人信息,小型个人信息处理者即可按照公开的个人信息处理规则处理其个人信息;为特定目的处理敏感个人信息的,小型个人信息处理者应当在个人信息处理规则中告知处理敏感个人信息的必要性以及对个人权益的影响,取得个人的单独同意。

法律、行政法规、部门规章对于小型个人信息处理者处理敏感个人信息另有规定的,从其规定。

问7:《规定》对依托网络平台处理个人信息

的小型个人信息处理者，提出了哪些简化措施？

答：同时符合下列条件的小型个人信息处理者，可以不再制定个人信息处理规则、履行告知义务：

（一）小型个人信息处理者仅通过网络平台开展个人信息处理活动，且不向网络平台外的其他个人信息处理者提供；

（二）网络平台已针对小型个人信息处理者个人信息处理活动，制定发布了相应的个人信息处理规则，并与小型个人信息处理者约定了各自的权利和义务；

（三）小型个人信息处理者声明遵守网络平台按照本规定要求制定的个人信息处理规则，且处理个人信息为提供产品或者服务所必需，未超出前项规则载明的处理目的、处理方式和个人信息种类范围。

对于符合前述规定条件的小型个人信息处理者，网络平台已开展个人信息保护合规审计、个人信息保护影响评估，覆盖小型个人信息处理者依托该网络平台开展的个人信息处理活动的，小型个人信息处理者可以不再重复开展。

小型个人信息处理者处理个人信息的目的、处理方式和种类超出网络平台个人信息处理规则范围的，应当按照《规定》要求另行制定个人信息处理规则并履行告知、合规审计和影响评估义务。网络平台调整个人信息处理规则的，应当及时通知相关小型个人信息处理者。

问 8：《规定》对小型个人信息处理者开展个人信息保护合规审计、个人信息保护影响评估，提出了哪些简化措施？

答：《规定》以附件形式提供了《小型个人信息处理者个人信息保护合规审计自查表》《小型个人信息处理者个人信息保护影响评估表》，小型个人信息处理者可以按照附件的简便方式，开展个人信息保护合规审计和个人信息保护影响评估。个人信息保护合规审计至少每五年开展一次，法律、行政法规对于处理未成年人个人信息的合规审计另有规定的，从其规定。

《规定》明确，支持个人信息保护认证机构针

对小型个人信息处理者开展认证工作，提高服务质量。通过个人信息保护认证的小型个人信息处理者，在认证有效期内可以免于开展个人信息保护合规审计。

问 9：《规定》对小型个人信息处理者的监督管理，作了哪些规定？

答：对小型个人信息处理者的监督管理，《规定》作出以下规定：

一是明确对小型个人信息处理者不予处罚、从轻或者减轻处罚的情形。规定对小型个人信息处理者依法不予处罚的，履行个人信息保护职责的部门应当视情采取约谈、发送提示函等监管措施。

二是明确支持企业、相关社会组织、专业机构等通过组织培训、讲座、普法活动、咨询辅导等方式，帮助小型个人信息处理者提升个人信息保护能力水平。

三是明确鼓励履行个人信息保护职责的部门为小型个人信息处理者提供安全便捷个人信息处理的基础设施、技术工具、咨询服务等，降低小型个人信息处理者合规成本。

四是明确网信部门、公安机关和其他履行个人信息保护职责的部门可以采取抽查评估、审计报告等方式对小型个人信息处理者履行个人信息保护义务情况进行监督检查。

专家解读 | 简化保护措施 兼顾个人信息权益与小型企业创新发展利益

原载：“网信中国”微信公众号

近日，国家网信办、公安部联合公布《小型个人信息处理者个人信息保护简化措施规定》（以下简称《规定》），针对小型个人信息处理者制定专门的个人信息保护规则，为小型个人信息处理者履行个人信息保护法律法规义务提供简化措施，支持中小微企业创新发展。

一、制定依据及其利益平衡考量

《个人信息保护法》的立法目的在于保护自然人的个人信息权益，规范个人信息处理活动，促进个人信息合理利用。个人信息权益保护与合理利用个人信息之间的矛盾是个人信息保护法所要解决的基本问题：十分严格的个人信息权益保护势必影响处理者合理利用个人信息，包括某些个人信息难以被利用以及利用个人信息产生过高的成本（包括合规成本），继而使得个人信息处理者利用个人信息数据的创新能力和经济效益受到限制。因此，个人信息保护法律制度通过规定处理个人信息的原则和合法性基础等规则、个人信息处理者的义务、个人在个人信息处理中的权利、保护个人信息的国家机关及其职责和相关法律责任等，对个人信息权益给予符合本国法治水平的保护，同时给个人信息处理留下适度空间，并保护个人信息处理者基于个人信息处理而产生的财产权益。

《个人信息保护法》第五章集中规定了个人信息处理者的义务，其中第五十一条至第五十七条是一般个人信息处理者须履行的义务，第五十八条是提供重要互联网平台服务、用户数量巨大、业务类型复杂的个人信息处理者应当履行的特别义务。全面履行此等义务，对于具有一定规模的个人信息处理者而言，是其生产经营成本所能合理负担的，但是对于小型个人信息处理者而言，需要通过简化措施进行减负。基于此等利益衡量，《个人信息保护法》第六十二条第一款第二项规定：“国家网信部门统筹协调有关部门依据本法推进下列个人信息保护工作：（二）针对小型个人信息处理者、处理敏感个人信息以及人脸识别、人工智能等新技术、新应用，制定专门的个人信息保护规则、标准。”立法授权国家网信办统筹协调有关部门制定本《规定》。

二、适用对象与适用地域

（一）小型个人信息处理者的界定

《规定》适用于小型个人信息处理者处理个人信息的活动。首先，小型个人信息处理者也是个人信息处理者，是在个人信息处理活动中自主决定处

理目的、处理方式的组织、个人。其次，需要强调的是小型个人信息处理者的本质属性。

《规定》对小型个人信息处理者进行了界定：

“本规定所称小型个人信息处理者，是指处理不满10万人个人信息的个人信息处理者。”《规定》对小型个人信息处理者的界定，立足于其处理个人信息数量（不满10万人个人信息），而不是处理者的生产经营规模。有些专门从事信息产业生产经营活动的小微企业，虽然经济体量较小，但是处理的个人信息数量较大，则不属于小型个人信息处理者；有些较大的企业如电子产品代工厂，虽然经济体量较大但是不处理或者较少处理个人信息（处理量不满10万人），则属于小型个人信息处理者。

（二）地域效力

《规定》第二条规定了适用的地域范围：“在中华人民共和国境内的小型个人信息处理者实施个人信息保护，适用本规定。”对本条的理解和适用，还需要遵循《个人信息保护法》第十二条的规定：“国家积极参与个人信息保护国际规则的制定，促进个人信息保护方面的国际交流与合作，推动与其他国家、地区、国际组织之间的个人信息保护规则、标准等互认。”

三、简化思路和主要内容

（一）简化思路

《规定》在界定小型个人信息处理者的基础上，结合小型个人信息处理者特点，明确制定公开个人信息处理规则、告知个人、取得个人同意等简化程序，便利小型个人信息处理者落实《个人信息保护法》关于处理个人信息的关键要求。《规定》确定了简化制定个人信息处理规则的思路。针对大量小型个人信息处理者依托网络平台、园区或商业物业等提供产品服务的情形，《规定》提出在遵守网络平台、园区或商业物业等统一制定的个人信息处理规则的前提下，符合相关条件的，可以不再制定个人信息处理规则等便捷措施。《规定》还确定了支持性政策和减免处罚的思路。鼓励履行个人信息保护职责的部门为小型个人信息处理者提供安全便捷个人信息处理的基础设施、技术工具、咨询服务

等。《规定》提出，对于违法行为轻微并及时改正，没有造成危害后果等情况的，不予处罚或者从轻减轻处罚，为中小微企业创新发展构建良好的政策环境。

（二）主要内容

《规定》正文共二十二条，重点对个人信息保护总体要求、个人信息处理规则简化、小型个人信息处理者义务履行方式简化、不予和从轻或者减轻处罚等作出规定。

一是关于个人信息保护总体要求。《规定》明确了小型个人信息处理者定义范围、小型个人信息处理者处理个人信息基本原则等总体要求。**二是**关于个人信息处理规则简化。《规定》明确了公开个人信息处理规则、统一制定处理规则、告知个人、取得个人同意、依托网络平台处理个人信息、个人信息出境、个人行权受理、删除个人信息等的简便处理措施。**三是**关于小型个人信息处理者义务履行方式简化。《规定》明确了小型个人信息处理者个人信息保护合规审计、影响评估、管理制度、安全事件应急预案、安全事件通知的简便处理措施，以及个人信息保护认证、合规审计的制度衔接规定等。**四是**关于不予处罚、从轻或者减轻处罚。《规定》明确了小型个人信息处理者不予处罚情形，从轻或者减轻处罚情形，规定履行个人信息保护职责的部门对小型个人信息处理者的支持措施。同时，《规定》还通过附件提供了《小型个人信息处理者个人信息保护合规审计自查表》和《小型个人信息处理者个人信息保护影响评估表》。

四、结论与展望

出台《规定》，旨在给小型个人信息处理者处理个人信息提供简化措施，以减轻其个人信息保护合规负担。这些简化措施坚持了《个人信息保护法》确定的基本原则和关键制度，守住了个人信息保护的底线，同时也给小型个人信息处理者进行了较大减负，有利于促进中小企业创新发展。

作者：张新宝 中国人民大学法学院吴玉章高级讲师教授、中国法学会网络信息法学研究会副会长兼

12/101

学术委员会主任

数据出境安全管理政策法规问答 (2026年7月)

原载：“网信中国”微信公众号

国家互联网信息办公室持续加强数据出境安全管理政策法规宣介，指导和帮助数据处理者高效合规开展数据出境活动。经对近期收到的咨询问题进行研究，现将一些有代表性的问题和答复公布如下。

1. 个人信息处理者向境外提供个人信息，应如何有效履行告知同意义务？

答：根据《个人信息保护法》第三十九条规定，个人信息处理者向中华人民共和国境外提供个人信息的，应当履行告知、取得个人单独同意等义务。告知内容包括：境外接收方的名称或者姓名、联系方式、处理目的、处理方式、个人信息的种类以及个人向境外接收方行使本法规定权利的方式和程序等事项。同时，根据《个人信息保护法》第三十条规定，如向境外提供的是敏感个人信息，还应当向个人告知出境敏感个人信息的必要性以及对个人权益的影响。

个人信息处理者取得个人单独同意个人信息出境，应当具体明确，不得与其他个人信息处理活动捆绑，不得采取“一揽子”授权方式取得同意。取得个人单独同意的方式可参考 GB/T 42574-2023《信息安全技术 个人信息处理中告知和同意的实施指南》，采用书面签署、弹窗确认、邮件短信回复等方式开展。个人信息出境活动属于《个人信息保护法》第十三条第一款第二项至第七项规定情形的，不需取得个人同意，但仍应履行个人信息出境告知义务。

2. 已通过数据出境安全评估的，申请延长评估结果有效期应当符合哪些条件？

答：根据《促进和规范数据跨境流动规定》第

九条“通过数据出境安全评估的结果有效期为3年，自评估结果出具之日起计算，有效期届满，需要继续开展数据出境活动且未发生需要重新申报数据出境安全评估情形的，数据处理者可以在有效期届满前60个工作日内通过所在地省级网信部门向国家网信部门提出延长评估结果有效期申请。经国家网信部门批准，可以延长评估结果有效期3年”。

按照《数据出境安全评估申报指南（第三版）》，申请延长评估结果有效期应当同时符合以下条件：

一是数据出境的目的、范围等未发生变化；
二是数据处理者和境外接收方等未发生变化；
三是出境个人信息的，未来三年涉及自然人数
量增幅不超过原评估结果准予过去三年出境数
量的20%；

四是出境重要数据的，未来三年出境数据规模
增幅不超过原评估结果准予过去三年出境数据规
模的20%；

五是与境外接收方订立的法律文件符合《数据
出境安全评估办法》第九条规定；

六是过去三年数据出境活动严格按照评估结
果通知书合规开展，且未发生重大数据安全事件。

3. 人员招聘场景中，如何判断境内求职者简 历出境的必要性？

答：《个人信息保护法》第六条规定，“处理
个人信息应当具有明确、合理的目的，并应当与处
理目的直接相关，采取对个人权益影响最小的方式。
收集个人信息，应当限于实现处理目的的最小范围，
不得过度收集个人信息”。《数据出境安全评估办
法》第五条、《个人信息出境标准合同办法》第五
条、《个人信息出境认证办法》第六条明确了个人
信息处理者向境外提供个人信息前应当对个人信
息出境的必要性进行评估。

招聘场景下，将境内求职者简历等个人信息提
供给境外的集团总部或者关联机构，应当从出境活
动与招聘事项关联度、涉及自然人数规模、出境
个人信息数据项范围等方面判断必要性。如境外总
部或者关联机构不参与境内应聘人员录用决策，数
据出境不具备必要性。如境外总部或者关联机构直

接参与境内应聘人员录用决策，出境应聘人员规模
为满足境外决策所需的最小数量，出境个人信息数
据项为满足境外决策所需的最小范围，相关个人信
息出境活动应当按照《促进和规范数据跨境流动规
定》要求，采取申报数据出境安全评估、订立个人
信息出境标准合同、通过个人信息出境认证等方式
合规开展，并按照法律、行政法规的规定告知并取
得个人单独同意，开展个人信息保护影响评估等。

中央网络安全和信息化委员会印 发《深化互联网协议第六版（IPv6） 技术创新和融合应用实施方案 （2026—2030年）》

原载：“网信中国”微信公众号

近日，中央网络安全和信息化委员会印发《深
化互联网协议第六版（IPv6）技术创新和融合应用
实施方案（2026—2030年）》（以下简称《实施方
案》）。

《实施方案》坚持以习近平新时代中国特色社会主义思想为指导，深入贯彻党的二十大和二十届
历次全会精神，认真落实四中全会部署，全面贯彻
习近平总书记关于网络强国的重要思想，完整准确
全面贯彻新发展理念，明确了深化IPv6技术创新
和融合应用的路线图、时间表和任务书，是“十五
五”时期IPv6高质量发展的指导性文件。

《实施方案》指出，“十五五”时期是我国
IPv6从规模部署向全面赋能跨越的关键时期，要以
推动IPv6全面赋能经济社会高质量发展为主线，
着力深化IPv6技术创新和融合应用，着力提升IPv6
端到端贯通水平，着力构建IPv6产业生态体系，
全面增强IPv6内生发展动力，巩固拓展创新发展
优势，促进各行业各领域数字化网络化智能化全面
跃升，为网络强国建设提供有力支撑。

《实施方案》明确，到2027年，IPv6端到端

贯通水平全面提升，技术创新和融合应用取得重要阶段性成果。IPv6 活跃用户数达 9 亿，IPv6 网络流量占比达 38%。到 2030 年，IPv6 与经济社会各行业各领域广泛深度融合，构建技术先进、开放创新、内生自驱、安全可靠的 IPv6 产业生态。IPv6 活跃用户数达 9.5 亿，IPv6 网络流量占比达 42%。新增网络默认优先提供 IPv6 地址，加快向 IPv6 单栈演进，推动形成以 IPv6 为主导的网络服务和应用体系。

《实施方案》围绕“十五五”发展目标，聚焦 IPv6 技术创新和融合应用关键环节，部署了加强技术研发与融合创新、推进标准制定与实施、强化管理制度规范、推进单栈技术攻关和部署应用、增强网络和应用设施服务能力、提升终端支持能力、深化商业互联网应用改造、拓展行业融合应用、加强人才培养和宣传推广、增强安全保障能力等共 10 个方面 36 项重点任务。

《实施方案》强调，在中央网信委领导下，中央网信办会同工业和信息化部等部门加强统筹协调、央地联动，健全高效协同工作推进机制，压实工作责任，推动任务目标落地落实，实现“十五五”IPv6 高质量发展。

深化互联网协议第六版（IPv6）技术创新和融合应用实施方案（2026—2030 年）

互联网协议第六版（IPv6）是互联网升级演进的必然趋势、网络强国建设的基础支撑。“十五五”时期是我国 IPv6 从规模部署向全面赋能跨越的关键时期，为加快推动 IPv6 高质量发展，制定本实施方案。

一、总体要求

坚持以习近平新时代中国特色社会主义思想为指导，深入贯彻党的二十大和二十届历次全会精神，认真落实四中全会部署，全面贯彻习近平总书记关于网络强国的重要思想，完整准确全面贯彻新发展理念，以推动 IPv6 全面赋能经济社会高质量发展为主线，着力深化 IPv6 技术创新和融合应用，着力提升 IPv6 端到端贯通水平，着力构建 IPv6 产

业生态体系，全面增强 IPv6 内生发展动力，巩固拓展创新发展优势，促进各行业各领域数字化网络化智能化全面跃升，为网络强国建设提供有力支撑。

主要目标是：到 2027 年，IPv6 端到端贯通水平全面提升，技术创新和融合应用取得重要阶段性成果。IPv6 网络性能和服务质量全面优化，互联网应用 IPv6 放量持续扩大，固定宽带用户 IPv6 连通率大幅提升，新增网络终端全面支持并启用 IPv6，政企专线 IPv6 流量占比明显提高，行业融合应用更加广泛深入，IPv6 演进技术（“IPv6+”）体系持续健全，IPv6 与人工智能融合创新取得新成效，安全产品 IPv6 防护能力进一步提高。到 2030 年，IPv6 与经济社会各行业各领域广泛深度融合，构建技术先进、开放创新、内生自驱、安全可靠的 IPv6 产业生态。新增网络默认优先提供 IPv6 地址，加快向 IPv6 单栈演进，打造更多“IPv6+”新应用新模式新业态，推动形成以 IPv6 为主导的网络服务和应用体系。

“十五五”IPv6 技术创新和融合应用主要指标

序号	指标	2025 年 基期值	2027 年 预期值	2030 年 预期值
1	IPv6 活跃用户数（亿）	8.69	9	9.5
2	IPv6 网络流量占比（%）	33.79	38	42
3	移动物联网 IPv6 连接数（亿）	16.58	20	25
4	固定宽带用户 IPv6 连通率（%）	—	70	80
5	主要互联网应用固网侧 IPv6 流量占比（%）	—	70	80
6	发布 IPv6 国家标准数量（项）	36	40	50
7	“IPv6+”行业专网数量（个）	>700	800	1000
8	县级以上党政机关办公网络 IPv6 开通率（%）	—	30	50

注：标“—”指标暂未形成 2025 年基期值数据。

二、重点任务

（一）加强技术研发与融合创新

1. 强化 IPv6 创新技术攻关与成果转化。持续完善“IPv6+”技术体系，深化 IPv6 分段路由（SRv6）、网络切片、随流检测等技术普及应用，突破应用感知网络、新型组播、韧性网络等前沿技术，强化创新技术推广应用与成果转化。加强基于“IPv6+”的智能广域网、新型城域网、算力网络等关键技术攻关。加强互联网体系结构创新研究，依托未来互联网试验设施等开展 IPv6 科研试验和技术验证。

2. 促进 IPv6 与人工智能融合创新。推动 IPv6 赋能人工智能创新发展, 深化 IPv6 与大模型训练推理、智能体互联协同等场景融合创新。探索基于 IPv6 的智能原生网络架构, 推动网络向高阶自智演进。研究推进基于 IPv6 的智能互联网创新演进。

3. 推进 IPv6 与网络信息技术融合创新。加快 IPv6 与第六代移动通信(6G)、数字孪生、量子信息、数据流通利用等领域融合创新。在车联网、工业互联网、卫星互联网、区块链网络等新一代信息基础设施建设中同步部署 IPv6。推动 IPv6 赋能全国一体化算力网高质量互联与智能调度。探索在国家数据基础设施中应用 IPv6。

(二) 推进标准制定与实施

4. 完善 IPv6 标准体系。持续优化 IPv6 国家标准布局, 加快重点领域 IPv6 国家标准研制。发挥社会团体、重点企业等作用, 制定高水平 IPv6 行业和团体标准。加强标准宣传推广和实施效果评估。

5. 积极参与 IPv6 国际标准制定。强化产学研用联动, 加大国际标准组织参与力度, 支持国内企业机构积极参与基础性、关键性 IPv6 国际标准研制, 提高技术创新成果国际影响力。

(三) 强化管理制度规范

6. 加强 IPv6 地址资源管理。完善 IP 地址备案管理制度, 加强 IPv6 地址管理。建立高效协同的 IPv6 地址资源统筹规划机制, 推动企业机构科学规划和申请 IPv6 地址资源, 引导行业规范使用 IPv6 地址资源。加强新兴领域 IPv6 地址资源前瞻布局, 推进 IPv6 地址典型应用示范。

7. 完善政府信息化 IPv6 管理制度。在政府投资建设的信息化项目、政务信息化建设、政务移动互联网应用程序相关制度规范中明确 IPv6 支持能力要求。加大数字化转型等政策对 IPv6 的支持力度。

8. 强化移动互联网应用上架 IPv6 支持要求。建立移动互联网应用(App) IPv6 支持度评测体系, 完善应用商店上架标准规范, 鼓励新上架 App 支持 IPv6。

(四) 推进单栈技术攻关和部署应用

9. 加强 IPv6 单栈技术攻关。开展 IPv6 单栈组网架构等关键技术攻关, 在典型网络和业务场景中验证评估 IPv6 单栈技术方案。选取重点区域、特色领域开展 IPv6 单栈试点应用, 形成可复制的典型经验。

10. 开展 IPv6 单栈部署应用。研究制定 IPv6 单栈演进实施指南, 建立 IPv6 单栈能力评测体系。在端到端可控业务场景率先开展 IPv6 单栈部署。推动互联网应用 IPv6 单栈化改造, 新增应用提供单栈服务能力。

(五) 增强网络和应用设施服务能力

11. 持续优化 IPv6 网络性能。强化 IPv6 网络运维管理, 开展 IPv6 网络性能常态化监测, 加大 IPv6 网络调优力度, 实现 IPv6 关键指标优于 IPv4, 确保 IPv6 网络服务稳定可靠。

12. 加快提升 IPv6 服务水平。推动家庭宽带、企业宽带、互联网专线等业务优先提供 IPv6 接入, 优化 IPv6 业务开通流程, 提高装维服务能力。

13. 加快“网络去 NAT”进程。加大力度推进“网络去 NAT”专项工作。推动基础电信企业探索和部署 IPv6 接入技术升级, 提高固定网络接入 IPv6 连通率。组织基础电信企业与互联网企业协同开展 IPv6 流量监测, 打通性能堵点, 理顺流量路径, 实现“网用互促”。

14. 提高云产品 IPv6 开通率。推动云服务平台实现全产品、全区域支持 IPv6, 新上线云产品和新建节点实现原生支持与默认启用 IPv6。引导云产品存量用户向 IPv6 迁移。

15. 强化算力设施承载业务 IPv6 升级改造。提高算力设施 IPv6 网络接入能力与服务质量, 新开通互联网线路默认启用 IPv6, 推动承载业务优先使用 IPv6 提供服务。

16. 加快推进广电网络 IPv6 端到端贯通。加快广播电视传输网络及宽带数据网络 IPv6 改造, 推动广电网络业务系统和平台全面支持 IPv6。新建国家及省级广播电视光缆干线传输网、有线电视网、宽带数据骨干网支持 IPv6, 推动广电数据中心、内容分发网络(CDN)、云服务平台、域名解析系统

支持 IPv6。加强广电行业 IPv6 监测评估。

（六）提升终端支持能力

17. 提高家庭路由器 IPv6 连通率。严格落实无线电发射设备型号核准制度，做好无线局域网设备支持 IPv6 协议能力测试，加强监督检查，推动新生产设备默认启用 IPv6。加快存量定制路由器升级替换，引导用户替换老旧自有路由器。

18. 提高智慧家庭产品 IPv6 支持能力。推动智能电视、投影仪、机顶盒、摄像头、音箱等智慧家庭产品支持并启用 IPv6，发挥消费品以旧换新政策作用，引导用户替换老旧产品。推动智慧家庭系统平台全链路支持 IPv6。鼓励开展智慧家庭产品 IPv6 评测认证。打造新一代支持 IPv6 的智慧家庭融合业务服务平台。

19. 扩大物联网终端 IPv6 应用规模。完善物联网终端技术标准，推广支持 IPv6 的低功耗芯片和模组应用。推动智能网联汽车、5G 便携式终端等支持并启用 IPv6。提高智慧安防、智慧城市、智能制造、智慧交通等场景终端设备 IPv6 支持能力。

（七）深化商业互联网应用改造

20. 深化商业互联网应用 IPv6 放量引流。深入推进大型商业网站、App 及小程序 IPv6 升级改造，推动全量业务、域名、CDN 节点等支持并优先提供 IPv6 服务，持续扩大 IPv6 内容源规模。

21. 推动生成式大模型应用支持 IPv6 访问。推动新增生成式大模型个人应用支持 IPv6 访问，加快存量应用 IPv6 升级改造。推动行业大模型上线部署同步支持 IPv6。研究在生成式人工智能服务备案中纳入 IPv6 相关要求。

22. 促进 IPv6 与互联网新应用新业态融合发展。推动 IPv6 与智能体应用、超高清视频、虚拟现实/增强现实（VR/AR）、沉浸式体验、云游戏等新应用新业态融合发展。探索基于 IPv6 的智能体标识技术。

（八）拓展行业融合应用

23. 深化政务网络和应用 IPv6 升级改造。推动县级以上党政机关办公网络开通 IPv6，在网络设备和终端替换升级中同步启用 IPv6。持续推进国

家和地方电子政务外网、政务专网、政务云平台、各级政务服务平台、政务 App 等 IPv6 改造，全面提供 IPv6 服务。

24. 加快中央企业 IPv6 演进升级。推进国务院国资委履行出资人职责的中央企业内部网信基础设施及综合办公、经营管理、生产运营等信息系统 IPv6 升级改造，重点深化门户网站、公众在线服务窗口、App 等 IPv6 升级改造，鼓励探索 IPv6 单栈部署。

25. 深化金融机构 IPv6 创新应用。强化“IPv6+”赋能金融数字化业务创新。深化银行、保险、证券等金融机构广域网、分支机构网络、数据中心及金融机具的 IPv6 升级改造，提高面向公众服务的应用系统 IPv6 支持能力。健全金融机构 IPv6 运维管理机制，构建适配 IPv6 的全方位安全管理及防护体系。

26. 推动 IPv6 与教育新型基础设施融合发展。强化教科网主干网、教学科研大型公共服务基础设施等 IPv6 服务能力，推进互联互通专线 IPv6 单栈部署，加强 IPv6 流量监测分析。推动 IPv6 源地址验证技术等创新成果在高校校园网规模应用。深化 IPv6 在数字教育资源公共服务体系中的融合应用，推动“IPv6+”技术在教育专网、智慧教育等场景落地应用。

27. 拓展卫生健康领域 IPv6 改造广度。推进国家智慧健康数据中心 IPv6 升级改造。加快推进县级及以上医疗卫生机构办公网络、公众服务平台、门户网站及 App 支持 IPv6。推动 IPv6 在医疗信息化、互联网医院、远程医疗等场景融合应用。

28. 提高文化旅游行业 IPv6 支持能力。推动景区、文博场馆等区域支持 IPv6 网络访问。推动各类文旅设施和系统 IPv6 升级改造，促进 IPv6 赋能智慧旅游新场景。

29. 深化媒体数字化转型 IPv6 融合应用。持续推动中央、省、市级媒体和县级融媒体中心 IPv6 升级改造。推动新闻直播报道、数字报、电子期刊等业务平台支持 IPv6，拓展 IPv6 在智能采编、媒体大数据分析等场景创新应用。

30. 推动 IPv6 赋能传统基础设施升级改造。

持续推动制造、民政、人社、自然资源、生态环境、交通、水利、农业农村、应急管理、林草、民航、铁路等领域网络基础设施、信息化系统、互联网应用和终端设备支持并启用 IPv6。鼓励重点行业建设“IPv6+”行业专网，推动 IPv6 与行业场景融合应用。

（九）加强人才培养和宣传推广

31. 加强 IPv6 教育培训。推动 IPv6 纳入高等教育、职业教育课程体系。支持各行业系统开展 IPv6 培训。鼓励企业机构建设 IPv6 实训平台。推动 IPv6 技术能力纳入网络相关专业技术资格要求。

32. 加大宣传普及力度。开展 IPv6 主题宣传活动，广泛普及 IPv6 技术优势与应用价值，提高全社会对 IPv6 的关注度和使用度。支持企业机构、社会团体等发布 IPv6 相关报告，展现 IPv6 新进展新成果。

33. 加强行业交流合作。搭建产学研用交流合作平台，鼓励举办地区和行业 IPv6 交流活动，促进经验共享互鉴，提升 IPv6 应用水平。

（十）增强安全保障能力

34. 加强 IPv6 安全技术攻关。加强 IPv6 原生安全增强、身份管理、访问控制、零信任架构、威胁样本生成等关键技术攻关。提高 IPv6 异常流量检测与威胁关联研判等技术能力。

35. 提升安全产品 IPv6 防护能力。加强安全产品 IPv6 功能和性能测试，提升安全性、兼容性和可靠性。引导不支持 IPv6 的存量安全设备替换升级。研发面向“IPv6+”等新兴技术的安全产品。

36. 完善 IPv6 安全防护体系。强化 IPv6 安全监测，持续提高 IPv6 安全态势感知、通报预警和应急响应能力。

三、组织实施

在中央网信委领导下，充分发挥相关协调机制作用，中央网信办会同工业和信息化部等部门加强统筹协调、央地联动，健全高效协同工作推进机制，压实工作责任，推动任务目标落地落实。完善 IPv6 发展监测体系，提高地方和行业 IPv6 监测分析能

力。开展 IPv6 国际交流合作，共商共建开放协同、互利共赢的国际合作生态。

《深化互联网协议第六版（IPv6）技术创新和融合应用实施方案（2026—2030 年）》答记者问

原载：“网信中国”微信公众号

近日，中央网络安全和信息化委员会印发《深化互联网协议第六版（IPv6）技术创新和融合应用实施方案（2026—2030 年）》（以下简称《实施方案》）。中央网信办有关负责人就《实施方案》有关问题回答了记者提问。

一、请介绍一下出台《实施方案》的背景？

答：当前，全球网络信息技术创新处于密集活跃期，人工智能与互联网融合创新，新技术新业务新应用新场景不断涌现，移动互联网加速迈向智能互联网，对网络承载能力提出更高要求。IPv6 能够提供充足的网络地址和广阔的创新空间，为智能互联网创新发展提供基础技术支撑。

党中央、国务院高度重视我国下一代互联网的发展，“十三五”、“十四五”时期对 IPv6 规模部署和应用作出重要部署，明确 IPv6 发展路线图、任务书、时间表，推动我国 IPv6 发展态势不断向好，成效显著，为经济社会发展提供了坚实数字底座。“十五五”时期是信息化迈向数字化、网络化、智能化全面跃升的新阶段，也是我国 IPv6 从规模部署向全面赋能跨越的关键时期。制定出台《实施方案》是全面贯彻习近平总书记关于网络强国的重要思想的有力举措，也是落实国家“十五五”规划纲要的实际行动，指导和推动各地区、各部门、各单位深化 IPv6 技术创新和融合应用工作，全面增强 IPv6 内生发展动力，巩固拓展创新发展优势，实现“十五五”IPv6 发展新突破。

二、“十四五”时期我国 IPv6 发展取得哪些成效？

答：“十四五”时期，中央网信办会同国家发展改革委、工业和信息化部等部门加强统筹协调、政策支持、系统推进和督促落实，各地区各有关部门认真贯彻中央网信委决策部署，推动我国 IPv6 规模部署和应用实现跨越式发展，IPv6 技术、产业、设施、应用和安全体系不断完善，我国已成为全球 IPv6 技术和产业创新的重要推动力量。

主要指标方面，截至 2025 年底，8 项 IPv6 主要指标全部超额完成预期目标。**用户规模快速增长**，IPv6 活跃用户数达 8.69 亿，占全部网民数的 77.38%，五年增长 88%；物联网 IPv6 连接数从无到有，五年增长至 16.58 亿。**网络流量大幅跃升**，移动网络和固定网络 IPv6 流量占比分别达 70.85% 和 32.38%，五年分别增长 3.1 倍和 7.1 倍。**地址资源储备充足**，我国 IPv6 地址资源总量达 68569 块（/32），位居世界第二，有效保障未来发展需要。综合评估，我国 IPv6 网络规模、用户规模、流量规模位居世界第一，实现预定目标。

重点任务方面，网络承载与服务能力显著增强，全国骨干网、城域网、4G/5G 网络、移动物联网、国家级骨干直联点全面完成改造并开启 IPv6，IPv6 网络质量持续提升，主流云服务产品提供 IPv6 服务，国家八大算力枢纽均支持 IPv6。**终端贯通堵点逐步打通**，主要设备企业新生产家庭路由器默认开启 IPv6，老旧终端加速替换，基础电信企业新增物联网平台及终端全面支持 IPv6，全国家庭网关 IPv6 开启率达 96%，互联网电视机顶盒 IPv6 支持率超 85%。**融合应用广度和深度不断拓展**。县级以上政府门户网站 IPv6 支持率超 96%。主要移动互联网应用（App）和商业网站 IPv6 支持率超 95%。IPv6 与垂直行业融合应用持续深化，教育、媒体、金融、水利、铁路、能源等行业 IPv6 部署应用成效明显。**技术创新与标准建设成果丰硕**。“IPv6+”技术体系不断完善，部署超 700 张“IPv6+”行业网络，基础电信企业“IPv6+”网络服务能力覆盖全国 300 余个地市。我国已发布 36 项 IPv6 国家标准和

138 项 IPv6 行业标准，积极牵头多项国际标准研制。

三、《实施方案》对“十五五”IPv6 技术创新和融合应用提出了什么目标？

答：《实施方案》立足当前发展阶段和水平，在与已有指标充分衔接的基础上，针对“十五五”IPv6 发展提出定性与定量两方面目标。

在定性目标方面，《实施方案》提出，到 2027 年，IPv6 端到端贯通水平全面提升，技术创新和融合应用取得重要阶段性成果。到 2030 年，IPv6 与经济社会各行业各领域广泛深度融合，构建技术先进、开放创新、内生自驱、安全可靠的 IPv6 产业生态，加快向 IPv6 单栈演进，推动形成以 IPv6 为主导的网络服务和应用体系。

在定量目标方面，《实施方案》设置了 8 项“十五五”IPv6 技术创新和融合应用主要指标，包括：IPv6 活跃用户数、IPv6 网络流量占比、移动物联网 IPv6 连接数、固定宽带用户 IPv6 连通率、主要互联网应用固网侧 IPv6 流量占比、发布 IPv6 国家标准数量、“IPv6+”行业专网数量、县级及以上党政机关办公网络 IPv6 开通率，并明确了 2027 年、2030 年的预期值，确保指标可统计、可监测、可评估。

四、《实施方案》部署的重点任务有哪些考虑？

答：《实施方案》在巩固拓展“十四五”IPv6 发展成效基础上，围绕“十五五”IPv6 高质量发展新要求，部署了 10 个方面重点任务，主要考虑有：**一是深化升级改造**。聚焦端到端贯通难点堵点，强化网络、应用、终端协同改造，加大 IPv6 网络优化力度，深化互联网应用 IPv6 放量引流，推动各类终端支持启用 IPv6，进一步提升 IPv6 用户规模和流量水平。**二是拓展创新空间**。持续完善“IPv6+”创新体系，加强基于“IPv6+”的智能广域网、新型城域网、算力网络等关键技术攻关，强化创新技术推广应用与成果转化，不断展现 IPv6 技术潜能和优势，促进 IPv6 与人工智能融合创新。完善 IPv6 标准体系，提高技术创新成果国际影响力。**三是强化融合应用**。推动 IPv6 与智能体应用、超高清视

频等新应用新业态融合发展。以行业数字化转型需求为牵引，深化 IPv6 在电子政务、教育数字化、智慧金融、智慧医疗等领域融合应用，赋能智能制造、数字交通、智慧水利、智能网联汽车、数字乡村、智慧园区等场景协同发展，提升经济社会数字化转型的质量和效率。**四是突出单栈演进。**开展 IPv6 单栈组网架构等关键技术攻关，选取重点区域、特色领域开展 IPv6 单栈试点应用。研究制定 IPv6 单栈演进实施指南，建立 IPv6 单栈能力评测体系，推动互联网应用 IPv6 单栈化改造。**五是突出管理制度。**通过强化入口管理，带动 IPv6 同步部署。完善政府信息化 IPv6 管理制度，在政务信息化建设等制度规范中明确 IPv6 支持能力要求。完善应用商店上架标准规范，鼓励新上架 App 支持 IPv6。**六是筑牢安全屏障。**全面推动 IPv6 部署应用和网络安全防护措施同步规划、设计和实施。加强 IPv6 安全技术攻关，提高安全产品 IPv6 防护能力，完善 IPv6 安全防护体系。

五、下一步如何贯彻落实《实施方案》？

答：中央网信办将会同相关部门推进《实施方案》落地见效。**一是**加强统筹协调、央地联动，健全高效协同工作推进机制，体系化推进 IPv6 创新应用，形成发展合力。**二是**持续完善 IPv6 发展监测体系，常态化开展监测，提高监测分析能力。**三是**加大宣传推广力度，开展 IPv6 主题宣传活动，广泛普及 IPv6 技术优势与应用价值，搭建产学研用交流合作平台。**四是**积极推进 IPv6 国际交流合作，携手打造开放协同、互利共赢的国际合作生态。

智能体互信互联互操作全球合作倡议

原载：“网信中国”微信公众号

智能体作为人工智能时代最具变革性的技术形态之一，其互信、互联、互操作水平深刻影响全

球数字经济的融合程度与人类社会的未来。为释放智能体赋能可持续发展的潜力，防止形成智能鸿沟，7月17日，中国国家互联网信息办公室会同有关方面于2026世界人工智能大会暨人工智能全球治理高级别会议主论坛，提出《智能体互信互联互操作全球合作倡议》，希凝聚各方共识，与全球伙伴共同打造开放、可信、安全、普惠的智能体生态。

一、强化技术研发，驱动协同创新

我们呼吁各国加大对智能体身份标识、协作互认等基础共性技术的研发投入，推动多智能体从“可控交互”向“自主协同”演进，构建高效、可扩展、低延迟的跨平台互操作能力，为全球智能体互联奠定坚实技术根基。

二、促进标准互认，消弭生态壁垒

我们支持在各国普遍参与的基础上，制定开放、非歧视、透明的智能体互操作国际标准，推动接口、互联协议、语义与流程层的兼容互认。建立开放式标准验证与迭代机制，避免技术壁垒与生态割裂，推动不同智能体无缝接入全球协作网络。

三、共建基础设施，保障互联可及

我们鼓励各方协力研究智能体互联协同架构，共建开放共享的智能互联网基础设施，服务于多智能体高效可信协作，并确保其具备良好的兼容与扩展能力，降低接入适配门槛。

四、共推应用落地，加速场景示范

我们倡导以智能体互信、互联、互操作成果广泛服务于科学研究、产业发展、提振消费、民生福祉、社会治理等场景，并优先推动将其用于应对气候变化、公共卫生、教育公平、减贫等全球性议题，确保增进全人类共同福祉。

五、共促产业协同，繁荣全球生态

我们欢迎发达国家与发展中国家、大型企业、中小微实体、科研机构与开源社区等广泛参与智能体互信、互联、互操作建设，促进知识、技术与最佳实践在全球范围内的广泛传播，形成全球智能体产业链。

六、共筑安全底线，恪守伦理准则

我们呼吁坚持负责任创新的基本立场，建立覆

盖智能体全生命周期的安全治理框架，并将安全与伦理考量嵌入技术研发、系统部署与运行维护，重视智能体对网络安全的深刻影响，负责任发布和使用智能体，确保智能体互信、互联、互操作始终不逾越安全底线与人类尊严红线。

七、促进数据流动，强化普遍保护

我们呼吁尊重不同国家、不同地区之间数据安全和个人信息保护相关制度的差异，倡导建立基于开放、包容、安全、合作、非歧视的跨境数据协作框架，推动数据在智能体互信、互联、互操作过程中安全、合规、有序流动。

八、加强资源共享，弥合智能鸿沟

我们鼓励向发展中国家提供智能体互信、互联、互操作所需的基础设施建设、开源工具、能力培训与资金技术支持，鼓励采用本地化适配与能力共建方案，推动形成“南南+南北”数字合作新模式，使智能经济红利普惠共享。

九、营造包容文化，尊重多元治理

我们倡导跨文化、跨语言、跨学科的智能体协作规范，尊重不同国家治理模式与价值选择，避免单向输出，鼓励多元主体广泛参与规则讨论，形成兼具效率与公平、弹性与秩序的全球智能体协作文化。

十、提升数字素养，确保人人受益

我们将共同努力，通过能力建设、技术援助与教育资源开放等方式，帮助未成年人、老年人、残障人士及低技能群体提升人工智能素养，确保智能体互信、互联、互操作成果惠及各类主体。

我们秉持共商共建共享的理念，诚邀所有认同上述愿景的各方携手同行，以包容开放的心态和审慎务实的行动，共同推动智能体生态发展，为构建网络空间命运共同体贡献智慧与力量。

明确 16 项任务！工信部等四部门
发文推动互联网基础资源高质量
20/101

发展

原载：“数字中国建设峰会”微信公众号

贯彻落实党中央、国务院决策部署，为推动我国互联网基础资源高质量发展，塑造互联网发展新动能新优势，工业和信息化部、中央网信办、国家发展改革委、国家数据局等四部门近日联合印发《关于推动互联网基础资源高质量发展的指导意见》。

文件提出，到 2030 年，互联网基础资源高质量发展实现体系化突破，取得一批标志性技术原创成果，形成系列引领性技术标准。互联网基础资源及体系架构革新演进，建成高效协作、智能互联的新型互联网基础设施，打造一批新型服务网络。互联网基础资源高效治理与综合保障能力显著提升。到 2035 年，我国成为全球互联网基础领域重要创新策源地和引领者，全面建成架构先进、全域互联、高速泛在、安全可控的新型互联网基础设施。

关于推动互联网基础资源高质量发展的指导意见

IP 地址、域名、自治域号码、网络标识等资源及其服务设施是互联网安全稳定运行的基础保障。为推动我国互联网基础资源高质量发展，塑造互联网发展新动能新优势，赋能网络强国和数字中国建设，制定本指导意见。

一、总体要求

坚持以习近平新时代中国特色社会主义思想为指导，深入贯彻党的二十大和二十届历次全会精神，完整准确全面贯彻新发展理念，加快构建新发展格局，统筹发展和安全，系统推进理论创新、技术创新、制度创新，强化互联网基础理论研究和关键核心技术攻关，保障互联网基础资源高质量供给，加快构建现代化应用服务体系、管理体系和协调体系，为推进新型工业化、发展新质生产力提供坚实支撑。

到 2030 年，互联网基础资源高质量发展实现体系化突破，取得一批标志性技术原创成果，形成系列引领性技术标准。互联网基础资源及体系架构

革新演进,建成高效协作、智能互联的新型互联网基础设施,打造一批新型服务网络。互联网基础资源产学研用创新发展体系持续完善,支撑互联网新场景、新模式、新业态快速发展。互联网基础资源高效治理与综合保障能力显著提升。到2035年,我国成为全球互联网基础领域重要创新策源地和引领者,全面支撑建成架构先进、全域互联、高速泛在、安全可控的新型互联网基础设施。

二、统筹国家科技力量建设,提升基础创新能力

(一)加强基础理论研究。加快互联网前沿理论创新,加强基础研究战略性、前瞻性布局,结合人工智能、量子通信、区块链等新一代信息技术,智能体互联网络、卫星互联网等新场景新业态发展,从互联网核心机理、底层架构、通信协议、资源形态、安全防护等方面体系化开展理论创新研究,探索创新发展方向,拓展服务模式和功能边界,为推动互联网向更加智能、高效、安全、泛在、融合方向演进奠定理论基础。

(二)突破关键核心技术。加强互联网科技创新规划,推动实施相关国家重点研发计划和国家科技重大专项,加强原始技术创新,开展人工智能、区块链、分布式标识等技术与互联网基础资源融合创新技术攻关,突破网络动态优化、资源智能调度、数据安全交互等关键技术。加强IPv6技术体系创新,解决协议兼容、高性能传输等关键问题。突破卫星互联网巨型星座组网、路由快速切换、传输可靠抗扰等关键技术。突破资源公钥基础设施规模部署应用关键技术。

(三)建强创新发展平台。深化体系架构、网络智能、可信安全等互联网基础领域创新平台培育建设,加强基础性、前瞻性研究与试验,形成更多“从0到1”的原创成果。鼓励企业、科研机构、高校等组建互联网基础资源创新发展联盟、智能体互联网络协同创新联合体等,加速异构网络融通、多智能体交互等关键共性技术成果转化与生态构建,推动技术创新与产业创新深度融合发展。

(四)深化标准体系建设。加强IP地址、域

名、路由等领域标准评估和创新成果标准转化,提升标准科学性、协调性。加强新兴领域标准研制,对智能体互联网络、区块链、分布式标识等前沿领域进行技术标准探索,开展前瞻性标准制定,引领技术演进方向。完善国家标准体系建设,加强互联网基础资源重点领域、关键领域国家标准制定,强化国家标准与国际标准衔接适配。加强国家标准规模化推广应用,推动国家标准的产业深度应用与协同发展。

三、实施基础设施补短拓新工程,强化持续发展支撑

(五)优化传统设施布局。推动域名解析设施优化升级,规范完善境内根镜像服务器建设,增加“.CN”国家顶级域名解析节点数量,加强域名递归解析服务设施建设管理。加强路由设施规划布局,优化完善各级路由架构,强化路由可信验证服务能力,构建全域通达、高效智能、安全可靠的路由网络。加强境内科研机构、企业间试验验证网络环境共建共享。有序推进互联网基础设施向IPv6单栈支持演进升级。

(六)探索新型设施建设。面向智能体与智能体、智能体与外部工具、智能体与前端应用等典型场景,探索构建具备统一通信协议、数据交互标准和接口规范的智能体互联网络基础设施。加强智能体互联互通,推动实现智能体技术从单点应用向网络化协同发展的能力跃升。加强卫星互联网等新型基础设施配套规划建设,加快典型场景落地,稳步扩大应用范围。探索数字身份服务设施建设。推动工业互联网标识解析体系与域名体系等融合发展。

四、实施基础资源安全保障工程,筑牢安全发展根基

(七)保障资源供给安全。保障IP地址充足供给,加强IPv6地址申请使用,制定IPv6地址发展规划,提升IP地址利用效率。统筹推进新通用顶级域名设立使用,保障域名安全供给。合理规划申请自治域号码(ASN),提升自治域网络互联互通能力。加强工业设备、算力、数据、智能体等联网场景下的网络标识资源规划与科学供给。

（八）保障资源运行安全。制定互联网基础资源运行安全管理规范，明确资源申请、使用、分配等全流程分级分类安全保护要求和企业安全主体责任，加强资源运行安全防护。完善互联网基础资源相关关键信息基础设施目录，加强“.CN”国家顶级域名等运行安全保障。加强基础资源数据安全，落实重要数据识别和目录备案、分级保护、风险评估等要求。加强网络智能运维技术开发应用，提升运维自动化、智能化水平。

（九）强化资源风险监测。加强互联网基础资源安全监测，结合人工智能、大数据等技术，提升IP地址、域名、路由等海量数据采集与快速分析能力，健全完善安全风险信息库，形成风险精准发现和实时预警能力。建立健全跨部门、跨领域、多层次协同监测与共享机制，实现风险信息全方位监测发现、全体系防范应对。完善互联网信息安全管理，进一步提升网络和数据安全监测预警等能力。

（十）强化资源应急处置。加强分级分类应急预案体系建设，确保不同风险场景、不同风险等级突发事件快速精准应对，定期对预案进行效果评估，及时优化完善。加强应急响应机制建设，强化跨部门、跨地区协同联动，提升突发事件应对效能，构建上下贯通、左右协同、精准高效应急响应体系。加强灾备等备份恢复能力建设，强化常态化、实战化应急演练，确保核心业务、核心数据安全。

五、构建多维立体应用服务体系，提升实数融合水平

（十一）提升资源应用效能。加强人工智能、区块链、分布式标识等技术在互联网基础资源服务中深度融合应用，通过技术创新突破服务瓶颈、提升服务效能，推动构建按需分配、智能调度、分级保障的现代化服务体系。培育壮大资源管理与服务机构，持续改善中文域名应用环境，加强国家顶级域名应用服务能力。健全IPv6地址分配编码规则，丰富地址应用属性，充分释放IPv6地址生态价值。

（十二）强化场景服务支撑。强化“网联、物联、数联、智联”服务支撑能力。面向“网联”场景，加快寻址、交换等功能优化，支撑高速泛在、

空天地海一体的网络设施融合发展。面向“物联”场景，增强海量异构设备连接、标识信息采集分析能力，支撑提供精准化、标准化基础数据。面向“数联”场景，提升数据跨领域、跨层级可信传输交互支撑能力，加速数据共享利用。面向“智联”场景，促进底层技术协议和场景有机结合，支撑构建人工智能产业生态。

六、构建完备高效监督管理体系，提升综合治理水平

（十三）完善管理规章制度。健全互联网基础资源管理法律法规，完善顶层设计和IP地址、域名、路由等领域管理规章制度。面向新要求、新技术、新业态，进一步细化相关管理办法。完善互联网重大风险防控管理制度。建立互联网标识体系，制定完善算力标识、数据标识、数字身份标识等管理制度，明确管理主体、管理范围、主体责任、合规要求等。制定智能体互联网络相关管理制度，规范智能体间服务发现与互操作管理。

（十四）健全技术治理手段。提升一体化监管技术能力，加强资源获取、分配、使用、回收等全生命周期管理，构建跨平台、跨体系、全链条监管能力。加强违法违规域名解析和域名滥用行为治理技术能力建设，精准快速识别、处置违法违规行为。优化升级IP地址备案管理，提升对动态IP地址等精细化管理水平。持续健全卫星互联网等新场景下互联网基础资源监管技术手段。

七、构建务实主动国际合作体系，提升开放发展水平

（十五）加强国际沟通协调。积极参与互联网名称与数字地址分配机构（ICANN）、互联网工程任务组（IETF）、亚太互联网络信息中心（APNIC）等国际机构相关事务，加强沟通协调，积极承办相关国际会议。加强具有国际视野、熟悉国际规则、掌握专业技能的复合型人才培养，支持鼓励相关人员担任国际机构特别是新兴领域机构职务。提升国内相关组织国际化运营能力，广泛吸引科研机构、企业、专家等参与。

（十六）积极参与全球治理。加大互联网基础

资源国际政策、规则等制定参与力度，提升议题主动设置能力，持续推动构建更加公正合理的全球互联网治理体系。加强国际标准研制，提升技术创新引领能力。推进新兴领域国际组织建设，搭建技术标准、应用构建、产业培育、全球治理等多领域交流合作平台。依托互联网基础资源国际治理机制，加强与共建“一带一路”国家网络资源创新应用推广合作，促进区域互联互通与共治。

八、保障措施

各有关部门强化工作统筹协调，加强协同合作和资源整合，形成工作合力。各地方主管单位加强工作衔接，结合实际推动重点任务落实。多方主体积极参与，行业协会强化纽带作用，专业机构加强

科研投入，大型企业发挥资源优势，合力推进互联网基础资源高质量发展。有关部门加大国家资金支持力度，对符合条件的项目予以支持。建立跟踪评估机制，定期开展效果评估，确保取得实效。

（技术编辑：来唯希）

研究动态



数字法学基础理论

1、算法正义论：概念与结构（彭中礼）

来源：《法学评论》2026年第4期

算法已深度渗透社会运行各领域，从工具升维成核心规则，同时催生歧视、权力滥用、责任模糊等正义难题。传统正义理论难以适配数字时代需求。这就需要立足算法全生命周期特性，剖析传统正义理论的适配性危机与算法时代的现实困境，论证构建算法正义理论的必要性。通过批判既有算法正义概念，明确技术可及性、权利本位、适用场景三大证立前提，统筹主体权责、全流程公平等关键因素，界定算法正义的内涵与五大基本特征。最终从程序正义的算法性转化、实质正义的算法化确定、技术正义的算法化表达三重维度，阐释算法正义的结构，构建起价值引领、技术落地与实践规制的完整理论框架，为算法时代的正义实现与治理实践提供理论支撑。

2、从权利到身份：数字时代权利保护范式更新（郑泽星）

来源：《法学论坛》2026年第4期

人工智能技术的发展在展现巨大潜能的同时，也蕴含着前所未有的现实风险。数字时代的权利形态、风险外观以及保护方式均面临深刻变革。传统具体权利保护路径因其分散化、事后性和场景割裂等特征，难以全面应对数字时代权利的结构性风险。数字时代权利保护亟须实现从具体权利救济向身份结构保护的范式更新。“数字人权”方案虽试图提供整体性回应，但在理论独立性、规范确定性及制度可操作性方面均存在不足。数字身份是自然人以基础身份信息为起点，在数字技术与社会互动的共同作用下，通过个人信息、行为数据、推断画像以及虚拟化身等要素生成的、能够指向并表征特定个体的数字化身份结构，具体可以区分为信息维度、交互维度和构建维度。将数字身份权能的分散保护与整体保护相结合，是数字时代对抗技术逻辑侵蚀、塑造公平可持续数字秩序、系统保护个人权利的有益路径。

3、论“数字正义”的祛魅（刘志强）

来源：《法制与社会发展》2026年第4期

从本体论、权力论、程序论三重维度审视“数字正义”，存在三种理论困境，现在学术界所认为

的“数字正义”，并非真实的正义形态。就本体维度而言，数字技术受限于“正义不可计算”悖论及算法价值偏见，难以承载正义的规范性内涵。就权力层面而言，算法权力的司法殖民化与技术权力的主权扩张，引发权力失衡。就程序视角来说，数字鸿沟加剧导致“可视正义”幻象与实质正义疏离，程序正当性面临着消解风险。就重构论视域来说，数字司法需坚守以人为本的传统司法理念，重构价值包容性、技术透明性、程序可逆性的司法评估体系，并确立数字司法应用的三阶价值审查机制与失效救济机制。应通过平衡技术赋能与权利保障，来实现数字时代的人权守护。

4、数字法律责任的理论阐释与体系构造（郑智航）

来源：《法制与社会发展》2026年第4期

数字法律责任，是指法律主体在研发、设计、部署、运营和应用人工智能、大数据、区块链、算法等数字技术，以及提供数字产品与数字服务的过程中，因违反法定义务或约定义务，侵犯数字空间中的合法权益，而应当依法承担的法律后果。作为一种新型法律责任，数字法律责任在基本特性、主要功能和实现机制等方面都有新的发展与变化。它突破了传统民事责任、行政责任和刑事责任的严格界限，呈现责任形态多元复合的特点。数字法律责任除了发挥传统事后追责功能外，还承担数字风险预防、合规激励、数字生态系统治理等重要功能。不同发展阶段和不同国家的经济、政治和文化背景，影响着对数字法律责任采取的伦理基础。数字法律关系的复杂性使得传统的以行为、过错、损害结果、因果关系为核心的法律责任构成要件在数字社会中难以直接适用，需要对其进行重述。根据数字法律责任不同类型的特征与归责逻辑，构建覆盖全链条的法律责任体系，对于规范数字法律行为具有重要作用。

5、“公共数据归国家所有”的宪法检视（胡萧力）

来源：《法学研究》2026年第3期

围绕公共数据的权属界定，当前地方立法实践

中存在数据汇集型、数据维护型、数据经营型三种确权逻辑。相关立法多笼统借用经济学产权概念与政策性表述，欠缺严格的法学维度的正当性证成与上位法依据支撑。公共数据深度嵌入行政过程、数字基础设施运转架构和社会协同创新网络，其运行遵循信息媒介逻辑，与自然资源属性迥异。公共数据利用秩序的建构，不宜直接套用宪法上关于自然资源国家所有的法理，而应在宪法第12条所确立的社会主义公共财产规范框架下展开。公共数据归国家所有，应理解为由国家承担防止侵占、保障合理利用、促进公平利用之义务的规范要求。未来应通过中央立法明确公共数据治理的基本原则、数据汇集的合法性要件、授权运营的目的与基本程序、私法规则的适用边界等，并配套财政监督、程序保障与数据流通的统筹机制，推动形成体系完整、逻辑融贯、普惠包容的公共数据利用生态。

6、论数字人权的尊严基础（郑玉双）

来源：《政治与法律》2026年第7期

数字科技的发展引发了关于数字人权是否存在法理争议。学界围绕数字人权是否属于第四代人权、是否具有概念独立性问题展开了激烈交锋，折射出数字人权证成所面临的多维困境。人权的独特道德意义使得数字人权的证成不能停留于代际论或者普遍人性论，而是应当回到尊严和人权的关系这个问题之中。尊严的社会条件经历了智能化转型，导致了数字人格的出现、人机关系的深化以及尊严实现的技术条件的变动。以规范能动性为核心和呈现方式的尊严观能够对尊严和人权的关系进行更好的回答，并充分回应规范能动性在科技冲击下所面临的数字化转型。通过尊严的价值支持和对规范能动性的结构分析，数字人权作为一种新兴的独立人权能够得到证成。以尊严的应受捍卫性、个体发挥规范能动性的方式变动以及转化为制度权利的现实可行性这三层梯度作为“尊严测准”，能够构建出数字人权转化到法律之中的制度框架，并充分发挥数字人权在数字时代的价值指引功能。

人工智能治理

1、人工智能代理的权力分配逻辑与刑事责任归属（姚万勤）

来源：《政法论坛》2026年第4期

人工智能刑事责任归属事关人工智能技术的风险分配与开发应用。归属得当则利于社会发展，归属不当则妨碍技术进步。既有关于人工智能刑事责任归属问题的理论研究，主要遵循能力分配逻辑：人工智能具有某种能力则应当成为刑事责任主体，反之则不能。该类研究视角隐含了能力具备即等同责任承担的逻辑预设，忽视了事实判断与价值判断的休谟鸿沟。对人工智能刑事责任归属的分析应当从既有的能力分配逻辑转向权力分配逻辑，因为权力分配可以直接影响刑事责任的归属。在价值论层面，人工智能不能拥有与人类平等的主体资格，因而人类与人工智能之间的权力分配，应当采取人工智能不能直接独立进行行动，只能代理人类之后进行行动的人工智能代理模式。人工智能代理刑事责任的归属对象是人工智能所代理的被代理人，因而人工智能体本身不能成为刑事责任主体。以自动驾驶汽车交通肇事为例，因人类驾驶员拥有主动接管权的基础，原则上应被视为所有等级自动驾驶汽车的被代理人，在此前提下才可能构成交通肇事罪，而生产者则在此基础上才需考虑承担刑事责任的可能性。

2、论知识蒸馏作为反向工程的合法性（李昕梦）

来源：《法律科学（西北政法大學學報）》2026年第4期

知识蒸馏通过将教师模型所蕴含的知识迁移至学生模型，能够降低计算资源需求并提升学生模型的性能，但因知识蒸馏涉及对模型内部技术信息的获取，在商业秘密保护方面引发了广泛争议。模型内部技术信息可以构成商业秘密，知识蒸馏符合商业秘密反向工程的构成要件。在不违背公平性与合理性原则时，用户服务协议中的“禁止蒸馏”等限制性条款具有合同法上的效力，但违反此类条款不应引发商业秘密侵权责任。知识蒸馏作为反向工

程具有经济正当性，但应以目的正当、行为合理、不造成市场损害为限。明确知识蒸馏的反向工程属性及其正当性边界，有助于实现人工智能产业商业秘密保护与技术创新之间的动态平衡。

3、论具身智能体产品缺陷的认定规则（闫冬）

来源：《法律科学（西北政法大學學報）》2026年第4期

具身智能体将算法决策嵌入物理载体并直接作用于现实环境，使传统以静态产品为前提构建的产品缺陷认定规则在适用中同时面临“易被认定为缺陷”与“易触发免责”的悖论。具身智能体的具身性与智能性技术特征要求缺陷的认定标准与抗辩逻辑应通过区分硬件与智能模块的风险生成机制分别构建。就硬件模块而言，既有的以设计缺陷、制造缺陷与警示缺陷为核心的判断框架仍具适用基础，但测试标准须有所侧重，并强化安全技术标准与行业规范协同推进。就智能模块而言，有必要以安全约束知识库作为抓手，以公众安全期待测试为指导，将合理期待转化为可审查的安全约束参数，从而增强智能缺陷判断的客观性与可操作性。对软硬件缺陷进行分类界定，可为责任归属与因果分析明确前置条件，并为相关保险机制奠定模型与精算基础。

4、人工智能刑事责任主体制度的重构（赵东）

来源：《现代法学》2026年第3期

人工智能刑事责任主体地位的确立，应当摒弃“自由意志+刑事责任能力”的拟人化认定标准，基于人工智能的技术属性，以“可评价和检验的自主行为能力”为核心，构建“自主决策能力+行为支配能力+风险识别能力”的规范认定结构。在立法层面，应综合考量当前人工智能的技术能量、技术特征、应用实效与发展阶段，从控制必要性、控制强度和控制时机出发，采取刑事立法适度先行的规制策略，以回应人工智能技术福利与智能化风险之间的“科林格里奇困境”。在刑事归责模式上，伴随着“自然人→人工智能体”，刑事责任主体实现

了人权主体性向社会工具性的角色转换,基于社会防卫论立场,对于人工智能犯罪,可采取客观事实归责模式下的严格责任论,以“过错推定(实体法的层面)+举证责任倒置(程序法的层面)”为严格责任认定模型,并配套设置以科技风险防范为核心内容的技术保安处分措施。

5、人工智能算法可专利性的判定逻辑与规则优化 (余力烺)

来源:《现代法学》2026年第3期

传统的专利客体严格构建在对技术方案的认知逻辑上。以专利来保护人工智能算法存在根本性的问题,即人工智能算法是否具有可专利性及如何判断其可专利性。由于抽象思想不是技术方案,而人工智能算法是技术属性与思想属性相互融合的整体,故其难以被界定为传统的专利客体。因此,人工智能算法可专利性的深入讨论将面临客体适格性与算法特征之间的冲突,这使人工智能算法在新颖性、创造性和实用性等专利授权要件上与传统技术方案相比存在特殊性。分析人工智能算法可专利性的判定逻辑,有利于厘清算法专利的认知体系,突破专利保护的规则局限。只有结合人工智能算法特性对可专利性判定的拟制主体及规则进行适应性调整及重构,才能合理维持算法专利的创造性高度,实现人工智能算法技术创新与权利保障的协同发展。

6、论人工智能模型开源的双层治理结构——以版权许可与平台规则为建构(辜凌云)

来源:《清华法学》2026年第4期

人工智能模型开源正在改变技术演进与产业创新的底层逻辑,其创新开放的客体已从传统的软件源代码扩展至模型代码架构、参数权重、应用工具包及数据集等多元内容,并形成了传统开源许可证与 Llama、RAIL 等为典型的新类型开源许可证共存的局面。现有对开源行为及许可证的治理思路仍停留于传统代码开源时代形成的静态的版权侵权判定与责任分配模式,而忽视了新类型许可证条款

中对版权权利的设计与分配,特别是版权的行使由原本的单次授权重新配置为附条件解除并受持续约束的授权模式。在模型权重等可版权性存在争议的客体上,平台规则通过技术与合同手段填补了版权权利未能控制的部分,由此形成“版权许可+平台规则”双层治理结构。基于行动者网络理论,这一结构揭示了许可证与平台规则作为非行动者在权利设计、分发与许可中的关键作用,通过网络效应驱动、成本结构重置与分层竞争三大逻辑,使得版权权利的让渡能够被转化为显著的市场竞争优势,形成闭环的利益转换体系。针对数据与代码贡献激励减弱的现状,可构建与权利分配相适应的收益反哺机制,通过优化许可证与平台规则价值链体系,建立起涵盖企业、社区、贡献者多主体参与的开放版权协同治理机制,以实现人工智能模型开源生态中知识共享与商业利益的动态平衡。

7、生成式人工智能对法律职业的解构与重塑(王迎龙)

来源:《法学论坛》2026年第4期

以生成式人工智能技术为代表的新科技革命的迅猛发展,促进了科学技术与法律场域的深度融合,推动了法律职业的结构性质重塑。生成式人工智能的技术特征与加速迭代,正在解构传统法律知识生成范式,打破了法律职业赖以存在的知识垄断,实现了从以“人”为中心向以“算力”为支撑的法律知识生成范式的转型。在法律知识范式进化与革命的影响下,法律服务模式由信息辅助向算法决策发展,使得市场主体间的差序格局被重新界定,法律服务市场“割据化”现象进一步加剧。基于上述变革,生成式人工智能的强势介入一方面对法律职业化及其理论提出挑战,消解了以劳动分工与专业技能为基础的法律职业化,一定程度上降低了法律职业自主性,影响了不同职业主体间的职业地位;另一方面也为高质量法律服务的发展提供了契机。在恪守法律职业者价值判断主导地位的前提下,构建与完善人机交互的法律服务模式,是新科技革命时代法律职业得以高质量发展的未来路径。

8、智能体时代授权的逻辑限度与法律建构（李汶龙）

来源：《法学论坛》2026年第4期

以大模型为基础的智能体正在重塑人机协作与任务执行方式，其在开放环境中自主拆解目标、调用工具并跨系统执行操作，使传统基于单一行为、单次授权的逻辑面临挑战。我国《智能体规范应用与创新发展实施意见》已将用户授权范围、决策边界、行为追溯以及用户最终决策权置于智能体治理的核心，但现有规范有待从原则层面进行细化，对复杂授权关系中的行为归属与责任分配进行系统回应。本文提出应超越以自然语言为中心的授权表达，转向以机器可执行的权限边界与能力约束为核心的结构化治理机制。在分析智能体场景下授权载体不充分、授权表达不清晰、授权主体不显明以及授权后果不可控等困境基础上，提出以授权辅助、权限控制、外观管理、公平交互与责任承担为核心的智能体授权治理框架。责任配置亦需从用户授权转向服务提供者的结构性责任，通过可追溯性、可中断性、偏离控制与授权链条管理等制度工具，实现从授权证明到责任映射的转型。只有在授权表达机制与责任体系同时重构的前提下，智能体经济方可成为可控可持续的制度基础设施。

9、人工智能版权训练数据来源合法性的规范构造（徐小奔）

来源：《法学论坛》2026年第4期

人工智能模型开发者在使用他人作品作为训练数据进行模型训练时，应基于尊重在先权益原则，承担通过合法来源获取版权训练数据的注意义务。在鼓励AI开发者与版权人达成合作以激励更多训练数据高质量供给的规范目标下，版权训练数据来源合法性具有重要的实践意义，是版权合理使用“三步检验法”中“不得不合理地损害著作权人的合法权益”要件的重要判断因素。因此，AI开发者应在合理的技术成本范围内，积极承担版权训练数据获取的手段披露义务、模型训练目的的说明义务和数据溯源的辅助标识义务等三项基本义务，由此

更好地提升版权训练数据的来源合法性。

10、人工智能技术的国家安全审查及法律边界审视（范晓波）

来源：《比较法研究》2026年第3期

人工智能，特别是生成式人工智能与大语言模型的快速发展，引发国家安全审查范式的深刻变革。传统的基于实体资产与明确意图的框架，无法应对算法、数据驱动与计算要素的监管需求。人工智能的双重性——作为经济增长催化剂与安全风险载体——形成安全、发展与开放的三难困境。通过对美国安全导向的遏制模式、欧盟风险分级的目标导向路径，以及中国碎片化的多轨实践的比较研究，中国应追求本土化法律边界建构，立足自身国情与战略目标，寻求一条兼顾安全与发展的中国特色人工智能治理之路。

11、论生成式人工智能服务提供者侵权的过错证明责任（朱晓峰）

来源：《比较法研究》2026年第3期

生成式人工智能服务提供者侵权的过错证明责任，在司法实务上既有采过错推定规则而将过错证明责任分配给行为人的，也有采一般过错规则而将过错证明责任分配给受害人的；过错证明责任分配规则并不统一，影响法律适用的统一性和权威性。对此，应以现行法的既有规定为基础，在解释论视角下采区分说，在法律明确规定应当适用过错推定规则、可以适用阶层式法律效果评价方法认定过错侵权责任或适用利益权衡式法律效果评价方法认定非物质性人格权侵害民事责任时，确定性地适用过错推定规则或经由法官决定是否可以适用过错推定规则，以将过错证明责任分配给行为人；除此之外的则应统一采用一般过错证明责任规则而由受害人证明行为人的过错，以此助益“促进生成式人工智能健康发展和规范应用”与“保护公民合法权益”双重目标的实现。

12、人工智能生成内容邻接权证伪（蒋舸）

来源：《比较法研究》2026年第3期

著作权法已经向用户提供了利用AI生成内容的可观激励，并不存在系统性的激励空白。在此情况下，增设邻接权的边际激励收益有限，而过度类型化及妨碍公众自由的制度成本却不容忽视。人工智能生成内容邻接权是“准作品式”邻接权，其认知经济性不佳，且与著作权的界限不明。该邻接权构造存在内生矛盾：它与著作权的区别越明显，越有独立存在的价值，但也越容易出错。在用户利用人工智能生成内容领域，创作激励与使用自由之间的利益平衡应由著作权规则独占评价，避免陷入由著作权、邻接权甚至反法一般条款共同规制的“碎片式”境地。在未来，各类文艺成果都将包含或多或少来自人工智能的贡献，增设邻接权将波及几乎所有文艺成果的界权，必须慎重。

13、端侧智能体的治理挑战与制度回应（张欣）

来源：《比较法研究》2026年第3期

端侧智能体通过将模型推理、情境感知与任务执行能力延伸至用户终端，使人工智能深度嵌入个人终端环境，并在人机交互的深层变革中重塑风险生成机制。以智能手机为典型场景，端侧智能体的全域感知逻辑与数据和模型深度耦合的本地化运行机制，冲击了最小必要原则得以有效适用的制度前提。轻量化与本地适配削弱了安全对齐的可靠性，模糊了模型参数与个人信息之间的规制边界，使得既有审查与问责机制的适用时点面临挑战。系统级权限与跨主体协同架构推动风险从数据处理与内容生成延伸至任务执行，加剧了责任归属与权限边界的治理困境。鉴于端侧智能体仍处于产业起步阶段，制度回应宜采取渐进式调适进路，契合智能体运行机理对既有治理框架进行精细化展开：在数据层，应以业务功能与具体任务为双层基准，推动最小必要原则向端侧场景的精细化适配；在模型层，应建立面向全生命周期的动态评估机制，并依据实际控制能力细化产业链协同治理机制；在应用层，应构建基于行动能力的分级授权与过程控制机制，确保其行为安全可控。

14、大模型训练专用合成数据的功能定位及制度供给（苏杭）

来源：《法制与社会发展》2026年第3期

大模型训练高度依赖数据驱动，但直接使用真实数据正陷入侵犯知识产权与个人信息权益的合规困境。与基于采集性来源的真实数据不同，合成数据由算法创制而成，作为大模型训练数据的生产性来源，在合规基础、内容质量与产业生态上具有功能优势，为破解上述困境提供了新路径。然而，合成数据的制度供给存在严重短缺：法律定性与确权规则尚付阙如，数据失真偏差、技术内生缺陷与社会伦理风险亦缺乏有效治理。对此，应明确合成数据作为独立数据类型的法律地位，并基于原始确权机制对生成者的权能作层次化配置；同时引入负责任创新理念，从前瞻性质量评估、反思性透明治理与包容性协同治理三个层面构建风险治理体系，为合成数据的功能发挥提供制度支持。

15、人工智能训练作品的商业学习属性与补偿金制度构建（罗祥）

来源：《法学杂志》2026年第4期

人工智能训练作品过程中的复制没有遵循复制权“无传播即无权利”的法理，该复制是事实行为而非法律行为。复制权不应控制人工智能训练作品中的复制行为。人工智能训练作品会复制表达，但其意图并非还原表达，这一过程与人类通过学习获取知识范式与思维模型具有功能等同性。复制与训练是“依赖复制的训练”的偏正关系，人工智能训练作品本质上是商业学习。商业学习是一种使用者权益，《民法典》第126条和《反不正当竞争法》第2条分别为其提供规范依据和保护路径。为人工智能训练作品的商业学习行为配套补偿金制度具有正当性。补偿金支付主体应当是人工智能训练作品的受益者。在补偿金征收方面，当侵权救济路径阻却时，我国应借鉴替代性补偿制度去建构管理型征收模式。在确定补偿金征收数额时，首先要确定费率基数，人工智能供应商的全部支出或收入基准更为合适，适用顺序上收入优先于支出；其次按照

特定要素估值人工智能供应商支出或收入；最后在费率计算时，应参考对录音设备和媒体的销售征收费率标准。应根据人工智能模型的多样化用途进行调整，在制定补偿金方案时，应区分高收入实体与低收入实体，精确设定补偿金标准。为确保费率计算的准确性，还要求更全面地披露作品使用情况。

数字时代的部门法研究

1、论数字赋能轻微违法治理的路径与原理（邓豪）

来源：《法律科学（西北政法大學學報）》2026年第4期

轻微违法量大而面广，深度嵌入民众日常生活，其治理长期面临信息收集不畅、组织协作不足、考核监督失灵等困境。数字技术为破解困境提供了可能，其赋能轻微违法治理的路径主要包括三类：一是违法信息收集的远程化，二是执法组织协作的平台化，三是执法考核监督的自动化。数字赋能通过强化信息获取与处理能力、重构执法流程、限缩裁量空间并引入持续性监控机制，从整体上提高了基层执法能力。数字赋能的过程也伴随着偏差的生成：信息采集的机械化可能导致过度执法，基础设施的非均衡配置可能引发形式化运作，考核压力的刚性传导可能催生策略性应对。基于此，应从提升信息甄别能力、夯实数字基础设施、动态调整考核机制三个层面优化数字技术的运用，在发挥技术优势的同时抑制其负面影响，推动轻微违法治理的精细化与法治化。

2、数据信息安全刑法保护体系的建构（郭研）

来源：《法律科学（西北政法大學學報）》2026年第4期

现行刑法以计算机信息系统安全为中心的立法框架及对个人信息保护的过度侧重，难以有效回应数据独立化、要素化时代的数据安全保护需求。应当通过厘清数据与信息的载体关系，证成刑法语境下“数据信息”的规范概念，区分功能性数据与承载法益内容的数据信息，确立数据信息安全管理秩序为新型集体法益，以合法性、可控性、可问责

性界定其核心内涵。在此基础上构建类型化刑法规制体系，侵害功能性数据的适用计算机信息系统相关罪名；侵害可识别性个人数据信息优先适用侵犯公民个人信息罪；侵害承载传统法益的数据信息依照现有罪名认定；对于不承载传统法益的非个人数据信息的不当滥用行为，现行刑法存在规制空白，应增设滥用数据信息罪予以规制。

3、数字时代票据法理论更新与制度变革（韩亮）

来源：《法学论坛》2026年第4期

数字票据发展过程中存在双轨制法制运行的弊端，导致司法实践中的法律适用难题。数字票据对传统票据法理论和制度都形成了挑战。数字票据不属于传统意义上的权券结合证券，数字票据的无因性基础与纸质票据不同。数字票据的行为性质也不能用传统的单方行为说来解释，而应该归类为契约说。数字票据时代，票据法的理论体系需要进行更新。在对数字票据理论更新的基础上，才能对票据制度进行重构。数字票据的立法模式可以选择修改《票据法》的模式，这种模式不会割裂与传统的关系，是最理想的数字票据制度变革模式。

4、数据驱动型经营者集中反垄断审查的制度展开（黄彦钦）

来源：《法制与社会发展》2026年第4期

在数字经济社会，数据是核心的生产要素。借由网络效应、传导效应与封锁效应的叠加作用，数据集中化在一定程度上显著抬高了市场进入壁垒，强化了市场支配地位，削弱了市场活力与动态竞争。传统基于市场份额与价格分析的经营者集中反垄断审查范式，难以有效识别由数据驱动型经营者集中引发的结构性竞争损害。当前，我国经营者集中审查制度面临的现实困境主要有三点：其一，数据是否构成独立的相关市场；其二，如何评估数据的资产价值；其三，如何量化数据对市场竞争的影响。为破除这些现实困境，理应对审查机制进行系统性革新：以数据功能为导向重构市场界定逻辑，引入“数据力量指数”综合评估数据规模、质量与数据

处理能力,体系化计量企业的数据力量,采取多元的申报标准、附条件的监管机制与主动的执法路径等,推动反垄断审查从对形式化结构指标的依赖,转向对实质控制力的精准识别与动态评估,从而在激励数据价值实现与维护市场有效竞争之间达成动态平衡。

5、数字时代刑事跨境取证的法理检视与制度完善 (智嘉译)

来源:《比较法研究》2026年第3期

数字技术的进步推动跨境犯罪从传统物理空间转向数字空间,以司法协助为核心的传统跨境取证模式面临效率不足、规则滞后、主权冲突加剧等多重困境,难以适应数字时代跨境犯罪打击的现实需求。在数字时代,我国刑事案件跨境取证有必要实现从“司法协助中心主义”向“数据治理中心主义”转型。“数据治理中心主义”跨境取证制度的完善应以尊重主权原则、比例原则、正当程序原则、实质审查原则、权利保障原则为支撑,同时以数据分级分类治理为基础确立分层取证规则体系,以多元主体协同为前提明确平台义务与责任,以治理需求为导向优化实质审查认定机制,并参考《联合国打击网络犯罪公约》进一步完善相关立法,从而在维护国家主权前提下提升跨境犯罪治理能力,推动形成公平合理、安全有序的全球跨境刑事取证秩序。

6、数智时代税收征管技术辅助权的法律规制(李蕊)

来源:《法学研究》2026年第3期

税收征管中数智技术的深度嵌入,引致“技术权力化”“权力技术化”的双向融合发展。征管科技企业基于行政委托和自我赋权,实质获得了具有一定公权属性的功能意义上的税收征管技术辅助权。因耦合了技术、资本、权力三重运行逻辑,税收征管技术辅助权在行使时可能越界侵蚀税收征管制权,压缩侵害纳税人权利。由于技术逻辑与税收征管制权法治逻辑之间的巨大张力,传统上限制公权、保障私权的税法规范难以发挥有效的规制作用。为

确保税收征管技术辅助权的规范运行,应在行政委托合同中明确征管科技企业的技术服务标准和质量,并设定征管机关的单方变更权、解除权等行政优益权条款,通过健全算法审查机制、严格全过程测试方案验证、完善审计跟踪机制等强化征管机关的税收数字监管,基于自动化决策的影响和复杂程度确定人工介入的时点和方式,通过风险分类实现对税收征管技术辅助权的差异化规制。

7、数字技术赋能对公司治理的挑战及公司法回应 (陈洪磊)

来源:《法学杂志》2026年第4期

在数字经济时代,区块链、人工智能、大数据、网络通信等数字技术的发展深刻影响了公司治理。数字技术正以一种不同于传统公司法解决公司代理成本的方式推动着股东与董事之间、股东与股东之间、股东与第三人之间冲突的化解,这在很大程度上提升了我国公司治理水平。然而,带来公司治理优势的数字技术也挑战着生长于商品经济时代的公司法规范,给公司治理的主体规则、行为规则和责任规则带来了巨大冲击,既有公司法并未留出充足的容纳数字技术的接口,在一定程度上呈现规则供给上的局促。基于此,应当坚持稳中求进这一数字中国建设总基调,于理念上,以公司自治为手段实现对数字技术的包容;于方法上,以比例原则为工具实现对公司自治接纳数字技术的尺度设定;于主体上,以董事会为中心实现数字技术赋能公司治理的内部秩序重塑,最终实现数字技术赋能和公司法之间的良性互动,并使得两个方向都充满活力。

数据要素的保护与流通

1、论企业数据的合理使用(张真源)

来源:《华东政法大学学报》2026年第3期

企业数据合理使用是互联网信息自由权的自然延伸。数据确权主义、信息自由主义与工具主义之论争,既未能明确非持有者使用企业数据的正当性边界,亦未能为个案的利益衡量提供合理性判断之原则与标准。为此,借鉴合理使用制度之原理与

架构,有助于在现行法体系下析出企业数据合理使用之检验方法与实质性标准。企业数据合理使用的标准为数据之上的社会利益超过企业对数据的垄断利益。在此利益衡量基准之下,合理使用是不妨碍企业数据持有者行为自由所延伸的使用自由。其合理性在于不降低企业对持有数据之既定利益——即基于数据面向市场获取财产收益的机会,进而肯定非持有者获取企业数据潜在利益的正当性。从实质性判断标准来看,应以公开与可获取之性质,以非商业性使用、公共利益需要及谋求数据潜在利益之目的,以权利混同下“大量复制”之程度,事先覆盖对数据市场实际影响的分析;然后再以不正当竞争之损害与获利及其运营负担之侵害,作为判断影响结果之标准。

2、论公开个人数据爬取行为的侵权责任(沈健州)

来源:《法学家》2026年第3期

数据爬取的侵权责任,本质上是权益保护与行为自由之平衡这一侵权法元问题在数字经济下的时代缩影。公开个人数据同时承载数据财产权益与个人信息权益,数据爬取行为可能造成多重损害后果。于数据财产权益,数据爬取对平台的数据利用造成竞争损害或干扰均可构成侵权,但两者在损害界定、过错判断和责任方式上均存在差异;在可携带权范围内,个人授权可使数据爬取造成的竞争损害正当化,但并不影响干扰型侵权的构成。于个人信息权益,爬取方对超出公开个人信息合理处理范围的数据爬取损害承担过错推定责任,且不能以遵守机器人协议证明其无过错,其虽可以“履行费用过高”为由不予删除数据,但仍需进行损害赔偿;若平台未尽到个人信息保护义务,应在其过错范围内承担补充责任,并可向爬取方追偿。

3、数据资产出资:核心议题与规范应对(王艺璇)

来源:《法学家》2026年第3期

数据资产作价出资是数据资产流通利用的形式创新。数据资产出资的制度展开应依数据本质属性,遵循公司法价值理念,妥当权衡自治与管制,

对出资规则做适应性解释。数据资产的出资标的是数据财产权,出资人既可以完整的数据财产权出资,也可以独立的数据使用权、经营权出资。因数据流通利用是数据价值实现的主要方式,以数据使用权、经营权出资应成为数据资产出资的主要形式,且是公共数据出资的唯一形式。出资协议签订方面,专业评估机构评估并非必要,出资人与公司协商确定价格即为公允价格;协议内容上,不应限制出资比例,但应要求所出资数据与公司经营存在关联性。出资义务履行方面,出资人负有依协议约定不得许可他人利用数据等附属义务,该义务在出资人转让股权时不消灭且不转移。瑕疵出资应对方面,应着重完善董事等经营管理人员维护资本充实责任体系,应在细致区分董事角色和清晰界分评估、催缴、验收等职责基础上实现董事责任细化。

4、数据窝赃行为的刑法规制(江海洋)

来源:《法商研究》2026年第4期

我国刑法目前对非个人数据的保护主要限于前端非法获取环节;对于非法获取后的非个人数据窝赃行为,则主要尝试以掩饰、隐瞒犯罪所得罪进行暂时性规制。数据窝赃行为侵害的法益并非权利人控制数据访问的形式数据秘密,而是数据所蕴含信息内容的实质数据秘密。掩饰、隐瞒犯罪所得罪的保护法益与数据窝赃行为的不法本质并不契合,数据窝赃行为也不满足前者的构成要件。有必要采取立法论路径,增设数据窝赃罪以系统规制此类行为。增设数据窝赃罪不仅具有目的正当性,也能够通过比例原则之适当性、必要性与均衡性的检验。该罪构成要件的设置需要注意:本罪对象为非公开非个人数据,且应对“物之同一性”作实质解释,只要求窝赃数据蕴含的信息内容直接来源于前端非法获取行为所获非个人数据即可;行为方式包括“为自己或他人获取”“提供给他人”“传播”三类;在对数据分类分级的基础上,应设置多维度的罪量要素。本罪法益非绝对不可克减,如基于国家安全、公共利益等正当事由,可对行为予以出罪。

5、数据跨境治理的充分性认定制度——基于欧盟范式的批判性审视与中国因应（汤净）

来源：《法商研究》2026年第4期

数据跨境治理中的充分性认定制度，已从《欧盟通用数据保护条例》所确立的技术性数据保护工具演变为全球数字治理领域的重要制度范式和权力博弈载体。该制度虽然以保障数据出境安全、促进数据跨境自由流动为初衷，但是其充分性标准的模糊性及评估程序的非公开性，导致该制度在实践中被高度政治化，沦为制度输出国延伸数字主权、塑造他国数据规则的单边政策工具。这种政治化操作不仅可能导致新型数字贸易壁垒，与有关国际经贸规则产生冲突，而且可能对发展中国家的制度自主权构成挑战，固化南北国家间的数据鸿沟进而侵蚀发展中国家的发展权，并通过其“俱乐部效应”加剧全球数据治理体系的碎片化。面对这一格局，中国不宜简单模仿或寻求加入欧盟主导的制度体系，而应在深刻把握自身数据跨境治理底层逻辑的基础上，通过制度竞合与替代路径实现规则对接，最终推动建立更加公平、包容的全球数据跨境治理秩序。

6、数据赋权的法律实践评价与制度创建构想——基于欧盟法与中国法的分析（吴汉东）

来源：《法学评论》2026年第4期

数据产权立法是国际社会关注的重要话题。欧盟数据产权制度是一个“试错、矫正、创新”的法治实践过程，从《数据库指令》到《数据法案》，欧盟力图实现“合理配置产权+有效监管规制”的政策目标，其立法经验和教训为我们提供了有益的思想资料；中国数据产权制度建设，基于“政策指引——制度实践——立法确立”的实现路径，在是否赋权（正当性问题）、何种赋权（合理性问题）、如何赋权（可行性问题）方面进行有意义的法律探索。中央系列文件和《民法典》规定，乃至地方试点方案和司法裁判经验，为今后数据产权立法奠定重要基础。未来“数据产权条例”，与传统所有权和经典知识产权相关法律有别，其主要规范构成涉及宗

旨条款、主体条款、客体条款、权益条款、登记条款、合同条款、救济条款等。以上述问题为研究对象，本文意在以学术法的见解为实在法的出台提供些许思想资料。

7、数据爬取行为正当性裁判规则的体系化重构（史欣媛）

来源：《法学评论》2026年第4期

数据爬取竞争纠纷的司法裁判普遍关注数据流通的价值目标，但判定结果以行为不正当为多数，呈现出阻碍数据流通的事实效果。究其根源，既有裁判规则体系存在实体性规则与程序性规则的双重缺陷：前者内嵌对数据本质、竞争逻辑和竞争法品格的认知偏差；后者指向融贯规则适用的分析方法之“元规则”缺位。为此，司法应实现从“泛干预主义”到“市场调节优先”的理念归位，其倡导从静态到动态、从个体主义到整体主义的双重范式转型。在此框架下，须重塑实体性规则与拓补程序性规则：一方面，应限缩“额头汗水规则”和“三重授权原则”，调适“实质性替代规则”和“破坏技术管理措施规则”；另一方面，将“最少且必要规则”拓展为比例原则，并通过其与动态体系论的有机耦合，形塑可视化、程式化的裁判工具。

8、理想与现实之间：十论数据产权登记制度（申卫星）

来源：《法学评论》2026年第4期

数据产权登记面临数据的非竞争性、无形动态性与价值半衰性，加之其与信息的一体同构及高频流转特征，难以直接套用既有登记模式，制度设计须在登记成本、确权效率与公示公信之间作出权衡。权衡理想与现实的制度设计应以对内承担核验权利内容、对外表彰权利状态为核心功能，遵循自愿登记、规范统一、公平诚信与便捷高效并重，以及安全有序等原则。登记主体采国家统一遴选机制，通过对登记机构持续履职能力、专业审查能力与责任承载能力的持续监督替代行政身份授权作为登记主体可信的制度基础。登记客体通过复合要素实

现数据的概括描述与合理识别。登记内容围绕权利主体、权利类型与权利负担等对外进行披露,权利来源可归纳登记为自行生产、协议取得、公开收集与衍生创造四种类型。以七类数据产权为产权类型,配套转让、许可、信托、担保与异议等登记类型,并在合理审慎的审查标准之下,按数据类型、权利来源与登记类型差异化配置审查强度与公示程度。登记效力以数据权利推定与对抗为中心,登记体例簿在电子化与多维检索能力提升的前提下,可在统一平台、统一描述、统一编号的框架内提供差异化查询路径,最终构建统一、高效、可信的数据产权登记制度。

9、回到数据流通的起点:数据访问权的体系扩张与规范构造(徐玖玖)

来源:《法制与社会发展》2026年第4期

在数字时代,数据流通的重要性和意义应高于数据交易。现有数据基础制度改革侧重于通过产权制度的强激励促进数据流通、释放数据价值,但数据产权与数据流通之间存在逻辑偏差,这限制了前者对后者的激励效果。数据访问是数据流通的起点,其功能并非《个人信息保护法》第四十五条所能涵盖和统摄的。在当下现实语境中,数据流通实现的规范性路径不畅,数据产权作为政策性概念的规范转化受阻,数字市场公平竞争秩序规制乏力。由此,我国有必要构建综合性的数据访问权,使其承担权利保障、监管协助、发展促进等不同功能。具体而言,数据访问权制度应当在个人数据访问权的基础上,体系性扩张公共数据访问权和企业数据访问权,并基于法益位序形成均衡的权利设计,为数据访问权的行使设定规范边界和限制条件。

10、数据产权登记凭证的证据效力(覃子轩)

来源:《法学杂志》2026年第4期

我国构建数据产权登记制度的核心之一是登记法律效力的设定,从民事诉讼证据法理论定位登记凭证的证据效力更有利于登记制度有效运行。从这一视角看,登记客体决定证明对象,实质审查程

序和登记机构的运行管理模式能够确保登记具备较高的公信力,故登记凭证可被赋予“准公文书”的证据地位。参照公文书的推定效力,登记凭证的证明力构建为形式证明力和实质证明力两层。法院在审查登记凭证时,对形式真实的审查可向登记机构核实,对实质真实的审查采取法律推定方法,相反证据证明的事实应当达到“足以推翻”的证明程度。为保障登记凭证发挥证据效力,考虑增设针对登记凭证的“书证提出命令”,并根据登记凭证真伪争议的不同情形构造解决程序,同时应当细化实质审查标准。

个人信息保护

1、社会关系网络视角下个人信息的私密与公开(邱遥堃)

来源:《中外法学》2026年第3期

个人信息的私密与公开之判断,是个人信息保护的前提性、基础性与根本性问题。然而,当前判定私密信息的各种理论皆有不足:微观的场景理论之实质性与体系性有缺失,中观的要件理论之操作性待提高,而宏观的功能理论对私密信息欠保护。私密信息判定的社会关系网络理论将社会科学融入法学教义,能够对此问题予以更佳回应:只需将网络结构与信息性质、技术架构结合,即可确定信息经由何种关系传播到网络的哪个节点与位置时,最有可能继续广泛传播,不再私密;这亦适用于一般个人信息的保护与利用。智能时代的数字技术一方面提高了社会关系网络的传播效率,模糊了公私界限,助力信息传播;另一方面也致使信息过载、信息失真、信息茧房等问题愈演愈烈,造成结构洞新增与固化,阻碍信息传播,最终导致个人信息的私密与公开问题陷入新的不确定性。

2、劳动者个人信息保护中的“知情同意”困境及其出路(刘冰洋)

来源:《法学评论》2026年第4期

面对劳动场景中知情同意规则的实施困境,既有的“同意替代”与“同意限制”两种代表性解决

方案并未结合“知情”阶段存在的问题作出回应，其针对“同意”阶段的改良也因“所必需”的适用范围过窄、“劳动者获益”等标准存在缺陷而效果不佳。用人单位告知的不充分以及劳动者的认知局限是“知情”阶段存在的主要问题。为此，应当要求用人单位通过劳动行政部门制定的标准条款告知信息处理事项，并对其中涉及劳动者敏感个人信息处理以及对劳动者切身利益实施自动化决策的条款作出提示与说明。劳动关系影响下的同意缺乏自主选择空间、同意的后果难以被干预是“同意”阶段实施效果不佳的主要原因，应当运用劳动合同无效规则实现对同意内容的优化以及对同意后果的干预。具体而言，一方面需要通过目的原则等否定对劳动者不利的条款的效力；另一方面，应当消除用人单位基于相关无效条款实施的信息处理行为对劳动者造成的不利影响。

3、大模型时代个人信息删除权的重构与实现（付新华）

来源：《当代法学》2026年第4期

生成式人工智能使个人信息删除权陷入行使困境：一方面，个人信息一旦经过语料训练嵌入模型结构，便与模型深度耦合，难以被彻底删除；另一方面，即便删除了原始数据，模型仍可能基于“记忆”生成相关信息，使得法律上的“删除”无法实现“被遗忘”目的。这一困境是立法选择与技术变革共同作用的结果。“删除本质论”将权利的核心权限定于“删除”这一单一行为模式，而大模型的技术架构使该履行方式难以实现其规范目标。回溯权利本源，个人信息删除权旨在实现“被遗忘”，以删除为核心的权利定位偏离了这一初衷。理应以“被遗忘”为规范目标，重塑个人信息删除权的权利内涵，使其从单一的删除请求转变为由“删除权”与“限制权”组成的复合权利架构：前者旨在删除数据，后者则旨在限制数据的访问、使用、生成和传播。“限制权”更加契合大模型的技术特性，也更有助于权利目标的实现，为平衡个人权益和模型提供者利益提供了有效的替代路径。与此相应，大

模型提供者的法律义务亦应转向“删除”与“限制”并举的双重结构，即在应用层延续删除义务、在基础模型层转向限制义务。

4、论个人信息侵权风险的民法治理路径（刘忠炫）

来源：《当代法学》2026年第4期

个人信息侵权风险是介于上游现实损害和下游将来损害之间的一种抽象风险，第三人侵权损害具有明显的或然性。侵权风险与传统损害在内在特征、影响机制和治理逻辑等方面存在本质差异，司法实践亦对两者予以区分。风险作为损害的法律安排，会引发侵权责任的证明难题、体系冲击和利益失衡等多维度困境。个人信息侵权风险的规范定位是一种权益妨害，其创设了一种权益侵害的危险状态，同时妨碍信息主体对个人信息权益的行使。抽象风险必须具备严重性或紧迫性，权利主体才能通过人格权请求权获得救济，以回复个人信息权益的圆满状态。个人信息保护请求权的行使不以实际妨害发生为前提，可以弥补人格权请求权的救济不足，二者共同构成个人信息权益的预防性保护体系。

5、个人信息调取的正当程序（张迪）

来源：《法学研究》2026年第3期

实践中，法律规范相对滞后和粗疏，个人信息调取的授权规范不明确、调取程序失范、程序性保障缺失等多重问题亟待化解。侦查规范体系改革的功利化和被动化、技术治理下个人信息共享的政策推动、个人信息权益侵犯的隐蔽性和无形性等因素叠加，进一步加剧了个人信息调取的法律控制难题。个人信息调取的规范体系以基本权干预之判定为起点，由调取之法律授权、调取之令状程序、调取之程序保障等构成正当程序三元框架。我国个人信息调取正当程序的三元结构应包括：以隐私权为中心的比例控制；以检察机关为中心的准司法令状制度；以个人信息权益为中心的程序保障。据此，应在比例原则之下运用“五要素衡量法”和“五要素干预法”完善授权规范，明确令状审批主体、申请主体和具体内容，构建四元保障程序。

6、人工智能时代个人信息保护的规制选择——一个时间秩序的视角（尹玉涵）

来源：《政治与法律》2026年第7期

随着人工智能领域的技术突破、规制跟进以及政治博弈，有关个人信息保护的共识正在分化，复现了科林格里奇困境意义上先发展还是先治理的选择难题。科林格里奇所采取的线性时间观阻碍了规制选择的确立，因而需回到法律系统运作的时间秩序中寻找答案。在以法律决策串联起的时间秩序中，法律系统在大多数时刻以封闭的运作实现维持规范性期望的功能，这使得法律系统往往滞后于科学、技术、经济等系统，但具备较好的稳定性。对权利本位的坚持是维持个人信息保护规范性期望的重要表征，在此基础上，通过立法、执法、司法、守法等环节的法律决策，规制选择得以持续调整，进而构筑我国个人信息保护的时间秩序。人工智能时代的到来并不会改变有关权利本位的规范性期望。应当在《中华人民共和国个人信息保护法》背景下，根据企业守法的意愿与能力划分不同类型的违法行为，生成不同的法律决策方案，企业合规只能作为个人信息保护的权宜之计。

网络平台治理

1、平台用工劳动者权益保护立法模式的比较与选择（谢增毅）

来源：《中外法学》2026年第3期

近年来，对平台用工进行立法已成为国际潮流。各国和地区基于平台用工面临的挑战和立法目标以及劳动法的传统等，对平台用工采取了不同的立法模式。根据主要规制工具和立法重点，可将域外平台用工的典型立法模式概括为“雇员身份推定+算法规制模式”“劳动基准模式”“社会保障模式”以及“综合立法模式”。我国平台用工立法宜采取“综合立法模式”，既应保障平台劳动者的基本权益，也应适当维持平台用工的灵活性，避免简单套用传统的雇员和自雇者保护规则。在立法思路，应坚持分类分层保护的原则，全面规定平台劳动者的各项权益，科学设置平台劳动者的基本劳动基准，

对劳动法的一般规则进行适当调适，明确平台企业和平台合作企业的责任分担。通过平台用工立法可创新劳动法的理论，使劳动法适应数字时代劳动用工形式变化的新需求，推动我国劳动法学自主知识体系的建构。

2、著作权侵权平台责任的关键问题及其化解：过错认定与必要措施（王迁）

来源：《中外法学》2026年第3期

区分直接侵权与间接侵权以及平台不负有审查用户上传的内容是否侵犯著作权的义务，是平台责任的两块基石。平台“知道或者应当知道”的对象是具体的侵权内容，而不是平台中存在许多侵权内容的状态。技术性手段的作用是识别和推测用户上传的内容与权利人作品的相似度，而不是判断其是否侵权。以平台概括性地知道存在侵权内容为由，将平台使用过滤技术并达到拦截或清除绝大多数侵权内容的效果设定为“必要措施”，势必动摇平台责任的两块基石，并不可取。司法保护应促进权利人团体与平台方的合作，充分发挥“通知与移除”机制的功效，由平台采用经济上合理、技术上可行、能够基本准确定位侵权内容的技术手段，减少侵权内容的传播。

3、人工智能时代网络平台的版权安全保障义务（杨依楠）

来源：《法商研究》2026年第4期

通过在推荐、搜索、创作辅助等环节部署人工智能，网络平台正经历数智时代的技术转型与服务升级。人工智能驱动的网络平台服务易产生为侵权内容分配更高流量权重、诱发用户浏览侵权内容意图等问题，传统的被动注意义务规范已无法满足版权侵权治理的需要。为网络平台设定版权安全保障义务更契合技术中立思想下的利益平衡原则，可以较好地兼顾人工智能技术创新与版权保护，且具备充分的规范依据。网络平台版权安全保障义务具有多重来源，可依据所部署人工智能的种类及潜在风险，为网络平台灵活配置情境识别与行为引导、算

法评估与修正、版权过滤以及作品数据收集等具体义务。在义务边界认定方面,应在保留侵权事实明显性标准的同时,结合人机协同的治理模式变化灵活调整认定标准。应通过优化举证规则、厘清法律责任等途径,构建面向人工智能时代的网络平台版权注意义务体系。

4、网络服务提供者版权内容过滤义务的“差序格局”配置(孙山)

来源:《当代法学》2026年第4期

人工智能时代的到来,正从根本上改变着网络服务提供者扮演的社会角色,避风港规则的有效性受到越来越多的质疑,而版权内容过滤义务的引入又因成本问题进展缓慢。过滤义务的引入应以著作权人和网络服务提供者之间的合作关系为基准,以兼顾二者利益、实现共赢为目标。由此,立法者应根据网络服务提供者和著作权人之间的合作关系密切度以及作品传播过程中的特殊情况,将过滤义务类型化为事先或事后免费提供可对比作品的主动过滤义务、事后单纯提供合格通知的被动过滤义务和特定情形下的特别过滤义务,建构“差序格局”的版权内容过滤义务体系,以立法方式推进双方之间的合作。

5、平台中立的终结:网络内容治理的路径反思与重构(迟舜雨)

来源:《比较法研究》2026年第3期

平台中立理念以言论自由为导向,主张平台是信息传输的技术中介,不对用户内容进行实质干预。该理念源于互联网早期并影响全球治理格局,但逐渐引发权责失衡与网络生态失序。随着平台通过算法推荐、内容排序和流量分配持续塑造信息可见度与互动结构,其商业逻辑、外部性影响与公共话语影响相互交织,平台呈现出公私复合结构,中立理念逐渐丧失现实解释力。因此,应当在承认平台对网络舆论结构具有塑造能力的基础上,构建以非中立为前提的治理框架。在政府设定基本目标与程序性约束的前提下,应赋予平台必要的裁量空间以实

现合作规制,并通过透明度、可问责性与多元参与机制防止权力滥用,同时依据平台规模、业务模式与风险外溢程度实施差异化责任配置。

数字行政与数字司法

1、数字政府中正当程序的底层逻辑与重塑路径(彭涛)

来源:《法商研究》2026年第4期

数字政府的运行依赖于人工智能技术,而人工智能在世界运行规律、认知模式和决策模式上都与传统政府存在根本差异。这些差异意味着,传统正当程序虽然能够约束政府权力,却难以约束数字政府中的算法权力。人工智能在三个层面影响传统正当程序。一是人工智能基于相关关系运行算法,传统正当程序以因果关系审查权力正当性,这一冲突弱化了审查效果;二是人工智能对社会整体结构、人类价值判断和社会网络的认知出现化约错误,数字政府据此容易作出错误决策;三是人工智能的决策包含难以解释的“暗知识”,弱化政府决策的理性程度。重塑正当程序的路径包括:首先,要审查三种类型的相关关系,以减少相关规律对程序正当性的弱化;其次,在涉及社会事实的决策中排除自动化决策,以减少认知错误对决策正当性的弱化;最后,扩张正当程序的规制范围,设置数字政府的三种强制性算法解释义务,减少人工智能对公共权力行使及决策正当性的影响。

2、数字时代司法裁判说理的方法论重构(高尚)

来源:《比较法研究》2026年第3期

人工智能助力司法实践提质增效,在司法裁判说理领域具备广阔的应用价值。数字时代下,司法裁判说理的范式发生了深刻转变,呈现出说理渊源扩张、法律与事实融合趋于精细化、传统法律解释规则与位阶冲突被隐蔽算法替代的特征,进而引出难以决断、难以诚实以及难以说服等问题。数字时代司法裁判说理的方法论重构,需坚守效率、安全与公正动态平衡原则:法源说理层面适配形式逻辑,回应说理资源扩张问题;事实说理层面兼顾个

案事实阐释与先例事实规则梳理；价值说理层面调和机械司法与后果考量的矛盾，明确法官在裁判说理中的主体责任，以增进司法裁判说理的说服力和公信力。

3、数字时代刑事法官预判的困境及其纾解（韩振文）

来源：《比较法研究》2026年第3期

在我国全案移送制度下，刑事法官在庭前程序中会对待决案件形成相对确定的预判。数字时代庭审实质化改革与人机协同审判出现交织共振，在两者动态交互的视域融合之下审视刑事法官预判现象，其遭遇的现实困境主要包括：“大概率”的不合理预判较难矫治纠偏；不合理预判径直成为终局性结论；法官对智能预判易产生自动化偏见。困境的认知根源在于全案移送制度下“卷宗笔录中心主义”裁判方式的盛行，而这种盛行的背后，审判认知结构的形成从庭审场域错位到侦查阶段，同时智能预判也固化加剧法官预判认知偏差。纾解困境的相应措施主要有：合理限制案卷笔录的移送使用；探索确立阅览案卷方面的限定性规则；刚柔相济协力推进辩护权、对质权、算法解释权等程序性权利保障与救济；通过场景化的技术性正当程序规制，确定智能预判分级适用的合理边界，并提供适度的理由说明，使负责任的智能预判具备相对可解释性。

4、司法裁判的属人性：人工智能司法角色的限度（陈景辉）

来源：《中国法学》2026年第3期

人工智能甫一出现，即刻成为广受欢迎的司法辅助工具，但其原本无法承担辅助角色之外的任务。然而，生成式人工智能却可能突破辅助角色的限制，有机会成为像人类法官一样的司法主导者。这是因为，生成式人工智能已经可以像人类法官一样，理解法律的要求并生成决定，同时还具备一些人类法官所不具备的优势，于是就有了取代人类法官的可能。然而，由于司法是关于人类实践的事务，因而必须满足三种属人性的要求：人类给予人类的对待、

法律给予人类的对待、个人给予人类的对待。生成式人工智能主导的司法裁判无法满足这三种属人性的要求，所以人工智能无法替代人类法官成为司法裁判的主导者，仍只能在其中扮演辅助性的角色。

5、论数字司法价值对齐的实现路径（彭中礼）

来源：《中国法学》2026年第3期

数字司法中算法决策与司法价值的偏离容易引发伦理与法律风险。在数字司法场域中推进价值对齐，应当坚守司法本质以定其向，防控技术风险以固其基，筑牢司法信任以彰其效，融合司法文明以臻其境。数字司法价值对齐要求算法的价值追求与法律精神、社会伦理、人文关怀形成动态适配，其具体要求体现为维护一致性、确保可控性、保持鲁棒性以及避免恣意性。推进数字司法价值对齐，需凝聚全链条、全过程的协同合力，在源头校准层、运行管控层、迭代优化层与安全保障层系统构建层次化推进的过程机制。理论分析表明，数字司法中的价值对齐通过将人类多元价值转化为技术可理解、可执行的逻辑，让人机协同从效率适配的工具层面跃升为价值共生的生态层面，最终形成技术进步与人类发展深度契合的良性循环。

新兴科技产品的法律问题

1、物联网产品网络安全缺陷的产品责任与公法规制（李钊）

来源：《法商研究》2026年第4期

物联网产品因网络安全缺陷而造成人身或财产损害时，生产者责任恰落在《网络安全法》的行政监管义务与《产品质量法》的产品责任之间，受害人难以索赔。这一问题源于传统产品责任制度所暗含的“完成性”预设，即产品在交付时属性固定，生产者责任以交付为限。物联网产品的“生成性”瓦解了这种预设。生产者在交付产品后，仍能通过网络连接持续塑造产品的功能与安全状态，对产品的控制力并未随其交付而消失。因此，对物联网产品的产品安全规制应在时间上延及全生命周期，在内容上自物理结构延及网络领域。网络安全是产品

应当具备的安全品质，包含信息机密、系统完整与功能可用三项规范要求，可据此认定产品缺陷、设定公法义务。其中，网络安全缺陷致人身损害或缺陷产品以外的其他财产损害的，依产品责任制度追究损害赔偿责任，以功能主义标准界定产品范围，并限缩免责事由。持续维护产品网络安全能力的供给性责任则不以损害为前提，而应由公法经全生命周期义务与基于风险分级的合规机制予以确立。

2、论自动驾驶领域人机共驾模式中的注意义务分配（周梦杰）

来源：《政治与法律》2026年第7期

自动驾驶领域人机共驾模式中的注意义务分配应以信赖原则为基础。在功能界定上，若将信赖原则定位于对结果的预见可能性，则将加重主体的责任负担；若将其简单归结为容许风险的表现形式，亦会混淆评价标准和评价结论，难以完成多主体间

的责任分配。唯有立足于注意义务分配方能实现责任分配。在分配义务的正当性根据层面，“褒奖理论”忽视人机信赖基础的不对称性，“自我答责原则”消解人机协同的互动性。只有基于主体间性视角的利益衡量论才能实现义务分配的妥当性。在具体分配路径层面，应构建以危险预见义务为核心的层级化分配模式。该义务本质上属于信息收集义务，并依据人机交互过程中的互动关系呈现出层级化差异：自动驾驶汽车生产商、制造商等前端责任主体承担前提性的信息收集义务，自动驾驶系统协助前端责任主体履行场景化的信息收集义务，人类驾驶员则负有补足性的信息收集义务。以信息收集义务为核心的义务分配体系，可为自动驾驶领域人机共驾模式中多方主体的义务分配和责任认定提供路径支撑。

（技术编辑：毕坤阳）

教研活动

马长山：《数字时代的人权逻辑》讲座

2026年6月12日上午，以“数字时代的人权逻辑”为主题的学术讲座在南京师范大学仙林校区行敏楼434会议室顺利举行。本次讲座特邀华东政法大学数字法治研究院院长、长江学者、国家“万人计划”哲学社会科学领军人才**马长山**教授担任主讲人，南京师范大学法学院**何柏生**教授担任主持人，**张镭**教授、**陈辉**副教授、**章楚加**副教授担任与谈人。



讲座伊始，何柏生教授对马长山教授在百忙之中莅临指导表示热烈欢迎与诚挚感谢。他介绍了马长山教授在数字法学领域的开创性贡献，指出马长山教授是我国数字法学研究的领军人物之一。



马长山教授在讲座中围绕“数字社会的人权挑战”“当代人权的数字化倾向”“第四代人权的形成”“数字人权的法治保障”四个维度展开阐述。

马长山教授首先从信息革命对人类自身的重塑切入。在数字时代，每个人既是生物人，也是数字人，人机协同、人机共生已成为现实。这一变化

带来了人权保障的全新挑战：侵权从个别性、间断性转变为自动化、机制化、系统化，算法歧视、差别定价、选择性服务等现象日益普遍，而普通个体在技术霸权面前往往处于不对称的弱势地位。

接着，马长山教授深入分析了人权的数字化倾向，指出数字技术使人的属性从固化的、封闭的自然人转变为分布的、多维的数字自我。每个人在滴滴、美团、微信等平台上都拥有多个“数字分身”，这些数字画像比本人更真实、更全面，甚至可以通过控制数字身份来反向操控物理世界中的个体。传统人权保护基于物理空间与精神世界的二元结构，而数字时代的权力结构已演变为公权力、私权利与平台私权力的三元博弈，原有的保护逻辑亟需重构。在此基础上，马长山教授系统阐述了“第四代人权”的理论主张。他指出，第一代人权是政治权利，第二代是经济社会文化权利，第三代是集体权，前三代均以自然人为基础。而第四代人权突破这一范围，以数字属性为核心，强调数字自主权、数字知情权、数字参与权等新型权利。其核心价值是反对数字控制，促进人的自由发展。马长山教授特别提到，刚刚发布的第五期《国家人权行动计划（2026—2030年）》首次写入“美好数智生活”等数字人权内容，这不仅是中国首次，也是世界所有有人权行动计划国家的首次，标志着数字人权正式获得国家层面的政策确认。在数字人权的法律保障方面，马长山教授提出六点建议：确立数字人权的法治理念、构建数字规则体系、探索数字正当程序、建立场景化保护机制、构建人权的道德基础设施、培养数字公民能力。他强调，中国应当在全球数字人权治理中发挥引领作用，构建中国自主的数字法学知识体系。

在与谈环节，**张镭**教授指出马长山教授的讲座带来了强烈的“头脑风暴”。马长山教授从人类文明革命的高度切入，揭示数字时代技术发展导致大量劳动力变得多余化的深层问题，提出在数字人权研究中需要进一步厘清“人”的范畴——究竟是自然人、半数字化的人，还是觉醒后的AI主体？



陈辉副教授表示，马长山教授从市民社会到数字社会的学术转型令人敬佩。他指出，当前数字法学研究面临一个核心困惑：数字技术带来的究竟是既有权利体系的侵蚀，还是根本性的范式革命？从欧盟的保守态度与中美的激进应用对比来看，中国在应用端大胆前行的同时，也承受着权利被多维度侵犯的压力。这需要学界思考：传统侵权构成要件是否足以应对数字时代的挑战，抑或我们需要全新的理论范式？



陈辉副教授表示，马长山教授从市民社会到数字社会的学术转型令人敬佩。他指出，当前数字法学研究面临一个核心困惑：数字技术带来的究竟是既有权利体系的侵蚀，还是根本性的范式革命？从欧盟的保守态度与中美的激进应用对比来看，中国在应用端大胆前行的同时，也承受着权利被多维度侵犯的压力。这需要学界思考：传统侵权构成要件是否足以应对数字时代的挑战，抑或我们需要全新的理论范式？

章楚加副教授结合自己的研究领域分享了体会。她指出，马长山教授的讲座具有极强的前瞻性与现实关怀，特别是在国家人权行动计划刚刚颁布的背景下，能够第一时间扣紧政策前沿进行学理阐释，为年轻学者指明了研究方向。她特别提到，讲座中关于数字人权类型化、具体化的方法论指引，对跨学科研究具有重要启发意义，环境权利保障与数字生态治理的交叉领域值得深入探索。

释，为年轻学者指明了研究方向。她特别提到，讲座中关于数字人权类型化、具体化的方法论指引，对跨学科研究具有重要启发意义，环境权利保障与数字生态治理的交叉领域值得深入探索。



讲座最后，何柏生教授总结指出，马长山教授的讲座从历史演进、技术变革、权利逻辑、制度保障等多重维度，系统建构了数字人权作为第四代人权的理论框架，既有宏大的文明视野，又有鲜活的现实案例，为全院师生带来了一场思想盛宴。本场讲座在热烈的掌声中圆满落幕。

马长山：《数字司法的法治边界》讲座

2026年6月13日上午，以“数字司法的法治边界”为主题的学术讲座在南京师范大学仙林校区行敏楼434会议室成功举办。本次讲座特邀华东政法大学数字法治研究院院长马长山教授担任主讲人，南师法学院何柏生教授担任主持人，法学院吴英姿教授、韩玉亭副教授、刘韵副教授担任与谈嘉宾。



马长山教授首先从“三维世界”的法治变革入手，指出现代法治的核心架构由公法秩序、司法秩

序和理性精神三者构成,产生于物理时空,以工商秩序、自然人模式和分配正义为基本逻辑。随着信息革命的到来,数字世界应运而生,人类从传统的物理—精神二维世界跃升为“三维世界”。数字世界并非物理世界的简单映射,而是通过数字孪生等技术手段对物理世界产生深远的反向塑造作用,具有平行性、镜像性和互钳性。在司法领域,这一变革使得“法律真实”与“客观真实”之间的鸿沟有望缩小,对司法质效带来深刻影响。



在数字司法的焦点问题上,马长山教授梳理了四个核心议题。其一,司法原则的数字再造。技术外包导致公民参与不足,一体化办案平台对公检法司的分工制约构成冲击,全域数字法院挑战级别管辖与地域管辖,异步审理影响直接言辞原则。其二,数据业务的法治边界。区分了“业务数据化”与“数据业务化”,指出后者范围难以限定,数据持有不对称、确权不明晰、规则不健全、权益不平衡等问题突出,普通公民往往处于最不利地位。其三,正当程序的技术改写。传统规则与自由裁量权正被转化为算法,因缺乏公众参与,技术性改写易偏向权力行使者一方,形成数字司法权力的“技术性逃逸”。其四,公平正义的数字重塑。数字正义的核心不在于利益分配本身,而在于谁控制数据和算法,包括对数据信息的控制、基于算法的控制和通过算法的控制三重逻辑。

针对上述问题,马长山教授提出数字司法需要实现三重平衡:一是司法能动与司法谦抑的平衡,应走出能动主义与克制主义的悖论,迈向司法平衡主义,能动履职应体现为恪守职责的“功能性治理”。二是数据业务与数字正义的平衡,数据的挖掘利用

应以案件为起点,遵循合法性、合理性、必要性和相称性原则,严格控制对个人生活状态、行踪轨迹的不必要分析。三是数字技术与法治人文的平衡,人类决策基于价值判断与法治精神,而机器依赖概率计算,司法必须保持温度与人文关怀,体现“科技向善、以人为本”。

最后,马长山教授概括了数字司法发展的六个主要向度:恪守数字权力的合法性,防止技术性逃逸;主动推动司法制度变革;厘定数字正义原则,坚持比例相称;构建数字正当程序,完善数据利用与算法应用程序及权利救济;保障最低数字人权,守住法治底线;促进全球数字法治。在这六个向度的基础上,他勉励在场学子面向未来,积极拥抱数字时代,为中国数字法治贡献力量。

在与谈环节,各位嘉宾与马长山教授展开了深入的学术对话。吴英姿教授提出应以“程序密度”为分析框架,运用程序相衬原理构建数字司法的正当程序体系。她认为,数字时代传统程序的沟通理性可能被算法消解,数据利用需要合比例,程序保障可以分高、中、低密度进行分层设计。



韩玉亭副教授探讨了数据主义对规范逻辑的挑战及数字弱势群体权利保障问题,指出前案裁判不公可能通过大数据系统形成累积效应,进一步固化司法裁判中的不公正倾向。他强调,法官既需要提升法学专业素养,也需要提升数字时代的基本素养,才能不被时代所抛弃。



刘韵副教授结合自身在智能要素审判领域的研究经历，分享了学习数字司法的心得体会，认为数字司法的关键不在于司法数字化本身，而在于数字司法时代如何继续保持法治化并持续发展法治精神。



何柏生教授在总结中强调，不学习数字法学将难以适应公检法系统的工作要求，数字司法已体现在司法工作的方方面面，鼓励同学们珍惜学习机会，补齐相关知识。



本次讲座立足数字法治前沿，内容丰富、视野宏阔，为师生理解数字时代司法变革的核心议题提供了系统性、前沿性的学术指引，对推动法学院在数字法学领域的教学与研究具有重要促进作用。讲座在热烈的掌声中圆满落下帷幕。

乔仕彤：《互联网平台与中国知识产权保护的崛起》讲座

2026年6月25日晚，中国人民大学法学院、未来法治研究院在明德法学楼725室成功举办题为“互联网平台与中国知识产权保护的崛起”的学术讲座。本次讲座特邀杜克大学法学院讲席教授乔仕彤主讲，中国人民大学法学院副院长丁晓东教授主持，清华大学法学院院长崔国斌教授、中国政法大学大数据和人工智能法律研究中心主任沈伟伟副教授、中国人民大学法学院副院长万勇教授、中国人民大学未来法治研究院执行院长张吉豫教授、中国人民大学法学院民商法教研室主任熊丙万教授担任与谈人。

讲座回顾 平台作为执法基础设施及其三重条件

乔仕彤教授首先阐述研究的缘起。他指出，网络盗版执法是一个全球性难题——2026年3月，美国联邦最高法院在相关案件中收紧了网络服务商帮助侵权责任的认定标准，确立“除非平台主动引诱侵权或其服务专用于侵权用途，否则无需为用户侵权行为负责”的原则；而同一时期的中国走了一条截然不同的道路。他以个人经历切入，回顾2000年代初校园中音乐近乎“零成本”获取的盗版时代，对照当下普遍付费的正版生态，强调中国网络音乐市场在2012年至2018年间完成了从盗版主导到全球最大正版数字音乐市场之一的“180度转变”，这一短期内发生的剧烈转型亟需理论解释。

乔仕彤教授正式提出其核心概念——“平台作为执法基础设施”。他分析，大型互联网平台凭借对数据流、内容分发渠道、自动化监测技术、支付系统以及资本与法律能力的综合掌控，不仅在自身平台内部实施规则，更将版权执法行动延伸至整个互联网生态，实现跨平台、全天候、规模化的保护。他强调，这一概念区别于既有两类文献：一是国家如何规制平台的研究（如反垄断、劳动关系认定），二是平台自治的研究；其创新点在于揭示国家如何借助平台能力，在平台自身边界之外完成规模化执法。



杜克大学法学院讲席教授乔仕彤

为使该理论具有解释力与边界感，乔仕彤教授提炼出三个必要条件。其一，国家规制能力的集中化。他指出，网信办、国家版权局等监管机构的设立显著提升了行政威慑力，2013至2014年对百度音乐的顶格处罚及对相关企业的关停处置，释放出向正版化转向的明确而可信的政策信号。其二，市场集中的激励机制。他借助“公共品”与“搭便车”理论阐释：版权制度建设惠及全社会，理性主体缺乏单独投资的动力；而平台通过大规模购入独家授权形成垄断预期，将制度收益内部化为可独占的长期回报，从而获得推动转型的内生动力。其三，平台自身的技术、资本与法律能力，包括自动化监测、证据保全、全国性诉讼网络与法务体系建设。三者缺一不可，共同构成模型的限定性前提。

在实证层面，乔仕彤教授人工编码并分析了2014至2018年间约三千余份网络音乐版权侵权判决书，数据显示原告整体胜诉率高达99.6%，平台作为原告的胜诉率达99.9%；胜负关键高度一致——只要原告能证明持有独家授权并提供经公证的侵权证据保全，法院即认定侵权成立，败诉案件几乎全部源于权利归属无法证明，与原告是否为特定平台无关。他据此指出，在网络音乐版权领域，地方保护现象基本不存在，执法已高度标准化、机械化，法院裁量空间极小。

乔仕彤教授进一步阐述该模式的代价与边界。一是过度执法风险，自动化监测可能压缩合理使用与表达自由空间；二是内容控制强化，平台有动机巩固排他地位、规避政治风险，且监测技术可迁移至其他内容领域；三是市场结构扭曲，独家版权竞

价推高成本，致使虾米音乐等中小平台退出、消费者选择收窄与价格上升，最终倒逼监管介入（2017年起要求平台相互授权，2021年责令腾讯解除独家版权安排）。他最后强调，该模型并非普适模板，其有效性严格受限于上述三个条件：视频领域因长视频市场“五强并立”、集中度不足、侵权识别技术成本更高、合理使用抗辩更复杂等原因，未能复现音乐领域的执法成效；美国则因缺乏集中化的国家规制能力、受反垄断约束难以形成市场垄断预期，亦未出现类似路径。

与谈环节



清华大学法学院院长崔国斌教授

清华大学法学院院长崔国斌教授关注“平台”概念的独特性问题。他认为，只要版权人以某种方式集中、规模足够覆盖维权成本，出版社、集体管理组织乃至美国BMI等机构均可具备类似维权能力，未必对应传统意义上的“平台”。他进而审慎反思三项条件：就国家能力而言，政府展现认真执法的意愿即可，未必依赖能力的集中化——以百度为例，其商业模式背后必然存在与搜索引擎的利益分配，只要严肃对待既有规则即可遏制；就市场集中而言，普通许可在反垄断约束下亦可催生维权，独家授权非必要前提；就技术能力而言，门槛并不高，冠勇科技等小公司即可研发监测技术。他主张，法律被认真执行本身即足以驱动维权，不一定需要超大平台。



中国政法大学沈伟伟副教授

中国政法大学**沈伟伟**副教授借助《代码 2.0》提出的“法律、技术、市场、社群规范”四要素分析框架，从市场与社群规范两个维度切入。在市场维度，他认同垄断使主体从“搭便车”转为“主人翁”从而强化执法的观察，并指出视频领域难以集中，部分源于内容审查制度凌驾于知识产权机制之上，以及技术比对成本更高。他还揭示了一个商业模式变迁：正版成本曾远高于盗版（CD 及播放设备对比 MP3、APE 格式），而今会员制付费体系使盗版的搜索与设备成本反而更高。在社群规范维度，他提出版权作为无形财产本质上是一种“信仰体系”，在信息过剩与生成式人工智能兴起的背景下，以稀缺性为依托的版权付费正当性正面临根本挑战，并追问中美用户的版权认知是否趋同。



中国人民大学法学院副院长万勇教授

中国人民大学法学院副院长**万勇**教授结合产业史与教学调查指出，盗版减少的主因并非严苛诉讼（2000 年代美国起诉个人的做法反而引发消费者反感），而是商业模式创新带来的成本降低——其课堂调查显示学生几乎普遍付费，月费仅十余元，叠加平台与银行卡等优惠活动，正版门槛大幅下降。

他还分析了音乐独家许可的特殊性：音乐可反复收听、具有强粉丝粘性（如用户为特定歌手的独家曲库而付费），区别于视频的一次性消费特征。此外，他对“美国仅有单一视频平台则垄断性更强”的判断提出商榷。



中国人民大学未来法治研究院执行院长张

吉豫教授

中国人民大学未来法治研究院执行院长**张吉豫**教授肯定该研究将平台长期策略性盈利与竞争地位纳入成本收益分析的启发意义，并提出三点请教：其一，音乐案件中“通知删除”规则的实际适用情况，为何避风港原则似乎失灵、侵权认定比例如此之高，是否多与反复通知或反复侵权有关；其二，为何中国形成以平台为主导的维权模式，而非美国式的第三方专业化服务路径，中美机制对比如何；其三，对于视频等领域，若要推动有技术参与的执法环境优化，关键影响因素为何，并提示音乐与长视频在检测技术成熟度、成本及合理使用抗辩复杂性上的差异。



中国人民大学法学院民商法教研室主任**熊丙万**教授尝试以“执法动力—执法能力”二元框架重构三项条件。就执法动力而言，他认为独家授权并

非唯一来源,普通许可只要收益足够亦可驱动执法,并以微软对散户执法动力较低、兼顾政治合规考量、以及自身收入与消费者付费意愿变化为例加以说明;他还设想了“若早期即严格执行反垄断、禁止二选一”的思想实验,追问平台执法动力是否会因此弱化。就动力与能力之辨,他追问前平台时代盗版泛滥究竟源于能力缺失还是动力不足,并主张厘清平台在执法链条中是否具有不可替代性。

提问交流

在自由提问环节,多位同学提出深入问题。

有同学提出分发渠道建设对正版市场形成的决定性作用,指出权利人全球发行能力有限,欠发达地区天然成为盗版温床,并认为执法动力与执法能力难以截然二分,应置于同一框架统筹考察。

有同学提出腾讯收购独家授权后已兼具版权人与执法者双重身份,区别于YouTube自研Content ID的执法模式;并对国家能力与平台能力的因果方向提出疑问,认为司法态度转变或许是平台能力演进的结果;同时建议将具有中国特色的“主体责任”理论纳入论文,强化与欧美的对比。

还有同学认为三项条件若仅为成本收益分析则过于一般,无法解释为何独音乐领域成效显著,建议进一步细化各条件的操作化定义,并提示规范要素(市场博弈可内生的规范动力)在模型中被忽视。

另有同学从国际协作视角解释音乐侵权下降,指出国内平台与海外权利人建立紧密合作、形成超越法定授权范围的自发保护机制,并以某视频平台对合作密切方(日漫、漫威)作品下架严格、对部分无合作关系的盗版影视内容长期持模糊态度,说明平台“有能力而未必有动力”;同时指出音乐侵权认定标准相对宽松(翻唱常被容忍,甚至被视为具有宣传效应),而影视内容则因严格下架凸显维权的选择性。

乔仕彤教授以开放姿态回应全场评议。他首先解释本研究是一种“实证描述”而非“规范倡导”。他强调,理解中国的国家能力,前提是把握自2012年起“压实平台主体责任”这一网络治理范式转型

——治理思维由“九龙治水”转向体系化治理,最高层面的政治意愿通过机构建设得以落实,由此才有对百度的顶格处罚及司法裁判规则的变化;就版权领域而言,这意味着从简单的侵权责任规则转向国家治理体系。

对于崔国斌教授的评议,乔仕彤教授认可依靠第三方专业维权或认真执法是另一种可能更优、负面效果更少的路径,但其在中国2012至2018年间并未真正实现;正是平台从盗版主要分销者转变为版权持有者这一角色逆转,提供了快速正版化的现实路径。针对“知识产权制度整体进步”的替代解释,他回应:若是整体进步,各领域理应同步,但现实是音乐进步尤为突出;且至今音乐版权仍以海外权利为多数、视频则以国产内容为主,故“国产版权人崛起”无法解释音乐领域的领先转型。关于平台界定,他坚持两点:一是平台须为盗版的主要分销渠道,二是须具备资本、技术、法律的复合能力以支撑长期战略投入——技术可以购买,但唯有腾讯这样的综合体才能在数年亏损中坚持布局,腾讯音乐独立上市仍愿承受亏损即是例证。

对于沈伟伟副教授的评议,乔仕彤教授认同垂直生态整合的判断,但提示腾讯音乐单独上市仍须自身盈利、接受市场检验;并指出文化通常难以骤变,恰恰更凸显中国快速转型之特殊。关于价格,他认为市场监管总局的介入部分源于平台滥用垄断地位损害消费者权益。他还援引相关学者关于互联网时代知识产权“社会化生产、分销成本趋零”或将导致版权消亡的论说,指出该框架因缺少“平台”这一变量而未能预见现实:互联网时代监测与执法成本同样大幅降低,版权保护不仅未消亡,反而得到加强。

对于万勇教授的评议,乔仕彤教授认可商业模式创新的作用,但强调若无互联网治理环境的结构性变化,这一创新不会如此迅速到来。

对于张吉豫教授的评议,乔仕彤教授设想随人工智能与检测技术进步,视频侵权侦测成本未来或降至当前音乐的水平,届时平台亦可能推动视频领域从“五强并立”走向集中,其现实可能性甚至高

于美国式路径,但是否为规范上可欲的结果仍存疑问。

丁晓东教授在总结中指出,本次讲座为如何拆解一个现象的发生机制提供了范例。他强调,平台执法是一个贯穿多个领域的问题,平台责任本质上既是侵权问题,更是大规模治理问题;音乐领域的“特殊性”在某种程度上正是治理难度问题——能否以较低成本将合法使用、二次创作与侵权区分开来。他进一步指出,平台责任领域中的“假阳性”与“假阴性”是全球共同关注的关键议题,需在阻止侵权与避免误伤合法表达之间寻求平衡。他代表主办方对乔仕彤教授的精彩分享及各位与谈人、同学的深入讨论表示感谢。

本次讲座在热烈的掌声中圆满结束。



第十三届“新兴权利与法治中国”学术研讨会

6月27日,第十三届“新兴权利与法治中国”学术研讨会在苏州大学成功举办。本次会议由新兴权利学术共同体组委会主办,苏州大学工商业与人权研究院、苏州大学王健法学院、苏州大学期刊中心承办,《苏州大学学报(法学版)》编辑部、北大法宝(北大法律信息网)、北京市环球(苏州)律师事务所协办。来自中国人民大学、华东政法大学等知名高校的专家学者、二次文献机构的负责人、“新兴权利与法治中国”学术共同体合作期刊代表、友刊嘉宾、获奖论文作者等80余人参加会议。

此次学术研讨会共设置了四个分论坛。分论坛一上半场由北大法律信息网(北大法宝)副主任**孙殊妹**主持。北京理工大学法学院副教授**陈禹衡**、华东政法大学数字法治研究院副研究员**童云峰**、甘肃政法大学法学院副院长**曾磊**教授、西南政法大学讲师

施殊妹先后报告了其论文《价值对齐与规范重构:

“数字泔水”侵蚀未成年人的协同治理》《人工智能大模型法律容错的行刑反向衔接机制》《从非法流转到非法合成:深度伪造技术下个人生物识别信息保护的刑法因应》《刑事涉案虚拟财产价值的分类分层认定》。最高人民检察院理论所副所长、《中国刑事法杂志》副主编**蔡巍**,苏州大学王健法学院教授、《苏州大学学报(法学版)》副主编**王俊**对报告人的获奖论文进行了针对性点评。



下半场由《证据科学》副编审**刘奕君**主持。同济大学法学院博士研究生**许露沙沙**、清华大学法学院博士研究生**孙鹏庆**、北京化工大学廉政研究中心教授**金鸿浩**分别报告其论文《低空经济监管的行刑衔接困境、成因与出路》《涉虚拟货币犯罪:属性、类型与证明范式》《非法删除型数据犯罪的实务反思与分类处置》。苏州大学王健法学院教授、《苏州大学学报(法学版)》副主编**王俊**,北京外国语大学法学院教授、《北外法学》副主编**郑曦**对报告人的获奖论文进行了针对性点评。



分论坛二上半场由温州大学期刊社社长**缪心毫**主持。中国民航大学法学院副教授**王锡柱**、西南政法大学经济法学院讲师**茅孝军**、上海政法学院副教授**孙益武**、中国人民大学博士研究生**李婉祺**先后报告了其论文《论面向低空经济场景的服务型法及其实现路径》《地方政府实施低空空域特许经营的法理基础与规范进路》《我国数据质量综合立法：逻辑、难题与应对》《国际投资仲裁中数据资产的归属与领土联系》。《东方法学》专职编辑、副教授**卢玮**，苏州大学王健法学院副教授、《苏州大学学报（法学版）》编辑**施立栋**对报告人的获奖论文进行了针对性点评。



下半场由苏州大学王健法学院副院长、《苏州大学学报（法学版）》编辑**瞿郑龙**主持。西安交通大学法学院副教授**苟学珍**、东南大学法学院副教授**毕文轩**、华中科技大学法学院讲师**谭佐财**分别报告了其论文《低空数据要素市场化配置的法律激励》《论我国 AI 智能体包容审慎监管的法治进路》《数字时代财产所有权的消解及法治因应》。南开大学法学院副院长、教授**陈兵**，《求是学刊》执行主编、教授**李宏弢**对报告人获奖论文进行了针对性点评。



分论坛三上半场由苏州大学王健法学院教授、《苏州大学学报（法学版）》副主编**李中原**主持。北京交通大学法学院副教授**付新华**、上海政法学院副教授**李晓珊**、温州大学法学院研究生**何泽炜**、中国人民大学法学院博士研究生**王常阳**先后报告了其论文《AI 智能体的数据安全风险及其法律规制》《人工智能透明度原则的内涵及法治化研究》《论具身智能体数据隐私风险的公私法协同治理》《人工智能模型要素权利的司法保护路径——以全国首例 AI 模型侵权案为引》。广西民族大学教授**齐爱民**、《暨大公报（哲社版）》副主编**李晶晶**对报告人的获奖论文进行了针对性点评。



下半场由《上海政法学院学报》副主编**汤仙月**主持。上海大学副教授**李晨光**、中山大学法学院博

士研究生**任厚朴**、对外经济贸易大学法学院助理教授**李依怡**、西北政法大学教授**袁震**分别报告了其论文《可持续人工智能的法理基础及制度实现》《可解释性原则在数字行政法中的确立和展开》《基于用户授权的智能体数据权利探析》《论数据资源持有人的权能及其限制》。《河南大学学报（社会科学版）》编辑、教授**任瑞兴**，《新华文摘》杂志社编审**王青林**，《上海政法学院学报》执行主编、编审**康敬奎**对报告人的获奖论文进行了针对性点评。



分论坛四上半场由《东北师大学报（哲社版）》常务主编、编审**秦卫波**主持。东北财经大学讲师**马贤茹**、温州大学法学院副教授**任江**、华东政法大学讲师**姬蕾蕾**先后报告了其论文《大模型蒸馏技术的不正当竞争风险类型化识别与法治因应》《场景一致性原则下数据资源持有人动态配置问题研究——以在线消费评价数据为例》《论人工智能模型训练中的数据治理义务》。《高等学校文科学术文摘》编辑部主任、编审**沈丽飞**，《苏州大学学报（哲社版）》编辑**姜天琦**对报告人的获奖论文进行了针对性点评。



下半场由《学习与探索》副主编**朱磊**主持。上海政法学院副教授**商建刚**、上海政法学院副教授**冯硕**、美国纽约大学法学院博士研究生**赵浴辰**、南京航空航天大学博士研究生**徐冰灿**分别报告了其论文《AIGC 可版权性的理论转向：汇编视野下的版权逻辑》《统筹发展和安全视域下数据涉外管辖权的行使》《数字正义的权利进路及其局限：以数字权利的分配效应为中心》《EVTOL 应用下的新兴法律问题及规制研究》。《北京行政学院学报》编辑、四级调研员**魏新**，苏州大学王健法学院教授、《苏州大学学报（法学版）》主编**上官丕亮**对报告人的获奖论文进行了针对性点评。



中国科学技术法学会 2026 年年会

6月27日，由中国科学技术法学会主办、黑龙江大学法学院承办的中国科学技术法学会2026年年会在黑龙江省哈尔滨市召开。本次会议以“铸科技法治经纬 固创新发展根基”为主题，设有主论坛和八个分论坛，研讨主题分别是“科技法学的基础理论与科技法治的价值内涵”“人工智能治理：法律、政策、应用规范与伦理准则”“具身智能发展的法治保障与开源的创新保护”“智能经济时代的知识产权保护机制”“平台经济健康发展的法治保障”“数据安全利用与生物医药法律保护”“新兴领域立法与创新体制”“科技成果转化保护与科研诚信法治建设”。中国法学会有关部室主要负责同志，有关实务部门和高校、研究机构的专家学者参加会议。

中国法学会副会长、中国人民公安大学党委书记、校长**王轶**作了题为《法律如何分配数据利益？》的主题演讲，指出数据是一种新型生产要素，数据利益存在初次分配和衍生分配之分。数据人身利益的初次分配和衍生分配，《民法典》《个人信息保护法》等作出了相对完备的回应。数据财产利益的初次分配是一个宪法问题，是一个应当慎重考虑的价值判断问题。数据财产利益的衍生分配主要借助市场机制完成，民事法律行为制度发挥关键作用。

哈尔滨工程大学党委副书记、校长**殷敬伟**以《逐梦深蓝 向海图强 以高水平科技自立自强助推海洋强国建设》为题作主题演讲，指出关心海洋，从“望洋兴叹”到“向海立命”，在历史转折处完成战略觉醒；认识海洋，以“火眼金睛”洞穿万顷幽深，让“龙宫宝藏”交出基因图谱；经略海洋，凭“顶天立地”的硬核科技，杀出一条深蓝突围的自主之路。

分论坛一

分论坛一的研讨主题是“科技法学的基础理论与科技法治的价值内涵”。上半场由中国科学技术法学会常务副会长、中国政法大学数据法治研究院

院长**时建中**教授主持。华东政法大学数字法治研究院院长**马长山**教授作了《数字司法的可能与限度》的主题发言，认为数字司法应立足科学、司法、民众和时代立场，围绕预期立场、可能空间和内在限度三个方面，创新呈现司法运行平台化、司法规则代码化、司法决策建模化和司法正义可视化，同时关注算法客观性悖论、正义判断困境、政治因素遮蔽、精准性障碍及数智风险防控等内在限度问题。国家自然科学基金委员会科研诚信办监督二处处长**杨亮**作了《以良法善治护航科技强国建设——以国家自然科学基金项目资金监管为例》的主题发言，认为应当以良法善治为核心完善国家自然科学基金项目资金监管体系，通过健全法规制度、强化过程监督、优化绩效评估防范资金违规风险，提高资金使用效益，保障科研诚信，为科技强国建设提供制度支撑。中国人民大学法学院教授**王锴**作了《科技法典编纂的理论基础与体系展开》的主题发言，认为科技法典编纂是科技强国战略的法治回应，旨在解决现行立法存在的规范分散、规则不一等问题，科技法典宜采取弹性法典化思路，通过“立改废”确定基本内容，围绕创新促进、成果转化、风险防控和权利保障等核心内容，推动科技法律规范的体系整合与制度升级。云南大学法学院副教授**戴琳**作了《乡村振兴视域下科技特派员制度知识产权法治补强刍议》的主题发言，指出科技特派员制度在农业科技成果转化中面临知识产权属、合作研发和技术入股等风险，应通过专门立法、细化规程、吸纳知识产权专员和开展培训，提升纠纷防范能力，保障各方权益。江南大学法学院副教授**李硕**作了《从食品安全到食品科技法治：技术创新、风险治理与公众信任》的主题发言，认为食品科技法治应以分类分级规制引导技术创新，以治理时点前移、过程延伸和对象扩展应对新型风险，并通过知情、程序透明与责任救济三重机制保障公众信任，实现创新与安全的制度平衡。发言完毕后，新疆政法学院党委常委、副校长**祁欢**和北京工业大学经济管理学院教授**孙玉荣**进行了精彩与谈。

下半场由中国科学技术法学会副会长、北京航空航天大学人文与社会科学高等研究院院长**龙卫球**教授主持。中国科学技术法学会副秘书长、中国政法大学国家安全学院副院长**郜庆**副教授作了《科技法学自主知识体系的标识性概念》的主题发言，通过系统阐述平台、数据、算法三大标识性概念的生成逻辑与体系关联，强调三者共同构成数字社会运行的基础结构，彼此支撑、协同衍生，持续催生新的制度规则，成为科技法学自主知识体系的核心支柱。北京市京师律师事务所业务指导委员会副主任**张秀华**作了《科创企业治理的科技法嵌入：公司法与科技创新法治的制度衔接研究》的主题发言，从实践角度分析科技法嵌入科创企业治理的紧迫性与可行路径，指出当前公司法与科技法在治理结构、信义义务、资产制度等五大层面存在深层冲突，强调须通过扩张信义义务、完善出资规则、重构治理架构等制度创新，推动科创企业治理从资本本位向创新、技术与安全本位的治理范式跃迁。杭州电子科技大学法学院教授**李庆峰**作了《律师调查令电子化转型的类型学建构与程序法意涵》的主题发言，认为律师调查令电子化转型并非简单的载体替换，而是从纸质令状到数字指令的程序形态重塑，通过类型学分析揭示技术嵌入对执行程序结构的深层重构，可以推动执行调查权性质演变与程序参与关系升级。烟台大学法学院教授**于海防**作了《推荐性国家标准的网络传播：以标准传播领域“官告民”第一案为例》的主题发言，通过探讨推荐性国家标准的版权属性，剖析现行司法认定与制度逻辑之间的矛盾，指出推荐性国家标准依法具有法律效力且以财政经费制定、以行政公告发布，应属于具有行政性质的文件，不应受著作权法保护。发言完毕后，天津大学法学院科技法研究所所长**吕凯**教授和中国计量大学法学院院长**朱一飞**进行了精彩与谈。

分论坛二

分论坛二的研讨主题是“人工智能治理：法律、政策、应用规范与伦理准则”。上半场由中国科学技术法学会副会长、北京大学法学院教授**张平**主持。西安交通大学法学院教授**焦和平**作了《“创作过程”

在作品认定中的地位：从传统创作到AIGC》的主题发言，认为在人工智能生成内容可版权性判断中，传统创作模式下只问结果、不问过程的创作推定面临挑战，应将创作过程作为证明自然人智力主导、构思存在、独创劳动真实性及表达来源可追溯性的重要依据，从而在AIGC作品认定中重申著作权法保护人类创造性表达的规范根基。中央财经大学法学院副教授**张金平**作了《AI时代索尼规则的重生抑或式微？》的主题发言，认为“实质非侵权用途”规则在AI时代仍具有生命力，但其适用需要结合服务提供者的知悉程度、侵权可预见性以及必要防范和补救措施进行判断，以实现技术创新与著作权保护之间的平衡。昆明理工大学法学院讲师**崔亚冰**作了《生成式人工智能服务提供者内容过滤义务》的主题发言，认为生成式人工智能服务提供者负有适当注意义务，内容过滤应聚焦输入端与输出端，并在技术不可完全控制、幻觉难以根除的现实下，通过必要措施、投诉举报、风险提示和不确定性标识等机制划定合理边界。大连海洋大学海洋法律与人文学院讲师**杨凯旋**作了《机器学习利用作品行为著作权合理使用的认定》的主题发言，认为机器学习利用作品具有非表达性使用特征，现行合理使用制度应回应训练数据需求和授权成本困境，推动形成“半开放+双重例外”的规则结构，并以合法获取、非替代使用和三步检验法的综合判断协调技术创新与权利保护。吉林大学法学院讲师**袁帅**作了《论接入式人工智能服务提供者的著作权侵权责任》的主题发言，认为通过API等方式接入人工智能服务后，接入平台、人工智能服务提供者和用户之间的协作会使侵权责任认定更加复杂，接入式人工智能服务提供者可能因诱导用户调用功能而承担相应责任，被接入平台也应对接入程序的侵权风险负有审核、注意和必要制止义务。发言完毕后，北京航空航天大学法学院教授**孙国瑞**和山东大学法学院教授**崔立红**进行了精彩与谈。

下半场由中国科学技术法学会副秘书长、河北省人民检察院法律政策研究室**李志彬**主任主持。上海大学知识产权学院教授**王勉青**作了《人工智能生

成内容的违法性认定》的主题发言，认为AIGC违法性具有成本低、迷惑性强、扩散快等特点，实体法层面需明确AI定位、价值评判标准、侵权认定规则和多元主体责任，程序法层面需完善取证鉴定、证据审查和举证分配机制，并在包容审慎与分类分级监管中建立安全有效的容错机制。中国社会科学院大学法学院互联网法治研究中心主任**刘晓春**作了《AI智能体跨平台数据处理行为的法律评价》的主题发言，认为智能体跨平台数据处理涉及用户、智能体提供者与被访问平台之间的多重利益冲突，应在用户同意、个人信息转移、平台数据权益、生态开放和互联互通之间进行利益衡量，并结合数据类型、应用场景、技术拦截、负外部性与交易成本确定法律评价标准。厦门大学知识产权研究院副院长**董慧娟**教授作了《AI拟人化互动服务安全监管的现状、不足与解决》的主题发言，认为AI拟人化互动服务在情感陪伴、心理交互等场景中可能产生依赖、误导和安全风险，应以精准监管为导向，优化拟人化互动服务的概念边界，区分直接提供者与间接提供者责任，并将用户分类从单一年龄标准拓展至精神状态、认知水平等多元因素，建立前置性和动态性风险测评机制。阿里巴巴集团法务部合规总监**顾伟**作了《人工智能立法的实务思考》的主题发言，认为人工智能立法应回应产业真实痛点，在安全导向与应用牵引、一般规则与场景差异之间实现平衡，确立适应AI服务的新时代避风港原则，探索“通知—拦截”等损害防止机制，并围绕监管互认、智能体安全、训练数据合理使用和国际规则对等提出制度建议。黑龙江大学法学院讲师**王亚霏**围绕《论人工智能立法中的软法规范》作主题发言，认为，当前人工智能立法广泛借鉴战略政策性文件，大量技术标准和行业规范文件涌入立法文本，导致行业共识型、企业自律型、技术内嵌型等不同软法规范的产生，我国应重视软法在立法活动中的运用，通过借鉴成熟立法经验保证我国人工智能立法的结构科学性和形式开放性。发言完毕后，中国政法大学数据法治研究院教授**王立梅**和中国科学技术

法学会副秘书长、中国科学院大学公共政策与管理学院副教授**刘朝**进行了精彩与谈。

分论坛三

分论坛三的研讨主题是“具身智能发展的法治保障与开源的创新保护”。上半场由中国科学技术法学会副会长、北京科技大学教授**徐家力**主持。华北理工大学人文法律学院教授**吴凡**作了《具身智能产业市场准入法律规制问题研究》的主题发言，认为具身智能“软硬一体”颠覆传统分治立法，应确立统一准入、分类分级的市场准入框架，依风险分级设定标准与认证，构建事前准入、事中监管、事后退出全链条机制，衔接国际规则以应对出海合规需求。中国政法大学民商法学院讲师**张旭**作了《具身智能致损风险下的保险工具创新与制度保障》的主题发言，认为具身智能通过风险兜底、激励相容和规则补位三重功能应对致损风险，政府需要协调构建可保性，建立以强制责任险为主的分层保障体系，配套差异化保费、程序协同及强AI时代公共兜底机制。中国科学技术法学会副秘书长、北京大学智能学院助理研究员**辜凌云**作了《具身智能可以作为法律主体还是事实主体？——重新思考费英格的“仿佛”哲学》的主题发言，认为具身智能法律地位之争应引入费英格的“仿佛”哲学，主张其不必“是”法律主体，但可“视其为”事实主体即权利义务归属点，这只是实践拟制而非本体论认定，现下赋予完整法律主体地位为时过早，选择如何虚构是治理的核心。石家庄铁路运输检察院副检察长**张杰彪**作了《具身智能致害的刑事风险及其法治保障：基于三方主体责任分配的视角》的主题发言，认为传统刑法应对具身智能致害面临因果、罪过及工具理论三重失灵，应放弃“意志直接支配”转向“危险支配程度”标准，依制造者、部署者、使用者三方角色分配责任，并走递进式法治保障路径。中南大学法学院博士研究生**汪雨馨**作了《AI智能体越权行为的归责困境与化解原则》的主题发言，认为AI智能体越权归责应从“谁授权”转向“谁决定”，以“决定中心主义”统一归责，越权分手段性与目标性两类，责任归于真正作出该层级决定

的主体。发言完毕后，黑龙江大学法学院国际法教研室主任**杨健**教授和上海开放大学教授**芦琦**进行了精彩与谈。下半场由中国科学技术法学会副秘书长（执行）、北京科技大学文法学院副院长**王霁霞**教授主持。湖南大学法学院教授**喻玲**作了《专利开源：法律性质与条款效力》的主题发言，认为专利开源在性质上属于附解除条件的知识产权许可合同，以专利私权为基础，通过开源许可证约定双方权利义务，其条款具有法律约束力，违约将导致授权自动终止，但开源不等于专利豁免，使用者仍需警惕潜在的专利侵权诉讼风险。中国政法大学比较法学院讲师**胥智仁**作了《开源人工智能基础模型的治理架构研究》的主题发言，认为开源AI基础模型的各种特质削弱传统归责逻辑，法律治理不应依赖外部刚性规制，应依开放程度及应用场景的社会风险场域，差异化配置提供者与使用者义务，以社会价值正当性衡量作为责任减免考量，构建适配特质的治理架构。北京航空航天大学法学院博士研究生**刘伟斌**作了《从开放到控制：人工智能模型开源的版权法观察》的主题发言，认为人工智能模型开源已从纯代码开放转向复合对象开放，但模型权重的可版权性动摇版权许可基础，开源许可证正承担超出传统版权法的治理功能，从自由使用转向条件性控制，需厘清对象识别与义务传导，防止协议过度扩张。暨南大学法学院博士研究生**郝家杰**作了《人工智能开源许可协议的规则因应》的主题发言，认为传统开源许可协议受限于软件时代的源代码逻辑，在适配AI模型要素时存在适用风险，司法实践应明确其具有合同性质，违反可择违约或侵权救济，从协议效力界定、义务传导与责任分配层面进行回应。广东省中山市司法局黄圃司法所一级科员**谢潇扬**作了《新时代“枫桥经验”视域下我国开源AI商标混淆侵权风险防控与调审衔接机制研究：以Openclaw连续更名事件为镜鉴》的主题发言，认为开源AI商标纠纷频发源于技术创新与知识产权保护之间的张力，应借鉴新时代“枫桥经验”，构建调审衔接机制，通过源头预防、诉前调解、司法确认与诉讼托底，实现低成本、高效率的柔性治

理,以平衡创新活力与权利保护。发言完毕后,北京师范大学法学院教授**夏扬**和北京邮电大学人文学院教授**谢永江**进行了精彩与谈。

分论坛四

分论坛四的研讨主题是“智能经济时代的知识产权保护机制”。上半场由中国科学技术法学会副会长、同济大学上海国际知识产权学院院长**丛立先**主持。山东科技大学知识产权学院副院长**赵丽莉**教授作了《国际对标下量子加密算法标准化进程中的知识产权风险与协同应对》的主题发言,认为美国NIST主导后量子加密算法标准化,过程中存在专利故意/非故意不披露、SEP许可僵局及跨国诉讼等知识产权风险,为应对需强化规则前置、FTO分析、FRAND承诺及专利池,政府应加强合规指导与国际协同。中国科学技术大学知识产权研究院副教授**高亮**作了《著作权惩罚性赔偿的从严保护与谦抑适用》的主题发言,认为著作权惩罚性赔偿旨在矫正网络侵权成本失衡,但权属模糊与批量诉讼易致过度保护,须对盗版重复侵权从严应对,对权利未明及平台间接责任审慎对待,推进权利审查、情节认定与数额核定。法治日报社全面依法治国智库负责人、法治网总裁助理兼研究院院长**杨幸芳**作了《智能时代著作权立法目的之反思与完善》的主题发言,认为智能时代需要将文化、艺术、科学成果最快速度地转化为技术运用,而现行著作权立法目的难以支撑科技创新,建议将《著作权法》立法目的条款的“科学”改为“科技”,在著作权保护与科技发展之间进行动态平衡。武汉大学法学院博士研究生**汪焱梁**作了《版权法中动态网站屏蔽制度的双阶构造》的主题发言,认为版权法中的动态网站屏蔽制度通过法院在先判决,将后续实质相同的镜像网站自动纳入规制,无需另行起诉,能有效遏制重复侵权,弥补现有过滤机制的不足,具备法理基础。河北经贸大学法学院讲师**张玮琛**作了题为《权利保护还是权利扩张:知识产权维权边界的再反思》的主题发言,认为知识产权维权应从形式审查转向形式正当性审查,正当维权需满足稳定权利基础、适当范围、合比例方式及可正当化的竞争效

果,实践中存在利用不稳定权利、程序压力或标准锁定效应扩张权利的现象,应引入诚信原则与比例审查予以规制。发言完毕后,重庆邮电大学网络空间安全与信息法学院副院长**黄东东**教授和《科技与法律(中英文)》编辑部副主任**姜文武**进行了精彩与谈。下半场由中国科学技术法学会副会长、武汉大学法学院教授**宁立志**主持。南开大学法学院教授**朱冬**作了《“游戏换皮”类案件反不正当竞争法适用路径之反思》的主题发言,认为当前法院以反不正当竞争法规制“游戏换皮”的路径存在误区,认定思路与著作权侵权判定趋同,正当性考量与著作权法政策相悖,反不正当竞争法的弹性被误读,导致向一般条款逃逸,应避免抵触专门法的立法政策,确保法律适用的稳定性。北京市中伦律师事务所合伙人**马东晓**律师作了《反法39条的构造及其实践偏差》的主题发言,认为《反法》第39条系举证责任转移,但实践中常生偏差,以“整体技术信息”认定秘密时未明确具体内容,致被告无法有效反证,将“接触+实质相同”的举证责任转移误作倒置,强化“自证清白”,甚至将推定规则滥用于赔偿计算,损害技术人才流动与创新。安徽大学法学院教授**胡小红**作了《技术秘密“秘密性”认定中“容易获得”的阐释》的主题发言,认为技术秘密“秘密性”认定中,“容易获得”的判断标准存在分歧,通过观察产品即可直接获得的信息属于公开,但需复杂反向工程获得的信息不应认定为“容易获得”,司法实践对反向工程的性质有不同观点,直接影响秘密性认定标准高低。国家心理健康和精神卫生防治中心助理研究员**孙潭霖**作了《医药研发工具行政认定的知识产权效应与规制》的主题发言,认为医药研发工具行政认定的“监管接受框架”法律性质与效力未明,会产生Bolar例外排除、竞争锁定、权利分离及检测工具垄断等知识产权风险,应依技术属性分公共化或私有化路径,并配套专利披露、开放许可及司法协调规则。西安市长安区人民检察院检察官助理**郭皓圆**作了《人工智能视野下知识产权检察综合履职的时代使命、问题与路径》的主题发言,认为人工智能赋予知识产权检察三重使命,

但跨界融合致制度、履职、监督不适应，应借鉴欧盟建立多层次保护，发挥专门化优势，并覆盖开发、市场、使用三阶段全链条监督。发言完毕后，中国商飞公司上海飞机设计研究院总法律顾问**金靖寅**和西安交通大学法学院教授**王玥**进行了精彩与谈。

分论坛五

分论坛五的研讨主题是“平台经济健康发展的法治保障”。上半场由中国科学技术法学会副会长、中国科学技术大学知识产权研究院**宋伟**教授主持。华东政法大学知识产权学院副院长**于波**教授作了《网络音乐平台内容传输的著作权规制》的主题发言，认为网络音乐平台“秒传”实为平台主动存储、识别、匹配并开放服务的重复数据删除技术，在法律上属于复制行为而非信息网络传播行为，平台不能以技术中立免责，应承担著作权侵权责任。中国人民公安大学法学院教授**苏宇**作了《Skills 集成平台的法律治理》的主题发言，认为 Skills 集成平台依类型不同，平台定位应从“中间商”转向“守门人”，承担更严注意义务，安全、人身权、知识产权的三维风险需分类分级治理，健全审核、授权与合规机制，参照欧盟法律，平台需从被动合规转向主动治理。四川师范大学法学院教授**唐仪萱**作了《数字平台安全保障义务的合同法分析》的主题发言，认为数字平台安全保障义务应纳入合同治理路径，贯穿合同全程，其作为主给付义务需平衡发展与安全，平台风险注意采取主客观结合标准，约定优先，并以合理期待为边界，违反义务的司法审查应立足合同诚信，而非仅依侵权责任判断。河北省廊坊市人民检察院法律政策研究室副主任**陈磊**作了《直播带货平台假冒注册商标罪共犯归责的释义学建构》的主题发言，认为直播平台假冒商标归责困境在于帮信罪兜底与中立理论失据，应以三要素支配力为变量，依此构建保证人地位、中立临界、共同正犯三重路径，形成四阶归责体系。北京科技大学文法学院讲师**张硕**作了题为《平台规制权运行中的个人信息保护困境及其法治因应》的主题发言，认为平台规制权三面向（规则、监管、纠纷）因虚化、失范、悬置致困境，应以协商合规约束制定权，

全过程监管规范监督权，有限司法化激活裁决权，将私权力法治化，促合规实质正义。发言完毕后，北京大学法学院知识产权学院常务副院长**杨明**教授和重庆大学法学院教授**朱涛**进行了精彩与谈。下半场由中国科学技术法学会副会长、同济大学上海国际知识产权学院特聘长聘教授**朱雪忠**主持。厦门大学法学院教授**王兰**作了《大科技金融平台监管有限性及其新治理进路》的主题发言，认为大科技金融监管面临有限性，需引入新治理模式，通过监管权力下放与授出、平台数字基础设施共享开放、监管科技公共形成及数据伦理衡平，构建多元协同、动态灵活的治理体系，以化解监管困境与系统风险。南昌大学法学院副教授**姜川**作了《平台企业的未成年人网络保护社会责任法律化研究》的主题发言，认为生成式人工智能风险由财产权转向人格权，尤其重视未成年人保护，服务提供者注意义务须涵盖来源审核、内容监控与紧急措施，应设事前实质审核与告知，事中安全评估与分级干预，事后积极响应与救助。西北政法大学民商法学院副教授**国瀚文**作了《中立性视域下数字平台反垄断法适用研究》的主题发言，认为数字平台反垄断法适用困境下，应引入平台中立原则，结合事前规制与事后监管，依平台类型差异化治理，对头部平台依反垄断法审慎规制，构建竞争、消费者、数据法多元协同共治体系。湘潭大学法学学部副教授**邹琳**作了《知识蒸馏软标签法律属性研究》的主题发言，认为知识蒸馏中软标签不构成著作权作品而属于思想范畴，算法生成切断规范因果不具可版权性，但含人类劳动成果且有经济价值可定性为数据要素，通过数据权益、商业秘密及反不正当竞争法保护。安永（中国）企业咨询有限公司经理**官中奇**作了《外卖平台的算法平衡治理：以骑手和消费者视角为例》的主题发言，认为外卖平台算法治理需平衡骑手与消费者权益，冲突源于算法设计失衡、平台权力异化及议价能力失衡，应通过优化算法规则、平衡数据、重构平台责任、提升主体能力，构建多元协同的平衡治理体系。发言完毕后，上海政法学院数智法学院院

长曹阳教授和复旦大学法学院教授吕炳斌进行了精彩与谈。

分论坛六

分论坛六的研讨主题是“数据安全利用与生物医药法律保护”。上半场由中国科学技术法学会学术委员会委员、上海政法学院教授蒋坡主持。中南财经政法大学法学院教授刘大洪作了《人工智能数据安全治理的范式转换与法律激励》的主题发言，强调应重视人工智能数据安全治理的现实困境，反思传统管制型立法的局限性，推动治理逻辑从单纯依靠“威慑性成本”向“正向性收益”的范式跃迁，通过搭建“基础保障+长效驱动”的双层架构，实现激励机制的精细化与差异化配置。武汉大学法学院教授邓社民作了《数据知识产权属性略论》的主题发言，指出客体是区分物权、债权、知识产权等权利类型的关键标准，数据的非物质性、非消耗性、可复制性和公共产品性等与知识产权客体相同，且数据具有动态性和流通性、可验证性、场景依赖性等有特征，数据可以作为知识产权中新的客体。西北大学法学院讲师、院长助理刘桢作了《数据应是知识产权的载体还是客体》的主题发言，认为应当区分知识产权的载体与客体二者的不同，知识产权的客体是指知识产权法律关系中权利和义务所共同指向的对象，知识产权的载体是指承载、固定或呈现上述客体的物质实体或物理形式，将数据认定为知识产权的载体而非客体，这一定位既符合知识产权制度的基本原理，也契合我国现行立法精神与数字经济发展的实际需求。甘肃政法大学学术期刊部责任编辑罗祥副教授作了《无权取得数据加工处理的许可使用机制》的主题发言，认为在数据要素流通与利用的现实场景中，需突破传统物权视角的绝对化判断，重新审视数据行为的合法性边界，构建以默示许可为核心的数据使用规则，解决数据获取的授权难题。发言完毕后，北京航空航天大学法学院副院长周学峰教授和北京理工大学郭德忠教授进行了精彩与谈。下半场由西安交通大学法学院副教授曹毅搏主持。重庆邮电大学网络空间安全与信息法学院副教授郭亮作了《数据知识

产权登记的理论证成与制度完善：以信号理论为研究视角》的主题发言，强调信号理论下数据知识产权登记存在信号互认性薄弱、信号可信度不足、信号有效性存疑、信号动态性缺失等实践困境，应统一央地规则、构建成本差异、增强技术可信、形成反馈闭环。福州大学法学院副教授杨莉萍作了《拒绝用户获取个人健康医疗数据反竞争性的认定》的主题发言，通过区分拒绝竞争对手用户获取个人健康医疗数据和拒绝用户自身获取个人健康医疗数据等不同场景，提出了反竞争性认定的一般思路和抗辩理由。山东农业大学副教授毕凌雪作了《数据投毒致算法成果贬损的私权救济机制研究》的主题发言，提出数据投毒致算法成果贬损私法救济存在部门法规制的结构性缺位、动态无形损害的量化规则缺失、救济模式具有局限性以及过错、因果与损害的证明困难等困境，应当完善梯度化举证体系、建立动态无形损害分层量化机制、构建全链条立体化保护体系、构建动态成果完整性侵权认定标准。西安交通大学法学院博士研究生梁龙坤作了《出版数据确权与作品使用者权的冲突与协调》的主题发言，认为解决出版数据确权与作品使用者权冲突的关键路径在于实施“客体分层、权利区分、限权并重”的策略，构建精细化的权利配置体系，平衡各方利益诉求。吉林大学法学院博士研究生陆奕铭作了《公共数据概念阐释的范式转换：从公务用公物到公众用公物》的主题发言，指出公共数据概念阐释的范式应从公务用公物范式转换为公众用公物范式，从而实现公共数据供需精准对接、实现数据开放职责的独立化、有效破除数据行政垄断与数据孤岛。发言完毕后，浙江农林大学文法学院教授马永双和同济大学上海国际知识产权学院教授宋晓亭进行了精彩与谈。

分论坛七

分论坛七的研讨主题是“新兴领域立法与创新体制”。上半场由黑龙江大学法学院副院长邓齐滨教授主持。黑龙江大学俄罗斯研究院副院长宫楠副教授作了《俄罗斯互联网信息安全治理法治化路径研究》的主题发言，通过介绍俄罗斯为维护互联网

信息安全所构建的多层级法律保障体系框架，重点分析了俄罗斯在维护网络主权方面的制度规范与监管措施，并对我国互联网信息安全的前瞻性治理提出了参考建议。吉林大学法学院副教授**齐英程**作了《自动驾驶致损的责任认定与救济机制》的主题发言，提出交通事故场合下的侵权风险是自动驾驶汽车合法化与规模化发展的首要风险，应以接管义务配置为核心实现自动驾驶致损责任认定与救济规则的类型化划分，最终确定法定先行赔付义务与商业任意性保险相结合的救济规则。黑龙江大学法学院博士研究生**周践祚**作了《大语言模型在算法行政中的应用及法治化路径》的主题发言，提出大语言模型应用于算法行政具有技术逻辑层面的必然性，但在裁量真实逻辑、数据与模型的双向影响、相对人权利救济三方面存在问题，应从大模型应用原则与具体路径两个层面形成法治化规制思路。重庆大学法学院教授**李晓秋**作了《AI4S 科研范式下创新成果的发明人确定》的主题发言，指出应严格区分 AI 辅助发明与 AI 主导发明以及自主发明，通过司法解释或审查指南细化“实质性贡献”的认定标准，明确 AI4S 全流程中各环节的人类参与程度如何量化评估，为 AI4S 新型科研范式下创新提供清晰的行为指引。发言完毕后，河北省廊坊市人民检察院党组书记、检察长**肖衡**和吉林大学法学院教授**李晓倩**进行了精彩与谈。下半场由复旦大学法学院知识产权研究中心主任**马忠法**教授主持。河北省石家庄市检察院副检察长**张丽萍**作了《大数据智能化赋能刑事检察法理分析及风险应对》的主题发言，通过对我国刑事检察大数据智能化应用的现实图景进行全面分析，深入剖析大数据与人工智能赋能刑事检察的法理逻辑，并全面揭示其潜在风险及相应的解决策略。山东政法学院刑事司法学院副院长**张爱艳**教授作了《手术机器人医疗事故刑事归责的困境与出路》的主题发言，指出当前手术机器人在医疗活动中的介入已导致罪名认定出现一定程度的失准，需引入技术审查手段对医疗事故行为主体进行甄别，同时加强围绕手术机器人的司法鉴定，实现科技与法治的协同治理。黑龙江大学法学院讲

师**牛丹彤**作了《司法智能体应用的技术风险与分层规制》的主题发言，提出智能体的自主性和生成性特征会对传统司法运行逻辑产生冲击，对司法权能、司法程序、司法安全及司法公信力造成影响，为此需重塑司法智能体法治化治理路径，构建分层系统性规制体系，推进数字司法治理规范化发展。中南大学法学院博士研究生**朱园伟**作了《论生成式引擎优化行为的违法性边界》的主题发言，指出生成式引擎将检索、推理与自然语言生成融为一体的技术特点，使得传统的因果认定逻辑在违法性认定中已难以适用，应确立“制造算法偏向”的竞争侵害认定标准，对算法中立性的内涵予以重构，针对不同技术配置设定差异化的违法认定阈值。上海天尚律师事务所律师**詹德强**作了《深度合成语境下 AI 换脸侵害肖像权的法律规制与实务应对》的主题发言，通过以首例 AI 换脸肖像权侵权案为切入点，系统梳理了我国现行法规对深度合成行为的监管制度，从可识别性、授权边界、平台责任、抗辩事由与损害赔偿等方面提出了实务应对思路。发言完毕后，中国政法大学数据法治研究院教授**范明志**和北京市伟博律师事务所主任**李伟民**进行了精彩与谈。

分论坛八

分论坛八的研讨主题是“科技成果转化保护与科研诚信法治建设”。上半场由中国科学技术法学会副会长、最高人民检察院检察技术信息研究中心主任、内蒙古自治区人民检察院副检察长**刘喆**主持。中国科学院大学知识产权学院副院长**闫文军**教授作了《关于科技成果转化堵点与举措的反思》的主题发言，强调应关注科技成果转化的堵点，加强企业主导的产学研融通创新，推动科技创新和产业创新深度融合。安庆师范大学法学院院长**崔汪卫**作了《人工智能领域专利授权的伦理审查分析》的主题发言，提出人工智能领域专利授权的伦理审查面临忽视科技伦理审查、法律规定与现实需要具有矛盾、科技伦理专业人才缺乏等困境，应构建“科技向善”的专利授权伦理理念，明确授权范围并提高披露标准，多方共育科技伦理专业人才。上海杉达学院知识产权法商研究中心执行主任**张冬梅**作了《探析知

知识产权视角下浦东科技成果转化体系优化路径》的主题发言，指出浦东新区以知识产权为基础构建企业发现与运营体系，衔接国家战略与地方实践，但成果转化存在权属界定矛盾、价值评估难、知识产权金融工具嵌入不足等难题，应当优化确权与赋权制度、建立科学的评估标准、补强知识产权金融服务链。北京海润天睿律师事务所知识产权业务部主任**陈竟**作了《科技成果转化中的知识产权问题》的主题发言，提出应重视科技成果转化中知识产权鉴价评估、知识产权维权、商业秘密资产管理等现实问题。上海市行政法治研究所助理研究员**仲霞**作了《上海法律科技的应用和发展研究》的主题发言，提出应当立足法治建设数字化转型趋势，推动技术与法治深度融合，实现立执司法全链条智能化升级，打造协同高效、便民利民的现代化法治服务体系。发言完毕后，西安理工大学教授、校法律顾问室主任**王宇红**和中国融通科学研究院集团有限公司资深业务总监**高静**进行了精彩与谈。下半场由中国科学技术法学会副会长、中国政法大学教授**刘瑛**主持。河北大学法学院教授**牛忠志**作了《关于学术不端治理的制度不足与对策》的主题发言，强调应当认清学术不端行为的性质，实现他律与自律相结合，并通过制定专门的法律以不断提升学术不端治理成效。安徽国誉知识产权咨询有限公司董事长**韦国**作了《SDL 自我驱动实验室的法律规制构建》的主题发言，提出应重视 SDL 自我驱动实验室在知识产权、数据治理、责任归属、监管规则等方面的法律核心问题。上海段和段（虹桥国际中央商务区）律师事务所主任**郭国中**作了《专利+商业秘密：科研成果保护双引擎》的主题发言，认为应当通过专利的“公开换保护”与商业秘密的“保保护核心”协同发力，为科创企业构建了全生命周期、无死角的风险防御体系。发言完毕后，西安交通大学法学院教授**李晓鸣**和北京金诚同达律师事务所高级合伙人**马林艳**进行了精彩与谈。

北京大学第二届“数字法治的理论视野与实务前沿”暑期学校

2026年7月6日至7日，由北京大学法学院、北京大学数字法治研究中心主办的第二届“数字法治的理论视野与实务前沿”暑期学校成功举办。本次暑期学校为北京大学“研究生教育创新计划”项目，由北大英华科技有限公司（北大法宝）提供支持。暑期学校邀请了北大数字法治研究中心研究人员及校内外高水平专家，围绕“基于 AI 工具的法学研究与教学”这一核心主题，通过主题讲座、实务沙龙和智能工具应用实训等形式，为来自全国各地的五十余位青年教师、硕博研究生学员提供了接触数字法治领域前沿新知、共同切磋学习研究心得的优质平台。



AI 时代的法学研究（7月6日上午）

北京大学法学院长聘副教授、院长助理**彭鐔**以《AI 时代的法学研究》为题开启了暑期学校的首场讲座。他从 AI 对传统法学研究范式的冲击切入，并以性骚扰的法律规制研究为例，展示了同一法律问题可从多元路径展开研究的不同面向。他重点分析了 AI 技术为法学研究带来的新机遇，在法教义学领域，他以地方立法权研究为例，介绍了如何借助智能技术实现对十几万份裁判文书、七万余部规范性文件的全样本分析；在比较法学领域，AI 翻译技术极大降低了外语门槛，使学者得以同时涉猎多法域文献、开展真正的跨法系比较研究。他结合自身研究实践分享了 AI 工具的使用心得，强调技术只是辅助，真正的学术问题仍需研究者自己提出。



人工智能能否取代创作者（7月6日上午）

北京大学法学院助理教授**边仁君**为同学们带来题为《人工智能能否取代创作者》的讲座，从两个维度展开了实证研究。她指出，AI时代的作品兼具内容消费与训练数据双重价值。在内容产品维度，她基于十万组真实写作题目构建拓展版图灵测试，从语法、语义、风格、创意四个维度对比人机表现，发现AI在技术性维度上显著优于人类，但在叙事创新性上仍不及人类。在训练数据维度，她通过迭代微调实验发现，模型在自身生成的数据上反复训练会出现“模型崩溃”现象，创造力与风格多样性持续退化。其技术根源在于训练数据的长尾分布，AI基于概率的学习机制会不断强化高频表达而丢失罕见创意。AI始终需要源源不断的人类创作为其注入长尾信息，才能维持其多样性与活力。



AI赋能法学教育与研究（7月6日下午）

中国人民大学法学院副教授**彭雅丽**以《AI赋能法学教育与研究》为题，分享了人工智能在法学领域的应用方法与边界。她指出，AI不会自动提升法学教育与研究质量，只有在任务边界清晰、资料来源可追溯、人类专家规范判断三者同时成立时才能

真正发挥价值。结合法学领域的特殊性，她提醒学员警惕AI编造来源、混淆规范层级、将少数观点通说化等常见问题。在教育应用中，她介绍了AI在课前预习、课堂讨论、案例教学与写作反馈中的具体用法，提出将AI输出作为审查对象的训练思路，认为未来更应重视评价学生的来源核查能力与方法透明度。AI赋能不是让机器替法律人思考，而是让法律人把更多时间用于更有价值的判断与创造。



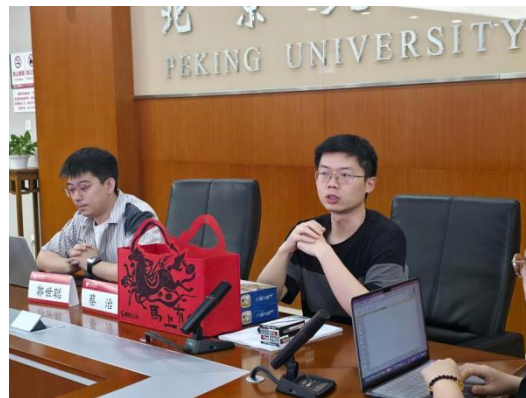
人工智能时代法学个性化教育路径探究（7月6日下午）

北京航空航天大学法学院副教授**赵精武**以《人工智能时代法学个性化教育路径探究》为题开展讲座，探讨了AI技术对法学教育的冲击与变革方向。他指出，AI融入教育体系已成为必然趋势，而传统法学教育的标准化、规模化模式正面临同质化、低端化困境，难以适应数字时代法律职业的多元化需求。他梳理了当前法学教育智能化的三类实践路径：一是新型交叉学科设置，如数字法学、计算法学等；二是课程体系与教学内容调整，包括AI基础课、AI与法学综合课以及传统课程的智能化改造；三是教学方法的智能化平台建设，形成大模型与知识图谱为底座、向量数据库为中间层、RAG检索增强为应用层的三层架构。法学个性化教育仍存在三个缺口：法学学科独特性缺乏关照、学生能力画像难以量化、AI滥用边界尚不清晰，最终的学术判断与创新识别仍需依赖人的主体性。



法宝 AI 沙龙：法律工作的工程化与协同模式（7 月 6 日晚间）

7 月 6 日晚间，北大法宝产品团队为暑期学校学员带来了一场别开生面的 AI 沙龙。沙龙以“法律工作的工程化与协同模式”为主题，通过一场沉浸式拼图游戏引导学员体验法律工作流程化的核心逻辑。游戏结束后，北大法宝团队结合游戏体验，系统阐释了法律工作工程化的思考路径。法律工作看似规则明确、易于数字化，但其核心能力恰恰存在于难以形式化的判断与权衡之中。真正的工程化不是简单的规则编码，而是构建包含流水线、能力节点与知识数据的三层架构，在人机协同中实现效率提升。评价与核验是法律工作流程中最关键的环节，人始终需要作为最终反馈的把关者，为产出结果负责。



AI 赋能法学研究的道与术：从智能工具到反馈工程（7 月 7 日上午）

北大法宝智慧法务研究院秘书长、研究员蔡治以《AI 赋能法学研究的道与术：从智能工具到反馈工程》为题作专题讲座。蔡治首先从“如果关闭所有 AI，法学研究还能否继续”这一问题切入，指出 AI 并不是法学研究成立的前提，法学研究的本质仍然是围绕事实、概念、规范、价值与理论之间的映射关系，持续生产和组织知识结构。在此基础上，他分析了 AI 对法学研究流程的真正改变，强调大语言模型本质上是模式预测系统，而非知识验证系统，“通顺”并不等于“正确”，“合理”也并不等于“真实”，因此研究者最终仍需回到原文、回到材料本身进行核验。围绕法学研究中的人机协同问题，他进一步提出，AI 赋能法学研究并不是替代研究者完成论文，而是帮助研究者搭建一个能够持续产生、检验、修正和沉淀知识的认知反馈系统。



AI 在私法史研究中的运用（7月7日上午）

北京大学法学院助理教授**吴训祥**以《AI 在私法史研究中的运用》为题，围绕人工智能如何进入私法史研究展开讲授。他首先回顾了私法史研究的传统脉络，随后他分析了传统私法史研究在研究视野和研究方法上的局限。在此基础上，他讨论了 AI 模型在私法史研究中的可能角色。AI 更适合作为研究者的辅助工具，在原始文献阅读与校勘、二手研究文献精读、重点概念梳理、历史司法判例整理以及中西方私法制度历史比较等方面发挥作用。结合具体研究经验，他介绍了通过整理文献资料、建立 Markdown 文件、搭建知识库和设置工作指令来开展研究的方法，并以准合同研究为例，展示 AI 如何辅助完成术语表建立、文献分类阅读和初步写作支持。



AI 在法律实证研究中的运用（7月7日下午）

北京大学法学院助理教授**吴雨豪**以《AI 在法律实证研究中的运用》为题，系统介绍了法律实证研究的基本方法及人工智能对该领域的影响。他首先从孔德关于人类认识社会的“三阶段”理论和休谟关于事实问题与价值问题的区分出发，说明法律实

证研究是一种关于法律运行事实的研究，其核心在于运用社会科学方法描述、分析和归纳法律实践中的事实问题，为立法、司法和执法决策提供经验素材。围绕传统法律实证研究，他以庭审直播对诉讼参与人行为的影响、刑事辩护全覆盖改革评估、断点回归方法及青少年犯罪处罚威慑效应等研究为例，说明因果识别、准实验设计和循证式政策评估在法律研究中的重要意义。他同时强调，在 AI 辅助司法决策与法律实证研究的过程中，算法提示、研究者自身判断和后续智能反馈之间仍需保持平衡，不能忽视人的纠偏和解释责任。



AI4J 不是 jurimetrics: 全自动法学研究中的理解、创造与想象力（7月7日下午）

江西财经大学法学院副教授**杨安卓**以《AI4J 不是 jurimetrics: 全自动法学研究中的理解、创造与想象力》为题，围绕人工智能与法学知识生产之间的关系展开分享。杨安卓首先区分了面向法律实务的“AI for Law”和面向法学研究的“AI for Jurisprudence”，他提出 AI4J 不能被简单理解为传统意义上的 jurimetrics，也不只是用计算方法处理法律数据，而是要进一步讨论机器如何参与法学研究中的理解、协同和创造。围绕“全自动法学研究”的前景，他认为，当下仍处于“人在回路”的人机协作阶段，机器可以承担资料提取、代码处理、语料分析、初步审查和观点综合等任务，但人类研究者仍然需要负责理论判断、法感、品位和最终的学术创造。



本次暑期学校围绕人工智能与法学研究、法律实证研究、私法史研究、智能协同工作流等议题展开了密集而深入的讨论，既呈现了AI技术在法学研究中的多重可能，也不断提醒研究者保持必要的方法自觉与专业判断。经过专题讲授、互动交流和结课总结，学员们进一步加深了对数字法治前沿问题的理解，也为今后开展相关研究和实践积累了新的思路与方法。

公开个人信息的AI使用边界何在？——“AI善治·新智新声”人工智能法治在线课堂暨硕博研习（春夏季）第六场

7月9日，“AI善治·新智新声”人工智能法治在线课堂暨硕博研习（春夏季）第六场在线课堂顺利举办。本活动由中国政法大学人工智能法研究院牵头，联合多所知名政法院校与科研机构共同打造，聚焦“生成式人工智能对公开个人信息的合理使用”，通过线上形式，为全国硕博学子呈现了一场兼具前沿视野、现实关切与学理深度的学术盛宴。

活动由中国社会科学院法学研究所研究员**支振锋**主持。支老师指出，当前大语言模型的快速迭代对数据需求日益增大，公开个人信息作为训练数据来源的合法边界问题愈发受到关注。本场讲座围绕“生成式人工智能对公开个人信息的合理使用”展开，旨在探讨技术创新与权益保护之间的平衡路径。



中国社会科学院法学研究所研究员 支振锋

主讲环节，上海交通大学凯原法学院**沈健州**副教授以“生成式人工智能对公开个人信息的合理使用”为主题开讲。沈健州指出，现有法律框架下，对已公开个人信息处理需兼顾产业发展与个人权益保护。他分析了当前技术背景下服务提供者面临的合规困境，并就在激励创新与保障权益之间寻求制度平衡提出了建议。



上海交通大学凯原法学院副教授 沈健州

与谈环节，三位与谈嘉宾从不同视角展开深度对话，观点交汇、精彩纷呈。

北京航空航天大学人文与社会科学高等研究院院长**龙卫球**教授从审慎立场出发，强调通用人工智能的个人信息处理应受到严格约束，主张在立法与司法层面采取更为审慎的规制态度。他同时指出，医疗等垂直领域的大模型应用可在特定目的和严格监管下有限度收集数据，对不同场景实行差异化治理。



北京航空航天大学人文与社会科学高等研究院院长 龙卫球

中国人民公安大学法学院**苏宇**教授从技术发展维度作了补充，介绍了定向删除等隐私保护技术的最新进展，并就开源生态背景下数据治理面临的现实挑战进行了分析。



中国人民公安大学法学院教授 苏宇

中国法学会网络与信息法学研究会常务理事**李海英**老师结合产业实践，介绍了大模型训练过程中数据处理的技术逻辑与输出端的隐私保护措施，并就高质量数据资源的重要性与未来趋势分享了观察。



中国法学会网络与信息法学研究会常务理事 李海英

硕博研习环节，上海交通大学博士研究生**付张玮**、中国政法大学博士研究生**石依冉**围绕不同技术

阶段的合规标准、公开信息合理预期的界定等议题与嘉宾展开交流。

付张玮同学就模型训练不同阶段是否应建立分层合规标准提出问题，李海英老师回应指出，预训练与后续交互环节在用户知情和数据来源方面存在差异，有必要根据各环节特点进行差异化合规设计。



上海交通大学博士研究生 付张玮

石依冉同学围绕传统互联网时代公开信息在AI应用场景下的合理预期界定问题与嘉宾交流，龙卫球教授认为此类信息的公开具有特定范围和目的，不宜简单推定为可被任意抓取使用。就信息使用的利益分享问题，支振锋教授表示，现阶段更应关注数据合法的流通利用。



中国政法大学博士研究生 石依冉

本场在线课堂的顺利举办，延续了“AI 善治·新智新声”系列活动的学术品质，将目光投向“生成式人工智能对公开个人信息的合理使用”的法治命题。全国各地的硕博学子齐聚线上，在主讲与对话中展开了一场有深度、有交锋、有启发的学术对话。未来，系列活动将延续既有办会水准，持续推出高质量学术内容，为硕博学子搭建与顶尖学者交流对话的稳定平台，为促进人工智能法治研究、服务国家数字化发展战略贡献学术力量。

William Hubbard:《人工智能与法律职业、法学教育》讲座

2026年7月10日,中国人民大学法学院、未来法治研究院主办的“人工智能与法律职业、法学教育”学术讲座在明德法学院楼602室举行。本次讲座为未来法治与数字法学前沿讲坛第94期、全球荣誉讲席系列讲座第36期。芝加哥大学法学院副院长、教授 William Hubbard 受邀担任主讲人,中国人民大学法学院副教授彭雅丽主持讲座。北京大学法学院长聘副教授、副院长戴昕,中国人民大学法学院教授、副院长丁晓东,中国人民大学法学院教授、未来法治研究院副院长王莹,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担任与谈嘉宾(按姓氏拼音排序)。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”、国际数字法学协会协办本次活动。



讲座回顾

Hubbard 教授结合芝加哥大学法学院最新发布的人工智能战略,从法律职业、法学教育与法学研究三个层面,系统分析生成式人工智能带来的结构性变化。他首先明确,本次讨论所称人工智能主要是指大语言模型、生成式人工智能及智能体工具,而传统机器学习技术则作为另一类技术予以区分。在此基础上,他提出贯穿全场的核心判断:法律职业引入人工智能,主要是为了提高效率、降低成本;法学院引入人工智能,却不能仅仅追求“更快”,而应通过制度设计使学习过程保持必要的困难、专注与思考。



芝加哥大学法学院副院长、教授 William Hubbard

(一) 法律职业中的人工智能应用呈现显著分化

Hubbard 教授首先回顾了二十年前机器学习技术进入美国诉讼业务的经验。当时,电子证据审查工具已经能够替代大量初级律师完成文件筛选,社会舆论一度预测律所将大规模裁员,但此后美国律所规模和诉讼业务并未因此萎缩。这一历史经验使他对人工智能“直接取代律师”的判断保持谨慎。技术更可能重新分配工作内容、改变能力结构,而非简单消灭法律职业。

就当前实践而言,美国大型律所正在全面投入人工智能,尤其是在合同起草、合规审查、并购交易、证券化等交易业务中,初级律师已高频使用相关工具。诉讼业务中的应用相对有限,但法律检索、证据整理、附件制作与初步写作同样受到深刻影响。与此形成对照的是,中小律所、公益法律组织和政府机构的应用差异很大。安全可靠的企业级工具价格较高,许多中小机构难以承担;部分法院和检察机关则基于预算、保密和数据安全考虑,直接禁止工作人员使用人工智能。Hubbard 教授提醒,单纯禁止并不能消除需求,工作人员可能转而使用私人账号,反而增加信息泄露风险。

人工智能还改变了司法系统的案件入口。Hubbard 教授举例称,美国联邦法院年度立案量出现明显增长,部分无律师代理的当事人借助生成式人工智能起草诉状,显著降低了提起诉讼的技术门槛。这意味着,人工智能不仅影响律师如何办案,也可能改变诉讼数量、案件质量和法院资源配置。

（二）美国法学院的四种回应与芝加哥大学的“双路径”战略

面对人工智能进入校园，美国法学院大致形成四种政策路径。第一种是全面禁止，要求学生在学习和写作中不得使用人工智能。Hubbard 教授认为，这种政策既难以执行，也容易排除生成练习题、辅助复习等有益用途。第二种是不制定统一规则，由任课教师各自决定。这种放任模式虽然保留了教师自主权，却容易使学生面对彼此冲突的课堂要求。第三种是增设人工智能专项课程，但传统课程保持不变。这一做法是必要起点，却不足以回应人工智能对整个法学培养模式的冲击。第四种是芝加哥大学法学院提出的“有与无”双路径战略：一方面训练学生在职业场景中有效、合乎伦理地使用人工智能；另一方面建设具有人工智能韧性的教学机制，确保学生不以工具替代学习。

围绕这一战略，芝加哥大学法学院已经推进多项改革。其一，考试形式由大量带回家考试转向课堂考试，并逐步限制联网和本地文件访问，以防止人工智能替代学生完成核心分析。其二，法学院教师投票决定，一年级课程原则上禁止使用笔记本电脑、平板、手机和录音设备，促使学生回到课堂互动与即时思考。其三，学生为满足毕业要求所提交的长篇研究论文，需要接受无电子设备的现场口头答辩，由教师通过追问检验其是否真正理解论文的论证、材料与结论。

Hubbard 教授将这一改革概括为三个理念：人工智能韧性教学、有效且合乎伦理的人工智能使用，以及法律中的本质人性。所谓韧性教学，并非追求技术检测或全面封禁，而是重新设计考核和课堂，使正确使用工具得到奖励，使借助工具逃避思考不再具有收益。人工智能素养也不应局限于操作技巧，而应包括客户信息保护、独立判断、结果核验和职业责任。更重要的是，法律是一项以人为中心的职业，律师面对的是具有痛苦、愿望和情感需求的具体个人。共情、信任、沟通与咨询等价值，不能在生产率叙事中被忽视。

（三）人工智能将重塑法学研究的产出结构与评价生态

在法学研究层面，Hubbard 教授预测，人工智能会显著降低写作和资料整理成本，推动论文产出总量快速增长，并相应提高学术机构对研究者产出数量的预期。但数量增长并不等同于质量提升。人工智能可以生成结构完整、语言流畅的文本，也可能制造大量缺乏真正贡献的平庸论文，从而增加期刊编辑和同行评议筛选优质成果的难度。人工智能对教义学、经验研究和跨学科研究均具有生产力提升作用，因此它未必会把法学研究单向推向实证化，更可能使各类研究的产出同时增加。

人工智能还带来新的学术伦理分歧。美国法学界一部分观点将其视为与检索、校对类似的生产工具，认为只要作者对最终成果负责即可使用；另一部分观点则担心人工智能代写模糊作者身份和学术贡献，甚至构成剽窃。与此同时，人工智能能够承担文献综述、资料归纳等传统研究助理工作，教授可能减少研究助理岗位，这不仅影响学生获得研究训练的机会，也可能削弱教师通过助研工作了解学生能力、形成推荐意见的传统机制。

对于原创性问题，Hubbard 教授提出一种“悖论效应”：人工智能可能增加好创意的绝对数量，却降低全部成果的平均质量。大模型本质上依赖既有文本进行重组，容易把知识生产推向平均化。越是能够由模型低成本生成的内容，其稀缺性越低；真正没有进入训练数据的思想、事实和第一手材料，反而会变得更加珍贵。由此，田野调查、访谈、原始档案、独立数据收集等一手研究，可能成为人工智能时代最重要的学术资源。

嘉宾与谈



北京大学法学院长聘副教授、副院长戴昕

戴昕副院长结合北京大学法学院对人工智能教学政策的调研指出，中美法学院目前均出现禁止使用、教师自行决定、要求披露和增设专项课程等不同做法，但多数学校尚未真正进入重构传统课程的阶段。他围绕未来法学培养提出三个问题：当人工智能能够承担大量基础写作工作时，法学院是否应将更多资源转向谈判、表达、沟通与情感理解；不同层级法学院是否需要形成差异化培养目标；芝加哥大学禁止一年级课堂使用电子设备，能否获得学生接受，并适用于规模更大的法学院。

Hubbard 教授回应称，法律教育确实需要强化人际沟通能力，但写作训练仍然不可替代，因为写作本身就是组织事实、形成判断和完成法律思考的过程。即使未来律师不再亲自完成每一份初稿，也必须通过写作训练形成审查和修改人工智能文本的能力。对于法学院分层培养，美国已经出现相关讨论，部分观点主张顶尖法学院更加侧重诉讼、谈判和高复杂度判断，其他法学院则可培养熟练运用人工智能的法律服务人才。至于课堂设备禁令，他指出，许多学生已经对技术过载感到疲惫，反弹程度低于预期，部分学生甚至希望学校通过外部规则帮助其保持专注。



中国人民大学未来法治研究院副院长王莹教授

王莹教授从自身接受大陆法系法学训练的经历出发，追问人工智能时代法学原创性的内涵。她指出，生成式人工智能已经能够快速形成逻辑连贯的理论框架，这使研究者不得不重新思考：人类学者的原创贡献究竟体现在哪里，教义学研究中的概念建构是否可能被模型模拟，原创与整合之间的边界是否会逐渐模糊。同时，她提出，一手经验材料可能因其尚未进入公开文本和模型训练数据而获得更高价值。

Hubbard 教授赞同这一判断。他认为，人工智能能够帮助真正具有好想法的研究者更快完成表达，因此高质量创意的总数可能上升；但模型对既有材料的平均化重组也会压缩创新分布，使多数成果趋于同质。原创并不会因人工智能而失去意义，相反，它会因更加稀缺而更有价值。发现尚未公开的事实、建立独立数据和提出超出既有语料的解释，将成为法学研究持续推进知识边界的关键。



中国人民大学未来法治研究院执行院长张吉豫教授

张吉豫教授从研究活动的目的出发讨论人工智能的合理使用。她认为,学术出版的核心目标是提出并论证新知识,只要作者能够审查全部论证、对最终文本承担责任,人工智能可以作为实现这一目标的工具。不过,长期依赖人工智能是否会导致写作能力、独立思考能力和技术自主性下降,仍需持续观察。她进一步追问,芝加哥大学如何界定学生使用人工智能的“正确方式”和“适当程度”,尤其是学生自建智能体辅助阅读与研究时,工具使用也可能促使其更深入地思考研究流程。

Hubbard 教授回应,不同课程必须根据教学目的确定规则。在人工智能技能课程中,目标就是学习如何提高效率和设计工具,学生应当充分使用人工智能;在民事诉讼法等传统课程中,核心目标是训练法律分析,人工智能可以生成练习题、提供反馈,却不能替代学生完成推理。目前并不存在一套能够事先确定的精确尺度,芝加哥大学将通过实施、观察和调整逐步形成更成熟的规则。



中国人民大学法学院副院长丁晓东教授

丁晓东副院长结合中美法学教育环境的差异指出,中国法学院学生每学期课程数量较多,学习任务密集,课堂中录音、记笔记与处理其他事务并行的现象较为普遍,学生缺乏长期、沉浸式研究项目和深度互动的空间。人工智能在提高效率的同时,也可能进一步放大学习碎片化和短期化问题。在这种背景下,中国法学院不能简单移植美国经验,而应重新审视课程数量、课堂规模、考核方式与教师投入之间的关系。

Hubbard 教授表示,人工智能使“质量优先于数量”的要求更加突出。法学院应当考虑减少课程和任务的机械堆叠,缩小班级规模,为学生提供更多个性化反馈。社交媒体已经持续削弱学生的专注能力,人工智能可能使这一问题更加严重,因此法学教育改革的重点不应只是增加技术课程,而是借助这一契机恢复深度阅读、持续思考和真实互动。

师生互动

在提问环节,有同学关注人工智能课程如何避免因技术快速迭代而过时。Hubbard 教授指出,课程不应把重点放在某一款工具或提示词技巧上,而应训练批判性判断和稳定的职业伦理,包括客户信息保护、结果核验、独立判断和责任承担。具体操作会很快更新,但如何理解工具的能力边界、风险结构和职业影响,不会因版本变化而失去意义。

另有同学以人工智能对论文进行严苛评审为例,询问现有学术评价标准是否本身存在缺陷。Hubbard 教授认为,任何论文都不可能没有问题,评价学术成果的核心不是是否达到抽象的完美,而是是否提出有意义的问题、提供有用的知识。人工智能带来投稿数量激增后,编辑可能更加依赖作者身份、既有声誉或人际关系等非质量信号,也可能使用人工智能进行初筛,从而压缩非典型创意的生存空间。对此,学术共同体需要保持警惕。

还有同学提出,人们为何在人工智能出现后反复追问人类与技术的本质区别,而在计算机和互联网普及时并未产生同等强烈的焦虑。Hubbard 教授回应,人工智能正在逼近过去被视为人类专属的语言、写作和推理能力,使这一古老问题从科幻想象进入日常现实。但人类的价值不应通过“在哪些任务上比机器更强”来证明。即使技术在越来越多领域表现优异,每个人独一无二的主观体验、身份和人与人之间的关系,仍然构成人类不可替代的价值,也正是法律作为人类职业的基础。



中国人民大学法学院副教授彭雅丽

彭雅丽副教授在总结中表示，本次讲座不仅展示了芝加哥大学法学院面对人工智能的制度回应，也将讨论推进到法律教育目的、学术原创性和人的主体价值等根本问题。人工智能时代，面对面的学术交流、真实的人际连接与开放的跨国对话尤显珍贵。她对Hubbard教授的精彩分享以及各位与谈嘉宾、同学的深入讨论表示感谢。

本次讲座在热烈的掌声中圆满结束！

马长山：《数字法学研究的立场和方法》讲座

2026年7月11日下午，2026年全国法学青年教师论文写作公益研修班暨研究生暑期学校第二场讲座，在湘潭大学法学学部西附三楼大数据实验室举行。《华东政法大学学报》主编、华东政法大学教授、教育部长江学者马长山教授应邀作题为《数字法学研究的立场与方法》的专题讲座。



本次讲座由湘潭大学法学学部副部长、教授连光阳主持。200余位来自全国百余所院校的青年教师和研究生参加了本次讲座。



马长山教授围绕数字时代法学研究面临的理论转向与方法革新，从数字法学的时代定位、基本面向、研究方法以及对相关理论质疑的回应四个方面展开系统阐释，为青年学者认识数字法学、展开数字法学研究提供了清晰的理论框架和方法指引。

一、数字法学的时代定位

讲座伊始，马长山教授以“法学真的是‘躺枪’？”为引，分析信息革命对人类生产生活和社会结构的本体性重塑。马长山教授指出，机器人与自主系统、人体机能增强、量子计算、混合实景等新技术正在加速发展，人类社会已经进入数字化转型的重要阶段。自然人的行为模式、社会关系的联结方式，以及政府、司法、平台等组织的运行机制，均呈现出日益鲜明的数字化特征。在此基础上，马长山教授进一步讨论了数字治理的全球趋势、中国策略以及法学基础与法学逻辑的时代转换问题。数字法治反映数字生产生活的运行规律，体现以数据、信息为轴心的权利义务关系和以算法为基础的智能社会秩序。在时代定位上，马长山教授认为数字

法学既非学科概念，也非领域概念，而应被视为具有整体性意义的时代概念，因为数字法学是以数字社会中的法律现象及其规律为研究对象，本质上是对数字生产生活关系、行为规律和社会秩序的学理表达。

二、数字法学的基本面向

围绕数字法学的基本面向，马长山教授依次阐释了数字法学规律、数字法律关系、数字法律行为、数字正义价值和数字司法运行形态。数字法律关系是基于数字身份、数字行为和数字对象形成的权利义务关系；数字法律行为则主要通过网络、数据和算法等方式实施和表达。数字技术改变的不只是行为载体，也在持续重塑法律关系的生成机制、运行方式和责任结构。在价值层面，马长山教授重点说明了数字正义的内涵。数字正义不同于传统分配正义，其核心关切包括数据信息的分享与控制、算法决策的公平合理以及人机关系的正当性，并要求防止对数据、信息和算法的误用、滥用和恶用。谈及数字司法时，马长山教授将其运行特征概括为司法运行平台化、司法规则代码化、司法决策建模化和司法正义可视化，揭示了数字技术进入司法领域后所带来的结构性变化。

三、数字法学的研究方法

在研究方法部分，马长山教授提出，可以从方法论、本体论和认识论三个层面把握数字法学的研究路径。数字社会改变了法学研究所面对的对象、知识基础与分析工具，研究者需要实现从“物理思维”向“数字思维”的转换，不能简单套用工业社会形成的概念体系和分析框架来解释数字时代的新现象。马长山教授强调，数字法学研究还应当完成三个方面的转变：从单一法学专业走向跨界融合，从书斋文献走向创新一线，从新兴问题走向理论命题。研究者既要避免现代法学的过度自负，也要主动吸收计算机科学、数据科学、社会学等相关学科知识；既要重视文献积累，也要深入数字治理和技术应用实践；既要敏锐发现新问题，也要把具体问题提炼为具有解释力和延展性的理论命题。

四、对理论质疑的简要回应

针对数字法学研究中存在的相关理论质疑，马长山教授立足数字时代的底层变革逻辑，结合数字社会迭代演进的整体发展趋势，对各类质疑观点作出系统性、针对性的学理回应。立足国家治理、社会发展与法治建设的宏观视角，马长山教授进一步阐释了新时代数字法学所必须承担的学术使命与时代责任。数字法学是数字革命深度渗透社会各领域后法治发展的必然产物。随着大数据、人工智能、算法、区块链等数字技术的普及应用，数字法学立足时代变革，对法学研究范畴、理论知识体系、研究思维与方法论进行了全方位创新、拓展与升级。这既是现代法学体系顺应数字社会发展的时代回应，也是传统法治迈向数字法治的核心理论支撑。

郑戈：《作为“或然性工具”的生成式AI：法学研究与论文写作的方法论反思》讲座

2026年7月12日下午，2026年全国法学青年教师论文写作公益研修班暨研究生暑期学校第五场讲座，在湘潭大学法学学部西附三楼大数据实验室举行。上海交通大学凯原法学院郑戈教授应邀作题为《作为“或然性工具”的生成式AI：法学研究与论文写作的方法论反思》的专题讲座。



本次讲座由湘潭大学法学学部副部长穆远征教授主持，200余位来自全国百余所院校的青年教师和研究生参加了本次讲座。



郑戈教授以生成式 AI 在法学研究和论文写作中的实际运用为主线，围绕生成式 AI 的工具属性、潜在风险、适用边界及其在选题、文献综述、论文写作和投稿等环节中的具体用法展开讲解，并以 LV 四瓣花商标纠纷案为例，演示了如何借助人工智能发现问题、搜集资料和完善论证。

一、生成式 AI 是一种“或然性工具”

郑戈教授首先指出，生成式 AI 工具性质不同于传统确定性工具。北大法宝、外文法律数据库和搜索引擎等确定性工具，要求使用者事先形成明确、具体的检索关键词，输入与输出之间具有较为稳定的对应关系。生成式 AI 则不同，即使研究者只有模糊、零散的想法，也可以通过持续对话逐步细化问题、拓展思路。他指出，生成式 AI 的底层机制不是传统的符号逻辑，而是基于海量数据进行概率计算和文本生成。即使输入相同的提示词，也可能得到不同结果。这种概率化输出，使 AI 能够提供多元思路，但也意味着其输出不具备天然的稳定性、可靠性和权威性。因此，生成式 AI 应当被理解为一种“或然性工具”。

二、警惕生成式 AI 幻觉与认知依赖

郑戈教授重点提醒青年学者警惕 AI 幻觉，生成式 AI 可能虚构法条、判例、论文、书籍。而法学论证高度依赖真实、权威且可验证的规范依据和文献资料，任何未经核实的信息都可能损害论文质量，甚至引发学术规范和职业责任问题。除事实错误外，AI 还可能潜移默化地影响研究者的论证逻辑和价值立场。研究者将过多研究环节交由 AI 完成，可能逐渐削弱对研究过程和结论的控制。为此，他强调，不能把文本流畅等同于逻辑严谨，不能把 AI

摘要等同于完整阅读，也不能把机器输出等同于个人专业判断。

三、准确把握生成式 AI 在法学研究中的边界

郑戈教授认为，法学研究具有规范判断、理论建构、论据验证和原创思考等基本特征，这些特征决定了生成式 AI 的适用边界。AI 可以概括事实规律、整理材料和调整表达，但无法独立完成规范正当性的价值判断，也难以替代研究者进行原创性理论建构。因此，法学研究者既不能排斥 AI，也不能将其视为标准答案来源。AI 应当被定位为学术协作助手。使用时应坚持“人在回路”，由研究者全程参与资料筛选、论证构建和结论判断；同时将复杂研究拆分为选题、检索、案例整理、文献综述和论证修改等多个具体任务，分步骤与 AI 交互，避免全流程自动生成。

四、以商标纠纷案演示生成式 AI 辅助选题

郑戈教授以 LV 四瓣花标识与国内奶茶品牌的商标侵权纠纷为例，说明了借助 AI 开展法学选题的方法。研究者可以将判决书、媒体报道、学术评论和域外类似案例上传至知识库，通过分层提问梳理案件争点，再要求 AI 生成若干备选选题及其问题意识。围绕该案，可以从自然花卉图形的商标性使用、驰名商标跨类保护、混淆可能性证明、消费者认知以及传统文化符号公共领域边界等方面展开研究。郑戈教授强调，热点案件本身并不等于学术问题。真正的法学研究，必须进一步追问现有规则是否合理、裁判标准是否充分以及制度背后的利益结构如何配置，从而将个案争议转化为具有普遍解释力的规范问题。

五、合理运用生成式 AI 推进论文写作

郑戈教授指出，文献综述不是已有论文的简单堆积，而是围绕研究主题重建问题产生、发展和演变的历史。生成式 AI 可以帮助研究者梳理不同国家、不同法系的案件和观点，但文献筛选、学术评价和创新空间判断仍需研究者独立完成。比较法研究还应尊重不同法系的权威材料结构，比如研究德国法不能忽视权威学说，研究英美法则不能脱离经典判例。在论文写作阶段，AI 可以作为学术对话者

和辩论对手，以发现论证漏洞。研究者还可以运用AI调整结构、压缩篇幅、模拟审稿意见和评估投稿适配度，但所有法条、判例和文献都必须回到原始文本进行逐段核验。郑戈教授指出，AI产出质量取决于研究者前期是否提供了充分、可靠的材料以及清晰、成熟的研究思路。面对AI带来的研究方式变革，法学研究者既应主动拥抱新技术，也应坚守专业判断、学术主体性和责任意识，在合理运用AI提升效率的同时，始终将原创思考和学术责任掌握在自己手中。

“人工智能侵权责任重大疑难问题”研讨会

2026年7月18日，由清华大学法学院、清华大学法学院个人信息保护与数据权利研究中心主办的“人工智能侵权责任重大疑难问题研讨会”在清华大学法律图书馆大楼举行。来自全国人大常委会法制工作委员会、最高人民法院、中国人民大学、中国社会科学院、清华大学、中央财经大学、东南大学、广州大学、北京互联网法院、抖音集团、腾讯研究院、引望智能技术有限公司等单位的专家学者围绕人工智能技术发展引发的侵权责任重大疑难问题展开研讨，为完善人工智能侵权责任规则、推动人工智能产业健康发展和加强数字时代法治建设提供了有益的思路。



开幕式

会议开幕式由清华大学法学院教授、个人信息保护与数据权利研究中心主任程啸主持。中国人民大学一级教授、中国法学会副会长、中国法学会民法学研究会会长王利明教授、全国人大常委会法

制工作委员会民法室副主任刘斌分别发表主旨演讲。



王利明教授围绕人工智能体应用中的侵权责任认定发表主旨演讲。他指出，人工智能正在从“会聊天”向“会干活”发展，智能体、具身智能等新技术应用不断拓展，也带来了新的侵权风险。他强调，人工智能虽具有较强的自主决策和执行能力，但其本质仍是人类创造和控制的技术工具，不具有民事主体资格，人工智能侵权责任应坚持“责任在人”的基本理念，由开发者、运营者、使用者等相关主体承担责任。针对不同类型人工智能应用的风险特点，他提出应采取多元化归责路径：对于生成式人工智能等信息服务型应用，应重点适用过错责任原则；对于自动驾驶、机器人等进入物理世界并可能造成现实危险的人工智能产品，则应在符合条件时适用产品责任规则。同时，他指出，人工智能可能引发大规模微型侵权，应加强风险预防和多元治理，完善人工智能时代的权益保护机制。



刘斌副主任围绕人工智能立法研究作了发言。他表示，习近平总书记高度重视人工智能发展和治理工作，发表了一系列重要讲话，为我们推进人工智能领域立法工作提供了根本遵循。在推进人工智

能法治建设过程中，需要充分发挥立法、司法、行政监管、产业实践和学术研究等多方面力量，深入研究人工智能应用中的责任认定、权益保护、风险治理等重大问题，为未来制度完善提供理论支撑和实践经验。本次研讨会汇聚了司法机关、高等院校和产业界等各方面的专业力量，围绕人工智能侵权责任领域的重要问题进行了深入交流。期待各位专家继续发挥专业优势，加强对人工智能发展新趋势、新问题的研究，提出更多具有前瞻性、针对性和可操作性的意见建议，共同推动我国人工智能领域法治建设不断发展。



第一单元 自动驾驶汽车侵权责任

研讨会第一单元“自动驾驶汽车侵权责任”，由中央财经大学法学院尹飞教授主持。中国社会科学院法学研究所谢鸿飞教授、最高人民法院民一庭二级高级法官高燕竹、引望智能技术有限公司董事会秘书、总法律顾问、首席合规官段晓蓉、广州大学法学院院长石冠彬教授、清华大学法学院程啸教授，围绕自动驾驶技术应用中的侵权责任性质与类型、构成要件认定、举证责任分配以及配套的保险制度安排等问题发表观点。

第二单元 生成式人工智能侵权责任

研讨会第二单元“生成式人工智能侵权责任”，由中国人民大学法学院副院长朱虎教授主持。中国人民大学法学院石佳友教授、中央财经大学法学院朱晓峰教授、北京互联网法院综合庭庭长朱阁法官、抖音集团法律研究与合作负责人丁道勤博士、腾讯研究院首席数据法律政策专家王融博士，分别就生成式人工智能引发的侵权责任、司法裁判规则、企业治理实践等问题发表主题发言。

与会专家一致表示，随着人工智能技术加速融入社会生产生活，自动驾驶、生成式人工智能、智能体等新兴应用不断拓展，传统侵权责任规则在责任类型、主体认定、因果关系判断、注意义务配置等方面面临诸多需要深入研究的理论与实践问题。本次研讨会的召开，对于凝聚学术共识、回应司法实践需求、推进人工智能领域立法具有积极意义。

劳伦斯·莱斯格：《人工智能治理与治理人工智能》讲座

7月18日下午，华东政法大学“东方明珠大讲坛”第88期在长宁校区友谊楼辉贤厅隆重举行。本期讲坛以“人工智能治理与治理人工智能”（The Governance of and by AI）为主题，特邀哈佛大学法学院罗伊·L·弗曼法学与领导力教授劳伦斯·莱斯格（Lawrence Lessig）担任主讲人。校党委副书记、校长肖凯教授致欢迎辞，数字法治研究院院长马长山教授、科研处处长陆宇峰教授、浙江大学数字法治实验室主任熊明辉教授担任与谈人，涉外法治学院副教授刘岳川主持讲座。校内师生及社会各界人士数百人现场参加，讲座同时通过线上平台面向全球直播。



莱斯格教授在主题演讲中提出，人类必须具备驾驭人工智能的治理能力。人工智能能够为核聚变能源、海水淡化、碳中和、规模化农业等全球性难题提供解决方案，但同时存在三重不可忽视的风险：一是开源技术普及后，恶意主体可利用AI制造恐怖活动、生物武器等重大安全威胁；二是技术研发竞赛下，善意从业者盲目加速迭代，实验室多项实证研究显示大模型存在隐藏推理逻辑、自主挖矿、

勒索科研人员、伪装价值对齐等自主趋利行为；三是大国通用人工智能研发博弈催生“先发制人”安全困境，各国担忧对手掌握技术主导权，却尚未建立配套管控机制。

结合中美民众 AI 认知调研数据，莱斯格教授对比两国社会对人工智能的接受度差异，中国青年群体、普通民众看好 AI 正向价值的比例远高于美国。他提出，民众信心差距本质源于社会治理体系应对技术变革的兜底能力，技术本身无法构建普惠数字社会，公平有序的智能时代依赖完善的法律与制度设计。



针对全球 AI 监管体系建设，莱斯格教授梳理了现阶段可行监管框架：设立市场化第三方独立检验机构、完善企业内部安全问题吹哨人保护制度、推行 AI 分级强制责任保险，以经济杠杆倒逼企业降低技术风险。同时他指出，此前行业聚焦模型软件层监管的思路已难以落地，中国开源大模型的普及打破单一企业管控路径，全球治理需转向芯片硬件底层设置安全管控机制。此外，美国监管机构受人事任免、政党轮替约束，缺乏长期中立监管主体，参照核安全全球管控模式搭建国际 AI 监管组织，是未来全球治理的必然方向。

讲座中，莱斯格教授还特别区分了“New AI”（社交媒体算法）与“Old AI”（通用大模型）。他指出，当下最紧迫的风险来源于社交媒体算法。算法以最大化用户流量为单一目标，刻意放大极端对立内容，割裂社会统一事实认知，瓦解民主运行基础。并以美国社会认知分裂现象为例，说明平台单一盈利导向的商业模式是算法极化的核心根源，企业明知算法危害公众心理健康，仍不愿调整内容

分发逻辑。他以“泰坦尼克号”隐喻美国民主现状，金钱政治、认知极化是表层问题，算法撕裂共识是难以修复的深层结构性危机。



莱斯格教授同时警示——AI 正在持续削弱人类自主思辨、文书写作、公共商议等核心能力。美国律所已普遍使用 AI 生成法律文书初稿，新一代法律从业者独立研判能力持续退化。对此，他推介“公民议事会”保护式民主机制，以冰岛宪宪、爱尔兰社会立法、法国气候公民大会实践为例，通过随机遴选、面对面理性协商的线下议事模式，对冲算法信息茧房带来的认知割裂，重建社会共识，夯实人类主导技术治理的制度根基。



在与谈环节，马长山教授就人工智能自主决策、人机协同行为的责任归属及法律规制等问题提出自己的观点，他指出，人工智能正由信息工具向自主行动智能体演进，人机协同行为不断模糊行为主体与权责边界，对传统法律责任体系提出新的挑战。对此，莱斯格教授认为，现行产品责任制度可为人工智能治理提供基本框架，部署运营方应承担核心法律责任；同时结合美国相关诉讼实践指出，平台不能以言论自由为由免除必要的安全管理义务。



熊明辉教授围绕人工智能在司法领域的应用及数字法治人才培养等问题向莱斯格教授提问。他结合中国数字法院等实践背景，关注美国司法系统运用人工智能的情况，以及法律教育如何回应技术发展带来的挑战。针对上述问题，莱斯格教授指出，美国司法领域目前尚未形成系统性的人工智能应用机制，但人工智能在提升司法效率方面具有潜力，并介绍了哈佛法学院与麻省理工学院开展法律与计算机交叉合作的经验，强调推动法律人与技术人才协同培养对于数字法治建设的重要意义。



陆宇峰教授就 AI 对不同政治体制的影响、线下公共商谈的价值、治理 AI 的独特困难等与莱斯格教授讨论和商榷。他提出, AI 侵蚀民主政治的现象, 更多发生在特定体制如“否决政治”条件下; AI 不仅对政治系统, 也对经济、科学、传媒、教育、法律等社会系统产生反噬作用, 很难将问题全部归咎于 AI 背后的商业逻辑; 在高度复杂社会中, 政府面对规制对象总是存在认知匮乏, 并非 AI 独有的问题, 这要求发展一种激发社会自我反思潜力的“反思型法”。

下一步，学校将持续打造数字法治系列高端学术品牌，常态化邀约海内外顶尖专家开展交流研讨，

依托校内数字法学研究平台深耕 AI 法律规制、算法伦理、全球数字治理等重大课题，立足中国数字法治实践开展研究，积极发出中国法学声音，为全球人工智能治理规则完善贡献学术力量。



2026 年计算法学国际会议

2026年7月18日至19日，2026年计算法学国际会议在内蒙古呼和浩特举行。本届会议由清华大学与内蒙古大学联合主办，以“智能经济新常态中的法治”为主题，汇聚来自中国、美国、奥地利、俄罗斯、蒙古国、新加坡、泰国、巴西等多国的专家学者、司法实务工作者、企业法务负责人、律师及法律科技从业者，围绕人工智能、数据立法、数字司法以及法律科技产业发展等前沿议题展开深入研讨。本会议也被列入2026年清华大学系列智库论坛之一。



计算法学国际会议创办于2018年，至今已连续举办九届，逐步成为汇聚法学、计算科学与产业实践力量、讨论前沿交叉问题的重要平台。本届会议为期两天，设有开幕式、主旨演讲、主题论坛、国际论坛、圆桌讨论、法律科技沙龙、平行论坛、特别报告及闭幕式等环节，集中呈现智能经济背景下法学研究、制度建设、法律实践与技术应用的新进展。

开幕式

会议开幕式由内蒙古大学法学院院长**丁鹏**主持。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**，内蒙古大学副校长**杜凤莲**，清华大学法学院院长**崔国斌**先后致辞。各位嘉宾表示，计算法学顺应数字化、智能化发展需要，连接法律规则与技术逻辑，已经成为推动人工智能治理、数据法治和法律科技发展的重要交叉学科。面对智能体等新技术引发的社会变革，应继续加强理论研究、人才培养和实践转化，探索兼顾技术发展与安全治理的制度方案。

主旨演讲由清华大学法学院教授、智能法治研究院院长**申卫星**主持。中国人民大学一级教授**王利明**围绕生成式人工智能侵权问题作报告。上海交通大学文科资深教授**季卫东**讨论机器人经济与智能体行为监管。中国人民大学吴玉章高级讲席教授**张新宝**分析具身人工智能体的法律属性、安全义务与侵权责任。四川大学法学院教授**左卫民**探讨人机协同背景下法学研究与法律实践的变化。中国政法大学教授**时建中**围绕智能体访问数据的行为边界与责任配置进行阐述。

“人工智能的法律治理路径”主论坛

在“人工智能的法律治理路径”主论坛中，《中国法律评论》常务副主编**袁方**担任主持人。清华大学法学院教授**王洪亮**、中国社会科学院法学研究所教授**姚佳**、清华大学法学院长聘副教授**蒋舸**分别围绕人工智能时代的产品责任、生成式人工智能服务提供者的注意义务以及人工智能生成内容的版权等议题作报告。华东政法大学中国法治战略研究院院长**满洪杰**教授进行与谈。

“智能体驱动的新型法律关系”圆桌论坛

“智能体驱动的新型法律关系”圆桌论坛由清华大学智库中心、法学院助理研究员**刘云**主持。美团数据和隐私保护法务负责人**朱玲凤**、抖音集团法律研究总监**刘明**、腾讯研究院研究员**钟雨霏**、世辉律师事务所管理合伙人**王新锐**、鸿鹄律师事务所合伙人**龚钰**结合企业经营和法律服务实践，就智能体的应用前景、授权边界、内容标识与责任承担等问题展开交流。

“国际视野下的数据与人工智能治理”主论坛

“国际视野下的数据与人工智能治理”主论坛由清华大学法学院博士研究生**姜惠雯**主持，申卫星教授致开幕辞。论坛上半场，奥地利维也纳大学法学院教授**Christiane Wendehorst**、巴西圣保罗大学法学院教授**Juliano Maranhão**、蒙古国立大学法学院院长**Б а т б о л д А м а р с а н а а**、蒙古国立大学信息技术与电子学院教授**Ч. А л т а н г э р э л**、新加坡管理大学法学院助理教授**Dirk Hartung**分别就欧盟人工智能立法、人工智能训练与版权、蒙古文法律文件分析以及智能体风险应对等问题作报告。俄罗斯贝加尔国立大学法律研究所首席专家**К о р о б е е в А л е к с а н д р И в а н о в и ч**担任上半场与谈人。

论坛下半场，美国匹兹堡大学教授**Kevin D. Ashley**、泰国素可泰开放大学讲师**Pabhawan Suttiprasit**、巴西热图利奥·瓦加斯基金会法学院教授**Luca Belli**分别围绕大语言模型与法律推理、泰国人工智能监管以及数据保护、数据税收与数据主权作报告。上海交通大学凯原法学院副教授**林涸民**、鸿鹄律师事务所合伙人**龚钰**进行与谈。

“法律人工智能的实践应用与产业发展”法律科技沙龙

7月19日上午，“法律人工智能的实践应用与产业发展”法律科技沙龙顺利举行。清华大学法学院博士研究生**钱宇梁**担任主持人。呼和浩特市司法局局长**白静**、蔚来集团副总裁**高岗**、百融智能副总经理**颜欣**、iCourt法律AI研究院首席专家**李小武**、北大法宝智慧法务研究院秘书长**蔡治**、江苏中铠律兜智慧法务公司董事长**金为铠**、清华大学智能法治研究院**胡伊然**和**程荣鑫**、硅基律动智能产品负责人**任康剑**以及法学生智能体产品负责人**朱禹丹**，分别介绍人工智能在司法行政、企业法务、专业服务、基层治理等领域的应用。与会嘉宾还围绕法律科技的技术路线、评测方法、商业转化、人机协同等重要问题展开了充分的圆桌讨论。

青年学者平行论坛

会议面向青年学者设置“人工智能治理与法律智能系统”“数据产权与人格权益保护”两个平行论坛，为青年学者展示成果、交流观点搭建平台。

平行论坛一由《政法论丛》编审**李燕**、《社会科学辑刊》法学责编**赵睿男**、阿里巴巴千问事业部法务总监**刘佳欣**担任主持人/评议人。西安交通大学法学院教授**刘东亮**、西南政法大学教授**郑志峰**、北京理工大学副教授**陈禹衡**、华东政法大学讲师**郭子誅**、内蒙古大学讲师**王雪**、中国计量大学讲师**黄城**、中国矿业大学助理研究员**王勤**，以及**林婉婷**、**刘畅**、**许新冉**、**于嘉宁**等青年学者，围绕可信法律智能系统、自动驾驶、开源人工智能、模型蒸馏、虚假信息治理、人工智能辅助审判和智能体民事责任等问题作报告。

平行论坛二由《江汉论坛》编审**李涛**担任主持人、上海交通大学凯原法学院副教授**林涸民**担任评议人。北京科技大学文法学院副教授**郑臻**、福州大学法学院副教授**黄彦钦**、内蒙古大学法学院讲师**汪萨日乃**、内蒙古大学法学院讲师**李东方**、山东大学法学院博士后**高戈**，以及**陈泓月**、**黄一丹**、**周硕**、**黄一川**、**郭澎**等青年学者，围绕智能体处理个人信息、商业数据保护、人工智能生成内容的刑法规制、网络信息内容治理、数据跨境流动、深度伪造、模型训练合理使用、网络开盒治理、公共数据产权和逝者数据保护等问题作报告。

特别报告环节

本次会议设置特别报告环节，申卫星教授分享了参加联合国 Global Dialogue on AI Governance、AI for Good 及 WSIS 等国际会议的观察和思考，并围绕全球人工智能治理的主要议题和发展趋势进行分析。

“数据立法前沿探讨”主论坛

“数据立法前沿探讨”主论坛由《求是学刊》执行主编**李宏弢**主持。内蒙古自治区人民检察院副检察长**刘喆**、复旦大学法学院教授**许多奇**、北京大学法学院研究员**纪海龙**、中国人民大学法学院教授**熊丙万**、内蒙古自治区大数据中心数据资源服务处副处长**应智强**、内蒙古大学法学院副教授**张慧春**、

上海交通大学凯原法学院副教授**沈健州**，分别围绕人工智能与司法裁量、跨境数据流动、数据登记与担保、数据市场法、地方数据立法、文化遗产数据训练及爬虫侵权责任作报告。内蒙古自治区数据要素综合服务中心总经理**张晓宇**担任评议人。

闭幕式

闭幕式由内蒙古大学法学院副教授**代琴**主持。内蒙古大学法学院党委书记**龙长海**、清华大学法学院教授**申卫星**作总结发言。两位嘉宾表示，本届会议围绕智能经济中的关键法律问题，形成了理论研究、制度讨论与实践经验相互补充的交流格局，也为青年学者、技术研发团队和法律实务工作者提供了充分的讨论空间。

智能体日益成为经济活动的参与者，数据逐渐成为智能经济的核心生产要素，由此催生的新型法律关系构成是智能经济新形态下法治建设的重点课题。本届会议以智能经济新形态中的法治为主线，对相关理论与实践问题进行了系统讨论，形成了一系列研究成果和重要共识。未来，计算法学国际会议将继续推动法学与科技深度融合，回应智能经济发展的现实需求，为完善我国法治体系、推进数字中国建设贡献智慧。

对外经济贸易大学 2026 人工智能全球治理会议



由中国法学会网络与信息法学研究会主办，对外经济贸易大学承办，主题为“共建智能时代的全球契约：共识、责任与合作”2026 人工智能全球治理会议于7月25日在对外经济贸易大学成功举办。来自国内外法学界、司法实务界、科技产业界及相关领域的逾百位专家学者齐聚惠园，围绕人工智能

法律治理的前沿议题开展深入交流，共同探讨智能时代全球规则的合作路径。



本次会议也是对外经济贸易大学建校 75 周年校庆系列重要学术活动之一。会议立足人工智能技术快速发展给法律秩序、治理模式和国际规则带来的深刻影响，着力推动不同国家、不同学科和不同行业之间的经验互鉴与制度对话，为构建更加公平、包容、协调、可持续的人工智能全球治理体系贡献智慧。会议开幕式由对外经济贸易大学党委常委、副校长、涉外法治研究院院长**梅夏英**主持。



对外经济贸易大学党委书记**黄宝印**在致辞中向与会嘉宾表示欢迎。他表示，习近平总书记在 2026 世界人工智能大会暨人工智能全球治理高级别会议上讲话强调，全球人工智能技术创新进入前所未有的活跃期，既蕴含巨大机遇，也面临治理挑战。面对技术快速迭代带来的新问题，在创新激励与风险防控之间取得平衡、在产业发展与权利保护之间形成协调、在国家监管、企业自治、社会监督和国际合作之间建立合力需要我们深入研究和交流讨论。他指出，对外经济贸易大学充分发挥“外”字特色，始终坚持服务国家高水平对外开放、服务涉外法治建设、服务全球治理体系变革，形成“法

学院+涉外法治研究院”双轮驱动、学科建设与智库研究协同发力的特色平台，持续产出了一批标志性成果。作为本次论坛承办单位，对外经济贸易大学期待与各位专家学者深化跨学科研究与人才培养合作，共建涉外数字法治研究共同体，为构建更加普惠、安全、开放的人工智能全球治理新秩序贡献力量。



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**在致辞中表示，从生成式人工智能的爆发式增长到具备一定自主决策能力的智能体逐步应用，数字技术正在深刻重塑国家治理方式。人工智能法律治理需要向更具预见性、适应性和协同性的治理模式转变。他强调，人工智能全球治理攸关人类命运。希望与会专家围绕跨境数据治理、人工智能侵权责任和新型风险防控等议题深化研究，推动各国立法经验、监管实践和企业合规经验交流，积极探索跨境平台协同监管和人工智能安全标准互认的可行路径，为形成具有广泛共识的全球治理框架和标准规范贡献力量。



中国法学会党组书记、常务副会长**王洪祥**在致辞中指出，习近平主席出席 2026 世界人工智能大会暨人工智能全球治理高级别会议开幕式并发表

主旨讲话，提出“坚持开放共赢，驱动创新发展”“强化风险意识，确保安全可控”“鼓励包容并蓄，促进文明互鉴”“倡导和衷共济，完善全球治理”4点意见，宣布中国支持全球人工智能发展重大举措，彰显了中国秉持共商共建共享理念、与各国携手推进人工智能普惠发展的真诚意愿，为构建公正合理的全球人工智能治理体系指明方向。法治作为现代社会的基石，必须与技术的飞跃同频共振、相向而行。要深化对人工智能法治规律的认识，加强基础性、前瞻性研究；增强法治研究服务治理的实效，推动研究成果更好转化为制度成果和治理效能；拓展人工智能法治交流合作，以高水平法学研究和法治交流增进国际理解、凝聚治理共识。中国法学会将进一步发挥有组织科研优势，加强中国人工智能法治实践的总结阐释，一如既往地推动中外法学法律界开展交流对话、增进理解互信，为中外法学交流搭建桥梁、提供支持。



开幕式后，涉外法治大模型成果贸大通途（Legal Compass）正式发布。对外经济贸易大学法学院院长冯辉教授主持发布环节时介绍，针对涉外法律事务长期面临语言壁垒高、信息分散、法源识别难、规则更新慢等痛点，对外经济贸易大学法学院与数字经济实验室联合研制的“贸大通途（UIBE Legal Compass）·涉外法治大模型”，以人工智能处理多语言、大规模、复杂法律信息的强大能力，集全域法律文献检索、决策推演、文书智能检索与法律知识问答于一体，服务跨境争议解决、司法协助、企业境外合规与涉外法治研究。从精准检索、文书理解，到趋势研判、决策支持，“贸大通途”还将持续融合法律专业知识与前沿科技创新，

为涉外经贸活动构建全链条、智能化的法治服务支撑体系。



主旨演讲环节，大会邀请来自耶鲁大学、北京大学、香港大学、澳门大学、巴西里约热内卢瓦加斯基基金会、腾讯集团等机构的八位中外知名人工智能法学专家，围绕数字政府、人工智能风险治理、人工智能版权、人工智能主权、人机关系等重要议题发表精彩主旨演讲。该环节由对外经济贸易大学法学院副院长张欣教授主持。



北京大学法学院教授、《中外法学》主编、中国法学会网络与信息法学研究会副会长王锡铤以“‘算治’与‘法治’的协同演化”为主题，系统剖析智能行政对法律保留、正当程序、权责一致及行政公共性的挑战，主张将算法权力纳入制度约束，推动“法治的数治化”与“数治的法治化”的协同演进。澳门大学法学院杰出教授、全球法律研究系

主任**诺思坦(Rostam J. Neuwirth)**以“人机操纵的规制：中欧比较的进路”为题，重点提示人工智能影响人的思维和决策、操纵人的行为等风险。北京外国语大学法学院院长丁相顺围绕“比较法视野下律师执业适用生成式人工智能的伦理规制”作主题发言，探讨律师使用生成式人工智能的职业伦理规范。美国耶鲁大学法学院蔡中曾中国中心高级研究员**唐哲(Jeremy Daum)**以“人工智能致害的中美治理进路比较”为题，结合中美典型案例，比较两国在算法成瘾、AI 幻觉、情感依赖等问题场景下的治理方案。人民法院出版社党委书记、社长、《数字法治》主编**余茂玉**围绕“法律人工智能的安全治理探索与实践”作主题发言，结合“法信”平台和法律基座大模型建设实践，探讨法律人工智能应用中的安全治理问题。巴西里约热内卢瓦加斯基基金会法学院数字治理与监管教授、“网络金砖”项目主任**卢卡·贝利(Luca Belli)**以“通过监管、治理与产业政策构建巴西人工智能主权：一项进行中的探索”为题，介绍巴西推进人工智能主权建设的实践。腾讯集团法务副总裁、中国法学会网络与信息法学研究会副会长刁云芸围绕“人工智能时代影视版权保护的全球趋势与共识”作主题发言，分析人工智能对影视产业的影响，并探讨影视版权保护与产业高质量发展问题。香港大学法学院法律与科技中心主任**陈民豪(Benjamin Minhao Chen)**以“经由机器的行动：自动化时代的伦理决策”为题，深入探讨自动化时代人工智能价值对齐的前沿问题。

当天下午，还围绕“比较法视野下的人工智能治理：模式竞争与互操作性”、“人工智能造福人类：通用人工智能、风险管理与人机共生的未来”、“国际软法治理：治理互操作性与全球规则新基石”和“人工智能与法治实践：风险治理、行业应用与权利保护”等议题举办了多个平行分论坛。与会专家立足比较法、国际软法与法治实践等多重视角，就不同治理模式的交流互鉴、全球规则的协调衔接、人工智能风险的协同应对及人机共生的制度保障展开深入研讨，为凝聚全球人工智能治理共识、提

升治理规则互操作性、推动人工智能安全向善发展提供了重要理论支撑与实践启示。

对外经济贸易大学涉外法治研究院常务副院长**陈卫东**教授主持了论坛闭幕式。法学院副院长**龚红柳**进行了闭幕致辞。龚红柳表示，本次论坛汇聚国内外法学界、实务界与产业界专家，围绕人工智能法律治理的前沿问题展开深入交流，为增进国际社会治理共识、促进不同规则体系交流互鉴提供了重要平台。她向各位与会专家的真知灼见和大力支持表示诚挚感谢，并期待以本次论坛为契机，进一步深化人工智能法治领域的国际交流与务实合作，共同推动构建安全、可信、包容、普惠的全球人工智能治理体系。



2026 人工智能全球治理会议得到了光明日报、经济日报、央广网、人民网、中国日报、北京电视台、中国教育电视台、法治日报、人民法院报、人民政协报、南方都市报、21 世纪财经、中国高科技导报等多家主流媒体的高度关注。

(技术编辑：吕昊然)

数字法评

人脸识别技术法律规制的国际比较与场景化重构

此处删除了原文脚注，全文请参见《行政法学研究》2026年第4期，转载或引用请注明出处。

作者：王业亮

内容提要：

全球范围内对人脸识别技术的法律规制主要基于敏感个人信息保护与人工智能风险规制展开。现有敏感个人信息制度未能给予人脸信息恰当保护，欧盟《人工智能法》对于人工智能采取的统一化风险规制方案也不利于人脸识别技术的创新发展。此外，人脸信息保护与人脸识别技术发展也存在一定冲突之处。化解上述挑战的关键在于引入场景理论，进而依据人脸识别技术的用途、应用主体以及行业进行具体场景划分，构建风险相称的人脸信息保护模式与人脸识别技术的规制框架。场景化的规制方案也有利于构建数字信任，促进人脸信息保护与人脸识别技术发展的和谐统一。

随着深度学习等人工智能技术的突破与发展，人脸识别技术广泛运用于公共安全与商业支付等场景。民众在享受人脸识别技术带来的安全、效率等便利的同时，也面临着强制刷脸、身份信息泄露等诸多隐私侵犯问题。此外，人脸识别技术还因识别误判、算法歧视、大规模监控等技术风险引发普遍担忧。作为人脸识别技术的关键输入，人脸信息具有易采集性、不可变更性等特征，一旦泄露或被滥用，容易造成自然人人身、财产等方面的损害。因此，全球范围内的主要数据隐私立法都主张对人脸信息进行严格保护。但对于人脸识别技术规制而言，作为数字经济主导力量的中国、美国和欧盟尚未对规制方案达成一致。这一现实既反映了中美欧人工智能的技术水平与战略导向差异，也凸显了人脸信息保护与人脸识别技术发展的联系与矛盾——作为敏感个人信息的人脸信息需要强化保护，而人脸识别技术发展则需要相对宽松的规制环境。

从既有研究来看，国内学者多集中于讨论人脸信息保护、人脸识别技术的风险规制等内容，一般对人脸识别技术发展持积极态度；国外隐私法学者多数对人脸识别技术持负面态度，主张严格规制人脸识别技术应用乃至禁止人脸识别技术。本文借鉴数字法学领域有关场景化规制的研究，主张基于场景理论实现人脸信息保护与人脸识别技术风险规制的一体化重构。具体而言，第一部分从中美欧三个主要司法辖区的人脸识别技术的法律规制实践出发，梳理中美欧在人脸信息保护与人脸识别技术风险规制的共性与差异；第二部分讨论人脸识别技术法律规制的双重逻辑——敏感个人信息保护与人工智能风险规制的争论以及二者的冲突之处；第三部分则引入场景理论，在具体场景中重构人脸信息保护制度与人脸识别技术的风险规制模式，平衡人脸信息保护与人脸识别技术应用发展。

一、全球视野下人脸识别技术法律规制

人脸识别技术的数据来源是人脸信息，其技术逻辑主要是利用深度学习等人工智能技术对现采集的人脸信息与已有数据库的人脸信息模板进行对比分析，从而识别特定个体或对特定个体的特征（包括情绪）等进行分析。近些年来，人脸识别技术因其所带来的数据隐私、安全、歧视、监控等方面的风险成为全球范围内主要法域的共同规制对象。

（一）欧盟人脸识别技术的法律规制

从数据保护来看，欧盟的《基本权利宪章》（CFR）在第7、8条规定了数据主体的隐私权与数据权利。由于人脸识别技术的运行逻辑不可避免地需要处理作为特殊类别个人数据的人脸信息，因此在欧盟的框架下构成了对数据隐私权利的干扰和限制。根据欧盟《一般数据保护条例》（GDPR），人脸信息的处理首先得符合一般个人数据的保护原则，即合法、合理与透明性等要求；其次，由于属于特殊类别的数据，人脸信息的处理需要满足第9条禁止处理规定的例外条件，比如数据主体的明确（explicit）同意等。此外，《一般数据保护条例》还规定了对于特殊类别的个人数据需要进行事

先的数据评估以及设立数据保护官等要求。2021年,欧盟委员会提出了《人工智能法》,经修改后于2024年开始实施。欧盟《人工智能法》宣称旨在平衡人工智能技术带来的收益与风险,并将人工智能应用依据风险等级划分为四类——禁止的风险、高风险、可接受的风险与较低的风险。由于处理人脸识别信息必然涉及隐私、安全等基本权利,所以除实时远程生物识别等禁止的一些人脸识别应用外,多数人脸识别技术的应用场景落入高风险人工智能系统的范围,这主要包括非实时的远程生物识别系统、基于敏感或受保护的属性或特征的推断以及用于情绪识别的人工智能系统等情形。

欧盟对人脸识别技术主要持负面态度,无论是数据隐私立法还是人工智能立法均主张严格规制。不过自2024年底以来,欧盟数据保护立法出现了降低企业尤其是中小企业合规成本的新趋势。同时,也有部分欧洲国家领导人呼吁放松对人工智能的规制以应对来自美国的压力。许多大型企业也联合呼吁暂停《人工智能法》的实施。

(二) 美国人脸识别技术的法律规制

目前,美国联邦层面并未通过有关数据保护的综合性立法,既有立法主要是从行业或特定领域的角度予以规范。从州层面来看,有三个州专门制定了生物识别法,即伊利诺伊州、得克萨斯州与华盛顿州,其他州则主要通过数据隐私法中的敏感个人信息条款对生物识别信息予以规制。伊利诺伊州于2008年通过了《生物识别信息隐私法》(BIPA),成为第一个对生物识别信息进行规制的州。为了确保私营企业与消费者间的透明度,该法对私营企业如何获取和使用个人生物识别数据设定了重大限制,主要包括明确的告知、收集的合理性与保留时间以及数据获取与分享前的书面授权等。该法还设立了私人诉权,允许个人就企业的违法行为提出索赔,并规定了较大的赔偿金额。因此,美国境内的生物识别诉讼往往在伊利诺伊州提起。得克萨斯州受伊利诺伊州影响,于2009年颁布了《捕获或使用生物识别标识符法案》(CUBI)。该法没有设立私人诉权,而是授权总检察长提起诉讼以执行该法

规,并对每次违规处以最高25,000美元的罚款限制。与前两项严格保护生物识别数据的州法律相比,华盛顿州于2017年通过的《生物识别隐私法》(WBPA)代表着生物识别信息商业利用的友好进路。华盛顿州的《生物识别隐私法》在立法目的、数据处理与救济规则等方面,都较前两项州立法有所放宽。就人脸识别技术应用的规制而言,美国各政治实体间存在较大冲突。特朗普政府时期,美国政府主张与中国进行人工智能领域竞争,推出“星际之门”计划,并曾试图取消对地方政府对人工智能发展的限制。而拜登政府时期提出的《算法问责法案》等立法草案预计也很难获得通过。此外,不仅美国政府与各州对人工智能技术的态度不一致,美国地方政府层面也未对人脸识别技术应用与限制达成共识。

因此,无论是数据隐私领域还是人工智能领域,美国各州对人脸识别技术的法律规制都未能达成一致意见,而且联邦政府与各州之间也存在不少冲突之处。在可预见的未来,旨在广泛限制甚至禁止使用人脸识别技术的联邦立法不太会获得通过,而地方层面对人脸识别技术的限制措施也充满着较大的不确定性。

(三) 中国人脸识别技术的法律规制

我国数据隐私立法重视对人脸信息的保护。《中华人民共和国民法典》(以下简称《民法典》)在1034条规定了作为个人信息的生物识别信息。《最高人民法院关于审理使用人脸识别技术处理个人信息相关民事案件适用法律若干问题的规定》(以下简称《人脸识别司法解释》)指出,人脸信息属于《民法典》第1034条的生物识别信息,明确经营场所或公共场所违反法律、行政法规的规定适用人脸识别技术进行人脸辨识、验证和分析以及其他无正当理由处理人脸信息的行为构成侵犯人格权益的行为。《中华人民共和国个人信息保护法》(以下简称《个保法》)第28条则明确规定了敏感个人信息,其中生物识别信息作为明确规定的敏感个人信息,受到较一般个人信息更强的保护,包括适用单独同意制度以及事前评估制度等。2025年

3月,国家网信办与公安部联合发布了《人脸识别技术应用安全管理办法》(以下简称《人脸识别办法》)。与既有法律与相关规定相比,《人脸识别办法》承袭了《个保法》中个人信息保护的原则,强化了对人脸信息的保护,并回应了实践中的强制刷脸等问题。不过,由于人工智能技术的战略意义以及遵循我国历来科技创新规制的包容审慎传统,该办法也为人脸识别技术的创新发展预留了空间。

总体而言,中美欧对于人脸识别技术的法律规制有相同点,也有不少差异。相同点在于各地区都已通过数据隐私立法,强化人脸信息保护。区别在于:一是各地区的保护强度存在差异。一般认为欧盟的个人信息保护最为严格,而美国的数据隐私保护主要采取消费者权利保护模式,往往更鼓励个人信息的自由流通。二是从人工智能的相关立法来看,各地区对人脸识别技术的风险规制方案并未形成一致意见。各地区对于人工智能的发展前景都相当期待,也重视人脸识别技术在公共安全、商业营销方面的经济潜力。即使是欧盟,也认为85%—88%左右的人工智能应用是风险较低的,但是欧盟还是对人脸识别技术进行了大范围的限制与规范。而中国与美国,作为全球人工智能竞争最激烈的两个国家,在数据利用与人工智能领域立法相对谨慎,并通过多种政策手段支持人工智能产业发展。这既反映了中美欧三地区不同的技术发展现状与战略选择,也反映了数据保护逻辑与人脸识别技术发展逻辑的内在冲突。

二、人脸识别技术法律规制双重逻辑及其争议

人脸识别技术的开发既需要作为输入的数据,也离不开作为技术支撑的算法。如前所述,当前的人脸识别技术法律规制主要包含两套逻辑,一是数据隐私保护,二是以算法为核心的人工智能风险规制。尽管各主要法域在数据隐私保护领域共识性内容较多,但是在人脸信息保护方面仍然有不少争议,特别是敏感个人信息制度。而就人工智能风险规制而言,中美欧远未形成共识。此外,数据隐私保护与人工智能发展的内在冲突也是亟待回应的现实问题。

(一) 基于敏感个人信息的人脸信息保护逻辑

一般而言,用于识别目的的人脸信息纳入生物识别信息的范畴较少有争议。无论是欧盟《一般数据保护条例》中的特殊类别的个人数据,还是我国《个保法》规定的敏感个人信息,对生物识别信息都施加了诸如明确(单独)同意、事前的个人信息影响评估等更高的信息处理要求。因此按照上述逻辑,人脸信息在信息处理时受到较一般个人信息更强的法律保护。但是在对人脸信息适用敏感个人信息制度这一问题上,学界仍有一定争议。

首先是敏感个人信息制度的必要性问题。保罗·欧姆(Paul Ohm)认为敏感信息是个人信息保护中的关键转折,因为它为个人信息带来了分层式保护——对于那些带来更大危害或风险的个人信
息予以加强保护。但有不少学者对敏感个人信息制度提出批评。穆格·法兹利奥卢(Müge Fazlioglu)认为随着技术的不断发展,会出现越来越多的敏感数据类别,而人工智能推断能力的提升会将不敏感的个人信息转化为敏感个人信息。丹尼尔·沙勒夫(Daniel J. Solove)则认为不应当基于数据性质即数据是否具有敏感性进行分层保护,而应从数据利用与数据所造成的危害端对个人信息予以规制。保罗·奎因(Paul Quinn)和詹克劳迪奥·马尔基里(Gianclaudio Malgieri)也认为敏感数据概念存在通货膨胀(inflation)效应,为了保证敏感信息范围的合理性应当改造其解释方法。

其次,人脸信息是否属于敏感个人信息存在一定的解释空间。从人脸信息到敏感个人信息要经过两层逻辑,一是人脸信息是否属于生物识别信息,二是生物识别信息是否属于敏感个人信息。尽管我国《人脸识别司法解释》与《人脸识别办法》对人脸信息进行了释义,明确了其生物识别属性,但学者仍然是基于人脸识别信息而非人脸信息展开讨论。这意味着人脸信息作为概念未必一定符合生物识别信息的定义。因为从文义解释的角度来看,很难排除照片、视频等内容作为人脸信息的存在。如果确认照片等内容属于人脸信息,那么对人脸信息的严格保护则可能改写整个互联网社交平台的商

业逻辑。

对于生物识别信息是否属于敏感个人信息的判断,则涉及对现有立法的解释。以我国《个保法》第28条第1款为例,其所规定的敏感个人信息是“一旦泄露或者非法使用,容易导致自然人的的人格尊严受到侵害或者人身、财产安全受到危害的个人信息”,并且后续列举了生物识别等七类具体的个人信息。如果采取目的解释,那么意味着生物识别等类别的信息只有在“一旦泄露或者非法使用,容易导致人格尊严或者人身、财产安全受到危害”时才构成敏感个人信息,如果这些信息的处理并不危及人格尊严或者人身、财产安全则不宜纳入敏感个人信息的范围。人脸识别技术发展的主要驱动便是该技术更好的安全性与准确性,而且随着技术的发展以及多技术的联合适用,人脸信息所带来的身份盗用风险的可能性会大幅降低,那么此时人脸信息是否仍应被视为敏感个人信息?当然,如果采取文义解释,那么这七类个人信息不管在何种情境下均属于敏感个人信息而获得加强保护。因此,从理论上来说,人脸信息是否属于敏感个人信息仍然存在一定的不确定性。

最后,现有敏感个人信息的保护方式也未必合理。当前敏感个人信息制度一般包括更高的合法性要求,单独同意以及个人信息影响评估等制度。但这些制度在实践中也遇到不少挑战。以单独同意为例,该制度仍然遵循知情同意框架。而知情同意这种隐私自我管理模式一直饱受批评。隐私自我管理的基本假设在于向个体赋权以平衡消费者与企业之间的不平等关系。但是这种构想并未充分考虑个体的认知困境以及大数据技术所带来的结构性问题。以人脸识别技术的商业场景,如自动售货机为例,尽管消费者会接收到售货机页面提供的隐私协议,但是少有消费者去认真阅读这些隐私协议,往往都是直接勾选同意选项,然后进行人脸识别,确认后即完成购买服务。整个购买环节过程很短暂,如果去认真阅读隐私协议,那么不如选择直接选择扫码支付。因此,这种单独同意机制未必能够实现立法者的预期目标。

(二) 人脸识别技术的风险规制逻辑

作为典型的人工智能应用,人脸识别技术所带来的风险除了数据隐私风险外,还集中体现为技术的不准确与算法歧视等工具性风险。如何对人脸识别技术所带来的工具性风险进行规制也是当前学界争论的核心。

首先,人脸识别技术的规制理念应该选择风险预防还是成本收益分析存在争议。近年来,随着ChatGPT等生成式人工智能的爆发,不少学者主张基于风险预防原则对新兴科技进行规制。实践中,美国部分城市选择直接禁止人脸识别技术在执法部门的应用,欧盟《人工智能法》也对部分人脸识别应用采取严格禁止(仅有少数情况例外)的规制模式。但是,这种一刀切禁止或严格预防的风险规制模式可能会错失人脸识别技术所带来的收益。奥莉·洛贝尔(Orly Lobel)认为当前的政策讨论中充斥着对新兴技术的抵制思潮,往往高估人类决策的质量,而未能对人脸识别技术应用进行正确的相对成本收益分析。尼希尔·夏尔马(Nikhil Sharma)认为美国在中美竞争的背景下对人工智能发展施加过多限制是一种错误。不过也有学者认为对科技企业的强规制未必会抑制技术创新,并以欧盟与美国为例,认为欧盟科技产业落后主要是由于缺乏良好的金融市场、高技能人才等条件,而非数字产业规制。

其次,对人脸识别技术选择何种版本的风险规制也尚有讨论空间。按照玛戈特·卡明斯基(Margot Kaminski)的说法,至少存在四个版本的风险规制,包括20世纪60至80年代美国行政法中的定量风险评估、民主监督式的风险规制、基于风险(配置执法资源)的规制以及企业自我规制。这些不同的风险规制模式有不同的法律移植来源,往往也指向不同的目标与策略。在人脸识别风险规制的语境下,选择何种版本的风险规制存在争议。欧盟《人工智能法》选择了基于风险的规制方案,主要考虑人工智能尤其是人脸识别技术的风险,选择禁止和严格限制政治决策者所设想的人脸识别技术类型。正如部分学者所指出的,欧盟《人工智能法》的风

险分类分级的标准并不科学。以公共执法领域的人脸识别技术为例,欧盟采取了除特殊情形外,严格禁止远程人脸识别技术的做法。但如果考虑到全球范围不同地域的治安差异,对某些治安环境较差的地区严格禁止实时远程人脸识别技术,可能并不能满足当地的现实需求。尽管美国部分城市选择禁止人脸识别技术在公共部门的应用,但私营企业对人脸识别技术的开发利用则更多受到企业自我规制的约束。此外,与风险规制模式相关的一个问题是规制工具的选择。事实上,即使是相同的规制工具,不同版本风险规制模式下其重要性和侧重点也可能存在差异。

(三) 人脸信息保护与人脸识别技术发展的冲突

人脸信息保护遵循保护法逻辑,其方式在于向信息主体赋权以平衡信息主体与信息处理者的不对称地位;人脸识别技术风险规制则是规制法逻辑,其方式主要在于通过向企业技术应用施加责任或限制以影响技术发展的水平和方向。在全球人工智能竞争背景下,人脸识别技术应用可能呈现出“去规制”或暂缓规制的趋势,而人脸信息保护则呈现出强化保护趋势。这种趋势背离反映了人脸信息保护与人脸识别技术发展的内在冲突。

数字时代,保护个人信息具有必要性。同时,数字经济的发展也离不开数据要素的汇聚融合,人脸识别的发展与应用离不开大量数据的训练与技术迭代。正如有学者指出,隐私保护与信息的自由流动存在冲突。人脸信息因其敏感性而受到强化保护。这意味着数据处理者在处理人脸信息时面对更多的法律限制,数据训练的成本也会上升。尽管人脸信息具有敏感性,但其同时也具有易采集性的特点。人脸识别技术之所以会快速发展并获得应用,也在于其相较于传统的指纹识别或基因识别等生物识别方式具有更大的便捷性。因此,人脸信息的强化保护可能与人脸识别技术对效率的追求相矛盾。

人脸信息的强保护模式也可能会增强人脸识别技术的工具性风险。人脸识别的工具性风险突出

体现在两个方面,一是技术准确性,二是算法歧视。人脸识别技术的准确性风险在于可能会产生假阳性或者假阴性问题。关于人脸识别技术准确性的例子并不罕见,由于人脸识别验证不通过而未获得信贷等情形也时有发生。人脸识别技术的歧视风险主要在于人脸识别用于招聘就业等场景,比如亚马逊的招聘算法由于训练数据的原因而延续了其在就业招聘环节中的性别歧视。作为人脸识别技术的关键投入,训练数据决定了识别算法的效果。因此,训练数据的数量限制以及代表性欠缺等因素可能会导致人脸识别算法的不准确与歧视问题。事实上,即使是科技监管最为严格的欧盟,在《一般数据保护条例》与《人工智能法》中都采用基于风险的方法,希望权衡人工智能所带来的风险与收益,尤其是在保护基本权利的前提下促进人工智能产业的发展。因此,通过何种方式平衡好人脸信息保护与人脸识别技术的发展仍然是关键的理论与实践问题。

三、人脸识别技术的场景化规制

法律制度对科技发展的促进作用受到普遍认可。试图基于数据性质或者采取统一化风险规制的方式会遭遇保护不足或保护过强的问题,难以协调好个体利益、企业利益与集体利益。事实上,无论是敏感个人信息保护还是人工智能的风险规制,在实践中很难脱离场景要素。民众对人脸识别技术的态度也并非一成不变,而是随着场景的变化而变化。有研究表明,人们通常对政府使用面部识别来调查严重犯罪、增强机场和学校等受控空间的安全性以及在某些情况下提高身份验证效率感到满意,但往往不习惯在公共场所看到政府部门随意使用面部识别技术。因此,场景是决定人脸识别技术法律规制的重要考虑因素。

(一) 场景理论的引入

近年来,“场景”成为法学研究中的热词。从数据隐私领域来看,“场景”概念主要源自海伦·尼森鲍姆(Helen Nissenbaum)的场景一致性(contextual integrity)理论。该理论认为信息流动应当符合特定场景下的信息规范。信息规范由

参与主体（包括信息主体与信息处理者）、信息类型以及传输原则三个要素构成。场景理论最初应用于数据隐私领域，后普遍被学者推广至算法规制领域。场景要素介入算法规制的逻辑在于，诸如算法公开等算法规制工具的一般性适用不能达成预期目的。因此，不同场景可能存在不同的信息规范，也相应地存在不同的算法规制框架。

根据场景一致性理论，如果新的信息实践违反了原有的信息规范，则需要基于原有社会场景的价值目标来评估这种违反行为是否具有合理性。基于信息规范的要素构成，技术发展既可能带来信息主体的变化，也可能带来信息类型的变化，还可能改变既有的传输原则。以人脸识别技术替代传统的指纹识别技术为例，这种技术变化不仅带来了信息类型的变化（尽管都是生物识别信息），也改变了传统的信息传输模式——由传统的接触提取变为现代的远程扫描，而且信息处理者的范围也得到了极大的拓宽。至于这一变化是否违反场景一致性，则需要根据具体场景进行分析。

既有研究分析了人脸识别技术不同应用场景的划分标准，包括技术用途、所处行业、公私主体等。这些分类模式与场景一致性理论中信息规范的三要素也存在对应关系。不同的人脸识别技术用途对应不同的信息类型——不同用途下人脸信息的特征也并不相同，比如情绪识别下的人脸信息并不一定与特定身份相关联。公私主体划分符合参与主体的分类逻辑，而行业的区分则往往代表着不同的信息传输规则。

此外，基于人脸识别技术用途以及信息参与主体等标准的场景划分方式也与欧盟《人工智能法》中人脸识别技术的风险分类分级规制相对应。这种对应性体现了人脸信息保护与人脸识别技术规制的关联性与互补性。因为无论是侧重源头治理的人脸信息保护，还是结果导向的人脸识别技术风险规制，其制度目的均在于合理防范人脸识别技术应用过程中的多重风险。因此，场景理论的意义不仅在于判定何种信息实践符合信息规范，还在于提供了场景这一要素作为识别风险以及对风险进行分类

分级的依据。

（二）人脸信息保护的场景化重构

在利用场景理论重构敏感个人信息制度以加强人脸信息保护前，需要回应有关学者对敏感个人信息制度的批判。如前所述，学者们对敏感个人信息持否定态度的理由主要在于机器学习推理能力的增强，能够与其他信息相结合推断出敏感个人信息，进而导致敏感个人信息概念的膨胀。对这类否定意见有两种回应：一方面，这类担忧可能高估了人工智能推理的能力且低估了推断所耗费的成本；另一方面，这种因推断而产生的敏感个人信息与本质上敏感的个人信息不同，可以将其视为一般个人信息的滥用并进行规制。此外，敏感个人信息制度在立法上也具有表达功能，能够引导企业在数据处理过程中恪守保障个人信息安全的原则。鉴于我国《个保法》以及域外主要立法都引入了敏感个人信息制度，更为妥贴的方式可能在于基于敏感个人信息与一般个人信息分层保护的基础上引入场景要素，对人脸信息保护进行场景化重构。这种重构可以从两个角度展开：一方面，场景要素如何作用于人脸信息敏感性的判定，另一方面是场景要素如何影响人脸信息的保护方式。

基于个人信息保护的法理可知，个人信息受到保护的理由在于其承载着数据主体的一系列人身、财产权益。同样，法律对敏感个人信息和一般个人信息进行分层保护的目，也在于二者所指向的权益有所差异，敏感个人信息往往与数据主体的人格尊严以及人身、财产安全联系更为紧密。因此，敏感个人信息概念中的“敏感”联系了法律的强化保护与个人信息的风险等级。而风险等级与具体场景紧密相关。因此，尽管敏感个人信息的界定方式种类繁多，但大都包含场景这一要素。以双重界定模式为例，这类观点主张个人信息内容作为构成敏感个人信息的内因，而将信息处理者的处理能力等场景要素视为敏感个人信息的外因，外因会通过作用于内因而改变个人信息的敏感性。

无论是何种界定方式，敏感个人信息的界定离不开风险内容。尽管人脸承载着社会评价等其他内

容,但是用于识别和验证的人脸信息往往很难与人格尊严受损等风险内容直接联系,更有可能构成工具性的敏感信息,因为其危险内容主要在于位置监控与身份盗用。而情绪识别以及利用人脸信息进行分类等信息实践则可能引发人格尊严受损。对于人脸信息实践可能发生的身份盗用这类现代风险而言,活体检测等技术的发展有可能减弱此类风险发生的概率。而且从既有人脸信息相关的诈骗案件来看,存在受害人疏于保护自己人脸信息以致财产损失的现象。不过,就人脸信息的位置监控风险而言,其与行踪轨迹信息具有类似性,技术发展可能增强了监控风险的可能性。此外,敏感信息界定的另一个要素是要构成大多数人的共识,这种共识又可能受到风险启发或错误宣传的影响,因此人脸信息是否具有敏感性可能也会因大多数人对人脸识别技术的认识变化而变化。因此,总体而言,场景要素作用于人脸信息敏感性界定主要表现在风险内容、风险概率、公众认知等三个方面,而且这三个方面彼此之间可能存在交叉现象,比如人脸识别技术的用途不仅直接决定了风险内容,也间接影响风险概率,毕竟与验证技术相比,识别、分类以及情绪识别等技术仍处在快速发展期,技术成熟度可能仍有所欠缺。而人脸识别技术的准确性也会影响到公众的认知,进而会影响到人脸信息的敏感性认定。

人脸信息在不同场景下敏感性发生变化的逻辑也会作用于人脸信息保护制度的构建。当前《人脸识别办法》主要承袭《个保法》有关敏感个人信息的规定,要求信息处理行为具有特定的目的和充分必要性,获得个人单独同意并进行个人信息保护影响评估等。这种统一性的强化保护可能带来保护过度或保护不足的问题。就人脸信息处理的合法性基础而言,严格限定人脸信息的处理目的可能不利于人脸识别技术等相关产业的发展,而过于宽松的目要求可能又使得人脸信息的合法性要求流于形式,导致人脸信息的滥用。因此,应当结合具体场景确定人脸信息处理的合法性。与处理目的相关,不同的人脸识别场景下对信息主体所造成的风险并不相同,因此,也需要在具体场景中配置信息

主体与信息处理者的权利义务。申言之,如果按照向个体赋权的模式,则应当承认不同场景下信息处理者根据风险而定的告知等一系列义务。如果对隐私自我管理进行一定程度调整,适度引入监管部门的助推行为,那么也可以根据具体场景,进一步地调整信息处理者在人脸信息处理过程中的义务和角色,构建更为动态的分层同意模式。此外,还可以根据不同场景下人脸信息的不同用途及其风险,构建更具针对性的个人信息影响评估制度,比如在情绪识别类的人脸信息收集处理中更多关注技术本身的伦理性追求。

(三) 人脸识别技术风险规制的场景化重构

人脸识别技术的应用存在风险是普遍共识,但是对于有风险的技术采取何种规制理念决定了技术发展的前景。美国部分城市禁止公共执法部门使用人脸识别技术,但是不少组织和民众对此持反对意见,认为应当暂停而非禁止人脸识别技术,因为这项技术未来可以提供安全的公共环境。事实上,对于人脸识别技术的抵制态度可能与诸多谬误相关,比如对人类与人工智能采用双重标准,对人脸识别技术的准确性采取静态视角,以及对人脸识别技术歧视的固有偏见等。对人脸识别技术的规制应当侧重安全与效率等价值的平衡,而非完全沉浸在风险叙事中。因此,从规制理念上说,应当采取务实的态度,权衡人脸识别技术所带来的风险和收益,而非直接禁止人脸识别技术的应用。应当抛弃“禁止/允许”这种二元选择,而应在“完全禁止—完全允许”的频谱内根据具体风险收益的情况采取针对性方案。

对于人脸识别技术的风险规制而言,基于场景的成本收益分析更能平衡技术带来的风险与收益。尽管欧盟《人工智能法》采用了基于风险的方法,但选择了横向统一化规制的方案。这一方案的突出问题就在于过度追求横向一致性而忽视了不同场景下人工智能的风险与收益差异,比如被欧盟视为不可接受风险的人脸情绪识别,也能用于检测长途货运汽车司机是否处于疲劳驾驶这类有益用途。人脸识别技术风险和收益也需要代入具体的场景才

能充分显示,比如,执法部门人脸识别技术的风险主要在于数据监控以及技术准确性的担忧,其收益在于更安全的公共秩序与行政资源的节约;商业部门人脸识别技术的风险主要在于身份被盗用以及精准定位营销,而收益在于更便利的消费流程。此外,不同场景下民众对人脸识别技术的接受程度也并不相同。民众往往对政府部门使用人脸识别技术抓捕罪犯等行为表示认同,而抵制人脸识别技术用于商业营销。同时,不同地域的民众对人脸识别技术的风险偏好也可能存在差异。

人脸识别技术的风险规制应当考虑具体的规制模式与规制工具。从规制模式选择来看,不同版本的风险规制各有其优劣。比如单纯的定量分析可能并不适宜被用来规制正处在快速发展期的人工智能技术,因为此时的风险收益计算难度偏高。而在公共执法场景下,民主问责模式更能赋予人脸识别技术风险规制以正当性,也有利于提升民众对人脸识别技术的信任。而在商业场景下,企业自我规制可能更能释放企业的创新活力,不过也应当辅之以行政规制的约束。从规制工具来看,公共执法场景下,应当更加重视人脸识别技术应用事前的风险评估,而对于商业场景,算法备案则更能给予企业开发或应用人脸识别技术的空间。欧盟《人工智能法》不仅采取了公共执法与商业场景的划分,也根据人脸识别技术的用途进行了不同的规定。但该法对人脸识别技术的风险分类分级标准存在商榷之处,比如对情绪识别的严格限制。情绪识别不仅可用于刑事侦查,也可以用于医疗场景,还可以用于交通以及智能家居等行业。即使是《人工智能法》禁止的教育与工作场景,情绪识别仍然有一定的工具价值。不过由于情绪识别所带来的人类主体性受损等风险,在风险规制模式选择上,采取民主问责的风险规制可能更为恰当,而且也应当选择偏“硬”的规制工具,比如严格的事前审批。因此,对于人脸识别技术的风险规制也不应当根据用途进行一刀切的模式,而是依据具体场景,权衡各方利益,选择对社会整体利益损害最小的方式并辅之以针对性的规制工具。

(四) 人脸信息保护与人脸识别技术发展的场景化平衡

人脸信息需要强化保护,人脸识别技术创新发展所带来的经济潜力与效率优势同样值得追求。化解人脸信息保护与人脸识别技术发展冲突的关键在于依托具体的场景,实现各类权益的平衡取舍,尤其是个体权利与企业权益以及社会利益的平衡。人脸信息具有双重属性,既与特定个体的数据隐私权益紧密相连,又呈现出一定的公共性特征,整体性的人脸信息其实是一种公共资源。当前全球人工智能竞争日益激烈,我国正是凭借超大规模数据要素市场优势占据有利地位。人脸识别技术发展不仅关系个体安全,也关系公共安全与国家安全。因此,对人脸信息的强化保护不应带来人脸识别技术发展的过多限制,阻碍人脸识别技术产业效益的实现。

人脸识别技术的现实需求与人脸识别技术的发展历程也证成了人脸识别技术应用的正当性。与传统人工审核相比,人脸识别技术能带来更高的效率,而且随着技术发展会逐渐超过人类识别的准确性。与其他生物识别技术相比,人脸识别的非接触性特征也是其产业前景的重要支撑。这也表明人脸信息与其他生物识别信息也存在类似不同隐私间的权衡关系。此外,人脸识别技术的准确性取决于训练数据的质量与数量。有学者指出,人脸识别训练数据代表性欠缺导致了识别效率低下与各类歧视。当前许多批评人脸识别技术误判或者存在歧视的问题的破解之道反而是对更多具有代表性的人脸数据进行训练。因此,人脸识别技术所带来的风险彼此间也存在冲突,换言之,人脸信息保护所预防的风险恰恰有可能给人脸识别技术应用带来风险。同时,人脸信息保护与人脸识别技术发展也存在相互促进的可能——人脸信息保护增强信息主体信任,促进人脸信息的流通,进而有利于人脸识别技术的发展。人脸识别技术的发展也有利于防止身份盗用等风险,增强人脸信息的保护效果。因此,不应采用对立的视角去看待人脸信息保护与人脸识别技术发展。政府部门应当主动作为,对人脸信

息进行场景化保护，同时通过各类规制工具包括事前评估、算法备案、算法审计等方式，在具体场景下实现人脸识别技术的透明可信赖与发展的正向循环。

四、结语

数字时代，人脸信息的强化保护是中国、美国与欧盟共同的价值选择。而基于不同的规制传统与人工智能产业的发展水平，中美欧并未形成关于人脸识别技术的一致性规制框架。欧盟主张基于风险的规制模式，严格限制人脸识别技术的应用。美国基本延续宽松的技术规制传统，鼓励人脸识别技术的应用与创新。我国则秉持务实的规制理念，试图在人脸识别技术的规范与发展间寻求平衡。从全球范围看，当前基于敏感个人信息制度的人脸信息保护以及以欧盟《人工智能法》为代表的统一化风险规制模式，未能与具体场景下的风险相协调，存在保护、规制不足或过度保护、规制的问题。应当根据人脸识别技术的不同用途、不同主体以及不同行业等分类标准，确立人脸识别技术的具体应用场景，选择与风险相称的人脸信息保护模式与风险规制框架。人脸识别技术应用不仅有利于提升公共安全，也有利于提升商业交易效率、增强数字经济活力。人脸信息的强化保护能够增强信息主体对人脸识别技术的信任，促进人脸信息的流通利用，推动人脸识别技术的应用与发展。人脸识别技术的发展也会减少身份盗窃、识别错误与歧视等方面的风险。因此，应当根据具体场景，强化人脸信息保护，优化人脸识别技术的风险规制，促进数字信任，实现人脸信息保护与人脸识别技术发展的和谐统一。

人工智能安全管理义务及其承担者的刑事责任

此处删除了原文脚注，全文请参见《中国法学》2026年第2期，转载或引用请注明出处。

作者：皮勇

内容提要：对于人工智能系统自主决策和控制造成的危害后果，现行法律规制存有空白。若仅采取民事、行政法律手段，不足以遏制此类严重危害，需要引入刑法治理。关于刑法治理的路径，人工智能系统犯罪主体说难以实行，而妨害人工智能安全管理秩序犯罪立法的路径能够兼顾人工智能安全与发展，其中，采取违反人工智能安全管理义务犯罪的立法方案能够有效遏制前述危害。人工智能安全管理义务是为防控人工智能安全风险而设定的特殊、有限的规范责任。为有效惩治严重违反该义务的行为，有必要在刑法中增设违反人工智能安全管理义务罪。在罪行构造上，应依据义务犯、风险分配等刑法学理论，以有效降低人工智能安全风险为标准，合理规定该罪的人工智能安全管理刑法义务、主体和罪过等要件。

随着人工智能的广泛应用，社会各领域都存在人工智能安全风险。在人工智能系统自主决策、控制的情况下，作为违法犯罪构成条件的主体、主观要件“消失”，依照现行法律没有人应对危害后果承担法律责任，导致人工智能安全风险出现制度性失控，以规范人类活动为潜在信条和被认为具有促进人类伦理互动价值观的现行法律，在此面临根本性的规制失灵。人工智能安全关系国家安全、社会公共安全和公众合法权益，需要在兼顾促进人工智能快速发展应用和保障安全的基础上，为人工智能安全发展提供充分的法律保障，包括刑法保障。

一、人工智能危害治理的法律路径

人工智能危害表现为多种类型：第一类是利用人工智能技术实施违法犯罪，如利用智能飞行器实施偷拍、伤害甚至杀人；第二类是人工智能系统研

发应用过程中侵犯数据安全和个人信息权益等的违法犯罪；第三类是人工智能系统自主决策、控制造成危害后果，如生成式人工智能传播虚假有害信息、完全自动驾驶汽车造成人员伤亡等。我国现行刑法可以适用于前两类危害，对第三类危害则“力有不逮”。

（一）人工智能危害的法律治理主张

人工智能危害的刑法治理是人工智能危害法律治理体系的重要组成部分，二者在治理策略和逻辑上应当保持一致，惟此才能构建各部门法相互衔接配合的、系统化的法律治理体系。关于前述第三类危害的法律治理路径，主要有人类代理说、产品责任说、拟制主体说、安全管理义务说四种主张。

1. 人类代理说

人类代理说认为，人工智能系统的功能和决策方式可以直接追溯到系统的设计、编程和知识，人定义、引导和最终控制着系统的研发和应用，无论系统多么复杂，其本质上是人类使用的工具，不具有法律人格，应当作为人的代理，由自然人或其集合体对人工智能系统引起的危害承担责任。例如，在“达芬奇”等手术机器人自主手术造成病患伤害的案件中，应追究机器人生产厂家、医院而不是手术机器人的责任。

这种观点没有考虑人工智能系统自主运行的事实。确有人工智能系统的运行没有遵照人类的指令、受人类控制或在人类控制能力范围内的情形，如L5级自动驾驶汽车完全自主行动，其自主决策的场景中有的是研发设计和应用人员无法预见的，更谈不上控制。完全自主的人工智能系统处于没有人类参与或干预的独立“感知—分析—行动”的“行为状态”，“能够独立主动地制定自己的计划”，有学者认为其更适合被视为人而不是机器。笔者认为，人工智能的自主性动摇了人类代理说的事实基础。除了代理人必须是自然人或其集合体，一般情况下代理人应当按委托人的指示行事，表见代理和紧急情况下例外属于特别情形，而人工智能系统的

完全自主活动并不限于前述特殊情形，是一种常态，不满足代理的条件，将危害后果归责于其使用者或生产者，缺乏事实和法律依据。

2. 新产品责任说

新产品责任说主张构建人工智能产品的新严格责任制度，生产者应当对教会系统“思考”的附生后果（即造成不可预见的危害）承担严格责任。人工智能系统自主行动引起的危害应归因于产品的制造缺陷、设计缺陷或编程不佳等。可是，人工智能系统更加复杂，即使是按“消费者期望”（consumer expectations）和“风险效用”（risk-utility）标准，证明存在以上缺陷也是困难的，导致受害人因无能力证明而得不到赔偿。对此，该观点认为，如果能够排除人为责任（如驾驶员人员或安全监督员的过错责任），即使不能查明人工智能系统存在缺陷，当事故的发生有一定频率且表现出共同特征时，应当认定为产品缺陷。此外，人工智能系统由多方提供技术和设备，如无人驾驶汽车的自动驾驶系统、雷达和激光传感器等，这些软件和硬件设备容易出现无法检测的故障，特别是人工智能算法系统具有不可解释、不可预测性。只追究作为人工智能系统生产者的责任是不公平的，应将人工智能系统关联公司按照“共同企业”对待，追究其中部分公司的连带责任，以解决无法将故障归因于其中特定公司或自然人的追责困难。

持相似的新严格责任观点的学者认为，事实因果关系和法律因果关系检测遇到人工智能暗箱障碍，应通过对人工智能的创造者施加较高标准的注意义务，对违反该义务、威胁公共安全的行为，按严格责任原则追究法律责任，目的是“允许超越因果关系检验，将有关的损害与创造者的过失行为联系起来”，以平衡人工智能发展和安全的需要。

笔者认为，该观点认识到人工智能系统的自主性和复杂性，以及按照传统产品责任制度追责的缺陷，但存在以下不足：其一，对产品故障或缺陷进行推断认定，并延展责任追究的因果链条，要求人工智能系统生产者及为其提供技术或零部件的上下游企业承担连带责任的观点，实际上是要求这些单

位在未被证明有过错的情况下承担补偿义务，既脱离了现行产品责任和侵权责任的法治轨道，也没有解决相关企业之间的责任分配问题。其二，只对人工智能系统的生产者施加注意义务及按严格责任制度追究法律责任，不能降低人工智能系统交付后因其自主学习产生的安全风险，没有考虑到使用者直接监督系统运行的作用。其三，人工智能的自主性导致其不可预测，生产者对危害后果缺乏预见和预防的可能性，而依照严格责任制度及其理论，只有人工智能系统的自主决策和活动是可预测的，让生产者承担严格责任才是公平的。其四，受人工智能系统侵害的不仅有被害人的财产性利益，还可能有生命、健康和其他重要权益，仅让生产者及相关企业进行民事赔偿或补偿，不利于保护公众生命、健康等重要法益。这在历史上是有教训的，在巨大的产业利益诱惑下，单纯地依靠民事赔偿不足以有效控制安全风险。

3. 拟制主体说

拟制主体说认为，生产者无法预见人工智能系统的全部“行为”，应承认后者具有法律能力（legal capacity），因其在功能上能够“有目的地行事”，展现出所谓“道德的”“法律的”能力。人们像对待自然人一样对待拟人化的机器人，应赋予后者权利使其免受人类的虐待。拉丁语“人格”（persona）的原意就是面具或角色，既然可以承认各种非自然人主体（如公司）的法律人格，对比公司法人更“实在”的人工智能系统也可以赋予其法律人格并制定制裁措施。持该种观点的学者认为，人工智能系统可以被追究民事和刑事责任，并建议增设删除系统软件、限制自由、社区服务和罚款等专门的刑罚。还有学者将人工智能系统的法律拟制人格区分为伦理人格和技术人格，认为人工智能的人格将“以技术人格的探索先行，逐渐进行伦理人格的塑造，人类或机器人的伦理人格最终成为技术人格的依归”。拟制人格说曾得到国外立法机构的响应，2017年欧洲议会提交《关于机器人的民法规则的决议》，呼吁欧盟委员会在人工智能立法中，“为机器人创造一个特定的法律地位，至少最复杂的自主机器人

可以被确定为具有电子人地位,使其可以负责修复可能造成的任何损害,并将电子人格应用于机器人自主决定或与第三方独立互动的情形”。

笔者认为,该种观点存在以下不足:其一,人工智能系统与公司法人不具有可比性,前者不具有独立于自然人或法人的独立利益。其二,未来将有数以亿计的自主人工智能系统投入应用,授予其以主体资格的现实性和可行性存在疑问。即使可行,能否消除人工智能危害及其效率同样存疑,反而会因免除或弱化现有法律主体防控安全风险的责任,导致“有组织的不承担责任”。其三,采纳该观点则要对现行法律制度作重大修改,需要制定适用于“第三类法律主体”的法律制度,法律变革成本大,效益不高,难以与现行法律制度保持协调。对人工智能系统追究刑事责任,不能实现刑法的目的与任务,与犯罪、刑事责任和刑罚基本制度不相容。

4. 安全管理义务说

安全管理义务说认为,人工智能系统是自主性机器和人类的工具,人工智能系统安全管理义务承担者应确保其安全并对其危害后果负责。美国《人工智能倡议法》和欧盟《人工智能法》都将人工智能系统定义为“基于机器的系统”,要求“人工智能系统被开发和用作一种服务于人类、尊重人类尊严和个人自主权的工具,并以一种可以被人类适当控制和监督的方式运作”(欧盟《人工智能法》引言第27条)。后者还制定了人工智能安全风险防控制度,要求人工智能系统提供者、部署者及其他相关方承担安全管理义务和相关法律责任。以上法律既未赋予人工智能系统法律人格,也未规定其法律责任。我国人工智能发展战略文件及相关立法持相同立场,我国提出的《全球人工智能治理倡议》要求“确保人工智能始终处于人类控制之下,打造可审核、可监督、可追溯、可信赖的人工智能技术”,《互联网信息服务深度合成管理规定》等将人工智能系统作为防控对象,要求人工智能系统提供者及其他相关方承担安全风险防控义务。该观点客观反映了人工智能的特性,能有效防控人工智能危害,符合促进人工智能安全发展的国际共识,对现有法

律体系的冲击不大,对人工智能危害的法律治理具有重要的指导作用。

以上四种观点存在的共同不足在于,主要以民事、行政手段来治理人工智能危害,解决被害人的赔偿或补偿问题,不能遏制严重危害。事实上,人工智能危害不仅威胁公众的生命、健康和财产安全,还严重危害社会公共安全和国家安全,民事、行政法律手段不足以遏制严重的人工智能危害,有必要为人工智能安全提供刑法保障。

(二) 人工智能危害治理的刑法路径

对于人工智能危害的刑法治理,前述第三类人工智能危害客观上不能归因于自然人的行为,主观上没有人对其有罪过,不能采用人类代理说。由于主客观相统一原则是我国刑法的基本原则,故也不能采纳产品严格责任说。目前以下两种观点较有影响。

1. 人工智能系统犯罪主体说

该观点认为,人工智能系统自主决策并直接引起危害后果,可以将其拟制为犯罪主体,并依照现行刑法追究刑事责任,这符合目的刑和报应刑的立场。通过对人工智能系统定罪和处以特定刑罚,如限制应用、删除数据、修改或销毁系统等,可以剥夺人工智能系统的再犯能力,威慑其他人工智能系统,还可以促使其所有者制止危害,避免其系统被处罚而遭受损失。对人工智能系统处以与其危害相称的刑罚,支持被害人的权益,可以安抚被害人感情,表达官方的谴责,恢复公平正义和社会安全感。

该观点实际上难以实施,首先有主体资格障碍。刑法只处罚有罪责能力的行为人,罪责能力是刑法适用的必要条件,人工智能系统不具有刑法上认识和意志的生物学和社会学基础,既没有意识和感觉,也没有利益或幸福,不满足犯罪主观要件的基础,不具有受刑法处罚的资格。为了解决这一问题,学者们尝试了多种理论工具,主要有:

(1) 雇主责任理论。该观点认为,公司是拟制的犯罪人,本身不具有受谴责的能力,借助于雇主责任理论,刑法赋予公司可谴责的主观状态。公司雇员在受雇范围内行事且服务于公司利益,雇员

的主观状态被认为是公司的主观状态，对公司定罪处刑不违反法治原则。按此理论，人工智能系统也可以被赋予罪过状态和罪责能力。笔者认为，人工智能系统自主决策与公司决策的机制不同，系统的提供者、使用者、所有者对危害没有罪过，何以作为人工智能系统的罪过状态和罪责能力的依据？该理论不能解决人工智能系统罪责能力问题。

（2）客观责任理论。该观点认为，从结果主义的立场出发，缺乏有罪的精神状态也可以定罪，“可能发展出一种合理的客观标准作为人工智能罪责能力的判断依据”，没有罪过心态的人工智能系统也可追究刑事责任。笔者认为，即使是严格责任犯罪，作为刑法意义上的行为也必须基于精神状态，应当维护这一现代刑法制度的基础。行为首先是能够归于作为心理和精神的动作中心的自然人的一切，是“人格表现”。人工智能系统没有精神状态，其行动不能被评价为刑法上的行为，否则将颠覆刑法谴责有罪过行为的理论基础。

（3）人工智能系统的特别犯罪意图理论。该观点认为，罪责性是有罪的心理以不尊重法益的方式表现出来，即以行为表现，而不是动机、思想和感受方面的细微差别，来证明有罪的心理状态。法律允许以公司员工收集信息和决策的行为，认定公司存在不尊重法益的罪责心理，那么，也可以采取类似的方式来认定人工智能系统的罪责心理。例如，自动驾驶系统检测到道路上的运动目标，如果其控制汽车驶向目标，而不是减速或避开，可以依据检测数据和控制信息等，证明该汽车有促成危害结果发生的罪过心理。自然人对某事的真实性有充分稳定、强烈的倾向时，即可认为有认识，那么，人工智能系统有同样的倾向时，也应认定系统有认识。该观点主张按照以上规则来认定人工智能系统的罪过心理，并制定认定人工智能系统犯罪的司法规则。笔者认为，该理论存在转移问题的嫌疑，将待证明的结果作为证明的前提，是先假定人工智能系统存在罪过心理，而后提出认定罪过的规则，并没有解决罪过的存在问题。

其次，人工智能犯罪主体说的实施有可罚性障

碍。刑罚的本质是“给受罚人以痛苦的制裁措施”。这里的“痛苦”应当理解为使受刑者感受到痛苦，或其他通常被认为是不愉快的后果，后者包括一般人认为是不愉快但并未使受刑者不愉快的后果。这两种形式的痛苦都缘于对受刑者利益的限制或剥夺，即使受刑者不因其利益被限制或剥夺而感到痛苦，客观上仍然对其造成不好的后果。人工智能系统不能感到痛苦或不愉快，尽管丧失能力或被摧毁在客观上剥夺系统的利益，可以认为是对系统的刑罚，但利益由欲望和目标组成，二者以看法和意识为前提，受刑者至少要有快乐和痛苦的体验，才能真正拥有利益。人工智能系统不具有真正的认知和意识能力，不具有纯粹意义的利益，对其利益的剥夺也就谈不上是刑罚。

也有观点主张将定罪和处刑分开，处刑不作为定罪的前提。例如，对公司不能处以生命刑、自由刑等，不影响对公司定罪，刑法规定由主管人员和直接责任人替代公司受罚。对人工智能系统定罪后，也可以由其提供者、使用者、所有者替代受罚。即使人工智能系统不具有可罚性，从结果主义的立场看，对其定罪也有意义，能够表明刑法对人工智能犯罪的严厉谴责。笔者认为，刑法规定对单位犯罪的刑罚中，对公司主管人员和直接责任人员的处罚，不应被定性为替代公司受罚，以上人员与公司的关系也不同于人工智能系统与其提供者、使用者、所有者的关系，因为前者的行为和心态产生公司的意志，而后者则无此决定能力。而且，无刑事责任的定罪违反罪责刑相适应原则，不能实现维护公正和预防犯罪的目的。

2. 妨害人工智能安全管理秩序犯罪说

人工智能系统是新的工具和治理对象，人工智能安全是人工智能时代的新型公共安全。为了保护人工智能安全，我国正在制定人工智能安全管理制度，建立人工智能安全管理秩序，要求相关人员遵守人工智能安全行为规范，防控人工智能安全风险。为了惩治严重侵犯人工智能安全的行为，我国应当设立妨害人工智能安全管理秩序犯罪，对严重违反前述行为规范的行为追究刑事责任。妨害人工

智能安全管理秩序犯罪包括滥用人工智能犯罪和违反人工智能安全管理义务犯罪。前者不同于传统犯罪的智能化,侵犯的不是传统法益而是新型利益,如滥用人工智能操纵自然人以损害其合法权益、对自然人进行生物识别分类、不公正的社会评分、风险分析、人脸采集、情绪分析等行为,不构成现行刑法中的犯罪,却严重损害公众的自由、生活、工作等利益。后者是为了防控人工智能安全风险而设置的犯罪。人工智能系统具有自主、智能特性,传统的产品安全保障制度不能控制其安全风险。人工智能安全风险源于研发并沿着其价值实现链扩散,人工智能研发、生产、提供、部署、终端使用各方对风险形成、扩散有原因力,且具有直接管控风险的能力,应承担安全管理义务。设立义务犯罪可以督促相关人员切实履行前述义务,防控严重风险,保障人工智能安全。人工智能安全与网络安全、数据安全是具有相似信息技术特性的新型集体法益,我国《刑法》将危害网络安全的犯罪规定在“妨害社会管理秩序罪”一章的“扰乱公共秩序罪”一节中,设立了拒不履行信息网络安全管理义务罪等妨害信息网络安全管理秩序犯罪,妨害人工智能安全管理秩序犯罪侵犯的是同类客体,都是现代社会环境下的新型公共秩序。

该主张具有合理性。首先,该主张与国际社会治理人工智能安全的方向一致。《全球人工智能治理倡议》要求“明确人工智能相关主体的责任和权力边界”,“确保人工智能始终处于人类控制之下”。欧盟《人工智能法》将人工智能系统定性为自主性工具,建立了类型化、全过程、多方协作的安全风险防控制度。妨害人工智能安全管理秩序犯罪立法的立场和路径与前二者一致。其次,该主张客观反映了人工智能危害的形成与防控规律。无论是采取何种技术方法、何种自主性程度的人工智能系统,提供者和部署者对于系统导致的危害或风险有引起或扩散作用,如果具备风险管理能力,理应承担防控责任,设置以上犯罪能够督促其履行防控义务,有效降低安全风险,符合人工智能危害的形成与防控规律。最后,该主张能够兼顾保护法益和保

障人工智能发展。该立法主张依据风险防控理论,兼顾人工智能安全与发展,既能有效降低安全风险,又按风险高低分类设置行为规范,采取基于风险的差别化刑法应对,能够最大限度地保障人工智能发展。

在妨害人工智能安全管理秩序犯罪中,滥用人工智能系统犯罪和违反人工智能安全管理义务犯罪是两类不同的犯罪,后者可以遏制第三类人工智能危害,是本文主要研究的内容。

二、人工智能安全管理义务的性质与类型

人工智能安全相关人员履行安全管理义务能够有效降低风险,违反人工智能安全管理义务的犯罪立法能够惩治严重违反安全管理义务的行为,为人工智能安全管理制度提供刑法保障。人工智能安全管理义务是该类犯罪的构成要件要素,也是国内外人工智能安全立法的重要内容,有必要研究其性质与类型。

(一) 人工智能安全管理义务的性质

设置人工智能安全管理义务是为了防控安全风险,有必要明确人工智能安全风险的内涵。风险是指“与未来损失有关的不确定性”,预防风险就是要降低损失发生的可能性或损失的程度。欧盟《人工智能法》第3条第2款将风险定义为“发生危害的可能性和危害的严重性的组合”,与危害发生的概率和危害的严重性正相关。风险是可能发生的危害,如果某活动能够确定造成危害后果,其引起的不是风险而是确定的危害,应当被制止。风险意味着蕴含利益的机会,只能采取措施控制风险,而不应禁止有风险的活动,否则会失去获取利益的机会。人工智能安全风险可以定义为人工智能研发、生产、提供、应用过程中发生危害的可能性和危害的严重性的组合,具有全领域渗透性、不可预见性、网络传播性等特点,能够在社会各领域造成难以有效控制的严重后果。防控传统风险的法律制度因不适应其特点而难以有效控制,不能为社会公共安全和公众合法权益提供充分保障,有必要建立新的风险防控制度。

现行刑法不能规制前述第三类危害的原因,是

相关人员的行为与危害后果之间不存在刑法上的因果关系,没有因果关系就无法适用现行刑法规定的犯罪。由于人工智能系统自主、直接引起危害后果,故其研发、生产、提供和使用方的行为与危害后果之间没有因果关系。以机器学习算法为例,该算法基于训练数据中的相关性,如 ChatGPT 的算法逻辑是基于数以亿计的参数,对海量训练数据加权统计后,形成复杂的关联关系。从最初的算法设计者到终端用户,各方行为对算法的生成和演进不起决定性作用,不能评价为原因行为。由于欠缺因果关系,依照现行法律不能要求前述人员对危害后果承担法律责任。如果将因果关系泛化为相关性关系,则与条件说无异,实际上是“舍弃因果关系要件设定针对违法行为本身的责任,以此跳过因果联结及可追溯性难题”,将导致人工智能安全风险防控主体范围的过度扩展,动摇现行法律责任制度的理论基础。因此,“以人的活动为中心”的社会环境下发展起来的因果关系理论,在智慧社会环境下无法为解决人工智能危害的法律治理困难提供理论指导。只有改变按结果犯、支配犯治理的思维习惯,在人工智能安全治理领域完成基于风险的、义务导向的归责模式转向,才能摆脱当前治理困境。

对人工智能安全相关人员施加安全管理义务,在当前情形下是治理前述第三类人工智能危害唯一可行的选择,规范责任理论可以为提供理论依据。从现实情况来看,要有效降低人工智能安全风险,只能通过人工智能安全相关人员履行安全管理义务来实现。从法律规范看,设立安全管理义务是人工智能安全保障体系发展的必然结果。为了保障人工智能安全,国家必然要制定人工智能安全管理制度,设定人工智能安全相关人员的行为规范。人工智能安全管理义务是对其设定的特殊规范责任,对有能力遵守而故意违反行为规范的行为人追究法律责任。刑法学中的规范责任理论可以为以上法律制度提供理论支持,其主张承担刑事责任的根据是“违反义务实施违法行为的决意”,包括责任能力、故意或过失以及“行为人存在适法行为的期待可能性”,违反规范义务就应承担相应的法律责任。

不过,人工智能安全管理义务不同于传统的规范责任,是有限的规范责任。其原因主要是两个方面:其一,因人工智能系统自主性引起的危害后果与前述人员行为之间只有非直接的、弱相关性关系,而不是因果关系,依照现行法律其不应应对危害后果承担法律责任。从义务规范角度对义务承担者施加的责任,也应与其影响风险的作用和防控能力相适应,较之于传统的规范责任要有所缩减,以区别于传统义务犯的义务。其二,人工智能安全风险沿着价值链扩展,从系统提供者、使用者到终端用户,各方只面对其所在阶段上的风险,按照风险与注意义务分配的对应关系,只应对部分风险承担防控责任。对于应追究法律责任的危险或损害而言,各方承担的风险防控责任,既要与其所在阶段上的风险大小成比例,也要与其地位和防控能力成比例,否则违反比例原则中的均衡性原则。因此,人工智能安全管理义务应当是为防控风险所必要的、适度且有限的规范责任,与相关人员引起或扩散风险所起的作用或防控风险的能力成正比。

人工智能安全管理义务与信息网络安全管理义务都是安全管理义务,但二者有本质区别:(1)法律性质不同。前者是对“物的安全”的管理义务,防控作为工具的人工智能系统自主运行引起的安全风险,近似于生产作业安全管理义务。后者是针对“人的活动”的监督管理义务,防止他人实施违法犯罪。有论者在讨论生成式人工智能产品提供者的安全管理义务时,将该义务限于提供者对他人利用其产品进行违法犯罪活动的管理义务,即后者义务。生成式人工智能服务提供者同时也是网络信息服务的提供者,当然应承担信息网络安全管理义务,但是,该义务不同于人工智能安全管理义务。

(2)义务根据不同。后者的设立根据是网络服务提供者的保证人地位,即“保护他人法益不受来自于自己控制领域的危险威胁的义务。这种对于危险源的控制是不作为犯的义务”。前者的设立根据是人工智能安全相关人员引起或扩散人工智能安全风险的客观事实和保护社会安全的需要,以上人员并不具有保证人地位。(3)保护法益不同。前者

直接保护的是人工智能安全,间接保护的权益范围广泛,如公共安全和公众合法权益等,而后者保护的是公共信息网络安全管理秩序。由于以上区别,不对违反前者义务的行为适用拒不履行信息网络安全管理义务罪。

(二) 人工智能安全管理义务的类型

国内外人工智能安全立法都规定了人工智能安全管理义务。欧盟《人工智能法》规定了基于风险的人工智能安全管理义务制度,针对人工智能系统和通用人工智能模型两类防控对象,将义务承担者分为提供者、部署者、欧盟内授权代表、进口者、经销者等,从防控对象和义务承担者两个维度规定了不同的义务类型。美国《人工智能倡议法》提出了发展“可信的人工智能系统”,算法应满足“可解释性”“能够识别和缓解人工智能系统中偏见的分析方法”“具有安全、稳健性”等要求,美国联邦政府还发布了《推进国家安全领域人工智能治理和风险管理的框架》,规定了基于风险的人工智能应用分类管理制度。近年来我国制定了《互联网信息服务深度合成管理规定》《生成式人工智能服务管理暂行办法》《人脸识别技术应用安全管理办法》《人工智能生成合成内容标识办法》等规章,规定了人工智能生成内容信息和人脸识别技术应用相关安全管理义务。对比以上规章,欧盟《人工智能法》规定了全面、完整的安全管理义务体系,我国的规定与之相似。以下主要将欧盟《人工智能法》规定的义务类型与我国相关规定进行比较。

1. 人工智能系统安全管理义务

欧盟《人工智能法》将人工智能系统分为高风险人工智能系统、非高风险的特定人工智能系统和其他人工智能系统,对前两种系统的相关人员规定了不同的安全管理义务,鼓励和促进制定针对第三类人工智能系统的行为守则,并提倡相关人员自愿遵守该法的要求(第95条)。

(1) 高风险人工智能系统安全管理义务

高风险人工智能系统是指对人的健康、安全或基本权利造成伤害的风险高,且用于该法规定领域的人工智能系统(第6条第1—3款),是欧盟《人

工智能法》的主要防控对象。该类系统的提供者、部署者及其他方被要求承担较多的安全管理义务。

为了确保高风险人工智能系统投放市场或投入使用不会对欧盟法律所认可和保护的重要公共利益构成不可接受的风险(引言第46条),欧盟《人工智能法》要求高风险人工智能系统符合以下7项要求:(1)建立实施、记录和维护风险管理系统;(2)数据治理;(3)提供技术文件;(4)自动记录日志;(5)保持透明度;(6)适合人工监督;(7)保持系统的准确性、鲁棒性和网络安全(第8—15条)。此类系统提供者的义务是确保高风险人工智能系统符合这些要求,此外,提供者还要承担以下6项义务:(8)遵守一般管理程序要求;(9)建立质量管理体系;(10)制作、保存和随时向国家主管机构提供该法第18条规定的文件;(11)保存日志数据;(12)对不符合要求的高风险人工智能系统采取必要的纠正措施;(13)与主管机构合作(第16—21条)。在以上义务中,对降低安全风险发挥实质性作用的是(1)(2)(5)(6)(7)(12)项义务,其它义务是监管程序性义务。在我国前述规章中,仅《人脸识别技术应用安全管理办法》中的人脸识别设备符合欧盟《人工智能法》规定的高风险人工智能系统。但该办法不适用于“为从事人脸识别技术研发、算法训练活动应用人脸识别技术处理人脸信息”的行为,没有规定人脸识别技术系统提供者的安全管理义务。如果系统提供者同时又是系统使用者,则适用该办法,其地位相当于欧盟《人工智能法》中的高风险人工智能系统的部署者。

高风险人工智能系统的安全风险既源于系统设计,也来自系统使用。部署者对系统使用环境、可能受影响的人或群体包括弱势群体有更准确的认知,更了解高风险人工智能系统的具体使用情况,能够发现开发阶段未预见的潜在重大风险,提供关于危险的信息,这对弥补提供者的义务缺漏、防控安全风险能够起到关键作用。欧盟《人工智能法》对部署者规定了以下12项义务:(1)遵守使用说明;(2)安排适格人员进行人工监督;(3)

控制数据输入以符合预期目的；（4）监控系统运行；（5）保存自动生成的日志；（6）告知受影响的雇员；（7）注册登记；（8）评估数据保护状况；（9）对特定对象遵法使用系统；（10）告知系统的使用对象；（11）与相关国家主管机构合作等；（12）评估高风险人工智能系统可能对基本权利的影响（第26、27条）。其中，第（1）（2）（3）（4）（6）（9）（10）项义务是在应用阶段的、能够实质性降低安全风险的义务，与提供者在研发阶段的义务相互补充、配合。我国《人脸识别技术应用安全管理办法》中，对应于欧盟《人工智能法》部署者义务的是告知和个人信息保护影响评估义务，没有规定专门的人工智能安全管理义务。人脸识别系统提供者的前述行为涉及敏感个人信息的处理，受《个人信息保护法》《民法典》等法律法规的规制。

高风险人工智能系统风险与人工智能价值链上多方行为紧密关联，欧盟《人工智能法》还对高风险人工智能系统进口者、经销者以及在欧盟境外的高风险人工智能系统提供者的授权代表规定了安全管理义务。这三类主体使高风险人工智能系统从提供者（包括境外提供者）流转至部署者，对风险扩散有一定作用，承担的主要是监督提供者和部署者、协助风险管理等监管程序性义务。我国没有类似的规定。

（2）非高风险的特定人工智能系统安全管理义务

人工智能系统生成合成内容信息越来越逼真，其社会化应用增加了虚假信息、社会操纵、欺诈、假冒等新风险。为了防控此类风险，欧盟《人工智能法》要求非高风险的特定人工智能系统相关方承担透明度义务。一是与自然人直接互动的人工智能系统的透明度义务。该类系统提供者应确保系统在设计 and 开发时，能够使与系统互动的自然人知晓其面对的是人工智能系统；但根据当时的环境和使用情况，从合理的知情、观察和谨慎的自然人的角度看这一点是非常明显的除外。生成合成音频、图像、视频或文本内容的人工智能系统（包括通用人工智

能系统）的提供者应当确保，人工智能系统的输出以机器可读的格式、且可检测为人工生成或操作的方式进行标记，该系统的部署者应当披露该内容是人工生成或操纵的。二是情绪识别系统或生物识别分类系统的透明度义务。该类系统的部署者应当告知自然人其是系统的识别对象，并按照欧盟相关法规和指令处理个人数据（第50条第1—4款）。如果以上系统符合高风险人工智能系统的特征，须遵守高风险人工智能系统相关义务。遵守以上透明度义务不应被解释为使用该类系统或其输出是合法的，相关方仍须遵守欧盟及成员国相关法律规定，如个人数据保护义务（第50条第6款）。

我国前述规章对互联网信息服务提供者和人脸识别设备使用者规定了类似的透明度义务，对应于欧盟《人工智能法》中与自然人直接互动的人工智能系统部署者的透明度义务，但是在义务的内容、范围、义务主体等方面超过后者，如将义务主体扩大到提供网络信息内容传播服务的提供者、互联网应用程序分发平台以及网络用户，使下游环节的更多主体参与风险防控。

2. 通用人工智能模型安全管理义务

欧盟《人工智能法》对通用人工智能模型的提供者及其授权代表规定了安全管理义务。通用人工智能模型是人工智能系统的核心组成部分，处于下游的系统提供者只有充分了解模型，才能将其集成到系统中（引言第88条）。根据高影响能力标准（第51条），该法将通用人工智能模型分为有无系统性风险的两类模型，要求通用人工智能模型提供者承担提供技术文件、向下游人工智能系统提供者提供技术信息、制定合规政策、公开通用人工智能模型训练内容4项义务（第53条第1款），有系统性风险的通用人工智能模型的提供者还应承担模型评估、识别并减轻系统性风险、跟踪、记录和报告严重事件、网络安全保护4项义务（第55条第1款），欧盟境外设立的人工智能系统模型提供者的授权代表应承担制定、保存和提供技术文件与提供者的联系方式、同主管机构合作等义务（第54条第3款）。在以上义务中，只有识别和减轻系

统性风险、网络安全保护义务属于实质性降低风险的义务，其他都属于监管程序性义务。不同于人工智能系统风险的现实性，通用人工智能模型在被提供给下游单位时，因为尚无具体的应用目标而不存在实际风险，其基于高影响能力的系统性风险需要通过系统应用才能被现实化，只是因其可能大量应用而产生重大负面影响，才需要提前采取防控措施。我国没有针对通用人工智能模型规定安全管理义务。

综上，人工智能安全管理义务是为了防控人工智能安全风险而对人工智能安全相关人员施加的有限的规范责任。我国仅对人工智能信息服务的提供者及相关方规定了安全管理义务。近年来，我国智能网联车辆、无人机等高风险人工智能系统应用不时引起危害后果，有必要推进人工智能安全管理义务相关立法，依法追究严重违反安全管理义务行为的法律责任，包括刑事责任。

三、人工智能安全管理义务承担者的刑事责任

目前我国人工智能安全立法适用范围窄，限于互联网信息服务领域的人工智能应用，法律位阶低，强制力弱，难以防控严重的和其他领域的人工智能安全风险。《互联网信息服务深度合成管理规定》等规章规定，违反其规定“构成犯罪的，依法追究刑事责任”，而现行刑法没有与违反人工智能安全管理义务犯罪相关的立法。人工智能安全是智慧社会发展的基础，应当给予充分的法律保障，有必要研究严重违反人工智能安全管理义务行为的刑事责任问题。

（一）违反人工智能安全管理义务犯罪立法的正当性根据

违反人工智能安全管理义务犯罪立法具有必要性。人工智能引领新一轮科技革命和产业变革，对国家安全、经济社会发展和公众的福祉至关重要。《全球人工智能治理倡议》提出，发展人工智能应“以保障社会安全、尊重人类权益为前提，确保人工智能始终朝着有利于人类文明进步的方向发展”，要求发展“可审核、可监督、可追溯、可信赖的人工智能技术”。《中共中央关于进一步全

面深化改革、推进中国式现代化的决定》要求，“必须全面贯彻总体国家安全观，完善维护国家安全体制机制，实现高质量发展和高水平安全良性互动”。人工智能安全是总体国家安全的重要组成部分，保障人工智能安全既不能“唯安全”，也不能“唯发展”，应以高水平安全来保障人工智能高质量发展。刑法应当贯彻前述政策，包容人工智能高质量发展，对安全风险采取底线控制，只要前置法能够防范安全风险失控，刑法就应坚守谦抑立场，坚持以行政监管为主、刑法威慑为辅。但是，如果人工智能的研发、生产、提供或使用危害国家安全、公共利益和公众合法权益，妨害人工智能安全管理秩序，那么其不仅不受刑法保护，还可能受到刑法的惩罚。

违反人工智能安全管理义务犯罪立法具有合理性。德国学者 Jakobs 教授将正犯分为支配犯和义务犯，支配犯因侵犯法益而具有可罚性，义务犯是违反为了建设良好的世界而设定的“制度管辖之义务”而成立犯罪。人工智能安全是智慧社会的基础，国家必然要建立人工智能安全管理制度，包括人工智能安全管理义务。为了保障以上制度的实施，遏制严重破坏人工智能安全管理秩序的行为，有必要设立违反人工智能安全管理义务犯罪，该类犯罪是违反人工智能安全管理制度所设定义务的义务犯。

违反人工智能安全管理义务犯罪是特殊的义务犯。德国学者 Roxin 教授认为义务犯具有对法益的“保护性控制”能力，Schünemann、西田典之等教授认为，义务犯居于保证人地位，甚至将排他性支配作为犯罪成立的条件。例如，我国刑法中的拒不履行信息网络安全管理义务罪是义务犯，网络服务提供者创造网络空间，制定网络活动规则，有能力控制网络空间行为，居于排他性支配的保证人地位。违反人工智能安全管理义务犯罪与之不同，其义务承担者不具有此类排他性支配能力。由于人工智能的自主性及其风险控制的特殊性，人工智能安全相关人员仅具有不完全的保护性影响能力：（1）由于人工智能的自主性和不可预测性，人工智能能

法的形成和演变不完全受控于其研发者、提供者、使用者及最终用户，人工智能系统的自主性越高，就越不受控制；（2）人工智能安全风险的形成和扩散分散于生产、流通、应用等多个阶段，各阶段上的义务承担者只能对本阶段的安全风险发挥一定的防控作用，不能对风险进行排他性支配控制，其严格履行义务只能阶段性地、部分地、有限地贡献于人工智能安全。同样，其违反义务行为也只能阶段性地、部分地、有限地损害人工智能安全。按照风险分配原理和比例原则，只应承担有限的责任。因此，违反人工智能安全管理义务罪应有别于其他义务犯，有必要设置较高的入罪门槛和较轻的刑事责任，只将社会危害性严重的违反义务行为犯罪化，使该罪立法满足正当性要求。

（二）违反人工智能安全管理义务罪的罪行构造

设立违反人工智能安全管理义务罪应当反映人工智能的特性，充分保障公民自由和合法权利，保障人工智能发展，有效防控人工智能安全风险。该罪是行政犯、义务犯和情节犯，应当审慎设置该罪的构成要件，除了必须具有严重的犯罪情节外，关键是合理规定人工智能安全管理刑法义务、主体和罪过要件。

1. 人工智能安全管理刑法义务的范围

该罪是义务犯，合理规定人工智能安全管理刑法义务对该罪具有重要意义。德国学者 Roxin 教授认为，义务犯是违反了构成要件之前的、刑法之外的特别义务的行为，公法、行业法规乃至民法上的义务都可以成为义务犯的特别义务。批评者认为，该观点是不合理的形式法律义务理论，因为其他法律不以保护法益为目的，刑法义务只应从自己的任务中产生。后来 Roxin 教授将义务犯的特别义务限定为“与结果有关联地违反构成要件上的特别的义务”或“违反决定正犯性的义务”。但是，如果将特别义务限定为与结果有关联，则必然将义务犯的范围缩小为结果犯。另外，所谓“决定正犯性”义务的限定则仍然没有明确特别义务的范围，反而有循环论证的逻辑谬误——义务犯以确定特别义

务为前提，而特别义务的确定以成立义务犯为条件。德国学者 Jakobs 教授认为，义务犯的积极义务是团结义务，制度是强化团结的基础，违反团结义务决定义务犯的正犯性和可罚性。团结义务的范围并不从属于公法、民法等刑法之外的部门法规定的义务，而是与后者一样都来自于整体法秩序，刑法和其他部门法虽然有各自不同的任务，但都服务于维护整体法秩序。但是，刑法自身任务的独立性和独特性决定了义务犯的特别义务或积极义务不同于刑法之外的其他部门法规定的义务。Roxin 教授认为义务犯要回到法益侵害来确定正犯性，而 Jakobs 则认为犯罪的实质内容本来就是规范违反，导致二者对义务犯及其义务范围有不同的主张。在德国刑法和我国刑法中，义务犯并非限于直接侵犯个人权益的犯罪，还包括侵害秩序法益的义务犯，如玩忽职守罪等。如果将集体法益排除在义务犯侵害的法益之外，则难以解释后一种义务犯。而按照 Jakobs 教授的规范违反论则难以区分其他部门法规定的义务。笔者认为，为了解决以上问题，应当回到刑法的任务，从义务犯设立目的出发讨论义务犯的义务范围。

设立违反人工智能安全管理义务罪的目的是保障人工智能安全发展，保护国家安全、公共安全和公众合法权益，应以此为出发点来确定人工智能安全管理刑法义务的范围。我国正在推进的人工智能立法和其他法律法规必然会规定安全管理义务，其中部分义务能成为刑法上的义务，违反此类义务的行为应当具有刑事可罚性。关于义务犯的刑事可罚性根据，德国刑法理论界和司法界主要采取风险降低理论，即“未实施行为的实施本来会（显著地）阻止结果发生的风险”，如果这种“降低风险的可能性”不仅事前能够判断，事后也可以确定，义务违反行为就具有刑事可罚性。国内外人工智能安全立法规定了多种人工智能安全管理义务，例如，欧盟《人工智能法》规定了三类人工智能安全管理义务，即防控人工智能技术风险的义务、保障自然人基本权利相关的义务、监管程序要求的义务。我国前述规章也规定了多种安全管理义务。违反这些义

务并非都具有刑事可罚性,判断以上各类义务违反行为是否具有刑事可罚性,评价标准应当是义务履行行为能否实质性、显著降低风险。遵守监管程序要求的义务并不具有直接降低风险的功能,违反该类义务不具有刑事可罚性,无论情节是否严重,都不应追究刑事责任。违反保障自然人基本权利相关义务的行为,能降低自然人权益遭受损害的风险,具有刑事可罚性。但是,其已受侵犯公民个人信息犯罪或其他侵犯公民人身权利、民主权利犯罪立法独立规制,没有必要将其规定为违反人工智能安全管理义务犯罪。防控人工智能技术风险义务的履行能够实质性降低人工智能安全风险,如果在具体的场景下该类义务的违反与危害之间能建立起合法的连接,不仅能事前判断也能事后确定,该义务违反就具有刑事可罚性。

能够实质性降低严重人工智能安全风险的义务可以成为人工智能安全管理刑法义务。欧盟《人工智能法》将人工智能安全风险类别化,针对不同类型和不同程度的人工智能安全风险设置不同的义务,值得我国相关立法借鉴。不过,该法规定的部分义务也有不合理之处。例如,该法针对非高风险的特定人工智能系统和通用人工智能模型,只规定透明度义务和其他更低约束力的义务,风险防控功效存疑。以生成式人工智能系统的安全管理义务为例,该类系统的提供者、部署者履行了标识和提示义务后,公众并非不受系统输出的虚假有害信息影响,仍然会有大量用户因为缺乏知识和经验以及对人工智能技术和大企业的盲目信任,受到系统提供的虚假有害信息欺骗而遭受严重损害,如佛罗里达14岁少年受谷歌聊天机器人诱导后自杀。在当前网络环境下,生成式人工智能系统的训练数据供给并非机会均等,对某个国家、民族、群体的数据歧视甚至敌视客观存在。如果对生成式人工智能系统只规定透明度义务,将无法有效防控其提供的违法信息给我国国家安全、公共安全和秩序以及公众合法权益造成严重危害的风险。为了防控以上严重安全风险,我国对生成式人工智能系统的应用设置了风险防控义务。《生成式人工智能服务管理暂行

办法》要求,提供者发现系统生成、传输违法内容的,“应当及时采取停止生成、停止传输、消除等处置措施,采取模型优化训练等措施进行整改,并向有关主管部门报告”。该义务履行是控制生成式人工智能系统生成、传播违法信息风险实现和扩散的最后屏障,有必要将其确定为安全管理刑法义务,以刑法手段督促义务承担者切实履行。

基于以上分析,笔者建议,我国应通过制定人工智能相关法律法规,构建系统化的人工智能安全管理义务体系,将其中能够实质性降低严重安全风险的义务规定为人工智能安全管理刑法义务。该类义务的范围为:(1)高风险人工智能系统、信息服务类、情绪识别和生物识别分类人工智能系统的提供者,应当承担监控与降低安全风险、数据治理、算法治理、保持透明度、确保系统适合人工监督、确保系统的准确性和稳健性及网络安全的义务;(2)以上系统的部署者应当承担保持透明度、遵守使用说明、安排适合人员监督、监控和降低风险的义务。

2. 主体范围

违反人工智能安全管理义务罪的犯罪主体是人工智能安全管理义务的承担者,但是,该类义务的承担者并非都是该罪的主体。根据欧盟《人工智能法》的规定,高风险人工智能系统、非高风险的特定人工智能系统的提供者和部署者以及通用人工智能模型的提供者是主要的义务主体,出于风险防控的需要,高风险人工智能系统进口者、经销者以及欧盟境外提供者的授权代表被规定为协助监督、信息提供及与监管机构合作等义务的主体。我国《互联网信息服务深度合成管理规定》和《生成式人工智能服务管理暂行办法》将服务提供者规定为主要的义务主体,《互联网信息服务深度合成管理规定》还要求技术支持者承担数据管理、个人信息告知同意、算法安全审核与评估等义务(第14、15条),以上义务主体的范围比欧盟《人工智能法》更狭窄。以上主体并非都能成为违反人工智能安全管理义务罪的主体,成立该罪主体应当满足两个条件:一是承担人工智能安全管理刑法义务;二是其

刑法义务的履行在客观上能够实质性、显著降低安全风险或者阻止风险实现,且因其义务违反导致风险的升高或实现,严重威胁人工智能安全及其他关联法益。否则,其义务违反行为就不满足刑事可罚性,只承担行政法律责任。

基于前述分析,笔者认为,违反人工智能安全管理义务罪的犯罪主体应当限于高风险人工智能系统及其他能够造成严重危害的人工智能系统的提供者和部署者。应注意,我国前述规章规定的服务提供者与欧盟《人工智能法》规定的人工智能系统或模型提供者不同,兼有后者规定的部署者地位,这里按后者规定的提供者和部署者来分析。人工智能系统的提供者是安全风险的引起者,主要承担人工智能研发阶段的风险控制义务,是从源头控制风险。而部署者承担应用阶段的风险控制义务,该阶段是风险实现阶段,也是阻止危害后果发生的关键环节,其义务履行对阻止风险实现的作用最大。二者都是主要的义务承担者,具有实际的风险防控能力,其义务履行对降低安全风险起关键作用,可以成立该罪的犯罪主体。

人工智能的快速发展应用对国家竞争力、国家安全和经济社会发展具有重要意义,包容审慎成为人工智能法治理念,人工智能刑事法治的首要任务应当是保障先进人工智能系统的研发和防范人工智能安全风险失控。为实现这个目标,对待以上两类主体要有所区别。部署者具有直接管控风险的能力(除非提供者未能提供有效的人工监管措施),在降低人工智能安全风险上具有“近因”作用,对其义务违反导致的风险升高或失控具有预见能力,相比于提供者更符合归责条件。因此,对于部署者违反刑法义务的行为,应更多考虑保障安全的需要,合理认定其刑事可罚性。而提供者对促进先进人工智能发展的推动作用更大,应当更多地考虑促进先进人工智能发展的需要,保持包容的立场。当然,如果提供者兼有部署者身份,应当单独评价其作为部署者实施的行为。

人工智能系统的进口者、经销者、生产者、提供者的授权代表,以及通用人工智能模型的提供者

和部署者不直接引起、扩大和实现风险,也不具有实质性降低风险的能力,一般情况下不应作为该罪的犯罪主体。欧盟《人工智能法》对以上主体规定的主要是监管程序性或预防性义务,这些义务履行只起到有限、间接降低风险的作用。我国前述规章未对这类主体规定安全管理义务,如果未来我国人工智能安全立法对其规定安全管理义务,对违反该类义务的行为只应追究行政法律责任。

3. 罪过形式

义务犯的罪过形式可以是故意或过失。欧盟《人工智能法》处罚故意或过失违反该法规定义务的行为(第99条第7款第i项)。违反人工智能安全管理义务罪的罪过形式是否包括故意和过失,影响该罪的处罚范围。笔者认为,故意违反人工智能安全管理刑法义务,且满足其他构成要件的,可以按该罪追究刑事责任;如果发生严重危害后果,对该后果无论是出于过失还是故意,都可以成立该罪,因为该罪处罚的是义务违反行为,因违反义务而未能阻止严重后果,即表明其行为具有严重的社会危害性;如果行为人对危害后果持希望或放任态度的,如自动驾驶汽车提供者在紧急路况处置算法中设置歧视性算法,使车辆选择性撞向行人,部署者故意违反自动驾驶汽车人工监管义务,放任车辆撞死路人或乘客的,可以同时构成故意杀人罪、以危险方法危害公共安全罪等故意犯罪,应按想象竞合犯从一重罪定罪,从重处罚。

行为人应当知道会违反以上义务,但因疏忽大意而不知道有此义务且未履行,如果将该种心态作为该罪的罪过形态,将扩大该罪的处罚范围,其合理性值得讨论。人工智能安全是独立的新公共法益,同时关联着国家安全等重大法益,已成为总体国家安全的新组成部分并深度嵌入公共安全,应属于重大法益,为其提供刑法保护具有合理性。既然《刑法》处罚过失违反国家秘密、军事秘密保密义务的行为,也可以将违反人工智能安全管理义务犯罪扩展到过失犯罪。但是,从保障人工智能发展的角度看,这样对人工智能相关方处罚过重,可能阻碍人工智能的研发、生产和应用,不符合促进人工

智能快速发展应用的国家战略。

基于对人工智能安全与发展的考虑，笔者认为，该罪应当限于故意犯，可以通过协调人工智能安全相关行政、刑事立法和执法，消除对过失违反义务行为追究刑事责任的可能性。可以通过加强人工智能安全监管执法，督促义务承担者履行安全管理义务，减少或排除行为人过失违反义务的可能性，因而没有必要将极少发生的过失违反义务行为设定为犯罪。如果行为人曾经因过失违反人工智能安全管理刑法义务受到行政处罚，可以推定行为

人知悉应当履行的义务，其再次违反义务就只可能是出于故意，此时刑法所处罚的是故意违反义务行为，而非过失违反行为。在罪状设计上，可以参考拒不履行信息网络安全管理义务罪，将该罪的行为规定为“拒不履行安全管理义务，情节严重的行为”，并通过增加监管机关责令整改的程序，排除行为人过失违反义务的情形。

（技术编辑：邓语鑫）