

我国 GAI 数据治理的多元协同模式研究

——新加坡治理经验的启示

胡裕岭,姚浩亮

(华东政法大学刑事法学院,上海 201620)

摘要:数据治理是现代国家发展生成式人工智能技术所关注的焦点议题。党的十八大以来,党中央、国务院高度重视数据治理工作,相继出台了《中华人民共和国数据安全法》等一系列法律和规定。然而,面对生成式人工智能带来的数据风险,我国当前的数据治理模式偏向以国家监管和法规施行为主的硬法模式,尚未形成国家、社会和个体协同参与的多元格局,且备案、审查、调查等静态治理手段难以应对技术发展变化带来的不确定性挑战。新加坡的数据治理模式充分吸收了欧盟、美国等人工智能数据治理优势国家及地区的治理经验,并在全球范围内率先推出技术治理方案,形成了集多元主体、多元规范、多元举措于一体的数据治理体系。在此模式下,AI Veirfy、Project Moonshot 等技术治理工具提升了治理规范的确定性;《生成式人工智能模型治理框架》等软法与《个人数据保护法案》等硬法保障着治理手段的可执行性;AI Verify 基金会架起了公私协作的桥梁,深化了治理空间的层次性。对此,我国应辩证地借鉴新加坡的治理经验,构建以权力、权利、义务为指引的多元主体参与路径;完善技术与规范相结合的多元治理举措,丰富数据治理的工具箱;充分发挥硬法的有效执行优势与软法的灵活治理优势,实现数据治理法律规范体系的内在融贯,以此构建符合我国基本国情、满足我国治理需求的多元协同数据治理模式。

关键词:生成式人工智能;数据治理;多元协同;层级治理;新加坡

中图分类号:DF37;TP18 **文献标志码:**A **文章编号:**1671-4970(2024)05-0086-13

生成式人工智能(Generative Artificial Intelligence, GAI)的出现不仅促进了人类社会与智能技术的深度融合,而且寓意着数据资源的充分开发与利用。数据不仅是拉动 GAI 前行的“马车”,更是决定 GAI 能力与“意识形态”的重要养料。“人工智能是引领新一轮科技革命和产业革命的战略性的技术,也是新质生产力的重要引擎。”^[1]然而, GAI 技术也加剧了数据泄露、网络攻击、虚假信息宣传等安全风险。囿于治理成本、治理及时性、专业性等因素,传统

的数据治理模式已然难以应对前述风险。因此,如何及时调整并优化 GAI 数据治理模式正成为世界范围内现代国家所共同面临的重要课题。

在探索 GAI 数据治理新模式的全球浪潮中,新加坡构建的颇具特色的多元主体参与、多元手段治理、多元规范并举的多元协同模式已经取得一定成效。根据《2021 年网络安全指数分析评估报告》,新加坡的个人数据保护能力及网络危机应对能力均得到满分,这意味着其

引用本文:胡裕岭,姚浩亮.我国 GAI 数据治理的多元协同模式研究——新加坡治理经验的启示[J].河海大学学报(哲学社会科学版),2024,26(5):86-98.

基金项目:国家社会科学基金重大项目(20&ZD199);上海市法学会国家安全法律研究会 2024 年度课题项目(2024SNSL04)

作者简介:胡裕岭(1987—),男,讲师,博士,主要从事数据治理研究。E-mail:2557@ecupl.edu.cn

在数据泄露等危急情况下仍处于完全准备状态^[2]。由于定位与策略匹配度较高,新加坡在《2024 年全球 AI 指数排名》中取得了“强度”第一、科研指数排名第三、总排名第三的成绩^[3]。当前,学界已经就数据治理模式展开了较多研究,提出了敏捷治理^[4]、合规治理^[5]、风险治理^[6]等方案,但大都以欧盟、美国为镜鉴对象,强调积极发挥公共监管机构的治理职能,而以新加坡为研究对象,或针对多元主体具体参与路径的研究较少。《中华人民共和国数据安全法》(以下简称《数据安全法》)第 9 条在我国确定了与新加坡近似的多元主体参与数据治理的基本理念,但在实践中尚未落实。因此,分析我国传统数据治理模式的局限,镜鉴新加坡多元协同模式的治理经验,进而构建贴合我国国情的数据治理新范式,落实多元协同治理理念具有极为重要的现实意义。

一、我国传统数据治理模式的局限

1. 多元主体尚未形成有效治理合力

《数据安全法》第 9 条将我国数据治理的多元协同理念具体表述为“推动有关部门、行业组织、科研机构、企业、个人等共同参与数据安全保护工作”。“共同参与”应当包含公权主体之间、社会主体之间、公权主体与社会主体之间共 3 个层次。然而,实践中不同主体或缺乏沟通协作机制或出于各自利益考量等原因尚未形成有效治理合力。

首先,公权主体之间可能会基于部门利益竞相治理或推诿。当前,国家网信部门、公安机关、国家安全机关等公共监管机构对数据治理均具有管辖职能,且能够以自身专业性为依托,实现多元领域的专门治理。然而,GAI 带来的“AIGC+”效应深化了数据的集中与叠加,可能会使不同公共监管机构陷入“管辖竞合”的泥潭,即不同监管机构可能在具有潜在利益的领域竞相治理,而对利益不足的事项避而不谈^[7],这不仅无益于实现治理效能,而且可能造成治理资源分配不均。

其次,社会主体之间出于经济利益考量往往缺乏协作意愿。智能行业中的社会主体作为最接近数据风险的主体,理应最适宜承担风险

防治第一责任人的职责。但由于缺乏有效的指引,这些社会主体大都受经济利益的驱使而将更多资源用于技术研发,缺少自我规范的动力和手段。当前,人工智能领域已呈现“寡头垄断”局面,即数据资源主要集中于有限主体手中,少数主体掌控着更多的话语权,智能行业主体之间难以在真正意义上达成凝聚最大共识的有效协作。此外,智能行业主体与其他社会主体之间存在天然的技术壁垒与透明度壁垒,缺乏实现治理合力的信任基础。

最后,公权主体与社会主体之间处于相对对立位置。传统数据治理模式主要依靠公共监管机构以“家长制”的方式进行干预、调控,缺乏多元主体的协商与互动,具有单向性、强制性^[8]。如此情势之下,智能行业主体向公共监管机构提供虚假数据或是隐瞒部分数据的现象屡禁不止,更枉谈“协助治理”。治理本身是一个过程性的结果,不仅需要通过制定与执行规范来实现,而且需要多元主体的互动、合作、协商才能发展^[9]。所以,亟须构建多元主体协同参与的数据治理模式,以更好地实现 GAI 数据治理效能。

2. 静态治理举措难以应对技术挑战

GAI 的发展放大了数据风险的突发性、广泛性,并使数据风险呈现瞬息万变的动态性。传统模式下,备案、审查、调查等治理举措往往基于既定的规则和流程,这种相对静态的治理模式在较为稳定的技术环境中有效,但面对 GAI 数据风险特征时,其已然无法满足当前的治理需求。

人工智能通常以确定性算法或非确定性算法作为运算的底层逻辑。其中,确定性算法根据预设规则输出预设结论,输出结果相对固定^[10],而 GAI 往往需要处理开放的、变化的数据,因此更倾向于采用动态调整的非确定性算法。非确定性算法能够通过组合非监督学习、生成对抗网络(GANs)、深度强化学习等技术,不断调整输入内容的权重分布,以生成不同的输出内容,从而应对多样化和不可预测的情况。这一特性使得 GAI 呈现高度复杂且动态的学习过程。如,GAI 大多基于深度学习模型,尤其是神经网络技术运行,这类模型的非确定性主

要体现在随机初始化和训练过程中,其能够通过自主学习和调整决策规则,依赖试错机制不断优化,而并非依靠固定规则进行决策,这进一步加剧了 GAI 输出内容的不确定性。正是这种非确定性,使得 GAI 与传统人工智能存在显著区别,也使人类往往只能感知输入与输出数据,而无法完全理解算法内部的运作机制和逻辑,这也为 GAI 的可解释性和治理带来了新的挑战,静态治理难以奏效。

此外,《网络数据安全条例》虽然补足了我国数据治理领域的部分空白,明确“支持相关行业组织按照章程,制定网络数据安全行为规范,加强行业自律”,但是难以满足数据治理的迫切需求,尚未走出静态治理的思维误区:一方面,该法以数据安全为核心强调“监管”,内容整体上延续了《数据安全法》等法律设置的传统治理举措,实则以前办法应对新问题。另一方面,该法第 47 条意在厘清不同机关的监管责任,但“在各自职责范围内”的表述并未明确不同主体间的责任分配问题。所以,我国需要一部能够高度适应技术不确定性、回应本土独特性需求、有效应对国际竞争的人工智能法律^[11],且数据治理思维不仅要关注管理的风险,更要关注如何最好地管理这些风险(包括利用技术工具)^{[12]43}。进言之,法律治理需要与技术治理结合,公共监管机构应拓展动态治理举措,深入技术前端。

3. 硬法主导治理难以回应发展诉求

法治包含硬法与软法两部分:以国家强制力为保障手段的法律形态,通常被称为硬法;软法意指“那些难以或不能运用国家强制力保证实施的具有公共规范性质的规范性文件或者惯例。”^[13]我国软法与硬法缺乏通畅的衔接渠道,两者在诸如生物识别特征信息等基础概念上存在明显差异^[14]。相较于软法,硬法在价值取向上更为保守,可能会为了保障安全而过度遏制发展,进而忽略科技的发展诉求。

在 GAI 数据治理中,硬法的治理局限更加明显。如维克多·迈尔-舍恩伯格所言,“如果没有足够的数 据,人工智能应用的质量将会降低,进而影响交易效率以及消费者福利。”^[15]

GAI 自动爬取数据能够有效提高模型学习效率,但数据爬取行为很有可能触及《数据安全法》第 32 条、《中华人民共和国网络安全法》第 27 条等法律所预设的合法性边界。如果一味强调硬法治理,以严格禁止换取数据安全,那么无疑会牺牲 GAI 的发展环境。硬法与人工智能的紧张关系如此深层,以至于任何协调两者的努力均可能削弱数据保护义务,干扰人工智能带来的既得利益,或是兼而有之^[16]。然而,GAI 数据治理并非一场零和博弈,可以兼顾安全与发展的双重需求。为实现这一目的,需要明确数据治理领域的哪些问题应当由硬法管理、哪些问题可以由软法管理,这是一个随着技术进步与治理经验丰富的动态过程^{[17]390}。对此,需要发挥软法的治理功能,寻求软法硬法并举的协调路径。

二、新加坡多元协同数据治理模式的范式镜鉴

当前,域外 GAI 的数据治理模式包括以欧盟为代表的严格监管模式、以美国为代表的软法治理模式、以新加坡为代表的多元协同模式。相较于欧盟对人工智能进行严格的分级风险治理,新加坡侧重于同时发挥软法的指引作用与硬法的兜底作用,为 GAI 的研发活动预留了相对宽松的发展空间;相较于美国“联邦-州”政府的双边治理形式,新加坡采用单一政府主导、多元主体协同的集中治理形式,通过政府牵头组建公私协作机构,统筹国内协作与国际协作的综合发展,在此模式下,公权主体与社会主体之间的联系更为紧密。此外,新加坡率先在世界范围内探索政府主持的技术治理方案,提升了治理规范的确 定性。

1. 新加坡多元协同数据治理模式的主要特色

作为亚洲地区较早开展人工智能数据治理的国家,新加坡确立了在 2030 年成为人工智能广泛应用的智慧国家、全球人工智能部署与创新解决方案的引领者等目标。近年来,新加坡政府创制了一系列治理规范、指引和工具(表 1)。

表 1 新加坡部分人工智能的数据治理规范、指引、工具汇总^①

时间	发布机构	规范/指引/工具
2019 年		Model Artificial Intelligence Governance Framework First Edition
2019 年		Trusted Data Sharing Framework
2020 年		Model Artificial Intelligence Governance Framework Second Edition
2020 年	新加坡信息通信媒体发展局 (IMDA) 个人数据保护委员会 (PDPC)	Companion to the Model AI Governance Framework-Implementation and Self-Assessment Guide for Organizations
2020 年		Compendium of Use Cases; Practical Illustrations of The Model AI Governance Framework (Volume1) (Volume2)
2020 年		A Guide to Job Redesign in The Age of AI
2022 年		A. I. Verify
2022 年		Invitation to Participate in Privacy Enhancing Technology Sandbox
2024 年	新加坡 AI Verify 基金会 (AIVF) 新加坡信息通信媒体发展局 (IMDA)	Model AI Governance Framework for Generative AI
	新加坡 AI Verify 基金会 (AIVF)	Project Moonshot

从相应规范和举措来看,新加坡的多元协同模式具有如下特色。

(1) 人工智能检测技术提升治理规范的确定性

新加坡政府开发了诸多辅助治理工具,如隐私增强技术沙盒、数据监管沙盒、AI Verify 等,其中的 AI Verify 推广度最高。AI Verify 由测试框架与工具包两部分组成,测试框架参考了 11 项国际公认的人工智能治理原则,形成了 5 项内容(表 2)。

表 2 AI Verify 测试框架支柱及原则概览^[18]

内容	原则
确保人工智能及人工智能系统的应用透明度	透明度原则
了解人工智能模型的决策过程	可解释原则
	可重复/可再现原则
确保人工智能系统的安全性与恢复性(resilience)	(风险)安全原则
	(网络)安全原则
确保(决策)公平	鲁棒性原则
	公正性原则
	数据管理原则
确保对人工智能系统进行适当人工监督与管理	问责制原则
	行为力与监督原则 ^②
	有益于人类和地球原则 ^③

为确保原则的可执行性,测试框架为每一项原则提供了定义及评估方案。如,在透明度原则的评估上,使用者可以向数据所有者提供人工智能的预期用途、风险评估结果、限制使用情况等信息,从而实现数据所有者的知情权。同时,测试框架还提供了定量、定性分析参数指标和指标的临界值,以实现测试结果的对比参照。

AI Verify 工具包是一个“一站式”检测工具。基于测试框架确定的基本原则,该工具包能够对人工智能的透明度、安全性等性能展开评估,并自动生成测试报告。使用者可以通过测试报告了解人工智能的信息,从而实现技术透明、可知。当前,AI Verify 针对不同行业、不同应用、不同产品形态分别制定了评估测试方

① 资料来源于 <https://www.imda.gov.sg/business> 网站。表格所示内容仅是笔者认为相关文件,罗列于此。以“Artificial Intelligence”为关键词在该网站检索,总共可以获得 339 份相关材料。

② 对于“human agency and oversight”的翻译,有译者译为“人类机构和监督”,或是“人类代理或监督”。但是,该原则的解释内容强调人工智能应用不应削弱人类的决策能力,需要突出显示人工智能的涌现性。而且,“human”主要强调人类对人工智能决策的监督,而非将人工智能称为人类的代理,所以,以“行为力”释义人工智能领域的“agency”更为合理,对此,该翻译应为“行为力与监督原则”。

③ 本部分原文是“Inclusive growth, societal and environmental well-being”,译为“包容性增长,社会与环境福祉力”原则,具有理解歧义。根据对该原则的解释内容,笔者译为“有益于人类和地球原则”。

案,并辅以相应的测试工具集与数据集,以期实现针对性治理^[19]。然而,该工具包无法适用于 GAI、大语言模型等特殊对象,所以 AI Verify 基金会(以下简称“AIVF”)新近推出了 Project Moonshot 工具包,旨在通过技术手段实现 GAI 的有形安全^[20]。目前,该技术已经投入测试,成效有待评估。可见,新加坡的 GAI 数据治理技术方案为多元主体实现自我治理、参与治理提供了可能,也增加了数据治理的透明度。

(2) 软法硬法并举保证治理手段的可执行性

2002 年起,新加坡便通过颁布数据保护立法(如《私营机构数据保护守则》等),组建跨机构数据保护专业委员会等举措构建数据治理体系。2012 年出台的《个人数据保护法案》奠定了新加坡数据治理的硬法基础。虽然该法案及后续修订内容未明确涉及 GAI 数据治理,但为多元主体的个人数据处理活动设定了框架。此外,在 GAI 应用的关键领域,新加坡政府积极推进立法。如,新加坡金融管理局(MAS)颁布了《个人问责和行为准则》来规范金融机构的人工智能实践,该准则将公平、道德、问责和透明作为 FEAT 原则的构成要件,用以约束金融领域的人工智能应用^[21]。

在硬法难以触及或不便触及之处,软法在新加坡 GAI 数据治理中发挥着重要的补充功能。表 1 所列大部分规范均为软法,以《人工智能模型治理框架(二)》(以下简称

《治理框架(二)》)和《生成式人工智能模型治理框架》(以下简称《GAI 治理框架》)最为典型。

《治理框架(二)》(图 1)的数据治理条文集中于“运营管理”部分,旨在围绕数据生命周期全流程构建企业内部的“数据问责”制度;在数据来源方面,以数据溯源方向为数据分类标准,并构建起差异化的数据回溯方式;在数据质量方面,鼓励企业基于数据的准确性、完整性、真实性等因素对数据质量展开自查,并解决质量问题;在数据偏差方面,指出数据固有偏差是影响人工智能的主要因素,指引企业通过保证数据完整性、使用异构数据集等方式消除偏差;在数据训练方面,列举了人工智能训练活动中可能存在的风险问题,建议使用测试数据集检验模型等方式予以查明;在数据审查和更新方面,建议企业定期清洗、审查数据,防止数据的重复运用令智能模型形成数据偏见。此外,《治理框架(二)》还通过列举“Suade Lab”与“Pymetrics”等案例指引治理活动,为智能行业主体提供了实践范例。

《GAI 治理框架》是新加坡信息通信媒体发展局(以下简称“IMDA”)与 AIVF 在整合全球 70 余家政府机构、企业、研究机构对框架草案的意见后形成的正式规范。作为一部软法,《GAI 治理框架》围绕建立全面和可信的 GAI 生态系统,在问责制、数据、可信开发与部署、事件报告、测试和保证、安全性、内容来源、安全性与价值对齐、人工智能促进公益九个维度作出

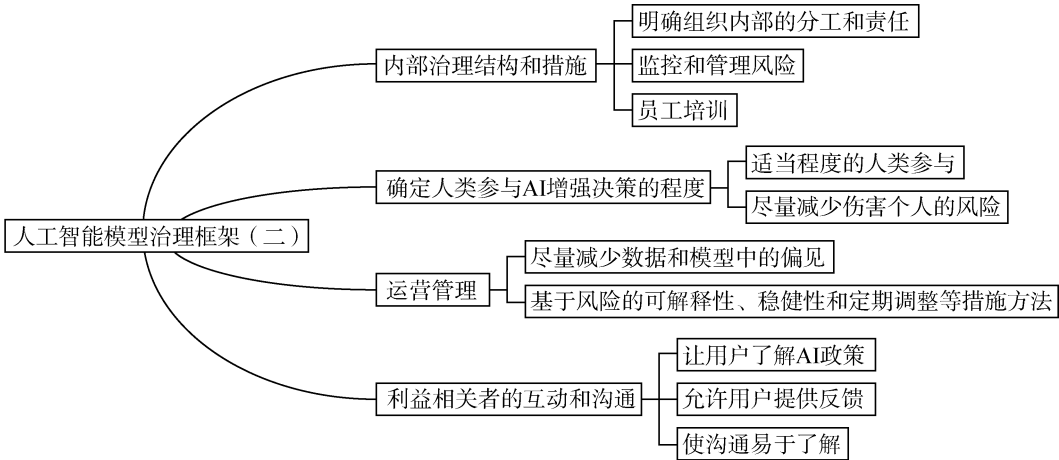


图 1 《治理框架(二)》概况^[22]

部署^[23]。数据治理内容在专章之余,遍布其他八个维度。一方面,《GAI 治理框架》将问责制置于首位,指导相关社会主体的责任分配。由于 GAI 涉及技术堆栈(tech stack)的多个层面,与云计算具有部分相似之处。所以,《GAI 治理框架》以云计算、软件开发治理规范为借鉴对象,制定 GAI 开发链上的主体责任分配规则^①。具言之,该规则包含事前与事后两个层次:事前层面,《GAI 治理框架》以开发链参与者对 GAI 的控制能力为分配基准,结合 GAI 模型的公开度与透明度分配责任。如,在开源模型的情形中,应用部署者应从正规平台下载模型以降低篡改风险;模型开发者需要为模型风险承担责任;事后层面,《GAI 治理框架》旨在构建以责任分担模式为基础,以赔偿、保险等为补充的安全网。在风险可视领域,《GAI 治理框架》要求模型开发者为可预见风险,如数据版权问题承担侵权责任。在风险半可视领域,其参考了欧盟的《人工智能责任指令》和《修订产品责任指令》,旨在降低终端用户的证明责任,确保用户与责任方在赔偿问题上的平等地位。在风险不可视领域,提出了将无过错保险^②作为解决方案,保证双方利益。另一方面,《GAI 治理框架》着眼于个人数据使用、版权数据使用与数据质量控制共 3 个方面:在个人数据的可信使用中,建议政策制定者、监管机构阐明个人数据保护法律在 GAI 数据治理中的运用路径,如规定知情同意规则的适用标准与例外;在版权数据使用中,指出各国均面临版权侵权问题,提倡通过多方对话寻求解决方案;在数据质量控制中,倡议设立世界范围内的可用数据库,建议各国政府与社会团体共建具有代表性的训练数据库来保证 GAI 的文化代表性。此外,《GAI 治理框架》还提出披露数据来源、提高模型透明度等多项举措以保证数据安全。

(3) 私权治理深化治理空间的层次性

新加坡的数据治理模式表现为“政府-社会-企业”的多元主体协同,以产学研融合共同应对技术与治理瓶颈^[24]。作为新加坡多元协同治理模式的重要产物,AIVF 是在 IMDA 牵头下成立的非营利性组织,旨在凝聚全球开源社会的集体力量以构建人工智能测试工具,为人工智能测试

框架、代码库、标准的研发及最佳实践作出贡献。在成员结构上,AIVF 在成立之初以 IMDA、微软、谷歌等 7 位成员作为高级成员,并吸纳了 Adobe、DBS、Meta 等 60 余名普通成员^[25]。同时,AIVF 集合多名具有世界影响力的人工智能专家、首席科学家组成全球人工智能治理小组,旨在剖析各国治理模式、研究治理对策、制定解决方案。可见,AIVF 承担着联系多元主体的重要功能,有助于寻求多方共同利益。此外,新加坡尚有诸多促进社会主体自治能力的举措:推出 GAI 监管沙盒,令中小型企业能够在 GAI 应用中受益^[26];发布《公私数据协作案例研究》,以指导社会主体的治理活动^[27]。足见,私权治理已经成为新加坡 GAI 数据治理领域的重要补充力量。

2. 新加坡多元协同数据治理模式的启示

首先,多元主体的利益诉求需要处于平衡状态。在 GAI 相关活动中,智能行业主体出于获取经济利益的目的,可能会将安全置于发展之后;公民出于数据利益、隐私保护的目的,可能会限制个人数据的授权使用;公共监管机构出于部门利益的需要,可能会导致治理的竞争与推诿。社会公共利益与个体利益至上存在天然矛盾,一味遵循社会公共利益或个体利益至上无益于实现治理效能。人类的本质属性是社会性,个体利益只有在社会公共利益得到满足的情况下,才能更好地实现。“以人为本”的价值理念始终居于人工智能价值体系的至高位置,在这一层面,多元主体存在利益一致性,对此,多元主体需要有序参与数据治理,让渡部分利益。公共监管机构需要积极履行职能为社会公共利益、社会主体合法利益的实现保驾护航;智能行业主体需要承担起数据风险治理第一责任人的职责,以更好助力数据治理活动;公民需要让渡部分数据利益,以实现社会整体利益最大化。

其次,智能行业主体的经济追求与个体德性需要维持相对稳定。如果智能行业主体在忽视

① 人工智能开发链上的参与者包括模型开发者、应用部署者、云服务提供者三类主体。

② 无过错保险(no-fault insurance),即无论谁有过错,利益相关者的费用均将得到承保,该措施当前被用于美国某些汽车事故索赔,该举措在人工智能问责制中的应用值得进一步研究。

安全、理性的情况下对经济利益过度追求,无益于社会整体安全与公民数据安全,明确 GAI 服务提供者的数据安全保护义务对于防范数据安全风险至关重要^[28]。个体德性促使智能行业主体回归理性人视角,承担应尽的治理义务。当前,大部分智能行业主体受外部压力,如媒体舆论、政府监管的影响,“被动”尝试化解权利不对称问题^[29]。在此情况下,智能行业主体虽然承担了部分治理义务,但是缺乏治理的积极性与主动性。个体德性应以自发为宜,智能行业主体需要强化主动治理意识,形成科学完备的内部治理体系与外部协作体系。只有在经济追求与个体德性达成一种相对平衡时,GAI 才能真正成为促进我国经济发展的新质生产力,数据治理的层次及体系方能更为完善。

最后,公共监管机构的治理活动需要兼顾发展价值与安全价值。每一项科技创新或多或少都会对已有秩序造成破坏,但亦会服务于秩序的重建。公共监管机构若一味遵循“家长制”的治理路径,仅偏重安全价值而忽视发展价值,无益于破坏的修复和秩序的重建。对此,公共监管机构需要将价值理性融入数据治理,防止安全价值对发展价值的绝对排斥,作出有利于实现社会公共利益最大化的决策。同时,公共监管机构的治理技术、治理思维尚存不足之处,引入社会主体参与治理具备必要性与正当性。

此外,新加坡的多元协同模式也有不足,如数据偏见预防规范尚未根治本国的 AI 决策歧视^[30]、隐私保护可能遭受挑战等,所以需辩证看待新加坡的治理经验。

三、我国多元协同数据治理模式的 范式构建

受“AIGC+”效应的影响,数据治理唯有经历一轮“破坏与创新”方可找到新的出路。新加坡的治理经验具有部分值得借鉴之处,可融入我国多元协同模式的构造中,以此落实我国《数据安全法》第9条承载的多元协同治理理念。

1. 构建有序分明的多元主体参与路径

(1) 以权力为指引的参与路径

随着 GAI 的深度应用,社会主体通过技术

赋权产生新型权力,由此,权力体系包括了国家权力(公权力)与社会权力(私权力):国家权力是主权者为维护统治秩序和社会秩序而由国家机构运用立法、行政与司法资源,对社会实施的“支配力、强制力和其他国家行为能力”;社会权力是社会主体运用拥有的社会资源对社会与国家产生的影响力与支配力^{[31]128}。通常情况下,社会权力易在国家权力的光辉下黯然失色,但仍不失为一股重要力量。

一方面,应当以国家权力指引公共监管机构参与的具体路径。公共监管机构主要通过“执法”实现治理目的,而数据治理的法律适用呈现出“民事-行政-刑事”的阶层性,以行为的影响程度差异为法律适用逻辑,所以,数据治理需要兼顾管理数据与规范行为两个方面。各公共监管机构的职权内容存在差异,如国家网信部门主要负责立法传播、内容管理等事宜,行使数据治理领域的统筹职权;公安机关主要负责追究数据犯罪、查获犯罪人等事项,行使侦查职权。对此,根据各公共监管机构的主要职权凝练出“内容-行为”的管辖治理模式,有助于形成治理合力。具体来说,可以遵循如下判断路径:其一,根据数据承载内容确定数据所属行业领域。通过领域分类,公共监管机构能够初步判断数据是否处于静态管辖职权范围,进而基于业务领域、流程环节等构建静态治理规范,实现初步管理。当前,国家标准《数据安全技术数据分类分级规则》已经确定了具体的分类思路及方法,有助于公共监管机构的静态治理。其二,根据数据的使用行为及其影响确定具体治理主体。对数据进行业务分类的意义在于判断数据使用行为是否属于业务领域的正常行为。如,银行职员利用职务便利倒卖内部数据显然属于业务领域的非正常行为,应由公安机关行使职权;银行职员制作数据报表供监管部门审查,属于业务领域的正常行为,由对应职能部门管辖;如果职能部门在审查过程中发现违法、犯罪事实,则需要移交公安机关等有权主体处理。所以,“内容-行为”的管辖治理模式能够明晰各公共监管机构的职权范围,充分落实“分工负责、互相合作”的治理原则,以此凝聚治理合力。

另一方面,应当以社会权力指引行业协会等自律组织的参与路径。社会权力的行权主体不仅包含行业协会等自律组织,而且包括公民、智能行业主体等。新加坡将部分治理任务赋权 AIVF 等社会主体,有效缓解了公共监管机构的治理压力。在我国,行业协会的权力主要来源于法律授权、政府委托、契约形成等 3 个方面,在长期实践中已然成为权力体系的重要组成部分。行业协会大都仅涉及行业领域内的管理活动,因此主要承担数据静态治理的职责。在具体的权力运作方面,行业协会具备的专业性、引领性令其得以在所属领域发挥引导功能。如,行业协会可以制定数据清单以落实分级分类管理制度,采用“定期走访+不定期抽查”的方式监督业内数据活动。充分发挥行业协会的社会权力,有助于形成权力层面的“公私协同”,但是社会权力的行使应当受到国家权力的监督与限制。所以,公共监管机构可以采用派员驻扎监督、定期交流报告等方式了解与规范社会权力的运作情况,实现精细治理。

(2) 以权利为指引的参与路径

作为数字社会的“原住民”,公民不仅是技术红利的享有者,而且是数据的所有者,当然可以参与数据治理以维护自身权利。但法律救济时常表现出滞后性:一方面,法律调整范围有限,无法及时解决新技术带来的权利侵害问题;另一方面,法律救济难以把握数据私权化与数据共享之间的界限,可能引发数据领域的新型“公地悲剧”^[32]。所以在立法保护之余,公民应当以数据所有者的身份直接参与治理。虽然公民大都缺乏对算法、法律的专业认知,但是对基础数据这种常识性的内容具有探讨能力,而且数据内容是否涉及隐私,这一判断标准应当具备可解释性,即需要公民结合自身情况予以说明。公民在向智能行业主体、公共监管机构反映问题的同时,亦实现了风险发现与修复的及时性。此外,删除权、更正权等救济性权利是公民维护自身权利的基础手段,需要实质的转化途径。在数据抓取的应用情形中,我国法律实践已经形成了“用户授权+数据控制者授权+用户再次授权”的认定路径^[33],对此需要稳固可行路径,指引公民实现合法权利。

(3) 以义务为指引的参与路径

智能行业主体需要对 GAI 应用产生的数据风险负责,这不仅是法律施加的义务,而且是其承担社会责任的表现。新加坡尚未将部分必要的 GAI 数据治理要素义务化,可能会导致智能行业主体执行不力的问题。对此,需将部分必要事宜“义务化”,引导智能行业主体积极实践,完成内部治理与外部协作两方面的任务。

一方面,需要提升智能行业主体参与协作治理的积极性。智能行业主体虽然具备技术性、灵活性等治理优势,但是终有力所不逮之时,需要与其他主体加强协作。在公私协作上,传统治理模式的对立格局压制了智能行业主体的协作意愿,需要调整主体定位来促发治理活力。新加坡的治理优势在于“少干预、多指引”,以此实现智能行业主体的自我治理效能。对此,我国公共监管机构可以借鉴新加坡的部分举措,将智能行业主体、技术人员、高校学者、行业协会负责人纳入案例汇编、政策制定等治理工作,充分吸收社会主体的合理需求,形成多元主体的平等治理地位。在社会主体之间的协作上,具备人工智能先发优势的行业主体需要主动承担起更多的治理责任,不仅应当持续完善自身的治理能力,而且需要加强与其他智能行业主体之间的交流合作,以自身影响力带动行业的规范发展。在具体实现路径上,可通过立法明确智能行业主体的协作义务:其一,重视《数据安全法》等法律中的协作条款的功能,要求智能行业主体积极履行公私协作义务;其二,形成“守门人”制度,充分发挥行业优势主体的引领功能,以此促进行业的内部协作与自律完善。与此同时,还需兼顾智能行业主体的利益诉求,为积极提供协作的智能行业主体提供声誉奖励、经济回馈等具体利益,以此激发协作积极性。

另一方面,智能行业主体需要坚守内部安全底线。新加坡的案例指引虽有助于智能行业主体的自我治理,但是忽略了不同主体的履行能力。对此,需明确自我治理的基础要素,并将其义务化。其一,建立数据溯源制度。溯源治理包括反向数据溯源、正向数据溯源、端对端数据溯源三种形式:反向数据溯源侧重从问题终

端回溯至数据源头,锁定问题成因;正向数据溯源的核心在于记录自数据源头至处理终端的全部数据活动,以此锁定问题环节;端对端数据溯源旨在通过数据源头与处理终端的双向溯源,对数据处理活动进行全面复查。通过3种溯源方式的综合运用,智能行业主体能够加强数据活动的透明度,实现有效的治理与追责。进言之,溯源治理具有双重意义:表面上,智能行业主体能够把控源头治理与过程治理,追踪问题来源与成因,以此降低数据质量、人工偏误等因素对人工智能的影响;深层上,溯源治理蕴含着数据来源需秉持合法、公开、可信的内在要求,有效保障了数据主体的合法权利,提升了数据传输活动的安全性。当人工智能的数据活动出现版权侵权、数据泄露等问题时,能够合理分配责任,保障权利主体得到及时救济。

其二,健全数据处理全流程的管理机制。在处理前端,智能行业主体应当进行数据分级分类管理,并对数据质量与数据偏差展开评估。一是对数据质量的评估。智能行业主体可以围绕数据的准确性、完整性、真实性、合法性、相关性展开评估,针对内容是否客观真实、来源是否合法可靠等进行分析。二是对数据偏差的评估。数据偏差产生于人工筛选、测量偏误、条件缺失等因素,智能行业主体需要保证数据的完整性,将数据筛选标准嵌入人工智能模型,以此降低数据的固有偏见。在处理中端,智能行业主体可以通过微调来降低数据安全风险,或是将人类社会的法律、价值伦理嵌入人工智能模型,实现模型的自我纠偏。同时,为防止价值偏误对人工智能决策的影响,智能行业主体可采用测试数据集检验、差异模型检验等方式验证决策内容。在处理末端,智能行业主体需要妥善管理数据,做好再利用、删除、存储工作。

其三,定期审查与筛除数据。数据定期审查需要明确审查的范围、方式、标准等内容。一是审查的范围应包含存储、训练、测试和验证数据。二是审查的方式。智能行业主体可以采用内部自评、专家审查或委托第三方机构评估等方式进行审查。实践中,第三方评估审查受到较多国家的认可,新加坡的第三方技术工具包审查方式亦有成效。由于外部评估能够切实增

强评估的公信力与透明度,可将此作为必要的审查方式。三是审查的标准,可围绕《数据安全法》等法律规范及各项国家标准确定宏观的审查要素,如来源合法性、内容真实性等,将具体内容交由行业协会、评估机构自行优化。经审查评估后,智能行业主体应该及时更新数据集。

其四,构建安全可行的存储机制与保护措施。为防止存储数据的完整性受损,智能行业主体可通过在内部云网络设立防火墙保护系统来实现入口控制,限制访问权限。对此,防火墙保护系统必须配备强大的入侵检测系统(IDS)与入侵防御系统(IPS),通过将网络活动与已知威胁数据库进行对比,拒绝有风险的访问行为。所以,智能行业主体必须定期更新威胁数据库及风险数据库,以此构筑外部风险的应对防线^[34]。此外,为防止内部人员的数据泄露行为,智能行业主体可采用分段技术将员工访问权限与工作权限挂钩,降低个体行为对整体数据安全的影响。

2. 完善技法结合的多元数据治理手段

(1) 公共监管机构的综合治理举措

一方面,公共监管机构需要创新治理工具。社会治理实践中,部分新型治理工具已经投入应用,如通过监管沙盒建立特定高风险情境中的责任追究制度等^[35]。监管沙盒制度能够有效缩小风险影响范围,兼顾治理与创新的双重需求,受到大部分国家的推崇应用。在我国,金融、科技、汽车安全领域已经开始监管沙盒的实践探索,GAI数据治理领域亦可以开展相应部署:适用对象上,监管沙盒可以优先向涉及公共安全等重大社会公共利益领域的技术开发者、中小企业和初创企业铺开;适用领域上,监管沙盒可以优先向医药卫生、金融创新等与社会公众联系密切的重点领域开放;适用标准上,公共监管机构可以结合GAI训练和部署的数据风险特性构建资格准入、运行监管、测试评估、申请退出的相应标准,与前述溯源制度形成内外追责的并行路径,实现对智能行业主体的有效监管。

另一方面,公共监管机构可以探索技术治理方法。当前,前瞻性评价、社会技术试验、技

术调节等多元化技术治理方法被渐次提出,治理举措更具灵活意义^[36]。为实现动态治理,公共监管机构可与智能行业主体共建智能监管平台,将法律规范嵌入系统,强化治理能力与效率。智能监管平台的主要功能在于日常巡视与异常预警:一是在日常巡视上,公共监管机构可从《数据安全法》等硬法中提炼与数据评估、备案审查、风险自评估相关的必备审查要素;从《新一代人工智能伦理规范》、各项国家标准、技术标准等软法中提炼参考审查要素,形成智能监管平台的实体标准与程序标准,对智能行业主体及人工智能数据处理活动中存在的泄漏风险、数据脱敏有效性、治理义务履行情况等予以监管。二是在异常预警上,公共监管机构需充分了解智能行业主体的运营情况、人工智能的数据处理机理、常见风险及突发风险的表现形式等内容,进而构建智能监管平台的预警数据库,由此实现事前风险预防、事中风险应对、事后风险分析的全流程治理。此外,我国可参照新加坡的技术治理路径,由国家网信部门牵头设立公私协作机构,致力于研发检测技术工具,便于公共监管机构及时掌握人工智能的系统指数,亦便于智能行业主体及时检测人工智能的安全性、实用性等属性。

(2) 社会主体的技术治理举措

当前,以智能行业主体为代表的社会主体已经开展了针对人工智能的技术治理研究,如探索价值对齐技术、用人工智能解释人工智能的方式破解算法黑箱、将隐私增强技术用于治理等^[37]。但是,技术治理的根源性问题,即研发可解释的人工智能,尚在进行中。

可解释的人工智能研究包括外源性研究与分解性研究^[38];外源性研究可以进一步分为以模型为中心的解釋研究与以主题为中心的解釋研究,前者侧重于解释人工智能的工作机理,如数据处理逻辑、训练数据的影响参数等;后者侧重于通过与其他人工智能模型的输出结果进行比较以解释某一特定输出结果,进而指出特定案例的输入属性变化如何影响人工智能的输出结果。分解性研究试图通过解释人工智能自身的运作,如披露人工智能的源代码或是重构另一个模型来寻找输入结果与输出结果之间的关

联性。为实现可解释的人工智能,部分智能行业主体尝试通过反向工程等技术方法,揭示人工智能神经网络的底层逻辑,以此明晰其运作机理,或是采用生成式对抗网络,通过循环往复地运算揭露人工智能的决策错误。此外,还有部分研究者专注于开发可信的人工智能,如通过知识提取、推理与内容的可靠应用等举措的综合运用来保证 GAI 数据生成物内容的可靠与可控^[39]。然而,大部分技术仍是通过事后分析方法来明晰不确定性算法的错误成因,难以满足事前预防、事中应对的治理需要,加之智能行业主体具有天然的逐利形象,以至于治理技术市场化缺乏足够的信任基础。所以,多元社会主体仍需不断推进技术研发,同时需要主动加强与公共监管机构的联系,共筑人工智能数据治理的技术生态。

3. 形成软硬并举的多元数据法律规范

在新兴技术尚未成熟前,软法在硬法不便触及的领域发挥着重要的补充功能;当条件成熟后,软法具备全部或部分转化为硬法的可能,最终形成更为完善的法治体系。新加坡充分发挥了软法的治理优势,但尚未完全协调软法与硬法的关系,以至于常常遭受重软法而轻硬法的诟病。因此,需要在秉持硬法底线的前提下,发挥软法的灵活治理优势,可以将软法执行情况与主体利益部分挂钩,以此增强智能行业主体的执行力。

(1) 以硬法作为数据治理的法律底线

硬法治理可能阻碍科技创新以及软法的推广应用,但其优势在于执行力度和执行效率。因此,数据治理离不开《数据安全法》等硬法规范,但同时需要为技术创新与软法适用提供一定自由度。一方面,硬法需要为软法提供现实土壤,实现法律体系的内在融贯。以数据分级分类管理为例,核心数据、重要数据的使用必须严格贯彻硬法规范,一般数据的使用可以部分交由软法细化。另一方面,硬法需要压实人工智能研发者、部署者、使用者的法律责任,防止人工智能应用逾越底线。在后续的人工智能法律制定过程中,需要明确三项内容:第一,人工智能应用的“受害者”可以在哪些情况下声称自己的权利受损,以及如何补救;第二,公共监

管机构如何有效保障多元主体的合法权利,以及职权范围是否需要调整;第三,开发者、部署者等智能行业主体需要遵循何种约束性规范,以及不同主体间的责任分配^[40]。只有通过硬法明确不同主体的权利与义务,“以人为本”的治理理念方能落实。

(2)以软法作为数据治理的操作框架

软法治理的优势在于风险应对的灵活性。当前,我国已经初步形成了自上而下的人工智能软法体系,具体包括国家制定的政策性文件、行业标准和技术标准、行业协会自律公约或合规指南等^{[8]102}。但是,在实践中软法常常面临执行力度不足、执行透明度不够的诟病。对此,有学者提出“软法 2.0 版”,并建议将企业道德委员会、供应链影响、政府采购业务、外部审计、专利许可、学者声明等 11 项措施作为增强软法执行度的工具箱^{[17]403}。强化软法的执行力度不仅要依靠智能行业主体的自觉,更要外部力量的推动。所以,可构建软法治理的政府引导和奖惩机制,如将政府采购业务与智能行业主体的治理评级挂钩,优先考虑治理成效优异主体;对于治理评级不合格者,则以消减声誉、减少交易机会等措施予以惩罚。此外,在硬法尚未出台前,科研机构、高校学者可以借助学术期刊等公开平台出台建议稿,进一步完善软法体系。

四、结 语

对规则的颠覆可以采取不同形式。然而,关键问题与其说是一般原则和规则的适用受到的挑战,不如说是回应的性质^{[12]28},这一点在 GAI 的数据治理中体现得尤为明显。当前,我国已出台了一系列数据治理法律法规,并在积极探索数据治理的多样化举措,但仍难以应对 GAI 带来的数据治理挑战。与此同时,《数据安全法》第 9 条确立的多元协同理念亦仅停留于理念层面,在实践中较为薄弱。本研究以我国与新加坡的数据治理模式为研究对象,旨在明确我国传统数据治理模式不足的基础上,借鉴新加坡的治理经验,以构建适应我国国情的多元协同数据治理模式。囿于篇幅,本研究尚未对新加坡的多元立法规范、多元治理举措实现完全发掘,主要以典型范例突出新加坡治理模

式的特色,且软法 2.0 的理论设想对于构建我国软法硬法协调并举的规范体系具有重要意义。在后续研究中,可以在着墨欧盟、美国等典型治理模式之余,对新加坡的治理模式有所着力,以此进一步深化 GAI 数据治理模式的理论研究。同时,还需要充分探索软法 2.0 模式下的软法硬法协同路径,设计更为完善的软法治理工具箱,用法治为 GAI 发展及数据治理活动保驾护航。

参考文献:

[1] 人工智能产业发展联盟 AIIA. 以治理促发展,推动智能向善——人工智能立法重大问题产业研讨会成功举办[EB/OL]. (2024-04-22) [2024-06-05]. https://mp.weixin.qq.com/s?__biz=MzU0MTEwNjg1OA==&mid=2247504627&idx=1&sn=cf59b86a1143cb8240e1e3361592622d&chksm=fb2c6536cc5bec2083e2d014765751d87036265a545f7db43e3988a50f8fcf90528bab202bba&scene=27.

[2] E-governance academy. National cyber security index [EB/OL]. (2021-10-18) [2024-10-13]. https://ncsi.ega.ee/country/sg_2022/?pdfReport=1.

[3] MOSTROUS A, WHITE J, CASAREO S. The global artificial intelligence index [EB/OL]. (2024-09-19) [2024-10-13]. <https://www.tortoisemedia.com/2024/09/19/the-global-artificial-intelligence-index-2024/>.

[4] 郑煌杰. AIGC 赋能新质生产力的数据风险及其敏捷治理 [J]. 河海大学学报(哲学社会科学版), 2024, 26(4): 89-102.

[5] 刘艳红. 生成式人工智能的三大安全风险及法律规制——以 ChatGPT 为例 [J]. 东方法学, 2023 (4): 29-43.

[6] 张涛. 人工智能治理中“基于风险的方法”: 理论、实践与反思 [J]. 华中科技大学学报(社会科学版), 2024, 38(2): 66-77.

[7] 毕文轩. 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 [J]. 比较法研究, 2023(3): 155-172.

[8] 曾雄, 梁正, 张辉. 人工智能软法治理的优化进阶: 由软法先行到软法与硬法协同 [J]. 电子政务, 2024(6): 96-107.

[9] COLEBATCH H K. Making sense of governance [J]. Policy and Society, 2014, 33(4): 307-316.

[10] HONG M C, HUI C K. Towards a digital government;

- reflections on automated decision-making and the principles of administrative justice [J]. Singapore Academy of Law Journal, 2019, 31(2): 875-906.
- [11] 张凌寒. 中国需要一部怎样的《人工智能法》? ——中国人工智能立法的基本逻辑与制度架构[J]. 法律科学(西北政法大学学报), 2024, 42(3): 3-17.
- [12] 罗杰·布朗斯沃德. 法律 3.0: 规则、规制和技术 [M]. 毛海栋, 译. 北京: 北京大学出版社, 2023: 28-43.
- [13] 马长山. 互联网+时代“软法之治”的问题与对策 [J]. 现代法学, 2016, 38(5): 49-56.
- [14] 龙龙. 敏捷治理理念下人工智能软法之治的问题与对策[J]. 江汉论坛, 2024(5): 128-134.
- [15] SCHONBERGER V M, RAMGE T. A big choice for big tech: share data or suffer the consequences [J]. Foreign Affairs, 2018, 97(5): 48-54.
- [16] KUNER C, CATE F H, LYNSKEY O, et al. Expanding the artificial intelligence-data protection debate [J]. International Data Privacy Law, 2018, 8(4): 289-292.
- [17] MARCHANT G E, CARLOS I G. Soft Law 2.0: an agile and effective governance approach for artificial intelligence [J]. Minnesota Journal of Law, Science and Technology, 2023, 24(2): 375-424.
- [18] LEE J. AI Verify: Singapore's AI governance testing initiative explained [EB/OL]. (2023-06-06) [2024-06-07]. <https://fpf.org/blog/ai-verify-singapores-ai-governance-testing-initiative-explained/>.
- [19] 季卫东. 何时真正迈入“人工智能治理立法时刻” [N]. 上海法治报, 2024-05-15(B3).
- [20] AI Verify Foundation. Project Moonshot——an LLM evaluation toolkit [EB/OL]. (2024-05-31) [2024-08-05]. <https://aiverifyfoundation.sg/project-moonshot/>.
- [21] CHANG D. AI Regulation for the AI revolution [J]. Singapore Comparative Law Review, 2023(1): 130-170.
- [22] IMDA. Model artificial intelligence governance framework (Second Edition) [EB/OL]. (2020-01-21) [2023-11-13]. <https://www.pdpc.gov.sg/-/media/Files/PDPC/PDF-Files/Resource-for-Organisation/AI/SGModelAIGovFramework2.pdf>.
- [23] AI Verify Foundation, IMDA. Model AI governance framework for generative AI: fostering a trusted ecosystem [EB/OL]. (2024-05-30) [2024-06-07]. <https://aiverifyfoundation.sg/wp-content/uploads/2024/05/Model-AI-Governance-Framework-for-Generative-AI-May-2024-1-1.pdf>.
- [24] SONDERLING K E, KELLEY B J. Filling the void: artificial intelligence and private initiatives [J]. North Carolina Journal of Law & Technology, 2023, 24(4): 153-200.
- [25] IMDA. Singapore launches AI Verify Foundation to shape the future of international AI standards through collaboration [EB/OL]. (2023-06-07) [2024-02-08]. <https://www.imda.gov.sg/resources/press-releases-factsheets-and-speeches/press-releases/2023/singapore-launches-ai-verify-foundation-to-shape-the-future-of-international-ai-standards-through-collaboration>.
- [26] Enterprise Singapore, IMDA. Singapore's first generative AI Sandbox to familiarise and help SMEs get head start in capturing new AI opportunities [EB/OL]. (2024-02-07) [2024-02-08]. <https://www.imda.gov.sg/resources/press-releases-factsheets-and-speeches/press-releases/2024/sg-first-genai-sandbox-for-smes>.
- [27] IMDA. Public-private data collaboration case study [EB/OL]. (2023-08-03) [2023-11-13]. <https://www.imda.gov.sg/-/media/Imda/Files/Programme/Data-Collaborative-Programme/Datathon-Case-Study.pdf>.
- [28] 王良顺, 李想. 生成式人工智能服务提供者的数据安全保护义务研究 [J]. 南昌大学学报(人文社会科学版), 2023, 54(6): 72-83.
- [29] MICHELI M, PONTI M, CRAGLIA M, et al. Emerging models of data governance in the age of datafication [J]. Big Data & Society, 2020, 7(2): 1-15.
- [30] KHOO S, CHOW E Z. Algorithmic fairness: challenges and opportunities for artificial intelligence governance [J]. Singapore Academy of Law Journal, 2022, 34(Special Issue): 736-767.
- [31] 肖梦黎. 从平台治理到治理平台: 平台自我规制的风险与问责分析 [M]. 北京: 中国政法大学出版社, 2023: 128.
- [32] 任颖. 数据立法转向: 从数据权利入法到数据法益保护 [J]. 政治与法律, 2020(6): 135-147.
- [33] 郭传凯. 数据抓取的合法性边界探究 [J]. 法学评论, 2024, 42(2): 122-132.
- [34] EN T M. Non-deterministic artificial intelligence systems and the future of the law on unilateral

mistakes in Singapore [J]. Singapore Academy of Law Journal, 2022, 34(1): 91-124.

[35] TRUBY J, BROWN R D, IBRAHIM I A, et al. A sandbox approach to regulating high-risk artificial intelligence applications [J]. European Journal of Risk Regulation, 2022, 13(2): 270-294.

[36] 贾向桐, 胡杨. 从技术控制的工具论到存在论视域的转变——析科林格里奇困境及其解答路径问题[J]. 科学与社会, 2021, 11(3): 26-39.

[37] 胡裕岭, 姚浩亮. 生成式人工智能的数据治理与“囚笼”——基于行业主体自我监管的回应型研究[J]. 征信, 2024, 42(2): 46-57.

[38] EDWARDS L, VEALE M. Slave to the algorithm? Why a “right to an explanation” is probably not the remedy you are looking for [J]. Duke Law & Technology Review, 2017—2018(16): 18-84.

[39] LI L H, ZHANG H P, LI C J, et al. Evaluation on ChatGPT for Chinese language understanding [J]. Data Intelligence, 2023, 5(4): 1-19.

[40] LANE L. Clarifying human rights standards through artificial intelligence initiatives [J]. International and Comparative Law, 2022(71): 915-944.

Construction of a Multi-Synergy Mode for GAI Data Governance in China: Inspiration from Singapore’s Governance Experience/ HU Yuling, YAO Haoliang (Criminal Law School, East China University of Political Science and Law, Shanghai 201620, China)

Abstract: Data governance is a topic of focus for the development of generative artificial intelligence technology in modern countries. Since the 18th National Congress of the CPC, the CPC Central Committee and the State Council have attached great importance to data governance, and have successively introduced a series of laws and regulations, such as the Data Security Law of the People’s Republic of China, the Measures for Security Assessment of Data Exit, and the Provisions for Promoting and Regulating Cross-border Flow of Data. However, in the face of the data risks brought about by generative artificial intelligence, China’s current data governance model is biased towards a hard law model of state regulation and enforcement of laws and regulations, and has not yet formed a pluralistic pattern of collaborative participation by the state, society and individuals, and governance means such as record-keeping, review and investigation are relatively static, making it difficult to cope with the challenges of the uncertainty brought about by technological development and change. Singapore’s pluralistic and synergistic data governance model has fully absorbed the governance experience of the EU, the U. S. and other countries and regions with advantages in AI data governance, and has taken the lead in launching technical governance programs globally, forming a data governance system that integrates pluralistic subjects, pluralistic norms and pluralistic initiatives. Under this model, technical governance tools such as AI Verify and Project Moonshot enhance the certainty of governance norms, soft laws such as the Generative Artificial Intelligence Model Governance Framework and hard laws such as the Personal Data Protection Act safeguard the enforceability of governance tools, and the AI Verify Foundation builds a bridge for public-private collaboration, deepening the hierarchy of the governance space. In this regard, drawing dialectically on Singapore’s governance experience, we should construct a path of participation by multiple subjects guided by power, rights and obligations, improve the combination of technology and norms of multiple governance initiatives, enriching the toolbox of data governance, give full play to the advantages of the effective implementation of hard law and the flexible governance of soft law, realizing the internal integration of the legal and normative system of data governance, so as to build a multifaceted and collaborative data governance model that is adapted to China’s basic conditions and meets its governance needs. In this way, a multifaceted and synergistic data governance model adapted to China’s basic national conditions and meeting its governance needs will be built.

Key words: generative artificial intelligence; data governance; multiple synergy; hierarchical governance; Singapore