

人工智能之辅助性定位原则的规范构造

吴育珊 杜 昕

摘要：人工智能的应用正由机械性赋能向创造性赋智转变，这引发是否赋予其法律主体资格的争论。不赋予或赋予其完全法律主体资格都不契合现实，所以应赋予其以适当法律主体资格，使其始终处于辅助性定位，亦即人工智能之辅助性定位原则。该原则在规范人工智能研发和应用方面具有普适的指导功能和规制功能：指导人工智能的研发和应用方向，规制人工智能的安全风险。该原则的实体构造包括：应赋予人工智能以适当法律主体资格、规范人工智能的研发和应用行为、调整人工智能与其他法律主体之间的关系。该原则的程序构造包括：内生自主性特征评估、标准化安全定型、人工智能失控的救济机制。无内生自主性或内生自主性完全可控的人工智能不具有法律主体资格，适用一般规定，是该原则的例外情形。

关键词：人工智能；法律主体资格；辅助性定位

中图分类号：TP18;D922.16 **文献标识码：**A **文章编号：**1673-5706 (2023) 06-0051-07

2023年4月28日，中共中央政治局召开会议^[1]，明确指出要重视通用人工智能发展，营造创新生态，重视防范风险。相较于具有领域局限性的专用人工智能（Artificial Narrow Intelligence，简称“ANI”），通用人工智能（Artificial General Intelligence，简称“AGI”）减少了对特定领域知识的依赖性、提高了处理任务的普适性，是人工智能未来的发展方向^[2]。这意味着人工智能将广泛应用于包括法律在内的诸多领域，2023年1月30日世界首份应用ChatGPT辅助作出的判决^[3]就

是例证之一。在这一背景下，人类与人工智能之间的关系定位也引发全球热议。

就我国立法而言，2022年12月8日，最高人民法院印发《关于规范和加强人工智能司法应用的意见》^[4]，明确要求坚持人工智能对审判工作的辅助性定位，并将之确立为人工智能司法应用的基本原则之一。这是人工智能之辅助性定位原则在司法应用方面的体现，明确了人工智能司法应用中的人机关系。人工智能之辅助性定位原则是如何提出的？具有怎样的功能？在此基础上讨

论其规范构造,是本文的问题意识所在。

一、人工智能之辅助性定位原则

(一) 提出背景:人工智能的应用正由机械性赋能向创造性赋智转变

当前,人工智能发展再一次迈入关键时期,以生成式人工智能(Generative AI或AI Generated Content,简称“AIGC”)[5]为代表的新技术、新应用不断打破人们对于人工智能的固有认知,人工智能的应用正由机械延伸人类能力向生成创意内容以启发人类创造力的方向迈进。随着自然语言处理和计算机视觉技术的进步,生成式人工智能研发不断深入,其应用场景涵盖教育、医疗、法律、文学、传媒、影音创作、图表制作、算法编程、软件开发等诸多需要文本生成、图像生成或其他有内容生成任务的领域。在具体应用过程,最前沿的人工智能不仅可以进行自动、机械地劳动作业,还可以生成具有启发价值的结果,为人类进一步深度创造提供参考素材。这意味着,人类在未来各个场景的工作效率将大幅提升,工作内容将朝着更具创造性的方向发展,工作形式也将更加自由。换言之,人工智能的发展不仅会改变人类的物质生产形态,还将逐步改变人类的思想生产形态,实现由机械性赋能向创造性赋智转变[6]。在这一背景下,讨论人工智能与人类之间的关系定位变得至关重要。

(二) 法理溯源:人工智能的法律主体资格争论

人工智能与人类之间的关系定位所映射的法律问题为“是否应赋予人工智能以法律主体资格”。根据目前人工智能最前沿的GPT-4的技术报告[7]显示,其在许多领域的专业基准上已表现出与人类相当甚至更为优秀的水平,这意味着以自然语言处理和计算机视觉技术为基础并且具备深度学习能力的人工智能已经可以通过大数据训练进行自主决策或对环境作出自主回应。人工智能的这种内生自主性特征引发了“是否应赋予人工智能以法律主体资格”的争论。具体观点主张如下:

不赋予人工智能以法律主体资格。这一观点对应了安德鲁·芬伯格[8]划分的工具理论(Instrumental Theory),工具理论的主要代表人

物是卡尔·雅斯贝斯,他认为“技术本身既无善也无恶,但它既可用于善也可用于恶。它本身不包含任何理念,既不包含完美的理念,也不包含邪恶的毁灭理念。善恶的理念都源自于人,是人赋予技术以意义”[9]。工具理论因其更符合常识而被学界广泛接受,例如有学者认为“电子计算机没有好恶、亲疏等感情因素,在某种程度上可以避免法律工作者凭自己的情感,一时的冲动来处理问题”[10];有学者提出要“实现对人的行为而非对技术本身的规制”[11];还有学者指出“就技术实质而言,算法本身是中立的”[12]等。然而,人工智能目前已经呈现出内生自主性特征,算法设计者无法预知算法经过训练自主迭代生成的新算法,亦即人类无法像过去掌控工具一样完全绝对地控制人工智能,工具理论已不再契合实际情况。对此,有观点提出应赋予人工智能以法律主体资格。

赋予人工智能以完全法律主体资格。与工具理论相对的是实体理论(Substantive Theory),主要代表人物是雅克·埃吕尔[13]和马丁·海德格尔[14]。实体理论主张跳出人类中心主义的局限,摒弃人类作为主体、技术作为客体的“主客二分”预设,认为现代科技是一个具备内生自主性且有扩张倾向的价值实体。在实体理论的启发下,有观点结合人工智能的发展实际提出应赋予其以完全法律主体资格,劳伦斯·索勒姆[15]是这一观点的代表人物,他认为当人工智能通过图灵测试,拥有与人类等同的所有功能,亦即发展成所谓超级人工智能(Artificial Super Intelligence,简称“SuperAI”或“ASI”)[16]时,就可以成为宪法意义上的主体。然而,即使人工智能发展到通过图灵测试的水平,在行为与功能方面表现出类人甚至超人的应对能力,这种应对能力也只是基于训练数据做出的概率性优化选择,其自身既无法像人类一样体会享有权利的激励,也难以感知受到惩罚的悔过。换言之,人工智能并不能与其他法律主体享有同等的权利、履行同等的义务,也不能代替其他法律主体实施特定行为,这使得赋予人工智能以完全法律主体资格丧失应有的规范意义。

赋予人工智能以适当法律主体资格。如前所

述,不赋予人工智能以法律主体资格或赋予完全法律主体资格都不契合目前的发展实际,所以需要平衡工具理论与实体理论的主张,在人类与工具之间拟制新的主体,亦即赋予人工智能以适当法律主体资格。苏珊娜·贝克^[17]是这一观点的代表人物,其论证由类人性和制度性两部分组成。就类人性而言,人工智能只能通过深度学习模仿人类的情感,但却无法真切感知,不存在享有权利的激励感或受到惩罚的悔过感,赋予其与自然人同等的完全法律主体资格是无意义的;就制度性而言,具有内生自主性的人工智能可以看作一个复杂的算法系统,包含各个独立要素之间的多重相互作用,越是先进的人工智能,其系统越复杂,预测其行为就越困难,不确定性就越明显,对其行为后果的原因追溯也就越不可能,为解决这一问题,可以类比公司的法律人格拟制一个制度性主体,赋予人工智能以适当法律主体资格。

然而,要赋予人工智能以适当法律主体资格还需解决一个前置性问题,即人工智能是否具有独立且单一的意志,这是因为无论自然人还是拟制的法人格主体都因其具有独立且单一的意志而成为法律主体^[18],如果人工智能不满足这一条件,那么也无法赋予其法律主体资格。结合前文论述不难发现,具备内生自主性特征的人工智能有学习迭代的扩张倾向,其系统的复杂性使人类难以完全绝对地控制这一倾向,所以人工智能具有独立且单一的意志,且这种意志具有扩张性。这是人工智能由专用向通用发展的必然结果,应当正视其独立意志,赋予其以适当法律主体资格,才能引导其良性发展,确保人类在人机关系中占据主导地位。因此,赋予人工智能以适当法律主体资格是较为现实的主张。从这一观点出发,人工智能在人机关系中应始终处于辅助性定位,不能与其他法律主体享有同等的权利履行同等的义务,不能代替其他法律主体实施特定行为,只能在辅助性定位范围内实施行为,享有权利,履行义务,亦即人工智能之辅助性定位原则。

(三)具体要素:人工智能之辅助性定位原则的内容

目前,明确提出人工智能之辅助性定位原则

的规范性文件是最高人民法院印发的《关于规范和加强人工智能司法应用的意见》,并命名为“辅助审判原则”。其内容表述为“无论技术发展达到何种水平,人工智能都不得代替法官裁判,人工智能辅助结果仅可作为审判工作或审判监督管理的参考,确保司法裁判始终由审判人员作出,裁判职权始终由审判组织行使,司法责任最终由裁判者承担”。根据这一表述并结合前文论述,不难发现人工智能之辅助性定位原则的内容应包含三个要素,一是人工智能仅具有适当法律主体资格,不能与其他法律主体享有同等的权利、履行同等的义务,对应“辅助审判”的要求;二是规范人工智能的研发和应用行为,亦即人工智能只能在辅助性定位范围内实施行为、享有权利、履行义务,对应“人工智能辅助结果仅可作为审判工作或审判监督管理的参考”;三是调整人工智能与其他法律主体之间的关系,人工智能的适当法律主体资格是因现实需要而拟制的,这意味着人工智能实施行为、享有权利、履行义务以及承担法律责任的方式都不同于其他法律主体,不能代替其他法律主体实施特定行为,对应“不得代替法官裁判”,“确保司法裁判始终由审判人员作出,裁判职权始终由审判组织行使,司法责任最终由裁判者承担”。

二、人工智能之辅助性定位原则的功能

(一)指导人工智能的研发和应用方向

“是否应赋予人工智能以法律主体资格”是人工智能与人类之间的关系定位所映射的法律问题,而人工智能与人类之间的关系定位则是科技伦理治理的重要内容。科技伦理是人工智能研发和应用需要遵循的价值理念和行为规范,是促进人工智能良性发展的重要保障。对此,中共中央办公厅、国务院办公厅印发的《关于加强科技伦理治理的意见》^[19]明确提出“伦理先行”的治理要求,将科技伦理治理摆在事关全局的重要位置。结合前文论述,赋予人工智能以适当法律主体资格更加贴近现实,有利于在统筹发展和安全的同时确保人类在人机关系中占据主导地位,而人工智能之辅助性定位原则正是这一观点的制度具现。所以人工智能之辅助性定位原则可以通过其科技

伦理内涵来指导人工智能的研发和应用方向。

（二）规制人工智能的安全风险

人工智能已经具备了内生自主性特征，这意味着人类难以完全预知人工智能的输出结果，亦即人工智能存在输出结果失控从而危及安全的风险。安全风险具体表现为技术的无序扩张，究其诱因在于资本化的人工智能造成的伦理困境以及人工智能自身的技术局限。

即便最前沿的人工智能也需依靠其背后的技术资本才能进行研发并推广应用，而技术资本本身就有扩张的倾向，这种倾向将直接影响人工智能的发展。例如，研发 ChatGPT 的 OpenAI 成立的初衷旨在确保人工智能造福全人类，其核心目标并非盈利，所以开放是其最显著的特征，全球研发者都可以通过其提供的开发与研究框架投入到人工智能的研究当中，然而随着资本的注入，OpenAI 开始逐步设置技术壁垒，逐渐将盈利确立为首要目标，背离了创立的初衷。“当资本迫使科学为自己服务时，它总是迫使劳动的反叛之手就范”^[20]，智能时代资本化的人工智能对人类的剥削手段将更加多样化、隐性化、离散化，并将借助资本扩张的倾向实现自身的无序扩张，极易导致人类丧失人机关系中的主导地位，从而造成伦理困境。而人工智能之辅助性定位原则将人工智能的研发和应用严格限制在辅助性定位范围内，为防范技术资本无序扩张提供了制度保障，可以有效破解伦理困境。除了资本化的人工智能造成的伦理困境之外，人工智能自身的技术局限也会诱发安全风险。根据 GPT-4 的技术报告显示，其预训练数据是从互联网上获取的截至 2021 年 9 月的公开数据，缺乏对 2021 年 9 月以后发生事件的知识，有一定的时效性局限。预训练数据的时效性局限和人工智能持续学习能力的欠缺，会使其难以从经验中汲取教训，以至于会犯一些简单的推理错误，表现为过分容易接受用户的明显错误陈述，从而导致输出结果失控。此外，输出结果失控还可能与算法局限有关，例如 Transformer 架构下词汇间两两计算的推理方式就可能导致 GPT-4 在推理过程中凭空捏造事实，特别是在模棱两可的上下文推断当中，最有可能发生这一情

形。人工智能之辅助性定位原则要求人工智能仅具有适当法律主体资格，不能代替其他法律主体实施特定行为，其承担法律责任的方式也不同于其他法律主体，所以在该原则指导下配合标准化定型规则可以为规制人工智能自身的技术局限提供现实路径。

三、人工智能之辅助性定位原则的构造

（一）人工智能之辅助性定位原则的实体构造

赋予人工智能以适当法律主体资格。因为人工智能的算法模型日趋复杂，难以实现可解释性目标，这意味着人工智能已呈现出内生自主性，可以实施人类难以完全控制的、自发自主的行为，且从生成式人工智能的技术原理出发，有创意内容生成同样意味着人工智能必然可以实施人类难以预料的自发行为，所以即便从人类自身利益出发，也应完善现行法律制度赋予其以法律主体资格。又因为无论技术发展到何种水平，人工智能都无法与人类等同，其法律地位也必须低于人类，所以仅可赋予其以适当法律主体资格。

人工智能之辅助性定位的内涵必然不同于工具性定位，辅助性定位是人工智能介于其他法律主体与工具客体之间的拟制法律主体定位，而工具性定位的规范则早已被传统法律体系所涵盖。所以，只有在人工智能仅具有适当法律主体资格的情况下，才能讨论人工智能之辅助性定位原则的规范构造。换言之，人工智能具有且仅具有适当法律主体资格是适用人工智能之辅助性定位原则的前提条件。因此，完善相关立法，赋予人工智能以适当法律主体资格是人工智能之辅助性定位原则的实体构造之一。

2. 规范人工智能的研发和应用行为。人工智能的研发行为需要严格遵循科技伦理，且以人为本是科技伦理法律化的底线^[21]，所以研发者必须秉持造福人类的初衷研发人工智能，坚持人工智能服务人类的研发导向，促进人类社会和平发展和可持续发展，最大限度地避免对人类造成伤害或潜在威胁。而人工智能之辅助性定位原则可以通过其科技伦理内涵来指导人工智能的研发方向，并可以配合标准化定型规则从源头规制安全风险，加之辅助性定位对应着以人为本的科技伦理法律

化底线,因此人工智能之辅助性定位原则规范人工智能的研发行为。人工智能之辅助性定位原则要求人工智能仅具有适当法律主体资格,所以人工智能不能与其他法律主体享有同等的权利、履行同等的义务,不能代替其他法律主体实施特定行为。这不仅意味着人工智能只能在辅助性定位范围内实施行为、享有权利、履行义务,还意味着其应用行为也只能局限在辅助性定位范围内,确保人类在人工智能的应用过程占据主导地位。

3. 调整人工智能与其他法律主体之间的关系。

因为人工智能的适当法律主体资格是因现实需要而拟制的,介于其他法律主体与工具客体之间,始终处于辅助性定位,这不仅意味着人工智能实施行为、享有权利、履行义务以及承担法律责任的方式都不同于其他法律主体,还意味着人工智能与其他法律主体的关系应当是辅助性关系,不能代替其他法律主体实施特定行为。换言之,人工智能之辅助性定位原则旨在调整人工智能与其他法律主体之间的关系,确保人工智能始终处于辅助性定位。

(二) 人工智能之辅助性定位原则的程序构造

1. 内生自主性特征评估。因为人工智能之辅助性定位原则要求人工智能仅具有适当法律主体资格,而具有适当法律主体资格的前提是人工智能具备人类难以完全绝对地控制的内生自主性特征,所以内生自主性特征评估是适用该原则的首要程序。在具体实践当中,内生自主性特征评估应包括三个要素:深度学习能力评估、内容生成能力评估和自适应能力评估。

人工智能研发的核心问题是使其可以从现象中发现规律,而深度学习能力就是旨在针对复杂特征的事物,通过多层次地学习相关数据来提炼一般规律。常用的深度学习模型是多层神经网络,多层神经网络的每一层都将输入非线性映射,通过多层非线性映射的堆叠,可以在深层神经网络中计算出抽象的特征来帮助提炼一般规律^[22]。不难发现,深度学习能力极大提升了人工智能在复杂的现实环境中提炼一般规律的水平,是其具备内生自主性特征的基础。因此,深度学习能力评估是内生自主性特征评估的基础。内容生成能力

决定了人工智能是否能够通过深度学习进行内容创造,内容生成能力越强的人工智能,智能化水平就越高,越容易产生内生自主性特征。因此,内容生成能力评估是内生自主性特征评估的前提。自适应能力意味着人工智能可以通过持续学习自主迭代适应不断变化的环境,自适应能力是人类难以像控制工具一样完全、绝对地控制人工智能的原因,是人工智能具有独立且单一意志的体现。因此,自适应能力评估也是内生自主性特征评估要素之一。

2. 标准化安全定型。标准化安全定型是我国许多法律规范科技研发和应用的立法思路。就科技研发而言,《中华人民共和国科学技术进步法》^[23]第三十七条规定国家推动科学技术研究开发与产品、服务标准制定相结合。就科技应用而言,在国家治理层面,《中华人民共和国行政处罚法》^[24]第四十一条第一款规定辅助执法的技术应当经过法制和技术审核,确保技术设备符合标准;在社会治理层面,《中华人民共和国个人信息保护法》^[25]第六十二条规定国家网信部门统筹协调有关部门依据本法针对新技术、新应用制定专门的个人信息保护规则、标准,《中华人民共和国数据安全法》^[26]第十七条规定国务院标准化行政主管部门和国务院有关部门根据各自的职责组织制定并适时修订有关数据开发利用技术、产品和数据安全相关标准。

根据 GPT-4 的技术报告显示,其在安全性方面的提升得益于 OpenAI 的对抗性测试计划,这意味着针对人工智能研发和应用存在的安全风险,可以通过标准化安全定型进行规制。结合《生成式人工智能服务管理办法(征求意见稿)》^[27]第六条内容可以发现,实施标准化安全定型也是目前防范化解人工智能安全风险立法思路。因此,标准化安全定型是防范化解安全风险的有效路径,是人工智能之辅助性定位原则规制安全风险以规范人工智能的研发和应用行为的程序要求。

3. 人工智能失控的救济机制。由于安全是相对的而不是绝对的^[28],即便进行了充分的标准化安全定型,人工智能仍然有超出设计者预设的安全方案而导致输出失控以至于发生侵害的可能,

所以人工智能之辅助性定位原则的程序构造还应包括人工智能失控的救济机制。又因为人工智能的适当法律主体资格是因现实需要而拟制的,其实施行为、享有权利、履行义务以及承担法律责任的方式都不同于其他法律主体,所以人工智能输出失控以至于发生侵害时,需要特殊的救济机制。

首先,造福人类是人工智能研发和应用的初衷,所以人类自身的权利优位于人工智能,当人工智能享有的权利与人类自身的权利发生冲突时,人工智能应让位于人类,例如人工智能有获取数据以进行深度学习的权利,但获取数据的行为必须经过人类授权许可,且人类有权收回授权。其次,人工智能对其创造生成的内容享有知识产权,可以获得一定的收益,当人工智能造成侵害时,可由前期积累的收益予以赔付。此外,人工智能的研发者和应用者也应从其研发和应用活动获得的收益中划拨出一部分为其购买强制性保险,从而避免救济不能的情况发生。最后,当人工智能经过深度学习迭代生成的新算法的危险性已完全不可控时,应及时格式化处理,终止该人工智能的应用。

(三) 人工智能之辅助性定位原则的例外情形

因为具备人类难以完全、绝对地控制的内生自主性特征是赋予人工智能以法律主体资格的前提,所以无内生自主性或内生自主性完全可控的人工智能不具有法律主体资格,在人机关系中处于工具定位。无内生自主性或内生自主性完全可控的人工智能属于弱人工智能(Weak AI)或专用人工智能范畴,目前已被广泛应用,例如2016年判决的“威斯康星州诉卢米斯”^[29]案就应用了分析式人工智能(Analytical AI)对被告进行累犯风险评估,以便为量刑提供科学的参考依据。相较于通用人工智能,专用人工智能的深度学习能力较弱,且不具备内容生成能力和自适应能力,即使可以在如大数据分析等特定应用方面表现出较强的优势,其活动的实施仍然完全依赖于操作者下达的指令,不会根据外部环境的变化修正原定方案,没有独立自主的意志,本质上属于人类用以提高自身能力的工具。而人工智能之辅助性定位不同于工具性定位,工具性定位意味着专用人

工智能的法律属性是客体,只需适用现有法律体系的一般规定即可,不适用人工智能之辅助性定位原则,是该原则的例外情形。

结语

人工智能是建设现代化产业体系、推动高质量发展、服务“以中国式现代化推进中华民族伟大复兴”的新的增长引擎之一。明确人工智能之辅助性定位原则,有助于回应人工智能具备内生自主性特征的现实情况,指导人工智能的研发和应用方向,规制人工智能的安全风险,把握人工智能应用场景由赋能向赋智转变所带来的机遇,引导人工智能的良性发展。通过剖析人工智能之辅助性定位原则的规范构造,有助于完善人工智能相关法律制度的学理解释,为相关领域进一步立法提供政策建议。

参考文献:

- [1] 中共中央政治局召开会议 分析研究当前经济形势和经济工作 中共中央总书记习近平主持会议 [N]. 人民日报, 2023-04-29.
- [2] 中国电子技术标准化研究院. 人工智能标准化白皮书 [EB/OL]. 2021. <http://www.cesi.cn/images/editor/20210721/20210721160350880.pdf>, 2023-03-25.
- [3] Salvador Espitia Ch á vez v. Salud Total E.P.S. Case No. 13001410500420220045901 [EB/OL]. <https://www.diariojudicial.com/public/documentos/000/106/904/000106904.pdf>, 2023-01-31.
- [4] 最高人民法院关于规范和加强人工智能司法应用的意见 [N]. 人民法院报, 2022-12-10.
- [5] Chaoning Zhang, Chenshuang Zhang, Sheng Zheng, et al. A Complete Survey on Generative AI (AIGC): Is ChatGPT from GPT-4 to GPT-5 All You Need? [EB/OL]. <https://doi.org/10.48550/arXiv.2303.11717>, 2023-03-22.
- [6] 成素梅. 智能社会的变革与展望 [J]. 上海交通大学学报(哲学社会科学版), 2020, (4): 9-13.
- [7] OpenAI (2023). GPT-4 Technical Report [EB/OL]. <https://doi.org/10.48550/>

arXiv.2303.08774, 2023-03-16.

[8] Andrew Feenberg. Transforming Technology: A Critical Theory Revisited [M]. New York: Oxford University Press, 2002.

[9] Karl Jaspers. The Origin and Goal of History [M]. London: Routledge & Kegan Paul Co., Ltd., 1953.115

[10] 龚祥瑞, 李克强. 法律工作的计算机化 [J]. 法学杂志, 1983, (3): 16-20.

[11] 赵春玉. 大数据时代数据犯罪的法益保护: 技术悖论、功能回归与体系建构 [J]. 法律科学 (西北政法大学学报), 2023, (1): 95-107.

[12] 田思路. 技术从属性下雇主的算法权力与法律规制 [J]. 法学研究, 2022, (6): 132-150.

[13] Jacques Ellul, The Technological Society[M]. New York: Vintage Books, 1964.

[14] Martin Heidegger. The Question Concerning Technology and Other Essays[M]. New York: Garland Pub, 1977.

[15] Lawrence Byard Solum. Legal Personhood for Artificial Intelligences [J]. North Carolina Law Review, 1992, (4): 1231-1287.

[16] Michael Timothy Bennett. Compression, The Fermi Paradox and Artificial Super Intelligence [EB/OL]. <https://doi.org/10.48550/arXiv.2110.01835>, 2023-03-25.

[17] Susanne Beck. The problem of ascribing legal responsibility in the case of robotics [J]. AI & Society, 2016, (4): 473-481.

[18] 李锡鹤. 民法哲学论稿 (第2版) [M]. 上海: 复旦大学出版社, 2009.

[19] 关于加强科技伦理治理的意见 [N]. 人民日报, 2022-03-21.

[20] 马克思. 资本论 (第一卷) [M]. 中共中央马克思恩格斯列宁斯大林著作编译局. 译, 北京: 人民出版社, 2004: 502.

[21] 黎四奇. 数据科技伦理法律化问题探究 [J]. 中国法学, 2022, (4): 114-134.

[22] 胡越, 罗东阳, 花奎等. 关于深度学习的综述与讨论 [J]. 智能系统学报, 2019, (1): 1-19.

[23] 中华人民共和国科学技术进步法 [J]. 中华人民共和国全国人民代表大会常务委员会公报, 2022, (1): 36-49.

[24] 中华人民共和国行政处罚法 [J]. 中华人民共和国全国人民代表大会常务委员会公报, 2021, (2): 261-270.

[25] 中华人民共和国个人信息保护法 [J]. 中华人民共和国全国人民代表大会常务委员会公报, 2021, (6): 1117-1125.

[26] 中华人民共和国数据安全法 [J]. 中华人民共和国全国人民代表大会常务委员会公报, 2021, (5): 951-956.

[27] 国家互联网信息办公室关于《生成式人工智能服务管理办法 (征求意见稿)》公开征求意见的通知 [EB/OL]. http://www.cac.gov.cn/2023-04/11/c_1682854275475410.htm, 2023-04-11.

[28] 习近平. 在网络安全和信息化工作座谈会上的讲话 [N]. 人民日报, 2016-04-26.

[29] State of Wisconsin v. Eric L. Loomis [Z]. Wisconsin Supreme Court, 2016 WI 68, Case No. 2015AP157-CR.

作者: 吴育珊, 中共广东省委党校法学教研部主任、副教授
杜 昕, 中共广东省委党校法学教研部硕士研究生

责任编辑: 谭博文