

人工智能法律政策图景研究报告 (2025 年)

中国信息通信研究院政策与经济研究所

2025年6月

版权声明

本报告版权属于中国信息通信研究院，并受法律保护。
转载、摘编或利用其它方式使用本报告文字或者观点的，
应注明“来源：中国信息通信研究院”。违反上述声明者，
本院将追究其相关法律责任。

前 言

习近平总书记指出：“人工智能是新一轮科技革命和产业变革的重要驱动力量，将对全球经济社会发展和人类文明进步产生深远影响。”当前，人工智能技术持续突破，创新产品快速涌现，正在以前所未有的速度、广度和深度改变生产生活方式。同时，人工智能发展也带来了隐私泄露、算法偏见、数据安全等一系列风险。在此背景下，世界主要国家和地区相继出台人工智能法律政策，加快相关立法与战略布局。国际组织也将人工智能发展与安全纳入核心议题，与各国政府共同构成全球人工智能治理的重要推动力量。

全球人工智能法治化进程不断提速。欧盟首部综合性人工智能立法引领全球，加快布局法案的实施与监督。美国进入特朗普 2.0 时代，更加强调维护自身产业技术领导地位，呈现人工智能立法监管放缓特点。中国发挥政策法规保障人工智能健康发展作用，有序推进人工智能相关立法工作。韩国、日本、巴西等结合产业实际和发展优势，也纷纷加快战略部署和立法进程，致力于打造具有本国特色的人工智能规范路径。此外，联合国等国际组织与平台加大对人工智能议题的关注，双多边治理成果不断丰富。

全球人工智能法律政策呈现新趋势。从总体格局上看，全球范围内初步形成以美欧为代表的两大主导模式，新加坡、马来西亚、泰国、沙特阿拉伯、阿联酋等新兴经济体加入治理行列，持续提升自身影响力。从政策实践来看，当前全球人工智能治理进入软硬法

接轨并行、综合立法和领域立法齐头并进的治理新阶段，促创新促发展成为全球政策法规的竞争重点，生成式人工智能、前沿大模型和先进算力的关注度与日俱增。

展望未来，伴随人工智能的新技术不断突破、新业态持续涌现、新应用加快拓展，人工智能立法将更趋理性且更加全面系统，人工智能风险治理与伦理考量不断深化。与此同时，面对渐趋白热化的全球人工智能竞争态势，各国支持人工智能创新与发展的政策力度或将只增不减。单边主义、地缘冲突、机制分歧等不利因素制约全球人工智能治理合作的广度与深度。打造一个全方位、多层次、汇聚广泛共识，具有真正的包容性、平等性、多元性的全球性治理框架仍是未来各国合作努力的重要方向。

目 录

一、全球人工智能法律政策图景	1
（一）欧盟：首部综合性人工智能立法全球瞩目	1
（二）美国：聚力推进基于创新应用的治理范式	6
（三）中国：以政策法规保障人工智能健康发展	13
（四）其他国家和地区：加快人工智能立法与战略布局	18
（五）国际组织：全球人工智能治理的重要推动力量	24
二、全球人工智能法律政策趋势动向	29
（一）人工智能治理总格局渐趋清晰，新兴经济体加入治理行列	29
（二）促创新和促发展成为竞争重点，以专设机构提升政策质效	32
（三）从软法规制到软硬法接轨并行，人工智能法治化进程加快	34
（四）综合立法和领域立法齐头并进，生成式人工智能关注剧增	35
（五）风险分级和分类规制不断完善，前沿大模型成为规范重点	37
三、全球人工智能法律政策未来展望	39
（一）人工智能立法更趋理性且更加全面系统	39
（二）人工智能风险治理与伦理考量不断深化	40
（三）人工智能创新与发展政策力度持续加大	41
（四）全球人工智能治理合作面临多方面挑战	42

一、全球人工智能法律政策图景

伴随人工智能技术的快速突破与变革性影响，世界各国竞相发力，加大人工智能立法和监管力度，各类法律政策不断涌现。截至 2024 年底，全球已有 69 个国家和地区制定了人工智能相关政策和立法，涵盖了从算法治理、隐私保护、数据监管到伦理规范等广泛主题，反映了各国对人工智能技术发展的重视及其潜在风险管理的关注。与此同时，国际组织也将人工智能发展与安全纳入核心议题，与各国政府共同构成全球人工智能治理的重要推动力量。¹

（一）欧盟：首部综合性人工智能立法全球瞩目

欧盟在人工智能方面布局较早，从科技伦理向法律监管稳步推进，旨在引领全球人工智能监管规则走向，从欧洲整体层面出发，建构协调一致的治理规则体系。

1. 出台全球首部综合性人工智能立法

欧盟《人工智能法》作为首部综合性人工智能监管立法受到全球关注。欧盟《人工智能法》由欧盟委员会于 2021 年提出，历经多轮立法修订和协商，于 2024 年 5 月 21 日获批通过，并于 2024 年 8 月 1 日正式生效。该法作为首部综合性人工智能监管立法，为全球人工智能的规则走向奠定锚点。欧盟《人工智能法》遵循风险分级的监管方法，参考功能、用途和影响将人工智能系统风险分为禁止性风险、高风险、有限风险、低风险或无风险四个级别，根据风险级别采取不同层次的监管措施，并为通用人工智能模型设置专有规则。该法的大

¹ 本报告主要统计数据截至 2024 年 12 月 31 日，对有关人工智能重要法律政策的汇总和梳理截至 2025 年 5 月 31 日。

部分规则将于 2026 年 8 月 2 日开始适用，但被认为具有“绝对不可接受的风险”的人工智能系统的禁令将在 6 个月后适用，而通用人工智能模型规则将在 12 个月后适用。此外，欧盟成员国必须在 2025 年 8 月 2 日前指定国家主管当局，负责监督人工智能规则的适用并开展市场监督活动。

《人工智能法》与数据等相关立法共同组成欧盟人工智能法律体系。欧盟《人工智能法》与《通用数据保护条例》（GDPR）、《数字服务法》（DSA）、《数字市场法》（DMA）、《数据治理法》（DGA）、《数据法》（DA）等共同构成欧盟数据战略框架下的重要治理规则，协同规范人工智能发展。目前，涉及数据处理的相关问题仍受 GDPR 等现行相关立法约束。此外，EDPB 近期陆续出台相关指导意见，旨在加强人工智能法与相关数据立法的协调。如，2024 年 12 月，EDPB 通过一项关于使用个人数据开发和部署人工智能模型的意见，即《关于人工智能模型处理个人数据的某些数据保护问题的第 28/2024 号意见》。该意见解决了在 GDPR 背景下人工智能模型的匿名性和合法性问题。²

2. 促进《人工智能法》落地实施

欧盟相关机构和部门加快布局《人工智能法》的监督实施。目前，该法已经进入分阶段实施阶段，欧盟各相关机构和部门正在为该法的实施作机构、人员以及文件指引等准备。

²https://www.edpb.europa.eu/news/news/2024/edpb-opinion-ai-models-gdpr-principles-support-responsible-ai_en

设置人工智能专门机构。一是欧盟委员会成立人工智能办公室，负责监督欧盟人工智能的发展、部署和监管，并将在实施欧盟《人工智能法》中发挥关键作用。人工智能办公室包括监管和合规部门、人工智能安全部门、卓越人工智能和机器人部门、人工智能造福社会部门、人工智能创新和政策协调部门。二是欧洲议会的内部市场和消费者保护委员会（IMCO）和公民自由、司法和内政事务委员会（LIBE）成立联合工作组，监督《人工智能法》的实施。三是欧洲数据保护委员会（EDPB）建议在多数情况下指定数据保护机构（DPA）为市场监督机构，增强监管机构之间的协调性，加强《人工智能法》和欧盟数据保护法的监督和执行。如葡萄牙国家通信管理局（ANACOM）已根据欧盟《人工智能法》指定了监管机构，以确保遵守旨在保护基本权利的欧洲立法。

制定《通用人工智能实践守则》。2024 年 9 月，欧盟人工智能办公室主办《通用人工智能实践守则》（Code of Practice on General-purpose AI，以下简称《实践守则》）的启动仪式。欧盟委员会任命多位学者讨论起草《实践守则》，旨在明确欧盟《人工智能法》中风险管理和透明度的实践规制。2024 年 11 月，《实践守则》初稿发布。同年 12 月，欧盟发布《实践准则》第二稿。该草案由独立专家团队编写，并基于对 2024 年 11 月初稿的反馈意见进行了修订。其核心目标是为通用人工智能模型的开发者和供应商提供清晰的指导，用于证明在人工智能模型生命周期各阶段遵守《人工智能法》的情况。第二稿侧重于澄清基础问题、增加细节规定、确保符合比例原则，在

后续迭代中，将确保该准则各部分的衔接性和易懂性等。2025 年 3 月，欧盟委员会公布《实践守则》第三稿草案，与前两稿相比，本版《守则》的结构更加精简，承诺和措施更加细化。草案前两部分详细说明了所有通用人工智能模型提供商的透明度和版权义务，并根据《人工智能法》对某些开源模型提供商的透明度义务进行豁免；第三部分仅与少数可能构成系统性风险的最先进通用人工智能模型提供商相关。后续还将进一步根据反馈意见和研讨进行修改完善。

发布相关指南加强协同治理。2025 年 2 月，欧盟委员会批准《关于人工智能法确立的禁止人工智能实践的指南》（Commission Guidelines on prohibited artificial intelligence practices established by Regulation (EU) 2024/1689 (AI Act)），旨在确保《人工智能法》在欧盟的统一实施。该指南提供了法律解释和实际事例，但不具有约束力，解释权属于欧盟法院。同月，欧盟委员会发布《人工智能法规定的人工智能系统定义指南》（Guidelines on the definition of an artificial intelligence system established by AI Act）。《指南》详细介绍了人工智能系统定义的七个要素，包括机器性、自主性、适应性、目标、推理能力、输出类型和环境交互。2025 年 3 月，欧盟委员会发布欧盟人工智能模型示范合同条款（EU AI model contractual clauses）的更新版本，具体分为适应于高风险人工智能系统的完整版本和适用于非高风险人工智能系统的轻量版本。示范合同条款为工作文件，专为采购人工智能解决方案的公共部门而设计，私营实体可以根据条款选择性

地调整适用。人工智能模型示范合同条款在欧盟《人工智能法》完全适用之前一直有效，并作为采购合同的附件使用。

3.构建统一协调的人工智能规范体系

欧盟自 2018 年以来陆续发布人工智能相关战略和文件，致力于构建统一协调的人工智能规范体系。2018 年 4 月，欧盟发布《欧洲人工智能战略》，旨在使欧盟成为世界级人工智能中心，确立以人为本和可信赖的人工智能发展方向。³2019 年 4 月，欧盟人工智能高级专家组发布《可信人工智能伦理指南》，提出以人为本的人工智能方法，并提出人工智能系统应满足“人类主导和监督”“技术稳健性和安全性”等七项关键要求，才能被视为可信。⁴2020 年 2 月，欧盟委员会发布《人工智能白皮书：欧洲实现卓越和信任的方法》，提出利用人工智能实现“单一数据市场”的愿景，建立区域协调的政策框架。⁵一方面，政策支持促进欧盟人工智能产业发展。2024 年 1 月，欧盟委员会推出人工智能创新计划，通过提供资金支持和超级计算访问权限等措施，支持人工智能初创企业和中小企业发展。此外，欧盟委员会还将与部分成员国建立语言技术联盟（ALT-EDIC）、CitiVERSE 两个欧洲数字基础设施联盟，解决欧洲语言数据的短缺问题，加快数字孪生技术应用，为人工智能发展提供基础设施支撑。⁶2024 年 12 月，欧洲高性能计算联合组织（EuroHPC）选定七项方案，将在欧洲各地建立并运营首批人工智能工厂。人工智能工厂将汇集人工智能发展所

³ <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=COM:2018:237:FIN>

⁴ <https://digital-strategy.ec.europa.eu/en/library/ethics-guidelines-trustworthy-ai>

⁵ https://commission.europa.eu/document/d2ec4039-c5be-423a-81ef-b9e44e79825b_en

⁶ https://ec.europa.eu/commission/presscorner/detail/en/ip_24_383

需的算力、数据、人才等关键要素，助力初创企业、产业和研究人员开发人工智能模型和系统。⁷

另一方面，治理实践推动重点领域监管走深走实。数据治理方面，明确人工智能数据隐私保护立场。2024 年 8 月，爱尔兰数据保护委员会对社交媒体平台 X 提起诉讼，认为其数据处理行为违反欧盟 GDPR，并寻求法院命令，停止或限制 X 处理用户数据以训练其人工智能系统。该裁决将对平台企业利用数据训练人工智能模型等相关实践产生重要影响。市场竞争方面，保护人工智能生态系统竞争力。2024 年 7 月，欧盟委员会与美国司法部、美国联邦贸易委员会、英国竞争和市场管理局联合发布《关于生成式人工智能基础模型和人工智能产品竞争的联合声明》，强调公平、开放和竞争性市场在开发和部署生成式人工智能方面的重要性。此外，欧盟当局还发起对虚拟世界和生成式人工智能市场竞争的研究调查，密切关注微软、谷歌、亚马逊和英伟达等大型科技公司投资人工智能的行为。

（二）美国：聚力推进基于创新应用的治理范式

作为人工智能领域的领跑者，美国高度重视人工智能发展，坚持创新为先和应用导向，强调维护自身的技术发展优势和领导地位，呈现人工智能监管立法“州政府先行、联邦稳步推进”的特点。与此同时，伴随特朗普重回白宫，美国人工智能立法模式与政策导向正在发生重大转变。

1. 积极探索人工智能立法方案

⁷ https://ec.europa.eu/commission/presscorner/detail/en/ip_24_6302

美国国会立法具有“雷声大、雨点小”的特点，正式签署成法的比例往往不足 2%。⁸尽管美国国会早在特朗普第一任期就开始了人工智能立法探索，但目前尚未出台一部具有强制效力的综合性监管法案。相比之下，美国州层面则进展较快领跑联邦立法进程，已有部分法律规范率先落地。

联邦层面，美国国会积极探索人工智能立法方案。综合性立法方面，美国国会于 2023 年提出《安全创新框架》（SAFE Innovation Framework）、《美国人工智能法的两党框架》（Bipartisan Framework for U.S. AI Act）等两党法案和规范框架，旨在帮助国会制定人工智能立法，提出维护美国在人工智能政策制定方面的领导地位等主要原则和措施。2024 年 11 月，美国众议院人工智能特别工作组发布报告，强调了人工智能立法的下一步优先事项，确保人类在人工智能技术及其应用中的中心地位，同时采取适度的监管方法。**重点领域立法方面，**就深度伪造、版权保护、算法歧视等提出一系列立法提案。**一是**加快规制人工智能深度伪造。如，2024 年 7 月，美国参议院一致通过《反抗法案》（DEFIANCE Act），允许人工智能深度伪造作品的受害者起诉其制作者并要求损害赔偿。又如，2025 年 5 月，特朗普签署通过《删除法案》（TAKE IT DOWN Act），全称为《应对网页及网站中深度伪造引发的性剥削工具法案》。该法案规定，在未经当事人同意的情况下，故意披露未经同意的私密影像（包括人工智能生成伪造内容）属于联邦刑事犯罪。**二是**加强人工智能知识产权保护。如，2024

⁸ 根据美国国会官网数据统计，第 118 届美国国会提出 23912 项立法，正式通过成法的有 274 项，通过率约 1.15%。

年 4 月，提出《生成式人工智能版权披露法案》（Generative AI Copyright Disclosure Act），要求人工智能公司披露受版权保护的生成式人工智能模型的训练材料；2024 年 5 月，提出《禁止伪造法案》（NO FAKES Act），禁止未经授权使用人工智能对个人肖像或声音进行复制，规范人工智能在音乐和电影行业的使用。三是注重人工智能安全治理。如，2023 年 11 月，提出《人工智能研究、创新和问责法案》（Artificial Intelligence Research, Innovation, and Accountability Act），鼓励人工智能创新，建立问责制，提高人工智能应用的透明度和安全性。又如，2023 年 12 月，提出《消除算法系统偏见法案》（Eliminating Bias in Algorithmic Systems Act），要求所有使用人工智能技术的联邦机构设立民权办公室，管理人工智能偏见和歧视。

州层面，多数州已提出或通过人工智能相关法案。截至 2024 年 12 月，全美至少有 43 个州通过人工智能相关立法，内容涉及选举、政府使用、就业、医疗健康、儿童色情等多方面。⁹例如 2024 年 2 月，夏威夷就政治竞选中的人工智能深度造假和虚假信息问题提出立法，禁止制造关于候选人或政党的虚假信息，并将散布虚假政治信息定为犯罪行为，授权夏威夷竞选开支委员会对虚假信息进行调查和罚款。2024 年 3 月，田纳西州通过了保护词曲作者、表演者和其他音乐行业专业人士免受人工智能滥用风险的立法。2024 年 3 月，犹他州签署通过《人工智能政策法》（SB0149），强调透明和非歧视地使用人工智能，保护消费者权益，并规定了人工智能政策办公室以及人工

⁹ https://hai.stanford.edu/assets/files/hai_ai_index_report_2025.pdf

智能学习实验室计划等内容。2024 年 5 月，**科罗拉多州**签署通过《关于在与人工智能系统的互动中提供消费者保护的法案》（SB205），旨在防止算法歧视，强化对消费者的权益保护。2024 年 9 月，**加利福尼亚州**针对人工智能深度伪造干预选举现象、好莱坞影视行业的人工智能克隆行为等，签署通过了 5 项人工智能法案。2025 年 2 月，**弗吉尼亚州**议会通过《高风险人工智能开发者和部署者法案》，该法案适用于自主决策或对决策有“重大影响”的人工智能系统，旨在保障消费者在就业、金融、教育等相关决策中免受算法歧视。2025 年 5 月，**阿肯色州**签署通过第 927 号法案《关于生成式人工智能工具生成的训练模型和内容所有权法案》，明确生成式人工智能工具的内容所有权和模型训练成果归属。

2. 稳步构建人工智能治理框架

自奥巴马政府时期，美国政府就开始陆续出台人工智能政策文件，加大对人工智能创新应用与安全治理的重视。目前，美国已经初步形成了以总统行政令、联邦机构标准指南以及行业自律为支柱的人工智能规范框架。

一是行政令统一推进人工智能创新应用。在出台《人工智能权利法案蓝图》《国家人工智能研发战略计划》等政策文件基础上，2023 年 10 月，拜登正式签署《关于安全、可靠和可信地开发和使用人工智能的行政令》（第 14110 号行政令），对美国下一步的人工智能发展和监管作出全面部署。内容涉及人工智能安全新标准、保护公民隐私、促进公平和公民权利、支持消费者和工人、促进创新和竞争、确

保政府负责有效地使用人工智能、加强美国的海外领导力等多方面。2024 年 10 月，美国政府发布《关于提升美国在人工智能领域的领导地位、利用人工智能实现国家安全目标以及促进人工智能安全性和可信度的备忘录》，旨在落实第 14110 号行政令，聚焦国防、情报等领域，为联邦政府使用人工智能提供进一步指导，推进美国国家安全利益。2024 年 10 月底，美国政府宣布已按时完成第 14110 号行政令要求的 100 多项行动，取得的重大成就涉及风险管理，工人、消费者以及隐私和民权，利用人工智能造福人类，将人工智能及相关人才引入政府，提升美国海外领导力五大方面。¹⁰

二是标准指南引导人工智能风险管理。行业标准作为美国人工智能规范体系的重要组成部分，持续发挥关键引导作用。2023 年 1 月，美国国家标准技术研究所（NIST）发布《人工智能风险管理框架》。¹¹该框架是一份非强制性的指导性文件，旨在更好管理与人工智能相关的个人、组织和社会风险，降低开发部署人工智能系统时的安全风险，避免产生偏见等负面影响，提高人工智能可信度，保护公民的公平自由权利。2023 年 3 月，NIST 启动可信和负责任的人工智能资源中心，促进管理框架的实施和国际协调。2024 年 7 月，NIST 发布《人工智能风险管理框架：生成式人工智能概况》，帮助识别由生成式人工智能造成的独特风险，并为生成式人工智能风险管理提出行动建议。此外，2024 年 10 月，美国国家安全委员会（NSC）发布《推进国家安全领域人工智能治理和风险管理的框架》，指导人工智能在美国军

¹⁰<https://www.whitehouse.gov/briefing-room/statements-releases/2024/10/30/fact-sheet-key-ai-accomplishments-in-the-year-since-the-biden-harris-administrations-landmark-executive-order/>

¹¹ <https://www.nist.gov/itl/ai-risk-management-framework>

事行动和国家安全中的使用。

三是行业自愿承诺提高人工智能安全和透明度。美国人工智能企业积极与政府开展治理合作。2023 年 7 月，微软、OpenAI、英伟达等 15 家科技公司已签署自愿承诺，确保人工智能产品安全、构建以安全为首位的人工智能系统以及获得公众信任。¹²2023 年 12 月，28 家医疗供应商也已承诺在医疗保健领域负责任地使用人工智能。¹³此外，2024 年 8 月，OpenAI 和甲骨文宣布与美国人工智能安全研究所签署合作协议，同意在模型发布前接受预测试。目前，OpenAI、谷歌、Anthropic 等正在通过组建专家团队、制定伦理原则和用户协议、开展对抗性测试和风险评估、添加内容标识等手段，主动落实自愿承诺，提高人工智能安全和透明度。如，2023 年 12 月，OpenAI 发布“防备框架”（Preparedness Framework），指导人工智能大模型的安全部署，包括风险监控与识别、建立安全基准、定期报告等。又如，2024 年 3 月，谷歌发布《在 Google Cloud 上安全部署人工智能的最佳实践》，分享在模型开发、应用程序、基础设施和数据管理四个方面确保人工智能安全的最佳做法。

3. 轻监管促竞争的立法转向持续加速

值得关注的是，特朗普政府上任后，美国政府人工智能立法导向发生重大转变，从风险审慎到竞争为先，意图以轻监管、去监管思路加速本土人工智能产业发展。2025 年 1 月，特朗普废除拜登第 14110

¹²<https://www.whitehouse.gov/briefing-room/statements-releases/2023/07/21/fact-sheet-biden-harris-administration-secures-voluntary-commitments-from-leading-artificial-intelligence-companies-to-manage-the-risks-posed-by-ai/>

¹³<https://www.whitehouse.gov/briefing-room/blog/2023/12/14/delivering-on-the-promise-of-ai-to-improve-health-outcomes/>

号行政令，并于进一步发布《消除美国在人工智能领域领导地位的障碍的行政命令》，指示相关机构制定人工智能行动计划，审查第 14110 号行政令下制定的所有政策文件，并酌情予以中止、修订或撤销。2025 年 4 月，美国白宫管理和预算办公室发布“关于通过创新、治理和公众信任加速联邦对人工智能的使用”和“关于推动政府高效利用人工智能”两份备忘录，并对应废除了拜登政府期间发布的关于“推进机构使用人工智能的治理、创新和风险管理”备忘录和关于“推进政府负责任地采购人工智能”的备忘录。特朗普政府两份备忘录以促进联邦机构对人工智能的采购与应用为重点，一定程度放松了拜登政府对风险管理的重视，同时保留了第 14110 号行政令中关于首席人工智能官、建立人工智能采购工具库等大部分内容。结合特朗普政府上一任期执政内容和新一任期举措，未来美国联邦政府或将采取轻监管促竞争的思路，以更大力度促进人工智能创新应用。

与此同时，美国新一届国会改变以往“府院分立格局”，高度配合特朗普政府执政理念，拟通过相关立法暂缓人工智能监管进程。2025 年 5 月，美国众议院通过名为“一项宏大美好法案”（One Big Beautiful Bill Act）的预算协调法案，其中包含了一项对州级人工智能监管的十年禁令，禁止未来十年内各州执行任何针对人工智能模型、系统或自动决策系统的监管法律。该条款旨在消除各州分散监管带来的法律障碍，为联邦政府制定统一的人工智能政策框架争取时间。结合美国国会同期出台通过的《删除法案》，美国或以小切口立法模式，针对人工智能具体用例和部分紧迫而明确的风险进行统一监管。

（三）中国：以政策法规保障人工智能健康发展

近年来，我国坚持发展与安全并重、促进创新和依法治理相结合的原则，不断加强人工智能法治建设，有序推进相关立法工作，构建形成了多层次、多维度的人工智能政策法规体系，促进和保障人工智能产业高质量发展。

1.持续推动出台人工智能产业促进政策

我国高度重视人工智能产业健康发展，积极部署人工智能产业发展促进政策，为人工智能科技研发、技术应用等注入强劲动力。2015年5月我国提出发展智能装备、智能产品和生产过程智能化后，人工智能相关政策进入密集出台期，国家发展改革委、工业和信息化部、科技部等出台了多个人工智能相关规划及工作方案。2017年7月，国务院印发《新一代人工智能发展规划》，首次将人工智能提升到国家战略地位，提出科技引领、系统布局、市场主导、开源开放的基本原则，部署构筑我国人工智能发展的先发优势。在此基础上，我国又陆续发布了《促进新一代人工智能产业发展三年行动计划（2018-2020年）》《国家新一代人工智能开放创新平台建设指引》《关于促进人工智能和实体经济融合的指导意见》等政策文件，夯实人工智能规模化应用和产业化发展基础，促进技术创新和模式创新，培育我国经济增长新动能。

随着人工智能发展进入新阶段，我国进一步创新政策文件，促进人工智能和实体经济深度融合，加快培育新质生产力，推动经济社会高质量发展。2022年7月，科技部等六部门出台《关于加快场景创

新以人工智能高水平应用促进经济高质量发展的指导意见》，着力打造人工智能重大场景，提升人工智能场景创新能力，加快推动人工智能场景开放，加强人工智能场景创新要素供给。2024 年 1 月，工业和信息化部等七部门联合印发《关于推动未来产业创新发展的实施意见》，提出围绕制造业主战场加快发展人工智能、人形机器人等未来产业，支撑推进新型工业化。2024 年 6 月，工业和信息化部等四部门印发《国家人工智能产业综合标准化体系建设指南（2024 版）》，进一步加强人工智能标准化工作系统谋划，加快构建满足人工智能产业高质量发展和“人工智能+”高水平赋能需求的标准体系。2025 年 3 月，政府工作报告指出，“持续推进‘人工智能+’行动，将数字技术与制造优势、市场优势更好结合起来，支持大模型广泛应用，大力发展智能网联新能源汽车、人工智能手机和电脑、智能机器人等新一代智能终端以及智能制造装备。”

2.初步构建多层次、多维度的人工智能法律体系

我国坚持统筹安全与发展，积极推进人工智能相关立法工作，初步构建起涵盖法律、法规、规章等多层次与数据、算法、应用等多维度的人工智能法律规范体系。

互联网领域综合立法为人工智能发展共性问题提供基础法律框架。我国先后制定出台了《网络安全法》《数据安全法》《个人信息保护法》等法律，明确了网络安全、数据安全、个人信息处理等方面的具体规则，为应对人工智能技术研发应用中的相关问题提供了明确指引。如，《电子商务法》《个人信息保护法》对非个性化选项进行

了规定，《民法典》明确禁止“利用信息技术手段伪造等方式侵害他人的肖像权”的行为。值得关注的是，2024 年 9 月，我国出台《网络数据安全管理条例》，以专条的形式对生成式人工智能训练数据的管理作出原则性规定。

围绕算法、深度合成、生成式人工智能等重点领域制定专门规则。

近年来，我国针对人工智能发展的重点问题制定“小快灵”规章，相关部门先后出台《互联网信息服务算法推荐管理规定》《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》等相关规定。2023 年 7 月，国家互联网信息办公室、国家发展改革委、教育部、科技部、工业和信息化部、公安部、广电总局联合公布《生成式人工智能服务管理暂行办法》。该办法是全球首部生成式人工智能专门立法，提出国家坚持发展和安全并重、促进创新和依法治理相结合的原则，采取有效措施鼓励生成式人工智能创新发展，对生成式人工智能服务实行包容审慎和分类分级监管，明确了提供和使用生成式人工智能服务总体要求，规定了生成式人工智能服务基本规范。此外，2025 年 3 月，国家互联网信息办公室、工业和信息化部、公安部、国家广播电视总局联合发布《人工智能生成合成内容标识办法》。该办法作为规范性文件，聚焦人工智能“生成合成内容标识”关键点，通过标识提醒用户辨别虚假信息，明确相关服务主体的标识责任义务，规范内容制作、传播各环节标识行为，进一步细化标识具体实施规范。

地方立足于促进法定位加快出台相关法规规章。目前我国部分地区结合本地人工智能产业发展实践，正在探索建立区域性的人工智能

规范体系，为人工智能健康发展提供制度保障。2022 年 9 月，深圳发布人工智能产业专项立法《深圳经济特区人工智能产业促进条例》，明确了人工智能概念和产业边界，提出了创新产品准入制度，对于国家、地方尚未制定标准但符合国际先进产品标准或者规范的低风险人工智能产品和服务，允许通过测试、试验、试点等方式开展先行先试。2022 年 9 月，上海发布《上海市促进人工智能产业发展条例》，同样立足于促进法的基本定位，注重创新性和引领性，充分发挥有效市场和有为政府的作用，重点促进基础硬件、关键软件、智能产品等高质量发展，聚焦打通产业发展中的堵点难点问题，立足更高起点推动人工智能产业、技术、应用加快发展。此外，部分地区也在探索通过立法释放“算力”这一人工智能发展关键要素。2024 年 2 月，安徽省芜湖市人民政府发布《芜湖市建设算力中心城市促进办法》，重点对算力基础设施建设、算力应用、算力产业培育、算力安全和服务保障等活动进行规范，发挥算力赋能高质量发展作用。

3.推动形成智能向善的人工智能治理主张

我国将伦理规范作为保障人工智能健康发展的重要手段，不仅重视人工智能的社会伦理影响，而且积极制定伦理规范和道德框架，确保人工智能安全、可靠、可控。

对内，我国出台相关指导意见和法规文件，明确人工智能科技伦理治理“红线”和“底线”。近年来，国家新一代人工智能治理专业委员会相继发布《新一代人工智能治理原则—发展负责任的人工智能》《新一代人工智能伦理规范》，为从事人工智能相关活动的主体提供伦理

指引和行动指南。2022 年 3 月，中共中央办公厅、国务院办公厅发布《关于加强科技伦理治理的意见》，成为我国首个国家层面的科技伦理治理指导性文件，为人工智能等新兴前沿技术的科技伦理治理作出顶层设计和系统部署。2023 年 9 月，科技部等十部门发布《科技伦理审查办法（试行）》，进一步落实《关于加强科技伦理治理的意见》中规定的科技活动主体的审查责任和程序，建立需要开展专家复核的科技活动清单制度，为各地方和相关行业主管部门、创新主体等组织开展科技伦理审查提供了制度依据。

对外，我国通过发表相关倡议或立场文件，积极宣介“以人为本、智能向善”的治理理念。2022 年 11 月，我国向联合国大会提交《关于加强人工智能伦理治理的立场文件》，从人工智能技术监管、研发、使用及国际合作等方面，明确了我国关于人工智能伦理治理的主张，积极倡导“以人为本”和“智能向善”的治理理念。2023 年 10 月，我国正式提出《全球人工智能治理倡议》，围绕人工智能发展、安全、治理三方面系统阐述了人工智能治理中国方案，为相关国际讨论和规则制定提供参考蓝本。2024 年 7 月，2024 世界人工智能大会暨人工智能全球治理高级别会议发表《人工智能全球治理上海宣言》，提出促进人工智能发展、维护人工智能安全、构建人工智能治理体系、加强社会参与和提升公众素养、提升生活品质与社会福祉等四方面内容，为全球人工智能治理提供新思路。2024 年 9 月，全国网络安全标准化技术委员会发布《人工智能安全治理框架》1.0 版，贯彻落实《全球人工智能治理倡议》，提出了包容审慎、确保安全，风险导向、敏

捷治理，技管结合、协同应对，开放合作、共治共享等人工智能安全治理的原则。2024 年 10 月，2024 年世界电信标准化全会（WTSA-24）期间，中国代表团主导并推动国际电信联盟首份人工智能新决议协商一致通过。

（四）其他国家和地区：加快人工智能立法与战略布局

除了中美欧三大经济体外，其他国家和地区也相继出台人工智能政策法律文件，在参考欧盟《人工智能法》等域外经验的基础上，结合本国产业实际和独特优势，加快人工智能立法与战略部署，力图实现发展与监管的二元平衡。

1. 韩国：“产业支持+适度监管”双着力的人工智能立法特色

韩国高度重视人工智能发展，近年来陆续推出多部人工智能政策文件，旨在跻身全球人工智能前三强。在人工智能规范治理方面，韩国也具有明显的产业支持导向，秉持宽严并济、管促结合、导主罚辅的立法思路，力求在日益激烈的全球人工智能竞争中，发挥立法对产业创新发展的促进保障作用。

韩国人工智能法律体系初步构建。一方面，2024 年 12 月，韩国国民议会最新通过《人工智能发展与建立信任基本法》（以下简称《人工智能基本法》），自 2026 年 1 月开始正式施行。该法使韩国成为继欧盟之后第二个制定人工智能综合立法的国家，旨在确保安全和负责任发展的同时，增强韩国在人工智能领域的竞争力。**另一方面**，韩国《个人信息保护法》《数据产业振兴和利用促进基本法》等现行立

法，也适用于人工智能应用场景下的数据使用、个人信息保护、自动化决策等相关问题。同时，韩国个人信息保护委员会等相关部门还发布了《人工智能时代个人信息安全使用指南》等指南与标准，从多维度规范人工智能发展。

韩国《人工智能基本法》具有明显的产业支持特色。重点关注三方面内容：构建国家层面的人工智能协同治理体系、支持人工智能产业发展、防范人工智能潜在风险。具体包括以下方面：**一是明确人工智能相关概念。**“高影响人工智能”是指有可能对人的生命、身体安全以及基本权利造成重大影响或带来风险的人工智能系统，例如能源供应、饮用水生产等领域。此外，还明确了人工智能从业者、生成式人工智能等相关概念。但与欧盟分级设置人工智能风险及义务不同的是，《人工智能基本法》主要针对高影响人工智能系统设置单独义务。**二是设立“人工智能技术发展和产业培育”专章。**《人工智能基本法》提出构建人工智能产业基础、支持引进和使用人工智能技术、促进人工智能产业融合发展等系列举措，旨在促进产业创新发展。同时，注重加强对中小企业的特别支持，包括在各类支持措施中优先考虑中小企业、促进中小企业参与人工智能产业，并在高影响人工智能运营者监管、人工智能影响评估等方面给予中小企业特别支持等。**三是设立人工智能专门机构统筹安全与发展。**国家人工智能委员会，直属于总统，负责审议和制定人工智能发展和信任基础构建等政策。**人工智能政策中心**，下设于科技通信部，负责整合资源、支持人工智能政策和国际标准制定以及分析人工智能对社会影响。**人工智能安全研究所**，下设

于科技通信部，负责保护国民免受人工智能相关风险，维护人工智能社会信任基础，开展人工智能安全研究和相关国际合作。**四是**规定高影响人工智能的认定方式及义务。高影响人工智能系统采取“自认定+官方认定”方式，由企业自己自行识别是否属于高影响人工智能，在不确定情况下，可以向科技通信部请求认定。高影响人工智能应履行制定运营风险管理方案、用户保护方案以及进行影响评估等义务。**五是**规定较低罚则。未能有效提示人工智能服务或未指定境内代表落实各项义务的境外人工智能服务方或不遵守停止、纠正命令者，将被处以 3000 万韩元以下罚款。相较于欧盟《人工智能法》高额罚款，此罚款金额相对较低，体现了韩国维护市场信心、引导人工智能健康发展的立法旨意。

2. 日本：技术发展为先的人工智能立法范式

日本坚持技术发展为先的人工智能立法范式，力求促进创新的同时解决风险。近年来，日本有序推动相关指引性文件和法律的制定，倾向以非严格监管的法律制度指导企业自我监管，具有高度的灵活性和引导性特征。

发布指导性文件保障人工智能安全可靠使用。2024 年 4 月，日本总务省和经济产业省联合发布《人工智能业务指南》，通过整合修改现有的《人工智能研发指南》《人工智能应用指南》和《实施人工智能原则的治理指南》，形成一体化的人工智能指导性文件。该指南无强制约束力，旨在帮助各行业的人工智能应用主体充分认识人工智能风险，并在整个生命周期内采取必要的应对措施，体现出日本人工

智能规范的“软法”特征。基本结构方面，该指南基于“Why-What-How”的逻辑框架，引导人工智能业务行为者明确管理任务以及具体落实的方法。主体责任方面，该指南将相关主体分为“人工智能开发者”“人工智能提供者”和“人工智能商业用户”。治理机制方面，该指南参考敏捷治理的治理模式，将根据人工智能发展需要进行动态更新。

出台首部人工智能法促进技术研发和应用。2025 年 5 月，日本通过《人工智能相关技术研发及应用促进法》，是日本首部全面规范人工智能的法律框架，旨在全面、系统地推进人工智能相关技术研发和应用，改善人民生活并推动日本国民经济健康发展。该法规定了人工智能相关技术研发及应用的基本原则，明确了政府、研发机构、经营者和公众等相关主体的责任义务，提出了促进人工智能相关技术研发及应用的基本措施。

一是规定基本原则。其一，以保持日本人工智能相关技术研发能力并提高人工智能相关产业的国际竞争力为目的；其二，以全面系统地推进人工智能相关技术的研发和利用为目的；其三，确保人工智能相关技术研发和应用过程中的透明度，并采取相关必要措施；其四，积极参与国际合作，并努力在人工智能相关技术研发和应用的国际合作中发挥主导作用。

二是明确主体责任。国家负责制定和实施相关促进措施，并积极推动政府部门使用人工智能相关技术；地方政府根据地区特点与国家机关适当分工；研发机构积极推动人工智能技术的研发以及成果的扩散，并培养专业人才；人工智能开发者、服务提供者以及使用者应当积极利用人工智能技术塑造新兴产业；公众应当加强对人工智能相关技术的理解。

三是提出基本措施。

该法明确国家应当为人工智能相关技术从基础研究向实际应用提供一体化支持，促进设施及设备的完善和共享，确保人工智能相关技术研究、开发和应用的适当性，并培养专业人才、推进国际合作等。**值得注意的是**，该法未专门设置违法处罚，体现出日本促进创新、避免过度监管的立法取向。主要通过两方面进行规制：一方面，依托现有监管框架规范人工智能，如，擅自使用个人信息训练人工智能将违反日本《个人信息保护法》。另一方面，授以主管机关一定指导权限。如，该法规定，国家应当收集国内外人工智能相关技术研发及应用的动态信息，针对人工智能技术侵害国民权益等案件启动调查研究，并对相关企业进行指导、建议。

3.巴西：技术主权话语下的人工智能规范进路

巴西在国内“技术主权”的呼声下，促进人工智能技术和产业发展的同时，高度关注人工智能相关风险和道德问题，加快推动人工智能立法进程，倡导人工智能的安全开发和应用，鼓励负责任的创新。

立足技术主权需求，明确人工智能负责任的创新原则。2021 年，巴西科技创新部发布《人工智能战略》，指导联邦政府人工智能相关政策的制定，包括人工智能解决方案的研发与创新，以及有意识、有道德、以创造美好未来为前提应用人工智能技术。《战略》还提出了未来开展的 74 项战略行动，其中将人工智能的立法、监管和道德使用作为首项行动。2023 年 11 月，巴西科学院发布《巴西人工智能发展建议》，强调了巴西面临的技术发展困境和国家主权忧虑。2024 年 7 月，巴西出台《巴西人工智能计划》，提出促进人工智能在巴西

的发展和使用，保障安全性、维护个人与集体权利、满足社会包容、确保信息完整性、保护劳动和就业、保障国家主权等目标，并明确将技术主权和数据主权作为指导原则之一。

人工智能立法加速，力图实现风险预防和创新激励双重目标。作为人工智能技术的主要应用者而非开发者，隐私侵犯、算法歧视、行为操纵等风险驱使巴西加快相关立法进程，制定人工智能监管工具和框架。2023 年 5 月，巴西提出首部实质性人工智能法案（第 2338/2023 号法案），旨在规范巴西人工智能的发展和应用。经过一年多的广泛辩论和修订，法案于 2024 年 12 月经巴西参议院批准通过。2025 年 3 月，巴西参议院将《人工智能法案》提交众议院审查，尚待众议院审议表决和总统签署。法案强调对公民基本权利的保护，一定程度借鉴欧盟风险分级监管思路，旨在促进负责任的创新，确保人工智能系统安全可信，推动社会、经济和科学技术发展。**一是**明确人工智能开发和使用的 17 项基本原则，包括包容性增长、自决权、人类监督、非歧视、公平正义、透明度和可解释性、尽职调查、可靠性和鲁棒性、互操作性、未成年人保护等。**二是**规定公民享有的权利，包括知情权、隐私权和个人数据受保护权、免受歧视和纠正歧视性偏见权等。**三是**建立更为灵活的分级监管框架。法案将人工智能系统分为“过度风险”“高风险”和“其他”三类。其中，禁止开发、部署和使用“过度风险”人工智能系统，包括使用潜意识诱导技术、利用自然人或群体缺陷、用于犯罪风险评估等；“高风险”人工智能系统必须履行严格的合规义务，包括自动化日志记录、算法逻辑可解释性、开展安全评估测试等。

此外，法案要求根据实际情况定期调整更新“高风险”人工智能系统清单，并创建高风险人工智能系统数据库。**四是**创建国家人工智能监管与治理体系（SIA）。SIA 由多个政府部门组成，其中巴西国家数据保护局将担任协调机构。**五是**纳入版权条款。法案力求在鼓励创新与保护版权之间取得平衡，规定科研组织和教育机构、博物馆、档案馆和图书馆等开发人工智能系统时自动使用现有作品不构成版权侵权。

（五）国际组织：全球人工智能治理的重要推动力量

1. 联合国发挥主渠道作用推进人工智能国际治理

制定首份全球人工智能伦理规范框架。2021 年 11 月，联合国教科文组织开创性制定全球首份人工智能伦理框架——《人工智能伦理问题建议书》。该文件以保护人权和尊严为核心，强调人类监督人工智能系统的重要性。建议书覆盖广泛的政策行动领域，使决策者能够将其核心价值观和原则，转化为数据治理、环境和生态系统、性别、教育和研究、医疗健康、社会福利等方面的行动。¹⁴

联合国高级别人工智能咨询机构发挥引导效能。为更好应对人工智能带来的风险和机遇，2023 年 10 月，联合国秘书长宣布成立高级别人工智能咨询机构，负责对人工智能近期和长远的发展与风险开展研究与建议。该机构由来自 33 个国家和众多行业的 39 位人工智能领域专家组成，是世界上第一个也是最具代表性的专家组。2024 年 9 月，联合国秘书长高级别人工智能咨询机构发布《治理人工智能，助力造福人类》（Governing AI for Humanity）最终报告，旨在消除人

¹⁴ <https://www.unesco.org/zh/artificial-intelligence/recommendation-ethics>

工智能相关风险，并向世界分享人工智能的变革潜力。报告提出成立国际人工智能科学小组、启动人工智能治理政策对话、创建人工智能标准交流中心等七项人工智能治理战略建议，为全球人工智能治理合作奠定了重要基础。此外，其他联合国机构也在研究人工智能对各领域的影响，例如世界卫生组织和国际电信联盟成立关于人工智能与健康交叉领域的联合小组。

通过联合国首项人工智能治理专项决议。在全球共识的基础上，2024 年 3 月，联合国大会一致通过《抓住安全、可靠和值得信赖的人工智能系统带来的机遇，促进可持续发展》决议。这是联合国大会历史上首次为针对人工智能治理问题确立全球统一规范所通过的专项决议，旨在阐明安全、可靠和值得信赖的人工智能系统的共同方法。此外，2024 年 7 月，联合国大会协商一致通过中国主提的加强人工智能能力建设国际合作决议，旨在帮助各国特别是发展中国家从人工智能发展中平等受益，弥合数字鸿沟，完善全球人工智能治理，加快落实 2030 年可持续发展议程。¹⁵

加大对生成式人工智能的关注。随着生成式人工智能的诞生与快速应用，联合国及其专门机构加大对其带来的伦理道德担忧以及信息完整性、隐私保护等问题的关注与研究。2023 年 9 月，联合国教科文组织发布《生成式人工智能教育与研究应用指南》，这是全球首份关于在教育研究领域使用生成式人工智能的指南，呼吁各国政府尽快就此问题实施适当的管制和教师培训，确保这项技术在教育中的运

¹⁵ 联大通过中国提出的加强人工智能能力建设国际合作决议，中国政府网，2024 年 7 月 2 日，https://www.gov.cn/yaowen/liebiao/202407/content_6960524.htm。

用遵循以人为中心的方法。¹⁶2024 年 1 月，世界卫生组织发布关于多模态大模型伦理和管理问题的指导文件《卫生领域人工智能的伦理和治理：多模态大模型指南》，列出供政府、技术公司和卫生保健机构考虑的 40 多项建议，以确保适当使用多模态大模型增进和维护人民健康。2024 年 6 月，国际货币基金组织发布《扩大生成式人工智能收益：财政政策的作用》，建议对人工智能碳排放征税，以应对人工智能导致的失业和收入不平等。

举办“人工智能向善”全球峰会。“人工智能向善”（AI for Good）全球峰会是联合国就人工智能开展全球包容性对话的主导平台，由国际电信联盟与 40 家联合国机构合作举办，并与瑞士政府共同召集。峰会自 2017 年以来已连续举办八届，旨在为人工智能领域政府、企业及专家学者搭建交流分享平台，推广“人工智能向善”应用案例和全球经验。2024 年 5 月，“人工智能向善”全球峰会在瑞士日内瓦举行，聚焦人工智能全球合作、监管与治理、标准建设，以及人工智能在教育、健康、通信、金融、气候变化方面的应用等议题，开展全球对话，推动全球对人工智能伦理、公平与正义的深入理解。中国、巴西、马来西亚、尼日利亚、阿联酋、坦桑尼亚、柬埔寨、土耳其、印度 9 个国家提交的 40 个案例入选“人工智能向善”案例集。中国、美国、津巴布韦、印度等 8 个国家的 13 名专家入选“人工智能向善”学者。

2. 更多国际组织将人工智能纳入关注议题

¹⁶ <https://news.un.org/zh/story/2023/09/1121282>

人工智能全球治理成果不断丰富。2019 年 5 月，经济合作与发展组织（OECD）通过《人工智能建议书》。这是全球首个关于人工智能的政府间标准建议，旨在指导人工智能参与者努力开发值得信赖的人工智能，并为政策制定者提供有效的人工智能政策建议。该文件目前经过 2023 年 11 月、2024 年 5 月两次修订，最新修订强调了人工智能虚假错误信息、透明度、安全等问题。2023 年 12 月，国际标准化组织（ISO）发布世界首项有关人工智能管理系统的国际标准《信息技术 人工智能 管理体系》（ISO/IEC 42001），旨在帮助企业等组织负责任地开发和使用人工智能系统。2024 年 1 月，世界经济论坛（WEF）发布《2024 年全球风险报告》，警告先进人工智能助长虚假和错误信息，可能会侵蚀民主进程并导致社会两级分化。2024 年 9 月，世界经济论坛发布《生成式人工智能与国际贸易分析》报告，梳理了全球人工智能治理进展，分析了生成式人工智能对现行国际贸易规则的影响以及世界贸易组织应当发挥的作用。

区域层面人工智能治理共识加快达成。2024 年 2 月，东盟发布《东盟人工智能治理与伦理指南》，旨在通过区域合作推动人工智能的负责任开发和使用。2024 年 5 月，欧洲委员会通过《欧洲委员会关于人工智能与人权、民主和法治的框架公约》。这是全球首部具有法律约束力的人工智能国际条约，旨在确保使用人工智能系统时尊重人权、法治和民主法律标准。该公约适用于公共和私营部门使用的人工智能系统，采用了类似于欧盟《人工智能法》的基于风险的方法，并于 2024 年 9 月正式开放签署。2024 年 10 月，七国集团发布内政

和安全部长公报，提出各国应合作制定和采用人工智能方面的趋同政策，营造有利于创新和数字安全的全球环境。

3.其他平台不断扩大人工智能治理影响力

人工智能安全峰会、世界人工智能大会等为全球人工智能治理“添砖加瓦”。人工智能安全峰会已连续举办三届，逐步探索人工智能全球治理。2023 年 11 月，首届人工智能安全峰会在英国布莱切利园举行，中美欧等 28 个国家达成了全球首份针对人工智能的国际性声明《关于人工智能安全的布莱切利宣言》（Bletchley Declaration on Artificial Intelligence Safety）。该宣言旨在关注对未来强大人工智能模型构成人类生存威胁的担忧，以及对人工智能当前增强有害或偏见信息的担忧。2024 年 5 月，第二届全球人工智能安全峰会在韩国首尔召开。首尔峰会以“安全、创新、包容”为主题，与会各方就进一步加强人工智能安全性、推动人工智能可持续发展等议题进行讨论，迈出了人工智能安全国际合作的重要一步。2025 年 2 月，人工智能行动峰会在法国巴黎召开，会议发布《关于发展包容、可持续的人工智能造福人类与地球的声明》，强调了加强人工智能生态系统多样性的重要性，提出促进人工智能的可及性和减少数字鸿沟、促进人工智能创新发展、鼓励人工智能部署、加强国际合作等优先事项。世界人工智能大会（WAIC）持续赋能全球人工智能产业繁荣与进步。WAIC 自 2018 年创办以来，已连续举办六届，逐步成长为全球人工智能领域最具影响力的行业盛会和国际高端合作交流平台。2024 年 7 月，2024

世界人工智能大会暨人工智能全球治理高级别会议发表《人工智能全球治理上海宣言》，是全球人工智能治理领域的又一重要成果。

二、全球人工智能法律政策趋势动向

在技术浪潮的推进下，全球人工智能法治化进程明显加快，域外初步形成以美欧为代表的两大对比性模式，更多新兴经济体加入治理行列，持续提升自身影响力。从各国法律政策实践来看，当前全球人工智能治理进入软硬法接轨并行、综合立法和领域立法齐头并进的治理新阶段，生成式人工智能、前沿大模型和先进算力成为普遍关注焦点。

（一）人工智能治理总格局渐趋清晰，新兴经济体加入治理行列

主要国家和地区采取的人工智能规范进路与其自身产业发展情况紧密相关，随着各国治理路径逐步清晰，目前域外主要形成以美、欧为代表的两大对比性模式。

美国以确保技术、产业和生态优势为出发点，采取市场主导和创新驱动的治理思路，坚持政府引导和行业自律相结合的弱监管模式，构建联邦层面以总统行政令、政策法规、标准指南等为主的“软法治理”框架。一是坚持发展导向、产业优先的规制理念。美国联邦层面现有法律政策多以鼓励人工智能研发和负责任使用为重点。如，美国拜登政府《关于安全、可靠、值得信赖地开发和使用人工智能的行政命令》，多处强调要促进人工智能在联邦政府内部的应用，将确保美国人工智能领先地位作为重中之重。虽然该行政命令被特朗普政府废

除，但促应用促发展的主线不会发生变化。如，2025 年 1 月，特朗普签署《消除美国在人工智能领域领导地位的障碍》行政令，要求巩固其全球人工智能领域领导地位。**二是**内外兼修、以管促用，以政府内部管理带动整体行业自治。如，2024 年 3 月，美国拜登政府发布《推进机构使用人工智能的治理、创新和风险管理备忘录》，明确联邦机构使用人工智能的监管底线，加强机构内人工智能应用协调、创新及风险管理。10 月，拜登政府再度发布“人工智能国家安全备忘录”，针对军事、情报等国家安全风险，要求各机构部门协调开展前沿人工智能模型安全自愿测试与风险评估，依托政府采购激励企业将联邦监管要求纳入行业自治实践。**三是**释放算力、数据等关键要素，促进人工智能高质量发展。如，2024 年 4 月，美国商务部发布《人工智能和开放政府数据资产信息征集请求》旨在推进透明度、创新以及公共数据资产的开放和使用，促进数据驱动的人工智能技术更好发展。

欧盟在技术和产业硬实力消减的背景下，以规则塑造软实力为进路，**率先采取立法统一规制的强监管模式**，于 2024 年正式通过《人工智能法》。**一是**以风险预防为导向，构建全链条监管机制。欧盟《人工智能法》从系统开发、部署到投入市场后，对人工智能系统进行全生命周期的严格治理，并对不同阶段的人工智能系统运营者规定不同的义务。**二是**详细规定人工智能各主体的严格义务。欧盟《人工智能法》将适用主体划分为人工智能系统提供者、部署者、进口商、分销商、授权代表、产品制造商，监管对象覆盖了人工智能价值链各参与方，细化明确各主体的责任义务，全面管控人工智能风险。**三是**确立

广泛的适用范围和域外效力。法案规定，所有在欧盟境内投入市场或投入使用的人工智能系统均适用该法案，无论该人工智能系统设立于欧盟境内还是第三国。即使没有在欧盟市场投放、提供服务或使用，只要该人工智能系统处理内容和欧盟相关或输出结果意图在欧盟内使用，仍会落入《人工智能法》的监管范围。**四是明确人工智能产品适用严格产品责任。**2024 年 12 月，欧盟新的产品责任指令生效，明确涵盖软件、人工智能系统或产品相关数字服务等产品，同时加大制造商的举证责任，将进口商或外国制造商的欧盟代表增列为责任承担主体。

与此同时，印度、马来西亚、沙特阿拉伯、阿联酋等新兴经济体也逐步加入治理行列，持续提升影响力。**一是亚太地区逐步成为全球人工智治理的新高地。**近年来，除中日韩等人工智能头部经济体外，马来西亚、泰国、新加坡等东南亚国家在人工智能治理方面也加快取得进展。其中，新加坡、马来西亚并未专门针对人工智能颁布强制性法规，选择借助灵活的非强制性政策文件开展柔性监管。如，新加坡发布《人工智能治理模型框架》《生成式人工智能治理模型框架》等多份指导性文件，马来西亚于 2024 年 9 月发布《国家人工智能治理与伦理指南》。泰国正在着手制定人工智能立法，《人工智能皇家法令草案》等多部法案已进入听证讨论环节。**二是中东地区积极争取人工智能发展和治理自主权。**沙特阿拉伯、阿联酋等通过积极主办或承办国际多边活动，力图在全球人工智能治理中谋得一席之地。如，沙特从 2020 年起，连续举办全球人工智能峰会，积极输出本国治理主

张。其中，沙特在 2024 年 9 月第三届全球人工智能峰会发起开创性倡议，推动符合伦理的人工智能研究与应用。此外，沙特还利用第 19 届联合国互联网治理论坛（IGF）主办契机，高调发布《利雅得宣言》，吸引全球目光。阿联酋连续多年举办全球最大的人工智能峰会之一“AI 万象”峰会，并于 2024 年 10 月批准通过《国际人工智能政策立场》，提出进步、合作、社会、伦理、可持续性和安全六项核心原则。

（二）促创新和促发展成为竞争重点，以专设机构提升政策质效

促创新和促发展成为各国战略竞争重点。由于人工智能在地缘政治和国防军事领域具有颠覆性影响，为抢占技术先发优势，各国纷纷加快战略部署和顶层设计。新加坡分别于 2019 年 11 月和 2023 年 12 月发布《国家人工智能战略 1.0》《国家人工智能战略 2.0》，以提升新加坡的经济发展水平和社会发展潜力。2019 年 12 月，韩国发布《人工智能国家战略》，提出了推动韩国“从 IT 强国向 AI 强国发展”的愿景，计划在 2030 年将韩国在人工智能领域的竞争力提升至世界前列。¹⁷2021 年 1 月，英国人工智能办公室发布《人工智能路线图》，明确英国人工智能发展战略和重点，随后又推出国家层面的《人工智能战略》和行动计划，计划通过长期投资、应用推广、政府治理三大支柱支持人工智能产业发展。¹⁸2024 年 6 月，非洲通信和 ICT 技术部

¹⁷ 韩国公布“人工智能国家战略”，人民日报，2019 年 12 月，<http://industry.people.com.cn/n1/2019/1219/c413883-31513358.html>

¹⁸ 英国推进人工智能产业发展，新华网，2024 年 2 月 21 日，<https://www.xinhuanet.com/tech/20240221/6fd7671accb0475491228345c4c02507/c.htm>

长们共同批准了《非洲人工智能战略》和《非洲数字契约》，旨在释放新数字技术潜力、加速非洲数字化发展。¹⁹2024 年 7 月，塞尔维亚公布《2024-2030 年人工智能发展战略》，旨在为塞尔维亚培育充满活力的人工智能生态。2024 年 11 月，希腊发布《希腊人工智能转型蓝图》提案，概述利用人工智能造福希腊经济和社会的战略。2024 年 11 月，瑞典人工智能委员会发布《瑞典路线图》报告，详细阐述以安全、可持续的方式开发利用人工智能的战略。

设立人工智能专门机构统筹协调政策制定与实施。2024 年 1 月，阿联酋成立人工智能和先进技术委员会（AIATC），监督人工智能研究、基础设施、投资政策和战略的实施和发展，并与当地和全球合作伙伴制定计划和研究项目，以提升阿联酋在人工智能领域的地位。2024 年 5 月，欧盟委员会宣布在委员会内部设立人工智能办公室，负责推动人工智能的发展、部署和使用，促进社会和经济效益以及创新，同时降低风险。²⁰2024 年 9 月，韩国成立国家人工智能委员会，将带头推动韩国社会实现人工智能转型。该委员会是总统直属机构，由 30 名民间委员和 10 名政府委员组成，包括人工智能专家、部长级官员等。²¹2024 年 12 月，马来西亚总理宣布，正式设立马来西亚国家人工智能办公室（NAIO）²²。NAIO 将负责制定 2026-2030 年人工智能技术行动计划、人工智能应用监管框架、人工智能道德守则，加

¹⁹<https://au.int/en/pressreleases/20240617/african-ministers-adopt-landmark-continental-artificial-intelligence-strategy>

²⁰ https://ec.europa.eu/commission/presscorner/detail/en/ip_24_2982

²¹ <https://cn.yna.co.kr/view/Ack20240926000900881>

²² <https://www.mydigital.gov.my/malaysia-launches-national-ai-office-naio/>

速人工智能技术适应，研究政府人工智能影响，发布国家人工智能趋势报告，建设人工智能相关数据集等 7 方面工作。²³

（三）从软法规制到软硬法接轨并行，人工智能法治化进程加快

“软硬法”协同是现阶段全球人工智能治理的重要特征，混合治理成为多数国家选项。人工智能技术和治理经验的成熟推动各国从软法先行进入软硬协同的治理新阶段。欧盟强调保障人的基本权利，非正式的伦理指引和正式的法律规范双管齐下，发挥监管的规范性影响力。欧盟早在 2019 年就发布《值得信赖的人工智能伦理指南》，提出值得信赖的人工智能全生命周期框架，并于 2021 年就启动《人工智能法》的立法程序，及早迈入硬法治理阶段，仿照 GDPR 以期放大欧盟在规则领域的“布鲁塞尔效应”，弥补技术短板。美国联邦层面“让位”州立法先行先试，以为国家层面的统一立法奠定实施基础。尽管美国州层面已有部分人工智能立法落地生效，但是在促创新促发展的坚持下，联邦层面仍将以行政令、标准指南等软法治理为主。通过美国州立法与联邦软法治理相互补充和对照实施，从而为美国统一立法汲取经验。

更多软法逐步被转化纳入硬法范畴，人工智能法治化进程加快。一是国际软法向国内硬法渗透。如，OECD 2019 年通过《人工智能原则》，要求成员国按照 OECD 提出的人工智能原则制定政策和监管框架，为全球治理的可操作性奠定基础。目前，欧盟、美国等多个

²³ <https://ai.gov.my/about-naio>

国家已经在其国内立法中采用 OECD 提出的人工智能系统和生命周期定义，以及透明度、人权和民主价值观等原则。**二是国内软法被转化纳入国内硬法。**随着治理实践的检验，部分证明行之有效的政策和伦理性原则被转化为法律规制。如，美国“未来生命研究所”于 2017 年发布阿西洛马人工智能原则（Asilomar AI Principles），已经被美国加利福尼亚州直接写入州立法之中。**三是通过监督管理机制赋予软法强制效力。**部分国家通过对科技伦理施加违规责任，实现人工智能伦理强制性监管目的。如，对没有履行其人工智能自治承诺的科技企业，美国联邦贸易委员会可将其视为“不公平的或欺骗的”商业活动并采取相应的措施，赋予行业自治以强制性法律效力。

（四）综合立法和领域立法齐头并进，生成式人工智能关注剧增

欧盟、韩国、巴西等综合立法取得进展同时，自动驾驶、金融、医疗等高风险应用领域立法正加快推进。自动驾驶领域，2024 年 12 月，美国国家公路交通安全管理局（NHTSA）发布“自动驾驶车辆安全、透明度和评估计划”拟议规则，旨在简化汽车制造商提交的豁免申请的审查流程，允许在不配备方向盘或刹车踏板等传统控制装置的情况下部署自动驾驶车辆，并鼓励自动驾驶汽车企业提供更多数据，增强公众对自动驾驶技术的信任。**金融监管领域**，2024 年 10 月，中国香港特别行政区政府发布《有关在金融市场负责任地应用人工智能的政策宣言》，提出将采取双轨模式，以促进金融服务业采用及发展人工智能，并同时应对网络安全、数据私隐及知识产权保护等潜在挑

战。**医疗产品领域**，多数国家正在制定相关监管框架。2024 年 3 月，美国食物和药品管理局（FDA）发布《人工智能与医疗产品：CBER、CDER、CDRH 和 OCP 如何协同工作》报告，促进人工智能医疗产品负责任和合乎伦理的开发、部署、使用和维护。此外，美国 FDA 还出台了《机器学习医疗设备的透明度：指导原则》《最终指南：人工智能/机器学习支持设备软件功能的预定变更控制计划的营销提交建议》等文件，旨在完善人工智能医疗设备监管框架。

随着生成式人工智能技术的持续突破和颠覆性影响，各国对于生成式人工智能伴随带来的风险挑战加大关注力度。中国于 2023 年 7 月出台的全球首部生成式人工智能监管法规《生成式人工智能服务暂行管理办法》，规定了生成式人工智能的安全评估和算法备案制度，以促进生成式人工智能健康发展和规范应用。**美国联邦层面**，尚未出台针对生成式人工智能的专门立法，行业标准、协会指南等非强制性指引为企业提供务实指导。2024 年 7 月，国家标准与技术研究院（NIST）发布《人工智能风险管理框架：生成人工智能概况》自愿指南，重点关注生成式人工智能特有的 12 种风险，提出 200 多项建议措施，鼓励企业加强人工智能系统风险管理。**美国州层面**，围绕生成式人工智能数据、算法、内容等核心关注的相关立法提案数量大幅增加。如，加利福尼亚州签署多部生成式人工智能监管法案，对人工智能生成内容水印要求（SB 942）、生成式人工智能透明度（AB 2013）、政府使用生成式人工智能（SB 896）等作出规定。**欧盟《人工智能法》**对“通用人工智能模型”作出专章规定，生成式人工智能作为通用人工智

能模型的典型范例，一并受到调整。2024 年 11 月，欧盟委员会正式发布《通用人工智能实践准则》草案，针对透明度、系统性风险分类与治理、技术措施等，拟为通用人工智能模型提供者提供具体指导。此外，欧盟数据保护监督机构 EDPS 还发布了《生成式人工智能与 EUDPR》指南，旨在为欧盟机构、部门、办公室和机构（EUIs）在使用生成式人工智能系统处理个人数据时提供实际操作建议。新加坡于 2024 年 5 月，发布《生成式人工智能治理模型框架》，提出问责、数据质量、可信开发和部署、事件报告、测试和保证等九大维度，以促进构建生成式人工智能的可信生态系统。

（五）风险分级和分类规制不断完善，前沿大模型成为规范重点

当前，风险分级和分类规制成为各国人工智能规范的主要进路。美国 NIST 发布的《人工智能风险管理框架》（AI Risk Management Framework）就以风险作为人工智能治理的核心，新加坡的《模范人工智能治理框架》（Singapore Model AI Governance Framework）也将风险视为人工智能治理的主要议题，欧盟《人工智能法》更是成为风险规制的典型代表。欧盟《人工智能法》将人工智能系统划分为禁止性风险、高风险、有限风险、低风险或无风险四类。其中，禁止性风险主要包括利用潜意识技术扭曲个人或群体行为、利用人工智能危害特定弱势群体、进行社会信用评分、利用人工智能进行特定侵入性的执法等类型。美国拜登政府第 14110 号行政令进一步强调解决人工智能系统的安全风险，包括 CBRN（化学、生物、放射性和核武器）风

险、生物安全风险、关键基础设施和网络安全风险以及国家安全风险，要求联邦政府相关部门针对不同类型的风险采取缓解措施。美国科罗拉多州人工智能法案《关于在与人工智能系统的互动中提供消费者保护的法案》（SB205）针对高风险人工智能系统予以重点规制，要求开发者尽到合理注意义务，采取合理措施保护消费者免受算法歧视风险。巴西参议院近期通过的《人工智能法案》部分参考了欧盟的风险分级思路，要求人工智能系统在投放市场或投入使用前进行风险评估，按照风险程度分为“过度风险”“高风险”和“其他”三类。韩国《人工智能基本法》同样采用了基于风险的监管框架，对高风险或高影响的人工智能系统施加严格要求，包括事先通知、确保安全可靠等。

越来越多国家倾向将前沿大模型作为规范重点，并将模型算力作为风险评估的主要标准。如，欧盟《人工智能法》将模型的参数数量、数据集质量和计算量等作为模型能力或影响力评估的标准，明确认定计算能力大于 10^{25} 浮点运算次数的通用人工智能模型具有系统性风险，要求具有系统性风险的通用人工智能模型提供者履行评估测试、风险缓解、跟踪记录、安全防护等义务。又如，美国拜登政府第 14110 号行政令引入“两用基础模型”概念，即经海量数据训练、自我监督为主、包含百亿参数、能够广泛适用，并在执行国家安全事务中表现出高水平性能的人工智能模型。2024 年 9 月，为落实该行政令要求，美国商务部发布《建立先进人工智能模型和计算集群开发的报告要求》拟议规则，要求算力大于 10^{26} 浮点运算次数的两用基础模型开发者每季度报告模型权重、“红队测试”结果等信息。此外，2024 年 8 月，

美国加利福尼亚州议会提出具有争议的《前沿人工智能模型安全创新法案》，虽未能最终被州长签署成法，但是反映了美国州层面以算力作为评估人工智能系统风险标准并进行监管的部分做法。该法案规定，训练算力超过 10^{26} 浮点运算次数或训练成本超过一亿美元的人工智能模型被纳入“受监管模型”范围。

三、全球人工智能法律政策未来展望

当前全球人工智能发展已经进入了新阶段，技术的进步正在深刻改变人类社会生活。面向人工智能时代，如何确保人工智能技术的发展安全、可靠、可控，且始终运行在向善的轨道上，已经成为全球需要共同面对的挑战。

（一）人工智能立法更趋理性且更加全面系统

近年来，全球与人工智能相关的立法数量显著增长。随着人工智能技术的深入发展和广泛应用，各国政府愈发重视以全面且系统的法律来应对人工智能复杂性。面向未来，人工智能的立法将更趋理性且更加全面系统，不仅关注技术本身的规范，也关注技术带来的多方影响。

一是由“碎片式立法”转向“体系化构建”。尽管全球各方均将人工智能治理作为一项紧迫且重要的议题，但总体而言，人工智能立法目前还处于探索阶段，高位阶、综合性立法数量较少。人工智能涉及众多复杂的交叉领域，如数据隐私保护、知识产权归属、产品责任界定、劳动就业影响等。未来各国立法将趋向于构建起一套涵盖各领域相互衔接、协调统一的法律体系，避免因法律规定的冲突或空白导致的治

理困境。

二是由“单一取向”转向“综合考量”。前期，面对人工智能技术发展带来的风险挑战，世界主要国家和地区多从自身产业发展和国家利益出发，设计偏向单一价值取向的法律规范，如欧盟构建以风险为基础的监管框架、美国强调本国人工智能产业创新发展。然而，随着人工智能战略地位进一步凸显、影响更加深远，各国关于人工智能治理的思考更加深入，未来将偏向于综合考虑“发展”和“安全”等多元价值取向，构建更趋理性的人工智能法律体系。

（二）人工智能风险治理与伦理考量不断深化

由于人工智能技术的特殊性，既需要以法律规则为代表的强制性要求，规制与国家安全、社会公共利益和个人权益密切相关的治理问题，也需要以伦理指南为代表的规范性要求，引导产业、企业进一步履行社会责。面向未来，人工智能治理将转向风险治理“硬法”规制与科技伦理“软法”规范双轨并行的深度治理。

在风险治理方面，当前，已有部分国家或地区探索形成风险分级分类管理制度。随着人工智能技术日益融入经济社会发展各领域全过程，其安全风险面不断扩大，需要充分考虑该技术对不同行业的渗透程度，细化可能出现的风险，探索更加敏捷、精准的风险识别与分类机制，把治理落实到实践中。

在伦理规范方面，人工智能伦理治理国际交流合作将不断加强。我国在多个国际场合下都提出了“以人为本、智能向善”的主张，引导全球在人工智能伦理的轨道上发展智能技术。为应对人工智能技术发展全球化特征，全球各国家和地区、国家组织将不断深化人工智能伦

理治理交流，谋求形成共识。同时，人工智能伦理治理技术化、工程化、标准化趋向将不断深化，全球人工智能伦理治理经验将不断丰富。

（三）人工智能创新与发展政策力度持续加大

人工智能的新技术不断突破、新业态持续涌现、新应用加快拓展，已成为新一轮科技革命和产业变革的重要驱动力量。今年以来，为应对国际间竞争的复杂局势，各国相继推出支持人工智能发展的政策举措。面向未来，为充分释放人工智能对经济社会发展的红利，各国政府将进一步加大对人工智能领域的政策扶持力度。

一方面，战略部署和政策支持力度只增不减。加快发展人工智能是事关能否抓住新一轮科技革命和产业变革机遇的战略问题。在此背景下，各国将进一步扩大政策支持范围，利用体制、机制、资金、人才等各类资源全方位助力人工智能技术的研发创新和广泛应用。如，美国白宫科技政策办公室就《人工智能行动计划》征求意见，涵盖从人工智能技术研发到实际应用等多个重要领域，旨在通过一系列政策行动保障美国在人工智能领域的全球领先优势。未来，伴随全球人工智能竞争进入白热化阶段，各国关于人工智能的战略部署将不断深入，相关优惠政策将密集出台。

另一方面，通过立法扫清人工智能发展制度阻碍。人工智能技术的突破性发展对相关法律体系提出更高需求，但现有法律制度并及时进行相应挑战，如，数据领域存在数据“不能用”、“不够用”、“不好用”等问题，这些问题对人工智能技术产业的发展构成了制约。面对新一代人工智能的发展需求，各国需要通过不断完善相关法律制度、构建创新合作框架等一系列举措，促进人工智能的健康发展。

（四）全球人工智能治理合作面临多方面挑战

随着人工智能技术的不断成熟和应用场景的不断扩展，全球范围内的竞争与合作将更加复杂多样，国际合作在人工智能治理中的重要性日益凸显。然而，由于人工智能治理与地缘政治、经济竞争日益交织，全球人工智能治理合作进程仍面临多方挑战。

一是全球人工智能治理合作面临较大不确定性。近年来，主要经济体依托联合国、二十国集团、金砖国家等双多边机制与平台，围绕全球人工智能治理开展对话协商，取得了部分关键成果。但是，随着人工智能技术的迅猛发展和地缘政治格局的最新变化，发达国家与发展中国家之间的智能鸿沟问题日益凸显，贸易摩擦等因素不可避免地渗透到人工智能治理合作领域。从分歧中寻找共识、在竞争中寻求合作，已成为多国领导人和全球科技巨头的共同目标。

二是全球治理方案碎片化趋势恐将进一步加剧。目前全球人工智能治理总体仍处于碎片化、阵营化状态，缺乏统一的全球协调机制。以联合国为代表，各个组织发布的众多治理方案及监管措施能否协调落地，成为全球人工智能治理秩序有效运行的突出问题。与此同时，各国关于人工智能安全可信、隐私保护、伦理道德、军事应用等领域的法律政策存在较大差异，美欧两大经济体间的理念分歧进一步扩大，预计未来全球治理方案碎片化趋势仍将延续。因此，打造一个全方位、多层次、汇聚广泛共识，具有真正的包容性、平等性、多元性的治理框架将是未来全球人工智能治理合作努力的方向。

中国信息通信研究院 政策与经济研究所

地址：北京市海淀区花园北路 52 号

邮编：100191

电话：010-62302581

传真：010-62304980

网址：www.caict.ac.cn

