

生成式人工智能的应用与治理

生成式人工智能的发展使人工智能技术成为第四次工业革命中最具影响力的创新之一，已成为科技发展的新高地、产业的新赛道和经济发展的新引擎，具有巨大的发展潜力和广泛的应用前景。本文通过梳理生成式人工智能的主要应用场景，总结生成式人工智能的发展堵点与风险挑战，进而形成生成式人工智能发展的治理体系，为生成式人工智能的健康发展提供有益支持。

文 | 蔺欣 国家发展和改革委员会 国家投资项目评审中心工程师

生成式人工智能是一种能够生成新的数据或内容的人工智能（AI）技术，通过深度学习模型生成文本、图像、视频等内容，显著提升人机交互的效率，极大扩展人工智能的应用场景，已成为科技进步的新高地、产业升级的新赛道、经济发展的新引擎，是加快发展新质生产力的关键要素。

一、生成式人工智能的应用场景

在数字化时代，人工智能（AI）正迅



赛迪网官方微信



数字经济官方微信

速成为创新和生产力提升的关键驱动力。特别是生成式 AI，它通过模仿学习现有数据模式，创造出全新的内容，从而拓展了人类创造力的边界，目前已被广泛应用于智能办公、生产制造、数字教育、政务服务等多领域场景。

（一）办公场景

随着人工智能技术的飞速发展，生成式人工智能正逐渐成为办公室和会议室中的得力助手，通过文字、语音和图像处理能力的跃迁革新传统的工作流程，

使得办公和会议活动更加智能化和高效。

智能办公方面，基于大模型的解决方案能够覆盖日常办公中的多样化需求。无论是文案撰写、演示文稿美化，还是复杂数据分析，均可以通过简单的自然语言交由智能系统处理。

智能会议方面，从会议的策划阶段开始，大模型就能够根据会议的主题和其他提示词，自动生成包括会议环节、分论坛设置、时间安排到预算编制在内的完整策划方案。会议进行时，大模型支持准确及时的同声传译。

此外，通过 AI 处理后的会议记录结构清晰、要点明确，极大地提升了会后信息回顾和决策的效率。

（二）教育场景

生成式人工智能正以其强大的数据处理和学习能力，在教育领域引发革命性的变化。通过将大模型融入教学和科研过程之中，既提高了教学科研的效率，也极大地扩展了教育视野。

在科研领域，大模型能够快速分析大量实验数据，预测实验结果。在某些情况下，能够在短时间内提出比传统实验方法更为简洁且成本效益高的解决方案。通过自动化的设计和模拟过程，缩短科研项目从构思到实验的周期，降低研发成本，从而加速科研成果的产出和转化。在教学方面，生成式人工智能作用同样不容小觑。不仅可用于生成定制化的学习材料，帮助学生理解复杂概念，还能够辅助教师制定教学大纲、生成讲义、设计课堂活动。

（三）政务场景

生成式人工智能已成为改善公共服务的重要工具，能有效提升政府服务的效率和质量，优化公共治理体系、提升政务服务能力。

以大模型为代表的生成式人工智能技术可作为政民互动虚拟助手替代部分政务服务人员的工作，能够根据具体场景更加精准高效地提供个性化的回复和服务，大大降低人工成本。

在行政执法的过程中，借助生成式人工智能通用大语言模型的基本特质及其大数据分析能力，有利于对不同情境中的同类行为进行比对和辨识，进而根据法律规定做出相应的行政行为。

二、生成式人工智能的发展堵点与风险挑战

（一）算力需求与算力成本增加引发的算力瓶颈

算力是 AI 大模型的关键基础设施，算力决定了模型的“智商”。大模型通常包含数亿，甚至数千亿的参数，需要大量的算力来训练和运行。这些模型的训练过程涉及大量的数据和复杂的计算，导致对高性能计算资源的需求不断攀升。例如，华为预计到 2030 年，全球数据年新增 1YB；通用算力增长 10 倍到 3.3ZFLOPS，AI 算力增长 500 倍超过 100ZFLOPS。

据人民网财经研究院发布的报告，以 ChatGPT 为例，微软 Azure 云服务为其提供了 1 万颗英伟达 A100GPU，这个算力也正是国内云计算技术人士共识的 AI 大模型门槛。同时，训练和运行 AI 大模

型需要巨大的经济投入。例如，训练一个大型模型可能需要使用上万颗 GPU，单次模型训练成本可能超过千万美元。这其中不仅涉及昂贵的硬件成本，同时需要大量的电力和散热系统支持，初期投资巨大。为了解决算力瓶颈问题，行业仍需探索包括改进模型结构、训练方法的优化、使用新型计算架构如 DPU、大数据技术的应用等多种技术创新和优化途径。

（二）工具不当使用带来的虚假信息误导和欺诈问题

AI 大模型的火爆使其被广泛应用于生产生活之中，在带来便利的同时也使得部分虚假信息“以假乱真”，由于大模型尚不能完美辨别用户请求的合法性与正当性，当其被不当使用或恶意利用时，可能会带来一系列安全风险，扰乱公共服务秩序。例如，AI 换脸和拟声技术的发展使得视频内容真假难辨，增加了欺诈行为的隐蔽性，同时也使得诈骗行为“批量化”，通过生成虚假的电子邮件或短信等形式，伪装成信誉良好的机构或个人进行诈骗。

此外，生成式 AI 技术的滥用还可能对社会秩序和国家安全构成威胁。一些不法分子利用 AI 技术编造虚假事件，制造社会恐慌，或者借社会热点造谣炒作，干扰公众视听。例如，中国互联网联合辟谣平台对 2023 年 10 月网络谣言的梳理分析，当月网上数据监测和网民举报显示，臆测编造民生政策、利用 AI 虚构案事件、借社会热点造谣炒作等传谣现象较为突出，自媒体利用网民对社保、

医疗等民生话题的关切，编造涉公共政策、社会民生领域不实信息，这类迷惑性强的信息在“三人成虎”的传播效应下，严重干扰了社会认知和预期。这就要求我们在技术创新的同时，认识到其所带来的内外部风险挑战，最大化地发挥 AI 技术的正面作用，防范和减少其潜在的负面影响。

（三）现有法律体系下对内容所有权和责任归属的重新审视

现有的法律体系可能难以适应 AI 生成内容的所有权和责任归属问题。传统的版权法是为了保护人类创作者的智力成果而设计的。AI 作为一个非人类主体，其生成的内容是否应该享有版权，以及谁应该拥有这些权利（AI 开发者、用户、还是 AI 本身），目前尚无明确的法律规定。

在创作性质方面，许多法律体系要求作品必须是人类智力劳动的结果才能获得版权保护。AI 生成的内容虽然可能具有创造性，但其“创作”过程是由算法和数据驱动的，这是否符合法律对“创作”的定义存在疑问。

在责任归属方面，当 AI 生成的内容涉及诽谤、侵犯隐私或其他违法行为时，确定责任归属变得复杂，是由 AI 开发者负责，还是由使用 AI 工具的使用者负责，抑或是 AI 本身是否能成为责任主体，这些问题在现有法律框架下难以解答。

此外，现行合同法尚未能适应 AI 生成内容在商业交易中应用的新情况，即当 AI 生成的内容被用于广告或产品描述，而内容不准确或具有误导性时，合同中的责任和义务的分配问题。应对这

些挑战，尚需要法律专家、技术开发者、政策制定者和社会各界人士的共同努力，推动法律与时俱进，确保 AI 技术的健康发展，保护个人和公共利益。

（四）学术工作原创性、真实性、安全性的新挑战

生成式人工智能的快速发展和应用，虽为学术研究和知识生产带来了便利，但同时也对学术诚信提出了前所未有的挑战。学生和研究人员可以通过 AI 工具有生成论文或文章，这不仅导致学术成果的原创性受到质疑，而且可能引起广泛的剽窃和抄袭问题。

AI 工具能快速生成大量文本，这使得传统的学术不端行为检测方法变得不再有效，学术机构必须重新考虑如何确保学术工作的诚信和质量。面对这一挑战，全球范围内的学术机构对 AI 工具的使用持谨慎态度。例如，加利福尼亚大学伯克利分校明确表示 ChatGPT 等 AI 工具并非学校支持的工具，并强调教师在使用这些工具时需要自行审查其可访问性、隐私性和安全性。

为了应对 AI 工具可能带来的风险，一些学校已经采取了具体的行动。他们明确禁止将学生作业上传至 AI 工具进行分析，避免学生的作业成为训练数据，防止潜在的查重泄密风险。然而，仅禁止使用 AI 工具并不能从根本上解决问题，学术机构需要与技术开发者、法律专家等共同探讨研究如何在保护学术诚信的同时，合理利用 AI 技术促进学术研究的发展。

三、生成式人工智能的治理体系

治理体系和治理能力是制度和制度执行能力的集中体现。面对人工智能的技术涌现超越目前制度体系约束范围的新情境，需要强化大模型的场景牵引作用，通过构建完善的治理体系，疏通发展堵点、防范潜在风险，使生成式人工智能有效支撑我国经济社会高质量发展，创造新价值、适应新产业、重塑新动能。

（一）构建协同治理体系，促进多方治理参与

面对生成式 AI 这一应用场景广泛、深入生产生活的技术工具，对技术的规范运用要求多方治理主体协同并进。政府机构应制定涵盖数据隐私、知识产权、内容监管和伦理标准等方面的明确政策和法规，为 AI 技术的健康发展提供指导和规范。技术开发者和企业应建立行业自律机制，制定并遵守行业标准和最佳实践，确保 AI 技术的负责任使用。学术研究机构应开展 AI 伦理、法律和社会影响的研究，提供科学的决策支持，加强 AI 治理相关的课程和培训。鼓励公众参与 AI 治理的讨论和决策过程，提高公众对 AI 技术的认识和理解，增强社会对 AI 治理的监督和反馈。

（二）完善制度约束体系，有序加强技术监管

建立有效的 AI 大模型监管体系，首先需要法律与政策的支持。《互联网信息服务深度合成管理规定》作为我国首个针对 AI 生成内容的管理规定，为监管提供了法律框架，要求服务提供者进行备案、安全评估，并明确了内容标识、

监督检查与法律责任。同时，监管体系应包括技术规范和标准，确保 AI 技术的透明性和可解释性。政府应采取渐进式监管策略，平衡技术创新与风险控制，制定前瞻性措施以应对技术可能带来的不良后果。

此外，需要强化网络谣言的打击力度，通过中国互联网联合辟谣平台指导网站平台强化监测查证、开展排查整治，提升公众的信息识别能力，增强防范意识。增强 AI 技术的透明度，让外部监管者和用户能够理解 AI 的决策过程，使增进公众信任成为监管机构有效监管的基础。

（三）优化数据标准体系，推动数据流通共享

生成式 AI 的技术发展离不开高质量的数据集，数据标准与共享是构建有效治理体系的关键组成部分。为了确保生成式 AI 的广泛应用和跨平台操作，需要制定统一的数据标准，包括数据的收集、存储、处理和交换格式，确保不同来源和类型的数据能够被 AI 系统有效利用。同时，需要建立数据共享平台和机制，促进公共部门、私营企业和研究机构之间的数据流通，在确保数据共享过程中的隐私保护和安全的同时，打破信息孤岛，提升数据的利用效率。

此外，AI 训练过程中所使用的语料库中的专利权和版权问题也不容忽视，需要明确数据使用权、版权归属和利益分配机制，保护原创内容的权益，同时允许合理的数据使用以促进 AI 技术的发展。

（四）推进国际合作体系，加强政策对话协调

AI 技术的快速发展和广泛应用，尤其在跨境数据流动、国际电子商务、全球信息传播等方面的作用，使得任何单一国家都难以独立解决其带来的全部治理挑战。因此，生成式 AI 的治理不仅需要国家层面的努力，更需要跨国界的合作与协调。

各国政策制定者应加强对话，协调各自的 AI 治理政策，以减少国际冲突，促进 AI 技术的全球普及和应用。鼓励国际间联合研究项目的开展，聚焦于 AI 治理的关键问题，如算法透明度、偏见检测和纠正、责任归属等。制定和采纳国际认可的 AI 伦理标准和操作规范，确保 AI 技术的开发和应用符合全球公认的价值观念和原则。充分利用联合国、世界贸易组织等多边机构的平台作用，推动全球 AI 治理的多边讨论和合作。

责任编辑：金烨 jinye@staff.ccidnet.com