

生成式人工智能垄断风险及其合作监管

王 磊^{*}

内容提要：生成式人工智能领域的垄断风险日益显现，如何监管这一新型垄断风险成为智能时代全球共同关注的重大理论与实践命题。生成式人工智能垄断风险的核心生成逻辑在于，大型科技企业控制了数据、算力等生成式人工智能发展的关键投入，制造了市场进入障碍。实践中，生成式人工智能垄断风险的主要形态是关键投入获取限制、基础模型访问限制等。生成式人工智能垄断风险的属性与构造十分复杂，单纯依靠反垄断执法机构，无法保证监管的相称性、动态性与回应性，应当构建“反垄断执法机构—大型科技企业”互动合作的监管模式。大型科技企业可以在基础设施层、开发层、部署层分层开展自我监管。反垄断执法机构则可以从维护生成式人工智能市场竞争、促进关键要素自由流动的向度展开外部监管。自我监管与外部监管应当合作、协同与紧密互动。

关键词：生成式人工智能 垄断风险 自我监管 外部监管 合作监管

一、问题聚焦：监管如何跟进生成式人工智能领域的垄断风险

随着人工智能技术加速迭代、人工智能能级不断提升，以 ChatGPT、DeepSeek 为代表的生成式人工智能模型已经具备了自我学习、自我成长的能力，并通过技术赋能等方式实现了对经济、政治等复杂系统的深层次嵌入，推动生产力与生产关系持续变革。鉴于人工智能技术，尤其是生成式人工智能技术蕴含的巨大价值，大型科技企业都在积极争夺生成式人工智能技术的制高点，一些大型科技企业试图控制生成式人工智能发展的关键投入（key inputs），并借此强化自身

^{*} 王磊，中国政法大学民商经济法学院副教授。

本文为司法部法治建设与法学理论研究部级科研项目“生成式人工智能关键投入垄断的合作规制研究”（25SFB3031）的研究成果。

市场力量，排除、限制竞争对手，以 Google、Microsoft 为代表的“技术利维坦”〔1〕已具雏形。

生成式人工智能领域日益显现的垄断风险引发了理论界、实务界的高度关注。2023 年 10 月，时任美国总统的拜登签署《关于安全、可靠和可信赖的人工智能开发和使用的行政命令》，要求“负责制定与人工智能相关的政策和法规的各机构负责人应当依法、合理地行使其职权，促进人工智能及相关技术领域的竞争。此类行动包括应对涉及关键投入的集中控制所产生的风险，采取措施制止非法共谋，防止具有支配地位的企业迫使竞争对手处于不利地位，并努力为小企业和企业家提供新的机会”〔2〕。一些国家和地区对生成式人工智能领域的投资和并购活动启动了初步调查程序，如 2023 年 11 月 15 日，德国联邦卡特尔局对 Microsoft 和 OpenAI 之间的合作协议进行了审查，以判断两者之间的投资合作关系是否构成反垄断法意义上的经营者集中。随后，欧委会和英国竞争与市场监管局先后宣布对 Microsoft 和 OpenAI 之间的投资合作是否符合其反垄断法律制度展开调查。在中国，2025 年 9 月 15 日，国家市场监督管理总局发布公告称，经初步调查，英伟达违反《反垄断法》和《市场监管总局关于附加限制性条件批准英伟达公司收购迈络思科技有限公司股权案反垄断审查决定的公告》（国家市场监督管理总局公告〔2020〕第 16 号），国家市场监督管理总局依法决定对其实施进一步调查。可见，生成式人工智能领域的垄断风险已经成为世界主要国家（地区）反垄断监管的核心关切。

既有研究也已经关注到生成式人工智能领域的垄断风险，并主要从“本体论”、〔3〕“要素论”、〔4〕“场景论”〔5〕与“行为论”〔6〕等向度进行了富有洞见的探讨，但面对生成式人工智能技术飞速发展和垄断风险日益加剧的新形势，仍有待深化：其一，分层分级深度解构生成式人工智能垄断风险的研究相对缺乏；其二，既有研究对大型科技企业在反垄断监管系统中的角色定位与功能发挥关照不足；其三，针对监管究竟应当如何跟进生成式人工智能领域日益显现的垄断风险，适度、动态与敏捷的监管模式如何建构，监管方案如何设计等关键议题的研究严重不足。

鉴于此，本文在深入剖析生成式人工智能垄断风险的生成逻辑，并对垄断风险展开层级解构的基础上，尝试从“合作监管”这一新的监管视角出发，构建“反垄断执法机构—大型科技企业”互动合作的监管分析框架与监管体系，提升研究的系统性。生成式人工智能垄断风险监管不是一个区域性、阶段性议题，而是人类社会在智能化进程中必然要面对的全球性、普遍性议题，深入研究生成式人工智能垄断风险的监管应对，保证反垄断监管的适度、动态与敏捷，对于促进生成式人工智能健康发展具有重要意义。

〔1〕 参见范如国：《平台技术赋能、公共博弈与复杂适应性治理》，载《中国社会科学》2021 年第 12 期，第 139-140 页。

〔2〕 Sec. 5.3 of Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence.

〔3〕 一些学者对生成式人工智能垄断风险的内涵、形态、影响展开研究，相关内容参见孙晋、胡旨钰：《开源模式下生成式人工智能的反垄断监管跟进》，载《西北工业大学学报（社会科学版）》2025 年第 4 期。

〔4〕 有学者从数据、算力、算法等生成式人工智能的关键投入切入，考察预训练数据垄断、算力垄断以及模型垄断对市场竞争的负面影响，相关研究参见杨东：《人工智能的垄断风险及其规制》，载《中国市场监管研究》2024 年第 8 期。

〔5〕 一些学者关注模型开发、模型部署等不同场景下的垄断风险，主张依场景之不同对生成式人工智能垄断风险展开规制，相关研究参见宁立志、孙慕野：《复杂系统视域下生成式人工智能的平台化及其监管路径》，载《西北工业大学学报（社会科学版）》2025 年第 3 期。

〔6〕 此类研究聚焦基础模型访问限制、数据自我优待等新型垄断行为及其规制策略设计，具体内容参见王晓晔：《我国平台经济反垄断监管“欧盟模式”批判》，载《法学评论》2024 年第 3 期。

二、生成式人工智能垄断风险的成因探求

生成式人工智能的创新发展高度依赖数据、算力等关键要素的持续性投入，大型科技企业对关键投入的集中控制制造了巨大的市场进入障碍，网络效应的存在、市场力量的价值链传导进一步提高了进入壁垒，引发严重的竞争担忧。

（一）生成式人工智能关键投入的集中控制

第一，数据的集中控制。数据是生成式人工智能的基石。没有数据，就没有生成式人工智能。^{〔7〕}无论是在预训练阶段，还是在微调阶段，生成式人工智能都需要大量的数据用于训练。^{〔8〕}生成式人工智能的训练数据主要包括三类：公开可用的数据（publicly available data）、合成数据（synthetic data）以及专有数据（proprietary data），其中专有数据又分为自有数据（first party proprietary data）与第三方专有数据（third party proprietary data）。^{〔9〕}公开可用的数据主要包括网络爬取数据和互联网中对外开放的开源数据。合成数据是指生成式人工智能企业根据真实数据的基本特征、分布，通过算法自主生成的数据。自有数据是指企业自己持有的数据。相比之下，第三方专有数据是指其他企业持有的数据。上述三类数据的获取难度存在差异。一般而言，公开可用的数据最易获取，合成数据次之，第三方专有数据最难获取。^{〔10〕}

由于公开可用的数据获取门槛较低，很多企业利用公开数据进行模型训练，OpenAI、DeepSeek 均是如此，而且它们主要是通过网络爬取的方式获取数据。但是，爬取数据存在侵犯著作权、隐私权等多重法律风险。^{〔11〕}2023 年 2 月，Getty Images 起诉 Stability AI 公司，指控 Stability AI 公司未经许可复制其数据库中超过 1200 万张图片用于生成式人工智能模型训练。^{〔12〕}可见，即便是公开可用的数据，将其投入模型训练仍然会遭受合法性诘难，而且，公开可用的数据不是无限的，由于数据训练规模的持续扩大，公开可用的数据可能会枯竭，从而无法支持生成式人工智能技术的持续发展。^{〔13〕}就合成数据而言，合成数据存在数据生成的数量、质量不稳定等问题，^{〔14〕}对模型性能的影响较大。基于这一担忧，一些研究认为，生成式人工智能模型训练的主要数据来源可能会从公开可用的数据转向专有数据。^{〔15〕}但是，专有数据的获取难度较大，一方面，大型科技企业对自有数据的管控非常严格，另一方面，一些大型科技企业往往还会与数

〔7〕 See Christoph Groger, *There is no AI without data*, 64 Communication of the ACM 98, 98 (2021).

〔8〕 参见〔美〕凯特·克劳福德：《技术之外：社会联结中的人工智能》，丁宁、李红澄、方伟译，中国原子能出版社、中国科学技术出版社 2024 年版，第 83-86 页。

〔9〕 See CMA, *AI Foundation Models: Technical Update Report*, (16 April 2024), <https://www.gov.uk/government/publications/ai-foundation-models-update-paper>, visited on 7 May 2026.

〔10〕 Ibid.

〔11〕 参见许可：《数据爬取的正当性及其边界》，载《中国法学》2021 年第 2 期，第 177 页。

〔12〕 See Getty Images (US), Inc. v. Stability AI, Inc., No. 1: 23-cv-00135 (2023).

〔13〕 See OECD, *Artificial Intelligence, Data and Competition*, OECD Artificial Intelligence Papers (24 May 2024), https://www.oecd.org/en/publications/artificial-intelligence-data-and-competition_e7e88884-en.html, visited on 8 May 2026.

〔14〕 See Peter Lee, *Synthetic Data and the Future of AI*, 110 Cornell Law Review 1, 35-37 (2025).

〔15〕 See CMA, *AI Foundation Models: Initial Report*, (18 September 2023), <https://www.gov.uk/government/publications/ai-foundation-models-initial-report>, visited on 8 May 2026.

据持有者签订协议，排除、限制竞争对手对于第三方专有数据的访问，进一步强化自身数据优势。例如，2024年2月，Google与Reddit达成合作协议，约定Google每年向Reddit支付6000万美元，作为对价，Reddit向Google提供其5000万用户产生的实时新增数据。协议生效后，“Reddit阻止其他爬虫对其网站进行索引，授予Google搜索引擎对Reddit的排他性访问权”^{〔16〕}。综上，大型科技企业对数据的集中控制使得生成式人工智能模型训练数据的取得存在一定的限制，数据获取难度增加可能会成为初创企业进入生成式人工智能市场的重大障碍。

第二，算力的集中控制。如同工业时代的电力一样，算力已经成为智能时代的基础能源。^{〔17〕}生成式人工智能模型的训练、运行更是离不开强大、稳定的算力供给。但算力资源，尤其是高端训练算力极为稀缺，供求严重失衡，生成式人工智能企业获取、调用算力资源的能力在一定程度上决定了其市场竞争力。企业获取算力资源的途径主要有两条：其一，自建算力设施；其二，购买云计算服务。^{〔18〕}自建算力设施的成本十分高昂，通常只有少数大型科技企业能够负担，当自建算力设施费用过巨时，潜在进入者只能转向租用云计算服务提供商的算力资源以开发、训练模型。但是，云计算服务市场属于结构性垄断市场，算力资源集中控制在少数大型科技企业手中，“根据Synergy Research Group的测算，2026年第1季度，AWS在全球云计算服务市场的市场份额高达28%，领先于Azure的21%与Google Cloud的14%。三大云计算服务提供商合计占有持续增长的云计算服务市场逾60%的市场份额，其余竞争者的市场份额均停留在个位数低位”^{〔19〕}，AWS、Azure和Google Cloud等企业事实上掌握了在自己和竞争对手之间分配算力资源的能力。^{〔20〕}算力壁垒是潜在进入者参与人工智能市场竞争难以逾越的一道鸿沟。^{〔21〕}

第三，人才的集中控制。生成式人工智能产业是典型的人才密集型产业。模型开发、算法设计以及数据处理都需要大量的高水平专业人才。作为技术创新、产业发展的关键支撑，一直以来，高水平AI人才就是生成式人工智能企业竞逐的对象，AI人才向头部企业集聚的趋势十分明显。究其原因有四：一是AI人才缺口大，供求关系本身就很紧张；二是大型科技企业招贤力度大，能够提供良好的研发环境、交流平台以及富有吸引力的薪资待遇；三是为了维护自身竞争优势，一些大型科技企业通过经营者集中等方式直接将初创企业或者其他竞争对手的AI专家团队纳入麾下，建立人才护城河。Microsoft在收购Inflection后，“聘用了Inflection几乎所有的员工，包括其AI业务首席执行官。此类行为通过阻碍竞争对手获取在行业内立足或者发展所必需

〔16〕 Courtney C. Radsch, *Dismantling AI Data Monopolies Before it's Too Late*, (9 October 2024), <https://www.techpolicy.press/dismantling-ai-data-monopolies-before-its-too-late/>, visited on 8 May 2026.

〔17〕 参见王志勤主编：《算力：筑基数字化竞争力》，人民邮电出版社2024年版，第11-12页。

〔18〕 See OECD, *Artificial Intelligence, Data and Competition*, OECD Artificial Intelligence Papers (24 May 2024), https://www.oecd.org/en/publications/artificial-intelligence-data-and-competition_e7e88884-en.html, visited on 8 May 2026.

〔19〕 Felix Richter, *Big Three Hold Dominant Lead in Accelerating Cloud Market*, (5 May 2026), <https://www.statista.com/chart/18819/worldwide-market-share-of-leading-cloud-infrastructure-service-providers/>, visited on 10 May 2026.

〔20〕 See Office of Communications, *Cloud Services Market Study-Final Report*, (5 October 2023), <https://www.ofcom.org.uk/siteassets/resources/documents/consultations/category-3-4-weeks/244808-cloud-services-market-study/associated-documents/cloud-services-market-study-final-report.pdf?v=330228>, visited on 10 May 2026.

〔21〕 参见殷继国：《论人工智能时代反垄断法必需设施理论的适用》，载《财经法学》2025年第3期，第142页。

的人才，可能产生排挤竞争对手的效果”〔22〕。此外，大型科技企业还会通过签订竞业限制协议（通常包含严格程度远高于行业平均水平的竞业限制条款）等方式锁定 AI 人才，强化对人才资源的集中控制。〔23〕大型科技企业对 AI 人才的集中控制可能会抑制生成式人工智能市场的竞争。

（二）网络效应制造市场进入障碍

网络效应理论认为，网络效应是消费者数量的函数，加入网络的用户越多，用户使用网络产品（服务）的成本就越低，获得的价值越高。〔24〕生成式人工智能的技术特性决定了，生成式人工智能模型的部署、应用也具有十分显著的网络效应，模型的用户越多，“人一机”交互越充分，建立在高频“人一机”交互基础上的模型缺陷反馈就越及时，模型设计者根据用户的反馈可以精准改进模型性能，从而吸引更多的用户。

生成式人工智能场景下的网络效应具有极强的数据驱动特性。详言之，模型的用户越多，模型设计者收集到的用户数据就越多、越丰富。相应地，数据分析与提取也就越全面、充分。OpenAI 在其发布的 ChatGPT 用户使用规则中曾声明，ChatGPT 用户与模型交互产生的数据会被用来作为模型迭代的训练数据。其结果是，模型设计者有更多的机会与可能去改进模型性能，提供更好的服务。英国竞争与市场管理局（CMA）的研究证实了这一点，CMA 在《人工智能基础模型：初步报告》中指出：“生成式人工智能企业通过用户是否点赞或者复制模型输出内容来收集相关反馈数据（用户点赞或者复制意味着内容呈现符合需求与期望），进而优化、调整模型输出。”〔25〕这一过程也被称为“数据反馈循环”（data feedback loop），即“更多的用户—更广泛的反馈数据—更好的模型性能—更多的用户”的闭环结构。〔26〕对于用户规模庞大、性能优越的模型，用户通常会表现出较强的依赖性，难以转向其他竞争性模型。

网络效应导致新进入者难以获得初始用户、形成有效规模。在用户获取、数据迭代、成本分摊等方面，新进入者处于显著劣势，根本无法进入市场与大型科技企业竞争。受网络效应影响，生成式人工智能市场上的竞争呈现出强者愈强的态势，容易诱发结构性垄断风险。〔27〕

（三）市场力量的价值链传导

对于生成式人工智能而言，关键投入的准备、模型设计与开发、部署应用等核心环节由多元主体参与、共同创造价值，因而生成式人工智能已经具备价值链特征。〔28〕生成式人工智能价值链呈现出一定的层级结构，主要由基础设施层、开发层以及部署层构成，且各层联动紧密，共同

〔22〕 European Commission, *Competition Policy Brief: Competition in Generative AI and Virtual Worlds*, (September 2024), https://competition-policy.ec.europa.eu/document/download/c86d461f-062e-4dde-a662-15228d6ca385_en, visited on 9 May 2026.

〔23〕 参见〔美〕埃里克·波斯纳：《劳动力市场中反垄断的缺席》，胡寅译，格致出版社、上海人民出版社 2025 年版，第 147 页。

〔24〕 参见〔美〕赫伯特·霍温坎普：《数字平台企业反垄断救济新论》，李中衡译，商务印书馆 2023 年版，第 152 页。

〔25〕 CMA, *AI Foundation Models: Initial Report*, (18 September 2023), <https://www.gov.uk/government/publications/ai-foundation-models-initial-report>, visited on 9 May 2026.

〔26〕 参见〔美〕莫里斯·E. 斯图克、艾伦·P. 格鲁内斯：《大数据与竞争政策》，兰磊译，法律出版社 2019 年版，第 197 页。

〔27〕 参见王先林：《平台经济领域强化反垄断的正当性与合理限度》，载《苏州大学学报（哲学社会科学版）》2024 年第 2 期，第 74—75 页。

〔28〕 参见谢尧雯：《生成式人工智能价值链行政监管与侵权责任的匹配》，载《政法论坛》2025 年第 2 期，第 37 页。

对关键投入获取的限制可能会产生严重的封锁效应。

就数据获取限制而言，由于一些数据具有较高的质量与价值，大型科技企业可能会通过采取控制数据调用接口的频次、数据量、时长等技术措施，以限制此类数据流动。更为严重的，大型科技企业可能会拒绝开放数据接口、限制数据互操作性，借此完全阻断竞争对手获取数据的机会与可能。“Microsoft 就曾公开表示，如果竞争对手继续使用从 Bing 搜索引擎获得的数据进行生成式人工智能开发，将终止其数据访问权限。”^{〔33〕} 拒绝开放数据接口、限制数据互操作性可能会构成反垄断法上的拒绝交易行为。在生成式人工智能场景下，拒绝交易行为的认定需慎之又慎，但并非完全没有可能。对于大型科技企业而言，“选择交易相对人的权利，不是绝对的，更非不受监管的……反垄断法禁止将其作为实现垄断目的的手段。”^{〔34〕} 就算力获取限制而言，在生成式人工智能市场上，很多中小规模的模型开发商都依赖大型科技企业提供必需的算力资源。然而，一些大型科技企业本身也是模型的开发者。出于排除竞争对手的目的，它们可能会拒绝为某些模型开发者提供服务。

此外，大型科技企业还可能向竞争对手收取不公平的数据许可费用或者算力服务费用，大幅提高竞争对手的进入成本，排除、限制生成式人工智能市场上的竞争，此类行为可能构成不公平高价行为或者差别待遇行为。

（二）开发层：限制部分用户基于商业目的的模型访问

基础模型（foundation model）是生成式人工智能的“技术底座”，^{〔35〕} 可以执行广泛的下游任务，因而具备较为明显的通用属性。在生成式人工智能市场上，绝大多数基础模型由大型科技企业主导研发，下游企业以基础模型为依托开发智能教育、智能客服、智能医疗与智能金融等应用，对大型科技企业具有较高的模型依赖性。

而且，大型科技企业往往具有双重角色，它们既是生成式人工智能基础模型的开发者，又是生成式人工智能应用服务的提供者，这一角色定位赋予了大型科技企业在交易关系中的强势地位。为了扩大在生成式人工智能应用市场上的竞争优势，大型科技企业可能会限制部分用户，尤其是竞争对手或者潜在竞争对手基于商业目的的模型访问，^{〔36〕} 如 Meta 在 Llama 3.1 的许可协议中就声明：“自 Llama 3.1 版本发布之日起，被许可方或者其关联公司提供的产品（服务）在前一个日历月的月活跃用户数超过 7 亿，则必须向 Meta 申请许可。”^{〔37〕} 大型科技企业通过独家许可、非独家授权许可与限制访问速率等方式排除、限制部分用户访问基础模型，或者对用户访问模型进行区别对待，可能会阻碍生成式人工智能应用市场上的竞争与创新。^{〔38〕}

〔33〕 Carugati Christophe, *Antitrust Issues Raised by Answer Engines*, Bruegel Working Paper (13 June 2023), <https://www.bruegel.org/system/files/2023-06/WP%2007.pdf>, visited on 8 May 2026.

〔34〕 *Lorain Journal Co. v. United States*, 342 U.S. 143, 72 S.Ct. 181, 96 L.Ed. 162 (1951).

〔35〕 参见李润生：《人工智能价值链的法律型塑》，载《东方法学》2025 年第 2 期，第 73 页。

〔36〕 See OECD, *Artificial Intelligence, Data and Competition*, OECD Artificial Intelligence Papers (24 May 2024), https://www.oecd.org/en/publications/artificial-intelligence-data-and-competition_e7e8884-en.html, visited on 8 May 2026.

〔37〕 《Llama 3.1 社区许可协议》第 2 条，https://www.llama.com/llama3_1/license/，2026 年 5 月 9 日访问。

〔38〕 See Autoridade da Concorrência, *Competition and Generative Artificial Intelligence*, (November 2023), <https://www.concorrenca.pt/sites/default/files/documentos/Issues%20Paper%20-%20Competition%20and%20Generative%20Artificial%20Intelligence.pdf>, visited on 9 May 2026.

（三）部署层：捆绑生成式人工智能应用与核心产品

为提高自研模型或者合作伙伴研发模型的市场占有率，大型科技企业热衷于将模型集成到自身构筑的数字生态系统中，如 Microsoft 很早之前就将 ChatGPT 嵌入到了 Office 办公产品、Azure 云产品和 Bing 搜索当中，而 Google 搜索则搭载了其自研模型 Gemini。通过将模型与核心产品捆绑，大型科技企业不仅强化了其在搜索服务等既有市场上的支配地位，还可能将市场力量传导至生成式人工智能市场，损害生成式人工智能市场上的竞争。

同样是利用在某一细分市场中的支配地位，一些大型科技企业可能还会实施优待自身产品（服务）的行为。举例而言，一些大型科技企业将搜索引擎与生成式人工智能模型整合，升级为 AI 交互式搜索引擎，用户在使用 AI 搜索时，生成式人工智能技术可以自动整理优化互联网内容提供者的创作内容，从而生成新的高质量内容并输出。此时，集搜索服务提供商和互联网内容提供者双重身份于一体的大型科技企业有动机和能力将生成式人工智能整合的内容置于搜索结果顶端，这一“优先展示”行为将使大型科技企业获得相较于其他内容提供者而言的不正当的竞争优势，涉嫌自我优待。

总体观之，在生成式人工智能价值链的不同层级，垄断风险具有较为鲜明的异质性。在基础设施层，生成式人工智能垄断风险主要表现为拒绝交易、不公平高价等风险，具体而言，大型科技企业可能会拒绝向竞争对手提供芯片或者索取不公平高价。在开发层，生成式人工智能垄断风险主要表现为差别待遇等风险，大型科技企业可能会通过授权许可等方式对竞争对手访问基础模型设置歧视性条件。而在应用层，生成式人工智能垄断风险则主要表现为自我优待、搭售或者捆绑等风险。

值得注意的是，在与技术交织、融合的过程中，生成式人工智能垄断风险的复杂性与不确定性急剧上升。一些垄断风险呈现出了跨层级特征，例如，科技企业之间的投资合作行为。近年，为了规避经营者集中反垄断审查，一些大型科技企业开始以“投资合作”代替并购交易。大型科技企业与初创企业之间的合作形式十分多样，包括算力合作、数据合作以及分销合作等，^[39] 其中 Microsoft 与 OpenAI 之间的合作颇具代表性。过去几年间，Microsoft 累计向 OpenAI 投资逾 130 亿美元，并一度成为 OpenAI 的独家云计算服务提供商，为 OpenAI 研发模型提供强大的算力支持。作为对价，Microsoft 获得了 OpenAI 49% 的股权和 75% 的利润分配权，Microsoft 在获得 1500 亿美元的利润之后，其持有的股权将回转至 OpenAI。此外，Microsoft 还获得了 ChatGPT 模型的独家授权。从上述公开披露的合作细节观之，Microsoft 与 OpenAI 是以一种非控股的方式进行合作。相较于直接收购或者控股，这一方式显然可以更好地规避反垄断审查，还可以在必要时保证风险隔离。但是，这类合作却隐藏着多重垄断风险：其一，OpenAI 的竞争对手若想训练同等规模的模型，可能面临来自 Microsoft 的“算力封锁”（如拒绝交易、不公平高价等），或者被迫加入 Microsoft 的生态系统；其二，基于合作协议，OpenAI 的领先模型深度捆绑并优先服务于微软生态，第三方应用（如 WPS、Slack）难以与 Office 等 Microsoft 的核心产品

[39] See CMA, *AI Foundation Models: Technical Update Report*, (16 April 2024), <https://www.gov.uk/government/publications/ai-foundation-models-update-paper>, visited on 8 May 2026.

展开竞争；其三，从外观上看，投资合作属于开放式合作关系，合作各方在经营上保持独立，但一方仍有可能对另一方施加影响。2023年11月，在OpenAI董事会作出解雇其总裁和CEO的决议后，Microsoft向OpenAI董事会施压，要求撤销该决议，最终OpenAI宣布恢复奥特曼的CEO职位并重组了董事会，并且Microsoft获得了一个无投票权的观察员席位。^{〔40〕}这一事件表明，在投资合作关系下，Microsoft对OpenAI的经营决策仍可能施加重大影响，以纵向一体化的形式排除、限制生成式人工智能市场上的竞争。

作为生发于智能社会的新型风险，生成式人工智能垄断风险兼具多元面向与多重属性，监管的难度非常大。鉴于此，反垄断执法机构应当在熟悉生成式人工智能垄断风险特质的基础上，革新监管理念与监管范式，并就各类风险之间的一体化监管进行系统考量。

四、生成式人工智能垄断风险的合作监管

生成式人工智能垄断风险的成因与构造十分复杂。然而，隐匿在复杂成因与构造背后的底层逻辑却异常朴素——大型科技企业是生成式人工智能垄断风险的唯一源头。因此，抓住大型科技企业这个关键主体，充分发挥其监管功能，构建反垄断执法机构主导、“反垄断执法机构—大型科技企业”互动合作的监管模式是实现生成式人工智能垄断风险有效监管的可行路径。

（一）合作监管目标的确立

1. 生成式人工智能垄断风险监管目标的文本分析

目标是行动的灵魂与驱动力，决定着行动的内容、性质、方向以及后果。因此，选定什么样的监管目标是开展生成式人工智能垄断风险合作监管首先应当解决的问题。

世界主要国家（地区）立足本区域生成式人工智能技术发展的实际情况，结合技术演进与风险转化的一般规律，针对生成式人工智能领域的垄断风险，确立了不同的监管目标，由于生成式人工智能领域反垄断监管专门立法尚未完成，因此监管目标散见于与人工智能相关的法律法规、政策文件中（详见表1）。

表 1 世界主要国家（地区）生成式人工智能风险监管目标列示

序号	国家/地区	法律、政策文件	监管目标
1	中国	《互联网平台反垄断合规指引》（2026）	第1条：“维护市场公平竞争秩序”“促进平台经济创新和健康发展”
		《网络安全法》（2025年修正）	第20条：“加强风险监测评估和安全监管，促进人工智能应用和健康发展”
		《生成式人工智能服务管理暂行办法》（2023）	第3条：“促进创新” 第4条：“不得利用算法、数据、平台等优势，实施垄断和不正当竞争行为”
		《科技伦理审查办法（试行）》（2023）	第1条：“促进负责任地创新”

〔40〕 2024年7月，Microsoft宣布放弃观察员席位，这一举动可能与反垄断执法机构加大对Microsoft与OpenAI投资合作的监管有关。参见文巧：《微软、苹果相继放弃OpenAI董事会观察员席位》，载《每日经济新闻》2024年7月14日，第8版。

续前表

序号	国家/地区	法律、政策文件	监管目标
2	欧盟	《人工智能法案》（2024）	前言 29：“保护消费者的自主决定权、自由选择权” 第 1.1 条：“改善内部市场的运行” 第 57 条：“促进创新和竞争，推动人工智能生态系统的发展”
		《数字市场法》（2022）	第 1 条：“确保欧盟范围内所有存在看门人企业的数字市场具备可竞争性及公平性，从而维护内部市场的良好运行”
3	美国	《确立人工智能国家政策框架》行政令（2025）	第 1 条：“促进创新”“降低企业合规成本”
		《赢得 AI 竞赛：美国 AI 行动计划》行政令（2025）	第一部分：“为企业主导的创新活动营造良好环境”
		《生成式人工智能问责法案》（加州法案，2024）	第 11549.65 条：“防范算法歧视，保护数据隐私”
4	英国	《人工智能（监管）法案》（2025）	第 1 条：“保护消费者权益”“保护隐私”“促进创新”
5	德国	《人工智能和数据保护指南》（2024）	第 2.5 条：“保护个人数据”“保护隐私”
6	日本	《人工智能相关技术研究开发及应用促进法》（2025）	第 3 条：“促进人工智能相关技术的研究、开发和利用” 第 12 条：“促进基础设施的开发和共享”
7	韩国	《人工智能发展与信任基础建立基本法》（2025）	第 3 条：“鼓励创新”“营造安全的人工智能应用环境” 第 15 条：“促进训练数据的生产、收集、管理、流通与应用”

在生成式人工智能垄断风险监管目标谱系中，“促进竞争”是世界主要国家（地区）共同追求的监管目标，也是他们在监管生成式人工智能领域的垄断风险时的核心目标。保证竞争机制的有效性、提升生成式人工智能市场的竞争效能本质上属于发展导向下的目标追求，^{〔41〕} 谋求的是生成式人工智能的技术创新、生态繁荣以及生成式人工智能产业的健康发展，生成式人工智能垄断风险的合作监管应当紧密围绕“促进竞争”目标展开。

2. “促进竞争”监管目标的设定

为推动生成式人工智能产业迈向更高水平，中国在监管生成式人工智能领域的垄断风险时应当始终秉持“促进竞争”的监管目标。作为新一轮科技革命和产业变革的重要推动力量，生成式人工智能关系一国的核心竞争力与发展利益，国家应当大力推动生成式人工智能技术创新和产业发展，提升国家创新体系整体效能，而这一切都建立在生成式人工智能市场有效竞争的基础之上。因此，国家必须直面生成式人工智能带来的挑战，着力预防“数据垄断”“算力垄断”等生成式人工智能发展过程中日益显现的垄断风险，保障用户隐私安全、保护企业竞争利益、维护生成式人工智能市场竞争秩序。

“促进竞争”监管目标的设定要求理论界、实务界在生成式人工智能这一新的场域中重新思考“竞争”“风险”与“反垄断监管”的关系。竞争从来都不是静态、空泛的概念，而是“一种

〔41〕 参见丁晓东：《全球比较下的我国人工智能立法》，载《比较法研究》2024年第4期，第63-64页。

• 169 •

使经济自由产生它所许诺的经济进步的机制”^{〔42〕}，竞争意味着经营者之间激烈角逐、持续冲突、彼此争胜，会引发各种类型的法律风险。因此，市场中竞争与风险具有紧密的伴生关系，即便是在竞争较为充分的市场中，垄断风险也是有可能存在的，垄断风险之存在不是反垄断监管启动的充分条件。更进一步说，反垄断执法机构在监管生成式人工智能领域的垄断风险时应当秉持包容审慎的监管理念，恪守“适度预防”而非“严格杜绝”的态度，^{〔43〕}允许生成式人工智能在风险可控的范围内自由发展，并努力为生成式人工智能技术创新营造相对宽松的外部监管环境。

（二）因应“促进竞争”目标的合作监管方案设计

为保证促进竞争监管目标的实现，生成式人工智能垄断风险的合作监管应当充分发挥大型科技企业、反垄断执法机构各自的监管禀赋，并注重不同监管系统之间的沟通与协同。^{〔44〕}

1. 大型科技企业的自我监管

作为被监管者的“内部式自律”行为，^{〔45〕}自我监管（self-regulation）的本质是基于管理的监管（management-based regulation），即由被监管者自己制定计划和内部规则，完善内部管理机制，以实现特定的公共目标。^{〔46〕}“当监管问题过于复杂，或某个行业存在异质性，或处于动态演进之中时，更适合选用自我监管。”^{〔47〕}就生成式人工智能领域的垄断风险而言，由大型科技企业开展自我监管具有两个明显优势：其一，大型科技企业是生成式人工智能技术的开发者和应用者，它们掌握更多的内部信息和专业知识，可以充分调动内部监管资源，因此更有可能设计出符合成本收益原则的监管方案。^{〔48〕}其二，大型科技企业比反垄断执法机构更接近垄断风险本身，容易获得垄断风险监管效果的及时性反馈，监管的效率更高。^{〔49〕}实践中，大型科技企业可在企业内部设立反垄断合规委员会，并根据生成式人工智能价值链分层开展自我监管。

第一，基础设施层中的自我监管。基础设施层中的主要垄断风险是大型科技企业限制竞争对手或者潜在竞争对手获取生成式人工智能的关键投入，如拒绝向竞争对手或者潜在竞争对手提供必需数据或者算力服务，或者对其收取不公平高价。鉴于此，反垄断合规委员会应当强化对拟投入市场的生成式人工智能产品（服务）的审查，排除产品（服务）设计中潜藏的限制互操作性、拒绝开放数据接口等垄断风险。产品（服务）投入市场后，反垄断合规委员会应当实时监测产品（服务）的运行，及时对垄断风险进行提示、预警。此外，反垄断合规委员会还应当建立垄断风险事后处置机制，一旦发生限制互操作性、拒绝开放数据接口等垄断风险，大型科技企业应当立

〔42〕〔美〕戴维·J. 格伯尔：《二十世纪欧洲的法律与竞争》，冯克利、魏志梅译，中国社会科学出版社2004年版，第44页。

〔43〕参见张守文：《数字经济发展的经济法理论因应》，载《政法论坛》2023年第2期，第43页。

〔44〕参见朱宝丽：《合作监管法律问题研究》，法律出版社2018年版，第15页。

〔45〕参见宋华琳：《人工智能立法中的规制结构设计》，载《华东政法大学学报》2024年第5期，第13页。

〔46〕参见〔英〕安东尼·奥格斯：《规制：法律形式与经济学理论》，骆梅英译，中国人民大学出版社2008年版，第112页。

〔47〕〔英〕罗伯特·鲍德温、马丁·凯夫、马丁·洛奇主编：《牛津规制手册》，宋华琳、李鸽、安永康、卢超译，上海三联书店2017年版，第169页。

〔48〕参见〔美〕凯斯·R. 桑斯坦：《权利革命之后：重塑规制国》，钟瑞华译，中国人民大学出版社2008年版，第93页。

〔49〕参见华文轩：《生成式人工智能的风险规制困境及其化解：以ChatGPT的规制为视角》，载《比较法研究》2023年第3期，第164页。

即暂停相关业务，并同步启动反垄断合规调查，经调查核实确实存在垄断风险的，应当对涉嫌违法行为进行整改。

第二，开发层中的自我监管。开发层中的垄断风险主要是大型科技企业限制竞争对手、潜在竞争对手基于商业目的访问基础模型。为避免歧视性许可等限制行为的发生，反垄断合规委员会可尝试设定并完善基础模型访问许可标准，依托标准实现对基础模型访问限制行为的“内部管理型监管”。^{〔50〕} 具体而言，其一，设定基础模型访问许可标准，明确基础模型访问的许可条件，如定价、用途限制等，避免在许可协议中设置地域限制、用户规模门槛等排他性条款，并且实现基础模型 API 接口标准化，制定统一的模型调用规范，降低基础模型访问限制行为发生的概率。其二，建立基础模型访问许可标准的内部自检机制。反垄断合规委员会应当定期评估基础模型访问许可标准的垄断风险防控效果，并对标准进行动态调整，如果发现标准未能得到有效执行，应当立即整改，相关标准不符合反垄断法律规范要求的，应当及时修订或者废止。其三，构建针对基础模型访问许可标准的争议解决机制。下游用户因基础模型访问许可标准与大型科技企业发生争议的，反垄断合规委员会应当及时受理申诉、举报，并提供相应的解决方案，持续优化基础模型访问许可标准。

第三，部署层中的自我监管。部署层中的垄断风险主要是大型科技企业实施捆绑、自我优待等垄断行为。以自我优待为例，自我优待本质上是大型科技企业基于算法实施的垄断行为，因此反垄断合规委员会应当以算法为中心展开自我监管。首先，反垄断合规委员会可以构建算法审计框架，利用代码审计方法，深入算法决策环路，评估和验证算法在模型部署和运行过程中是否存在自我优待等垄断风险，例如，Google 就专门设计了 SMOAT 算法内部审计框架，可从范围界定、映射、文件搜集、测试以及反思五个阶段进行系统的风险评估。^{〔51〕} 其次，根据《互联网信息服务算法推荐管理规定》第 12 条的规定，反垄断合规委员会可以建立增强算法透明度和可解释性的机制，在合理限度内依法向社会公开算法决策的基本规则、因素权重与运行逻辑，并进行必要的解释说明，提高算法的透明度和可信度。^{〔52〕}

综上，在自我监管场景下，反垄断执法机构将监管生成式人工智能垄断风险的权力部分地转移给了大型科技企业，大型科技企业不再是仅承受监管压力的被监管者，角色转变与利益激励给予了它较强的监管意愿与监管动力，^{〔53〕} 而这有助于实现“促进竞争”的生成式人工智能垄断风险监管目标。

2. 反垄断执法机构的外部监管

第一，禁止垄断生成式人工智能关键投入，促进要素自由流动。数据、算力等关键投入是生成式人工智能发展的基石，大型企业对关键投入的集中控制是生成式人工智能垄断风险产生的根本原因，因此，保证关键投入可及是遏制生成式人工智能垄断风险最有力的举措，这具体可

〔50〕 参见童云峰：《生成式人工智能技术风险的内部管理型规制》，载《科学学研究》2024 年第 10 期，第 2039 页。

〔51〕 See Inioluwa Deborah Raji et al., *Closing the AI Accountability Gap: Defining an End-to-End Framework for Internal Algorithmic Auditing*, Proceedings of the 2020 Conference on Fairness, Accountability, and Transparency (3 January 2020), <https://arxiv.org/pdf/2001.00973>, visited on 10 May 2026.

〔52〕 参见程雪军：《超级人工智能平台算法歧视的系统治理机制》，载《法治社会》2025 年第 3 期，第 60 页。

〔53〕 参见王磊：《论平台垄断的合规治理》，载《行政法学研究》2023 年第 4 期，第 130 页。

以从三个维度渐次展开：

其一，关键投入获取的个案实现。以训练数据的获取为例，在个案中，针对大型科技企业实施的拒绝开放数据接口、限制数据互操作性等行为，反垄断执法机构可以根据“数据最小化”原则确定训练数据开放的范围，^{〔54〕}要求大型科技企业向竞争对手开放生成式人工智能产品和服务开发所必需的数据，以消除竞争对手获取训练数据的障碍。数据开放、数据互操作涉及反垄断执法机构、大型科技企业、训练数据需求方以及用户等多方主体与复杂利益平衡，在满足训练数据需求方诉求的同时，为避免“搭便车”损伤大型科技企业的创新激励，大型科技企业与训练数据需求方可在反垄断执法机构的指导监督下，根据公平、合理、无歧视原则签订相关协议，如双方可签订训练数据互操作协议，并在协议中对数据互操作期限、数据格式以及数据互操作许可费用等事项进行约定。

其二，大力推动公共数据资源开发利用，拓宽关键投入获取渠道。除强制大型科技企业开放数据这一训练数据供给渠道之外，因公共数据具有规模大、质量高、潜能大等特性，公共数据资源开发利用正逐渐成为我国数据要素供给体系的重要组成部分。^{〔55〕}有学者提出，如果能够解决公共数据资源的开放使用问题，将可以在较大程度上满足生成式人工智能的训练数据需求。^{〔56〕}立足自身权责与定位，反垄断执法机构可以从两个向度促进公共数据资源开发利用，降低训练数据获取壁垒。一方面，严格开展公平竞争审查，政策制定机关应当加强对涉公共数据政策措施的公平竞争审查，避免政策措施中含有公共数据垄断等排除、限制竞争的内容，政策制定机关未进行公平竞争审查或者违反审查标准出台政策措施违反《反垄断法》的，反垄断执法机构可以向有关上级机关提出依法处理的建议。另一方面，鉴于公共数据授权运营是当下我国公共数据开发利用的重要方式，反垄断执法机构应当加强对公共数据授权环节、运营环节垄断风险的监管，以授权环节为例，反垄断执法机构应当密切关注独家授权引发的结构性垄断风险，遏制行政性垄断等潜在垄断行为，推动公共数据运营主体的竞争性遴选。^{〔57〕}

其三，推动生成式人工智能基础设施建设，实现关键投入共用共享。未来一阶段，反垄断执法机构应当立足自身维护市场竞争的功能定位，积极推动国家生成式人工智能基础设施建设。详言之，反垄断执法机构应当加强与国家数据局等部门的协作，构建反垄断监管与行业监管高效协同的跨部门综合监管机制，合力打破数据要素的区域壁垒和市场分割，并在此基础上推动建立全国统一的数据共享平台，实现全国数据共享一体化，为生成式人工智能模型的训练与应用创造良好的数据环境。此外，反垄断执法机构还应当加强对算力垄断的监管，打破算力市场壁垒，协助工业和信息化部等主管部门构建全国一体化算力网络，整合、优化全国范围内的算力资源，突破传统的单一算力供给服务模式，降低算力获取门槛和使用成本，以更好地满足生成式人工智能模型训练需求。生成式人工智能基础设施建设有助于实现不同行业、不同地域、不同系统之间的数

〔54〕 参见王磊：《反垄断视角下促进大型平台开放数据的三重路径》，载《地方立法研究》2023年第1期，第43-44页。

〔55〕 参见程啸主编：《数据权益与数据交易》，中国人民大学出版社2024年版，第388页。

〔56〕 参见赵精武：《论人工智能法的多维规制体系》，载《法学论坛》2024年第3期，第63页。

〔57〕 参见时建中：《数据概念的解构与数据法律制度的构建——兼论数据法学的学科内涵与体系》，载《中外法学》2023年第1期，第44页。

据、算力等关键投入的互联互通、高效配置，提高生成式人工智能关键投入的开发利用效率，充分激活、释放关键要素价值，有效纾解生成式人工智能垄断风险监管困境。

第二，革新生成式人工智能垄断风险监管工具，引入“监管沙盒”。生成式人工智能垄断风险由 AI 技术驱动，具有高度的复杂性与不确定性，为因应技术化、动态化的监管需求，反垄断执法机构可考虑引入“监管沙盒”，^[58] 通过构建一个具体、可控的监管框架，为生成式人工智能服务的提供者或者潜在提供者按照沙盒计划自由开发、验证、测试生成式人工智能产品（服务）提供试错空间，^[59] 以此来缓和垄断风险监管与生成式人工智能技术创新之间的紧张关系。“监管沙盒”以实用主义哲学为理论基础，^[60] 通过试验式地供给监管政策与措施，灵活调整监管方式，最终实现监管对象与监管路径、监管工具的适配，^[61] 在充满“持续变化与挑战”的智能社会中，这一摒弃标准化解决方案的监管思路与生成式人工智能垄断风险的监管需求高度契合。^[62]

就“监管沙盒”的适用而言，反垄断执法机构应当构建“监管沙盒”的“申请—测试—退出”机制，明确“监管沙盒”的适用范围、测试要求以及退出程序。在申请阶段，大型科技企业可以根据已公布的“监管沙盒”的目标和基本要求向反垄断执法机构提出书面申请，并提交包括测试计划（沙盒计划）、风险评估、企业信息在内的相关资料，由反垄断执法机构决定是否准予其进入沙盒进行测试。在测试阶段，大型科技企业应当根据事先商定的测试方案对拟投入市场的生成式人工智能产品（服务）进行测试，并定期向反垄断执法机构报告测试中出现的问题，做好垄断风险管控。为应对突发情况、提升测试的有效性，大型科技企业可以向反垄断执法机构申请变更沙盒，如延长测试周期、扩大测试范围等。测试结束后，大型科技企业应当向反垄断执法机构提交书面总结报告。在退出阶段，反垄断执法机构应当对大型科技企业提交的总结报告进行评估，如果评估结论为“通过测试”，反垄断执法机构可以为大型科技企业出具书面证明，该证明可作为产品（服务）反垄断合规的依据。如果大型科技企业没有通过测试，可在实质改进产品（服务）后再次申请测试。

第三，加强生成式人工智能市场竞争状况调查，构建反垄断梯次式执法体系。生成式人工智能垄断风险监管本质上是面向不确定性的监管活动。为避免过度监管，一方面，反垄断执法机构应当加强生成式人工智能市场竞争状况调查，深入了解生成式人工智能的发展趋势与市场整体竞争状况，提高对生成式人工智能垄断风险的感知和预测能力。另一方面，反垄断执法机构可以根据个案情况灵活运用柔性执法手段，通过柔性的、轻触式的执法引导大型科技企业有序竞争，^[63] 例如，反垄断执法机构可以向涉嫌实施垄断行为的大型科技企业发出提醒敦促函，要求其做好预防与整改工作，并书面反馈落实情况。除提醒敦促外，约谈制度也为生成式人工智能垄断风险监管提供了一种相对包容的监管方案。作为以“协商性”为基础的执法模式，约谈为反垄断执法机

[58] 参见张晨颖：《反垄断智慧监管的理据与图景》，载《探索与争鸣》2024年第5期，第109-110页。

[59] 参见解志勇：《高风险人工智能的法律界定及规制》，载《中外法学》2025年第2期，第299页。

[60] 参见〔美〕约翰·杜威等：《实用主义》，田永胜等译，世界知识出版社2007年版，第294页。

[61] 参见靳文辉：《试验型规制制度的理论解释与规范适用》，载《现代法学》2021年第3期，第124页。

[62] See Charles F. Sabel & William H. Simon, *Minimalism and Experimentalism in the Administrative State*, 100 The Georgetown Law Journal 53 (2011).

[63] 参见〔美〕理查德·赛勒、卡斯·桑斯坦：《助推》，姜智勇译，中信出版社2023年版，第XXXV页。

构与大型科技企业营造了一个平等对话、自由商谈的空间,这使得双方能够及时交换信息,充分表达自身利益关切,并在此基础上达成共识,保证反垄断执法的审慎、谦抑。实践中,如果大型科技企业涉嫌从事垄断行为,引发舆情或者造成不良影响,或者对前述提醒敦促事项逾期未整改或者整改不到位,反垄断执法机构可以约谈大型科技企业的法定代表人,并要求其提出改进措施。

在提醒敦促、约谈等柔性执法手段之外,责令停止违法行为、处罚等刚性执法手段亦不可或缺。“尽管制裁与执法行动不能够带来最适威慑,但是制裁与执法是反垄断法最适威慑实现的必要条件”^[64],如果大型科技企业在被约谈后逾期未整改、整改不到位或者整改后再次出现问题,或者有证据初步证明大型科技企业涉嫌实施数据垄断、算力垄断等行为,反垄断执法机构可以依法立案调查,经调查确有违反《反垄断法》的行为的,反垄断执法机构应当依法责令大型科技企业停止违法行为,并对其进行处罚。

综上,“提醒敦促—约谈—制裁”的梯次式执法机制克服了单一的“命令—控制”式监管模式的弊端,有助于实现对生成式人工智能垄断风险的“合比例性”监管。^[65]

3. 大型科技企业与反垄断执法机构的合作

根据监管空间理论,监管主体要实施有效监管,必须拥有权力、信息与财产等各种资源。这些资源分散在各个监管主体之间,因此,只有多主体合作监管才能实现最佳监管效果。^[66]就生成式人工智能垄断风险的监管而言,大型科技企业与反垄断执法机构的监管禀赋不同:大型科技企业具有知识与信息优势,但其监管权威不足;作为公共利益的代表,反垄断执法机构具有理性、权威等“官僚制的美德”,^[67]但其缺乏对生成式人工智能运行机理的深刻理解。大型科技企业的自我监管与反垄断执法机构的外部监管具有不同的监管侧重,发挥着不同的监管功能,二者必须合作才能够形成监管合力,^[68]并有效应对生成式人工智能领域的垄断风险。大型科技企业与反垄断执法机构的合作具体可从三个维度展开:

第一,合作双方应当建立监管互信,并基于交往理性深入开展信息交流。本质上,合作监管是政府将自我监管系统整合到政府监管框架中的结果,^[69]监管模式的再造、合作监管的运行需要不同主体发挥各自比较优势为监管活动提供权威、技术和信息等方面的支持。因此,反垄断执法机构与大型科技企业必须加强对话、交往以及合作,持续提升监管互信与合作意愿。一方面,双方可以建立常态化沟通机制,采取“常规沟通”“专项磋商”“紧急情况即时联动”多维交往模式共享信息、凝聚监管共识。其中,常规沟通聚焦生成式人工智能市场竞争状况、监管沙盒运行情况等合作监管重点事项;专项磋商针对训练数据使用、算力分配、API 开放策略等关键环节的潜在垄断风险展开研讨;紧急联动则应对企业并购、技术授权等可能引发市场竞争格局变动的突发情况。常态化沟通不仅有助于反垄断执法机构及时掌握监管情境的变化,也有利于大型科技企

[64] 喻玲:《从威慑到合规指引——反垄断法实施的新趋势》,载《中外法学》2013年第6期,第1209页。

[65] 参见苏宇:《风险预防原则的结构化阐释》,载《法学研究》2021年第1期,第51页。

[66] 参见〔英〕科林·斯科特:《规制、治理与法律:前沿问题研究》,安永康译,清华大学出版社2018年版,第31页。

[67] 〔美〕史蒂芬·布雷耶:《打破恶性循环——政府如何有效规制风险》,宋华琳译,法律出版社2009年版,第81-83页。

[68] 参见〔美〕拉里·A.迪马特奥、〔意〕克里斯蒂娜·庞西布、〔法〕米歇尔·坎纳萨主编:《剑桥人工智能手册:法律与伦理的全球视野》,郑志峰等译,当代中国出版社2024年版,第608页。

[69] 参见朱宝丽:《合作监管的兴起与法律挑战》,载《政法论丛》2015年第4期,第138页。

业深入理解监管的目的与导向，进而促使双方为实现促进竞争的合作监管目的而一致行动。另一方面，反垄断执法机构应当建立信息分级分类保密制度，对在常态化沟通过程中获取的诸如训练数据的来源、算力采购规模与渠道以及与生成式人工智能相关的大额投资、战略联盟等涉及企业核心商业秘密的敏感信息，通过加密、脱敏、专家保密审查等方式严格保护，仅向监管指定人员披露，避免信任崩塌导致合作监管失败。

第二，反垄断执法机构与大型科技企业应当加强在生成式人工智能反垄断监管规则制定方面的合作，统合双方利益诉求。面对飞速发展的生成式人工智能技术，由单一主体绝对主导生成式人工智能反垄断监管规则的制定容易导致监管脱节或者创新束缚，因此，必须推动监管规则从一元主体单向度供给向合作监管多元主体共同创设转变，^{〔70〕}在监管规则形成过程中增加合意的成分。具体而言，反垄断执法机构可邀请大型科技企业参与制定生成式人工智能反垄断合规指引，重点明确“拒绝交易”“自我优待”等行为的认定标准，厘清合规与违法的边界。同时，合理设置安全港规则，区分垄断行为与正当的商业行为，例如，将大型科技企业为保护商业秘密而采取的合理数据管控，与排除、限制市场竞争的数据封锁等行为区分开来，实现垄断风险防范与创新激励的统一。

第三，反垄断执法机构应当加强对大型科技企业的反垄断合规指导，双方携手迈向“指导型合作监管”^{〔71〕}。尽管合作监管是公私主体为了回应监管需求、实现共同监管目标而进行的合作，但究其本质，合作监管仍是一类监管活动，仍需满足监管活动本身对秩序的追求。^{〔72〕}为保证合作监管的秩序与效果，反垄断执法机构应当加强对大型科技企业反垄断合规工作的指导与监督。针对模型训练中的数据许可、算力资源分配中的非歧视原则、API 开放中的公平性条款等细分场景，反垄断执法可以指导大型科技企业制定具体的合规操作规范，并持续跟踪大型科技企业反垄断合规落实情况，对指导后仍存在的合规漏洞，开展专项督导整改。同时，将反垄断合规落实情况与经营者集中反垄断审查、合规激励等挂钩，对主动落实指导要求、合规成效显著的大型科技企业给予政策倾斜，对拒不配合、合规流于形式的企业强化约束，引导、激励大型科技企业积极落实反垄断合规责任。

综上，生成式人工智能垄断风险的合作监管是一个开放、协商、互动的过程。生成式人工智能垄断风险有效监管的基础不是单一主体的命令性控制，而是多元主体的合作、协同与良性互动。大型科技企业与反垄断执法机构之间全面、深入的合作，将进一步释放自我监管保障个体自由和自发秩序以及外部监管保障社会公共利益的监管效能，形成带有强回应性特征的“自我监管—外部监管”合作监管结构，并最终实现生成式人工智能技术的安全与创新发展。

五、结 语

随着生成式人工智能技术的快速发展及其对社会生活的深度嵌入，它所引发的垄断风险已然

〔70〕 See Dennis D. Hirsch, *The Law and Policy of Online Privacy: Regulation, Self-Regulation, or Co-Regulation?*, 34 Seattle University Law Review 439 (2011).

〔71〕 刘鹏、吴强：《平台经济中的政企合作监管类型及其生成逻辑——基于 A 市网络交易监管的多案例研究》，载《公共管理学报》2025 年第 2 期，第 161 页。

〔72〕 参见刘权：《分享经济的合作监管》，载《财经法学》2016 年第 5 期，第 50 页。

超出技术控制的范围，并逐渐波及政治、经济、社会等多维层面。生成式人工智能垄断风险监管应当解决的核心议题是风险的有效预防，即如何通过预防性行为与因应性制度来消解生成式人工智能可能产生的负面竞争影响，纠正大型科技企业的“技术专制”及其引发的市场竞争失衡。本质上，生成式人工智能垄断风险监管属于“决策于未知之中”的复杂活动，因此必须要保证监管的恰当时机与合理尺度，监管不应滞后于技术发展，也不宜走在技术创新的前面，而应当与技术创新同频共振。为此，应当摒除过度依赖反垄断执法机构的监管模式，充分发挥大型科技企业的自我监管功能，构建“反垄断执法机构—大型科技企业”互动合作的监管模式，提高生成式人工智能垄断风险监管的技术性、动态性以及回应性。大型科技企业可以在基础设施层、开发层、部署层分层开展自我监管。反垄断执法机构则可以从维护生成式人工智能市场竞争的向度开展外部监管：第一，禁止大型科技企业垄断生成式人工智能关键投入，促进要素自由流动；第二，革新生成式人工智能垄断风险监管工具，引入“监管沙盒”；第三，加强生成式人工智能市场竞争状况调查，构建反垄断梯次式执法体系。大型科技企业与反垄断执法机构的紧密合作、自我监管与外部监管的相容互促有助于实现“促进竞争”的生成式人工智能垄断风险监管目标，推动中国新一代生成式人工智能健康发展。

Abstract: The risk of monopolization in the generative artificial intelligence field is becoming increasingly evident. How to regulate this new type of monopoly risk has become a significant theoretical and practical issue of global concern in the intelligent era. The core logic behind the monopoly risks of generative AI lies in the fact that large tech companies control key inputs for the development of generative AI, such as data and computing power, creating barriers to market entry. In practice, the primary forms of generative AI monopoly risks include restricted access to key inputs and limited availability of foundational models. The nature and structure of these risks are extremely complex. Relying solely on antitrust enforcement agency cannot guarantee the proportionality, dynamism, and responsiveness of regulation. A cooperative regulatory model involving antitrust enforcement agency and large tech companies should be established. Large tech companies can engage in self-regulation at different layers, such as the infrastructure, development, and deployment levels. Antitrust enforcement agency, on the other hand, can implement external regulation by maintaining competition in the generative AI market and promoting the free flow of key elements. Self-regulation and external regulation should collaborate, coordinate, and interact closely.

Key Words: generative artificial intelligence, monopoly risks, self-regulation, external regulation, cooperative regulation

(责任编辑：范佳慧)