

欧盟人工智能知识产权监管升级新形势下企业合规风险防范及应对举措的研究

目 录

前 言	4
一、全球人工智能知识产权监管整体格局	4
（一）主要经济体 AI 知识产权立法体系概况	4
（二）全球监管核心议题存在显著分歧	5
二、欧盟人工智能知识产权完整治理体系	6
（一）《人工智能法案》风险分级制度与通用模型版权合规义务	6
（二）《数字单一市场版权指令》及文本数据挖掘例外规则	7
（三）训练数据强制透明度要求与公开摘要填报规范	8
（四）通用人工智能行业合规指引	8
（五）欧盟新版外观设计法规对 AI 生成设计的规制要求	9
（六）欧盟成员国典型司法裁判实践	9
三、其他主要经济体 AI 知识产权监管路径	11
（一）美国版权合理使用规则与司法判例演进	11
（二）日本宽松数据挖掘例外及产业友好监管导向	12
（三）英国现行版权制度安排与数字复制品立法规划	12
四、2026 年全球主要法域 AI 知识产权监管政策对比分析	13
（一）欧盟法案分步落地，透明度监管全面收紧	13
（二）美国司法诉讼持续发酵，数字人格权立法加速推进	14
（三）日本 AI 产业立法落地，版权补偿机制纳入规划	15
（四）英国维持现有法律框架，谋划数字复制品专项监管	16
（五）我国司法裁判持续完善，行政监管同步强化	17
（六）全球各法域监管核心要点对照表	18
五、境外新规对我国 AI 企业的多维度影响	19
（一）海外市场准入门槛大幅抬高	19
（二）全链条合规运营成本持续攀升	19

(三) 训练数据使用存在多重法律侵权风险	20
(四) 全球知识产权布局面临全新挑战	21
(五) 国际 AI 规则制定话语权有待提升	22
六、我国 AI 企业跨境合规应对工作举措	23
(一) 搭建适配全球业务的合规管理体系	23
(二) 落实训练数据全生命周期合规管控	24
(三) 统筹开展全球知识产权前瞻布局	25
(四) 完善合规技术手段，落实信息披露要求	26
(五) 主动参与国际 AI 知识产权规则制定	27
七、结论与发展趋势展望	29

前 言

人工智能技术的指数级演进正重塑全球知识产权治理的底层逻辑。生成式人工智能（Generative AI）的爆发式发展，使得模型训练数据的版权合法性、AI 生成内容的可保护性、以及跨境技术部署的合规边界，成为各国立法者与监管机关必须直面的核心议题。2024 年至 2026 年间，全球主要经济体密集出台或更新了人工智能知识产权监管框架，标志着这一领域从法律真空迈向规则竞争的新阶段。

在这一波立法浪潮中，欧盟以其特有的布鲁塞尔效应，通过《人工智能法案》（Artificial Intelligence Act, Regulation (EU) 2024/1689）与《数字单一市场版权指令》（Directive on Copyright in the Digital Single Market, Directive (EU) 2019/790）的联动适用，构建了全球最为系统、最为严格的人工智能知识产权治理体系。该体系不仅直接影响在欧盟市场投放 AI 系统的全球企业，更通过 extraterritorial 的适用逻辑与示范效应，深刻影响着中国人工智能企业的出海战略与全球知识产权布局。

本报告以欧盟人工智能知识产权监管要求为核心切入点，系统梳理全球主要法域的最新监管趋势，深入分析其对中国企业的影响，并提出具有可操作性的应对策略。

一、全球人工智能知识产权监管整体格局

（一）主要经济体 AI 知识产权立法体系概况

当前，全球人工智能知识产权监管呈现三轨并行格局：以欧盟为代表的规则先行+行政介入模式，以美国为代表的判例驱动+市场调节模式，以及以日本为代表的技术友好+例外宽松模式。三种模式在训练数据合法性、AI 生成内容保护、权利救济机制等核心议题上存在显著分歧，形成了复杂的跨境合规迷宫。

欧盟于 2024 年 8 月 1 日正式生效的《人工智能法案》，是全球首部综合性人工智能横向立法，采用风险分级方法，将 AI 系统分为不可接受风险、高风险、有限风险和最小风险四个层级，并针对通用人工智能模型（General-Purpose AI Models, GPAI）设置了专门的版权合规义务。该法案与 2019 年《数字单一市场版

权指令》中的文本与数据挖掘（Text and Data Mining, TDM）例外条款形成制度联动，共同构成欧盟 AI 知识产权治理的双柱结构。

美国则延续其普通法传统，未制定联邦层面的 AI 专门版权立法，而是依托《版权法》第 107 条的合理使用原则，通过司法判例逐步界定 AI 训练数据的合法性边界。2025 年，联邦法院在 Thomson Reuters v. Ross Intelligence 案与 Bartz v. Anthropic 案中作出截然相反的判决，凸显了该路径的不确定性。

日本于 2009 年率先在全球引入 TDM 例外，并通过 2018 年、2019 年的修法将其扩展至非享受目的的信息分析领域。2024 年 3 月，日本文化厅发布《AI 与版权通用理解》（General Understanding on AI and Copyright），进一步明确了训练阶段与生成阶段的版权责任分界，确立了 AI 友好国的政策定位。

英国在 2024 年 12 月至 2025 年 2 月就版权与 AI 问题开展大规模公众咨询，2026 年 3 月发布的政府报告决定维持现有法律框架，暂不引入 TDM 例外，但计划废除《1988 年版权、外观设计和专利法》第 9(3) 条关于计算机生成作品的保护条款，并着手规制数字复制品问题。

世界知识产权组织（WIPO）自 2019 年启动 AI 与知识产权政策对话（WIPO Conversation on AI and IP），并在 2024 年发布《Getting the Innovation Ecosystem Ready for AI: An IP Policy Toolkit》，试图为成员国提供政策工具箱，但尚未形成具有约束力的国际统一规则。

（二）全球监管核心议题存在显著分歧

训练数据的合法性构成当前全球 AI 知识产权治理的最大争议焦点。欧盟采取例外和选择退出（opt-out）机制，即在满足特定条件下允许 TDM，但权利人可通过机器可读方式保留权利；美国依赖灵活但不确定的合理使用四要素测试；日本则提供宽泛的非享受目的例外，且权利人无法选择退出。这种制度差异直接导致同一训练行为在不同法域可能面临截然不同的法律评价。

AI 生成内容的可版权性是另一关键分歧点。欧盟坚持人类作者原则，认为缺乏人类智力投入的纯 AI 生成内容不受版权保护；美国版权局 2025 年明确裁定，仅由 AI 生成的材料或人类对表达元素控制不足的材料不享有版权；英国则保留了独

特的计算机生成作品条款，但正计划废除；中国北京互联网法院在 2023 年菲林诉百度案中确立了人类参与是独创性必要来源的司法标准。

透明度义务方面，欧盟通过《人工智能法案》第 53 条建立了强制性的训练数据摘要披露制度，要求 GPAI 提供者公开充分详细的摘要；美国、日本、英国则尚未建立同等强度的强制性透明度框架，主要依赖行业自律或自愿性准则。

二、欧盟人工智能知识产权完整治理体系

（一）《人工智能法案》风险分级制度与通用模型版权合规义务

欧盟《人工智能法案》（Regulation (EU) 2024/1689）于 2024 年 7 月 12 日在《欧盟官方公报》发布，2024 年 8 月 1 日正式生效，采用分阶段实施机制：2025 年 2 月 2 日禁止性 AI 实践条款生效；2025 年 8 月 2 日 GPAI 模型义务与治理结构生效；2026 年 8 月 2 日高风险 AI 系统主要义务生效；2027 年 8 月 2 日产品相关高风险义务生效。该法案适用于所有向欧盟市场投放 AI 系统或将其投入服务的提供者，无论其设立地点是否在欧盟境内，具有显著的域外效力。

该法案将 AI 系统按风险层级分类规制。不可接受风险 AI 系统（如社会评分系统、利用潜意识技术操纵人类行为的系统）被绝对禁止；高风险 AI 系统（用于生物识别、关键基础设施、教育、就业、执法等领域）须履行全面的合规义务，包括风险管理系统、数据治理、技术文档、日志记录、透明度、人工监督、准确性、稳健性和网络安全措施，并在欧盟数据库中注册；有限风险系统（如聊天机器人、深度伪造内容）须满足透明度义务；最小风险系统可自由运营，但鼓励遵守自愿性行为准则。

对于通用人工智能模型，该法案第五章设立了专门义务框架。根据第 53 条，所有 GPAI 模型提供者必须：

第一，制定并持续更新技术文档，包括模型训练、测试过程及评估结果的文档（第 53(1)(a) 条）；

第二，向下游提供者提供必要的信息与文档（第 53(1)(b) 条）；

第三，制定并实施符合欧盟版权法的政策，特别是尊重《数字单一市场版权指令》第 4(3) 条规定的权利保留声明（第 53(1)(c) 条）；

第四，根据 AI 办公室提供的模板，编制并公开训练内容的充分详细的摘要（第 53(1)(d) 条）。

对于具有系统性风险的 GPAI 模型（如大规模语言模型），还须履行模型评估、风险缓解、事件报告和网络安全措施等额外义务（第 55 条）。

在处罚方面，该法案第 99 条规定了严厉的罚款结构：违反禁止性 AI 实践条款，最高可处 3500 万欧元或全球年营业额 7% 的罚款；违反高风险义务，最高 1500 万欧元或全球年营业额 3%；提供不正确信息，最高 750 万欧元或全球年营业额 1%。对于 GPAI 模型违规，罚款上限为 1500 万欧元或全球年营业额 3%。

（二）《数字单一市场版权指令》及文本数据挖掘例外规则

欧盟《数字单一市场版权指令》（Directive (EU) 2019/790，简称 DSM 指令）于 2019 年 3 月 26 日通过，为 AI 训练数据的版权处理提供了核心法律框架。该指令第 2 条将文本与数据挖掘定义为任何旨在分析数字形式的文本和数据以生成信息（包括但不限于模式、趋势和相关性）的自动化分析技术。

DSM 指令设立了两种 TDM 例外：

第一，第 3 条规定的科学研究目的 TDM 例外，适用于科研机构和文化遗产机构，允许其为科学研究目的进行 TDM，且权利人不得选择退出。该例外已被德国《版权法》第 60 条等国内法转化。

第二，第 4 条规定的一般 TDM 例外，适用于包括商业目的在内的所有 TDM 活动，但须满足三个条件：一是作品须为合法获取；二是复制和提取行为仅限于 TDM 所必需的范围；三是权利人未以适当方式保留其权利。该例外是 AI 商业训练的核心法律依据，但其适用以权利人未有效选择退出为前提。

所谓选择退出（opt-out）机制，是指权利人可通过机器可读方式（如 robots.txt 协议、TDM 保留协议、元数据标签等）明确表达其保留作品不被用于 TDM 的意愿。一旦权利人有效行使该权利，AI 开发者若继续使用其作品进行训练，则无法援引 TDM 例外，必须获得单独授权。这一机制在 2024 年德国汉堡地方法院审理的 LAION 案中得到司法确认，法院明确指出，非机器可读的权利保留声明无效。

（三）训练数据强制透明度要求与公开摘要填报规范

2025 年 7 月 24 日，欧盟委员会 AI 办公室发布了《通用人工智能模型训练内容公开摘要模板》（Template for the Public Summary of Training Content for General-Purpose AI Models），该模板是实施《人工智能法案》第 53(1)(d) 条的强制性工具，于 2025 年 8 月 2 日对新投放市场的 GPAI 模型生效，已上市模型须在 2027 年 8 月 2 日前合规。

该模板要求提供三类信息：

第一类一般信息，包括提供者身份、模型名称与版本、投放市场日期、训练数据模态（文本、图像、视频、音频等）、数据规模估计值、语言覆盖范围与人口统计覆盖范围。

第二类数据来源清单，要求区分公开数据集、商业授权数据集、网络爬取内容、用户数据、合成数据及其他来源。对于网络爬取内容，提供者须披露按内容量计算排名前 10% 的域名（中小企业为前 5% 或前 1000 个域名）；对于大型公开数据集，须单独标识并说明其在总训练数据中的占比（超过 3% 须披露）；对于商业授权数据，须确认是否存在授权协议及数据模态。

第三类相关数据处理方面，要求说明为识别和遵守权利保留所采取的措施、避免或移除非法内容的措施（如黑名单、关键词过滤、基于模型的分类器）、以及数据保护合规情况（包括是否进行数据保护影响评估）。

该模板明确允许提供者基于合法商业秘密保护目的对有限信息进行删减，但须向主管机关提供未删减版本，且公开摘要仍须达到充分详细的监管目的。这一设计体现了透明度与商业秘密保护之间的审慎平衡。

（四）通用人工智能行业合规指引

2025 年 7 月 10 日，欧盟委员会发布了《通用人工智能行为准则》（General-Purpose AI Code of Practice），作为自愿性合规工具，帮助 GPAI 提供者证明其符合《人工智能法案》义务。该准则分为透明度、版权保护、安全与保障三个章节，其中版权保护章节具有特别重要的实践意义。

版权保护章节要求签署方：

第一，制定、定期更新并实施单一版权政策，明确组织内部责任分工；

第二，确保网络爬虫仅复制和提取合法访问的内容，不规避有效技术措施，并将欧盟主管机关认定的盗版网站排除在爬取范围之外；

第三，通过机器可读手段（robots.txt 协议、其他适当的技术标准）识别并遵守权利保留声明；

第四，采取相称的保障措施降低生成内容侵犯版权的风险，并在服务条款中反映禁止侵权使用的规定；

第五，设立权利持有人投诉联络点并建立申诉机制。

AI 办公室承诺，在准则生效后的第一年（即至 2026 年 8 月），将密切与签署方合作，若签署方已尽力遵守准则，即使未完全立即实施所有承诺，也不会认定其违反法案义务。这一安全港机制为企业提供了宝贵的合规过渡期。

（五）欧盟新版外观设计法规对 AI 生成设计的规制要求

2024 年，欧盟通过了新的设计法规体系，包括《欧盟设计条例》（Regulation (EU) 2024/2822）和《设计指令》（Directive (EU) 2024/2823），前者于 2025 年 5 月 1 日直接适用（部分行政性条款延至 2026 年 7 月 1 日适用），后者须于 2027 年 12 月 9 日前转化为成员国国内法。这两项立法将设计和产品的定义显著拓宽，并引入从共同体设计到欧盟设计的术语转换。

对于 AI 企业而言，设计法新规的重要性在于：AI 生成的产品外观设计（如通过生成式 AI 设计的工业产品、包装、图形界面等）若满足新颖性和独特个性要求，仍可获得欧盟设计保护。然而，AI 生成设计的作者身份问题在欧盟层面尚未完全解决——与版权领域类似，欧盟设计法隐含要求设计者为自然人，纯 AI 生成的设计在权利归属和有效性挑战方面存在法律不确定性。中国企业在将 AI 生成设计投放欧盟市场时，须审慎评估其可保护性及潜在的权属争议风险。

（六）欧盟成员国典型司法裁判实践

德国汉堡地方法院于 2024 年 9 月 27 日就 LAION 案（Robert Kneschke v. LAION e.V.，案号 310 O 227/23）作出判决，这是德国乃至欧洲首例就 AI 训练数据集制作与版权例外关系作出裁判的案件，具有里程碑意义。

该案中，摄影师 Kneschke 指控非营利组织 LAION 将其照片纳入 LAION 5B 数据集（用于训练 Stable Diffusion 等生成式 AI 模型）构成版权侵权。法院最终驳回原告请求，认定：

第一，为 AI 训练目的制作数据集（包括下载、复制、比对、筛选等预处理行为）属于《德国版权法》第 60 条和第 44b 条项下的文本与数据挖掘例外，该等预处理行为与后续训练行为应作区分评价；

第二，LAION 作为非营利研究组织，其活动符合科学研究目的，且数据集免费向公众开放，满足 DSM 指令第 3 条和德国《版权法》第 60 条的适用条件；

第三，网站上的自然语言使用限制条款（如禁止使用自动化程序访问或下载内容）不构成 DSM 指令第 4(3)条意义上的机器可读权利保留，因此不能排除 TDM 例外的适用；

第四，AI 训练数据集的制作与后续 AI 生成内容的商业利用应作区分，前者本身并不直接产生竞争替代效应，因此不违反版权例外的三步检验法。

2025 年 12 月，汉堡高等地区法院（OLG Hamburg）在二审中维持原判，进一步确认：权利保留必须以机器可读方式表达方可有效，非机器可读的自然语言条款不能阻止 TDM 例外的适用。但同时，法院也指出，若权利人已部署有效的机器可读 opt-out 机制，则 AI 开发者必须遵守，否则将丧失 TDM 例外保护。

法国则采取了更为激进的立法路径。2026 年 4 月 8 日，法国参议院一致通过了由参议员 Laure Darcos 领衔提出的法案，拟修订《法国知识产权法典》，在版权民事纠纷中引入举证责任倒置机制——权利人仅需出示初步证据（如 AI 生成内容与原创作品高度雷同），法院即可推定 AI 提供者存在非法使用行为，由 AI 提供者自证其训练数据源自公共领域或已获合法授权。该机制旨在打破算法黑箱带来的举证壁垒，大幅降低创作者维权门槛。

然而，2026 年 5 月 12 日，法国国民议会十一位团体主席会议未将该法案纳入 6 月跨党派周议程，法案被实际驳回。尽管法国国务委员会已于 2026 年 3 月 19 日出具意见认为该机制符合宪法和欧盟法律，且法国文化界持续呼吁解冻该法案的审议，但截至 2026 年 7 月，该法案在国民议会层面仍处于停滞状态。若该法案最终通过，将成为全球首个在 AI 版权领域适用举证责任倒置的国内立法，对欧盟乃至

全球的 AI 知识产权执法产生深远影响。中国企业应持续跟踪该法案的后续进展，但目前无需基于该法案已生效的前提进行合规准备。

三、其他主要经济体 AI 知识产权监管路径

（一）美国版权合理使用规则与司法判例演进

美国未制定专门针对 AI 训练数据的版权例外，相关争议主要围绕《版权法》第 107 条合理使用原则展开。该原则通过四要素测试评估特定使用行为的合法性：使用的目的与性质（是否具有转化性）、被使用作品的性质、使用部分的数量与实质性、以及对原作品潜在市场的影响。

2025 年，美国法院在 AI 训练数据合法性问题上作出关键判例，呈现出高度分歧：

2025 年 2 月，特拉华州联邦地区法院在 Thomson Reuters v. Ross Intelligence 案中作出 summary judgment，认定 Ross Intelligence 复制 Westlaw 法律数据库的编辑性 headnotes 用于训练其 AI 法律研究工具，不构成合理使用。2026 年 5 月，陪审团进一步作出裁决（jury verdict），支持原告主张，确认侵权行为成立。法院认为，该使用行为具有直接商业替代性，而非转化性使用，且复制了原作品的核心表达性元素，因此四个要素均不利于被告。

2025 年 6 月 23 日，北加州联邦地区法院法官 William Alsup 在 Bartz v. Anthropic 案中作出 summary judgment，认定 Anthropic 使用受版权保护的书籍训练 Claude 模型构成合理使用。法院强调，训练过程类似于人类阅读学习，AI 输出与原始输入在表达上存在 spectacularly different（显著差异），且训练行为本身并不向公众传播原作品，因此具有高度转化性。

这种判例分歧表明，美国 AI 训练数据的合法性高度依赖个案事实，特别是训练输出与原作品的市场替代关系。对于中国企业而言，在美国市场部署 AI 系统时，须审慎评估训练数据的使用方式与输出内容的转化性程度，避免直接复制或竞争性替代原作品的核心表达。

（二）日本宽松数据挖掘例外及产业友好监管导向

日本在全球 AI 知识产权治理中占据独特地位。其《版权法》第 30-4 条于 2009 年率先引入 TDM 例外，2018 年修法后扩展为信息分析例外，允许在不以个人享受或使他人享受作品中表达的思想或情感为目的的情况下，以任何必要方式利用作品进行数据分析。这一非享受目的标准将 AI 训练行为纳入合法范围，且权利人无法选择退出。

2024 年 3 月 15 日，日本文化厅发布《AI 与版权通用理解》（General Understanding on AI and Copyright），进一步澄清了关键边界：

第一，训练阶段的版权例外：若 AI 开发者仅将受版权保护作品用于训练数据库和纯数据分析，且不以生成包含常见创意表达的素材为目的，则第 30-4 条例外适用，无需权利人同意；但若训练行为故意旨在复制原始表达，则例外不适用。

第二，生成阶段的侵权判断：AI 生成内容是否侵权，须同时满足相似性和依赖性两个要件。若 AI 从受版权保护作品学习，依赖性通常可被推定；但若技术保障措施确保训练作品无法被生成，则依赖性可能被否定。

第三，AI 运营者责任：若生成式 AI 模型频繁产生侵权内容，且运营者未采取预防措施，则其被追责的可能性增加。

日本自民党在 2023 年和 2024 年连续发布 AI 政策白皮书，明确将日本定位为最 AI 友好的国家，在版权与 AI 的博弈中明显倾向于支持 AI 产业发展。然而，2024 年《通用理解》也澄清，虽然日本对 AI 训练持友好态度，但如果开发者明知数据非法还故意使用，就不能援引非享受目的例外免责，必须承担侵权责任。这一澄清回应了国际社会对日本机器学习天堂的误读，表明其仍维持基本的版权保护底线。

（三）英国现行版权制度安排与数字复制品立法规划

英国在 AI 知识产权领域的立法进展相对审慎。2024 年 12 月至 2025 年 2 月，英国政府就版权与 AI 问题开展公众咨询，收到约 11,500 份反馈。2026 年 3 月，政府发布报告，决定暂不引入 TDM 版权例外，维持现有法律框架不变，但将继续观察欧盟等法域的发展态势。

报告明确排除了广泛的无 opt-out TDM 例外选项，并放弃了此前偏好的带 opt-out 的 TDM 例外方案。政府表示，将听取创意产业的关切，探索针对特定使用类型的聚焦型例外，并计划废除《1988 年版权、外观设计和专利法》第 9(3) 条关于计算机生成作品的保护条款。若该条款被废除，英国将失去对纯 AI 生成内容的独特版权保护路径，与欧盟人类作者原则趋于一致。

在数字复制品领域，英国政府表现出更强的改革意愿。由于深度伪造技术对表演者声音和肖像权的威胁日益加剧，政府计划在 2026 年夏季就引入专门的人格权或强化现有保护框架开展新一轮咨询。对于中国企业而言，若产品涉及 AI 生成的数字人、虚拟形象或声音克隆功能，须密切关注英国这一立法动向。

在司法层面，2025 年 11 月 4 日，英国高等法院就 Getty Images v. Stability AI 案作出备受期待的判决。法院驳回了 Getty 的次级版权侵权主张，认定 Stable Diffusion 的模型权重不构成英国法下的侵权复制件，因其并未包含或存储受版权作品的任何副本。但 Getty 在商标侵权（水印复制）方面获得有限胜诉。该判决是英国首例全面审查生成式 AI 训练和模型使用的版权案件。

四、2026 年全球主要法域 AI 知识产权监管政策对比分析

进入 2026 年，全球 AI 知识产权治理从立法框架搭建阶段加速迈向规则实施与深化阶段。各主要法域在 2024—2025 年确立的基本立场之上，通过司法裁判、立法推进、行政监管和战略部署，进一步细化和强化了 AI 知识产权规则。以下从五个维度进行系统对照分析。

（一）欧盟法案分步落地，透明度监管全面收紧

2026 年是欧盟《人工智能法案》实施的关键年份。根据法案第 113 条的分阶段适用机制，2026 年 8 月 2 日原定为高风险 AI 系统（Annex III）主要义务和 Article 50 透明度义务的生效日期。然而，2026 年 5 月 7 日，欧洲理事会与欧洲议会就数字一揽子简化方案（Digital Omnibus）达成临时政治协议，将 Annex III 高风险义务推迟至 2027 年 12 月 2 日，Annex I 产品嵌入型高风险义务推迟至 2028 年 8 月 2 日。值得注意的是，Article 50 的透明度义务未被纳入推迟范围，

仍按原计划于 2026 年 8 月 2 日生效，仅对已上市生成式 AI 系统的机器可读标记义务给予延至 2026 年 12 月 2 日的过渡期。

Article 50 的透明度义务具有独立于风险分级的普遍适用性，覆盖三类主体：一是直接与人交互的 AI 系统提供者（如聊天机器人、语音助手），须在首次交互时告知用户其正在与 AI 系统互动；二是生成合成音频、图像、视频或文本内容的 AI 系统提供者，须以机器可读格式标记输出内容，使其可被检测为人工生成；三是深度伪造内容的部署者，须清晰披露内容已被人工生成或操纵。对于中小企业而言，这意味着即使不训练模型，仅使用现成 AI 工具生成营销文案或客服聊天机器人，亦须履行披露义务。

在立法深化层面，2026 年 3 月 10 日，欧洲议会以 460 票赞成、71 票反对通过《版权与生成式人工智能—机遇与挑战》决议，要求欧盟委员会就 AI 训练数据透明度、opt-out 机制有效性、AI 生成内容标识义务等议题提出具体立法提案。该决议虽不具有直接法律效力，但向欧盟委员会和成员国发出了强烈的政治信号，预示着 2026—2027 年可能出台更为细化的 AI 版权实施条例。

2026 年 3 月 18 日，欧洲议会内部市场与消费者保护委员会（IMCO）和公民自由、司法与内政事务委员会（LIBE）就简化方案进行投票，5 月 1 日启动理事会—议会三方谈判。这一系列立法动态表明，欧盟正试图在强化监管与减轻企业负担之间寻求平衡，但 AI 知识产权保护的总体趋严方向未变。

（二）美国司法诉讼持续发酵，数字人格权立法加速推进

2026 年，美国在 AI 知识产权领域呈现出司法裁决与立法突破的双轨并进格局。

在司法层面，2026 年 3 月 2 日，美国最高法院拒绝受理 *Thaler v. Perlmutter* 案的上诉请求，维持哥伦比亚特区巡回上诉法院的裁决，正式确认 AI 不能作为版权法意义上的作者，纯 AI 生成内容在美国不享有版权保护。这一裁决终结了自 2023 年以来关于 AI 作者资格的司法争议，确立了人类作者要求（human authorship requirement）的终局性法律标准。

在训练数据诉讼领域，RIAA 诉 Suno AI 案持续推进。自 2024 年 6 月 RIAA 代表三大唱片公司起诉 Suno 和 Udio 大规模版权侵权以来，该案经历了多轮诉讼攻

防。2026 年 2 月，音乐界发布“Say No to Suno”公开信，进一步施压。该案是否以及在何时驳回合理使用抗辩，仍有待法院进一步裁决。该案涉及 AI 音乐生成公司 Suno 和 Udio 被指控使用海量版权音乐训练模型，原告寻求每件作品最高 15 万美元的法定赔偿。与此同时，2026 年 1 月，由 Concord Music Group 和环球音乐集团牵头的音乐出版商向加州北区联邦法院提交新诉讼，指控 Anthropic 非法下载超过 20,000 首受版权保护歌曲（包括乐谱、歌词及作曲作品）用于训练 Claude 模型，索赔金额超过 30 亿美元。该案是 Concord Music Group 诉 Anthropic 系列诉讼的最新进展。

在立法层面，2026 年 6 月 18 日，美国参议院司法委员会以全票一致通过 S. 4591《培育原创、促进艺术、保障娱乐安全法》（NO FAKES Act），这是联邦数字复制品权利立法首次从委员会层面推进至全院表决阶段。该法案将创设一项新的联邦知识产权权利，赋予每个美国公民（无论公众人物与否）控制其声音和视觉形象在 AI 生成数字复制品中使用的专有权。平台若明知（即已经收到有效通知或明显知道）托管未经授权的数字复制品且未遵守通知—保持下架机制，可能面临每件作品最高 75 万美元的处罚。该法案与田纳西州 2024 年《ELVIS 法案》形成联邦—州两级数字人格权保护体系。

此外，2025 年 9 月 5 日，Anthropic 就使用盗版书籍网站 LibGen 和 Pirate Library Mirror 的数据训练 Claude 模型，与原告达成 15 亿美元的和解协议并提交法院申请初步批准。该和解仅解决过去索赔，未授权未来训练，且 2025 年 8 月 25 日之后提起的索赔明确被排除。然而，随后美国联邦地区法官 William Alsup 对该和解协议提出严厉批评，认为该解决方案存在诸多缺陷，和解的最终批准仍存在不确定性。

（三）日本 AI 产业立法落地，版权补偿机制纳入规划

日本在 2026 年经历了从政策宣示到立法落地的质变，其 AI 知识产权治理呈现促进创新与权利补偿并行的双重转向。

2025 年 5 月 28 日，日本国会通过首部 AI 专门立法《人工智能相关技术研究开发及应用推进法》，2025 年 6 月 4 日公布施行。该法以促进 AI 研发为首要目

标，同时建立透明度和风险评估的基线义务，不设金钱处罚。日本内阁于 2025 年 12 月 23 日批准了依据该法制定的《AI 基本计划》。

2026 年 4 月 7 日，日本内阁批准了《个人信息保护法》修正案，并提交国会审议。该修正案为 AI 训练中的个人数据使用创设了额外的同意豁免情形，特别是经去标识化处理且用于统计分析或 AI 开发目的的个人数据，可在无需重新获取个人同意的情况下使用。该修正案预计在 2026 年通过，新规则最迟于 2028 年全面生效。

2026 年 6 月 12 日，日本首相高市早苗主持知识产权战略总部会议，正式通过《2026 年知识产权战略计划》，将生成式 AI 版权补偿框架和 AI 声音模仿规则纳入立法轨道。该计划明确呼吁建立权利人补偿机制，考虑为知识产权侵权受害者提供民事救济，标志着日本从无条件 AI 友好向有条件权利补偿的政策转型。在声音模仿方面，计划指示继续讨论修订《不正当竞争防止法》以规制演员和配音演员声音的未经授权模仿。

2026 年 4 月 16 日，日本公平贸易委员会（JFTC）发布生成式 AI 市场竞争研究报告，信号加强对主导型模型提供者的反垄断审查；4 月 17 日，法务省宣布成立 AI 生成滥用研究小组，聚焦深度伪造、声音克隆和图像操纵问题，预计 12—18 个月内出台执法指引。

在司法层面，日本法院尚未就生成式 AI 版权作出终局性裁判，但文化厅持续明确：纯 AI 自主生成的内容因缺乏人类创作贡献而不受版权保护；人类通过详细提示词工程、多轮迭代选择和后期编辑作出实质性创造性贡献的，可能满足独创性门槛。2026 年 1 月，日本文化厅进一步明确，简单描述性提示词属于功能性指令而非受保护的原创表达。

（四）英国维持现有法律框架，谋划数字复制品专项监管

英国在 2026 年延续了审慎观望的立法姿态，但在数字复制品领域展现出改革意愿。

2026 年 3 月 18 日，英国政府发布《版权与人工智能报告》，决定现阶段不引入新的版权法或监管监督机制，维持现有法律框架不变。报告明确排除了广泛的无 opt-out TDM 例外，也放弃了带 opt-out 的 TDM 例外方案，转而考虑针对特定使用

类型的聚焦型例外。报告同时提议废除《1988 年版权、外观设计和专利法》第 9(3) 条关于计算机生成作品的保护条款，若该条款被废除，英国将失去对纯 AI 生成内容的独特版权保护路径。

在数字复制品领域，英国政府明确表态将在 2026 年夏季开展专门咨询，探讨引入新的数字复制品权或更广泛的人格权。选项范围从仅针对 AI 生成合成输出的狭窄权利，到覆盖身份一般保护的广泛权利。目前英国法律不保护个人身份，个人无法阻止 AI 模型基于其过往作品训练后生成合成复制品。

在司法层面，英国高等法院在 Getty Images v. Stability AI 案中正在审理使用版权材料训练 AI 模型的合法性问题，上诉程序预计继续进行。值得注意的是，法院驳回了 Getty 的次级版权侵权主张，认定 Stable Diffusion 的模型权重不构成英国法下的侵权复制件，但 Getty 在商标侵权（水印复制）方面获得有限胜诉，却被判承担 Stability 69.4% 的诉讼费用，胜诉代价高昂

（五）我国司法裁判持续完善，行政监管同步强化

中国在 2026 年继续通过司法裁判和行政规制双轨推进 AI 知识产权治理，形成了全球最为活跃的 AI 版权司法实践。

在司法层面，2025 年 3 月 7 日，江苏省常熟市人民法院作出判决，确认 AI 生成图像在满足人类用户通过提示词工程、多轮输出选择和后期编辑作出有意义创造性贡献的条件下，可获得版权保护。该判决将北京互联网法院 2023 年菲林诉百度案和 2024 年李某诉刘某案确立的裁判标准扩展至省级法院层面，进一步巩固了人类参与是独创性必要来源的司法共识。

2026 年 1 月，上海市法院作出相反方向的裁判，裁定简单、描述性的 AI 提示词属于功能性指令，缺乏足够的原创性，不构成受版权保护的表达。该判决与北京、常熟法院形成保护性光谱：人类输入越复杂、越具创造性导向，AI 生成内容的版权保护越强；反之，仅输入 4K、高分辨率、写实风格等技术指令的，不产生版权。

在行政规制层面，国家网信办和版权局持续强化 AI 内容标识要求，平台须对 AI 生成内容采用可见标识和嵌入元数据双重标记。在版权登记实践中，申请人须

披露 AI 辅助情况，登记机构对纯 AI 生成内容不予登记，对 AI 辅助作品仅就人类创作部分予以保护。

在立法层面，中国《生成式人工智能服务管理暂行办法》（2023 年 7 月）第 7 条要求使用具有合法来源的数据，不得侵害他人依法享有的知识产权，但《著作权法》第 24 条规定的合理使用情形尚未明确涵盖 AI 模型训练。这种立法滞后导致中国企业在出海欧盟时面临国内法未授权、国外法要求合规的双重压力。

（六）全球各法域监管核心要点对照表

维度	欧盟	美国	日本	英国	中国
训练数据合法性	TDM 例外+opt-out 机制；2026 年 8 月 Article 50 透明度义务生效	合理使用四要素测试；2026 年 Suno 案驳回合理使用抗辩；Anthropic 15 亿美元和解	非享受目的例外；权利人不可 opt-out；2026 年拟建立补偿框架	无 TDM 例外；维持现有法律框架	合法来源要求；合理使用未明确涵盖训练
AI 生成内容版权	纯 AI 生成内容不保护；人类创作部分可保护	纯 AI 生成内容不可版权（2026 年 3 月最高法院裁决）	纯 AI 自主输出无版权；人类实质性贡献可保护	拟废除计算机生成作品条款；纯 AI 内容或失保护	人类参与是独创性必要来源；提示词复杂度决定保护强度
数字复制品/人格权	GDPR 框架下个人数据保护；AI 法案透明度要求 Article 50 强制披露	NO FAKES Act 推进中（2026 年 6 月委员会全票通过）	拟修订《不正当竞争防止法》规制声音模仿	2026 年夏季就数字复制品权咨询	深度合成管理规定；肖像权、声音权益民法保护
透明度义务	（2026 年 8 月）；训练数据摘要模板（2025 年 7 月）	无联邦强制透明度法；依赖行业自律	AI 法基线透明度义务；2026 年 4 月生效	无专门 AI 透明度法	生成式 AI 服务管理暂行办法；深度合成标识要求
执法机制	AI 办公室+成员国主管机关；最高 3500 万欧元或 7%营业额	联邦法院诉讼为主；州法补充；NO FAKES Act 拟创设联邦权利	PPC+PFTC+法务省多部门；行政指导为主	知识产权局+法院；无专门 AI 监管机构	网信办+版权局+市监局；行政监管与司法并行
2026 年关键节点	Article 50 生效（8 月 2 日）；Digital Omnibus 推迟高风险义务	Thaler 案（3 月）；Thomson Reuters 案陪审团裁决（5 月）	首部 AI 法生效（4 月）；IP 战略计划通过（6 月）	版权与 AI 报告（3 月）；数字复制品咨询（夏季）	上海法院提示词判决（1 月）；常熟判决扩展（2025 年 3 月）

维度	欧盟	美国	日本	英国	中国
		月)；NO FAKES Act 委 员会通过(6 月)			

五、境外新规对我国 AI 企业的多维度影响

(一) 海外市场准入门槛大幅抬高

欧盟《人工智能法案》第 2 条明确规定，该法案适用于所有向欧盟市场投放 AI 系统或将其投入服务的提供者，无论其设立地点是否在欧盟境内。这一域外效力意味着，中国 AI 企业即使未在欧盟设立实体，只要其产品或服务面向欧盟用户（如通过 API、云服务、应用商店等渠道），即须全面遵守该法案义务。

具体而言，中国 GPAI 模型提供者若向欧盟市场提供模型，须自 2025 年 8 月 2 日起履行第 53 条义务，包括制定版权合规政策、公开训练数据摘要、遵守机器可读 opt-out 机制等。2026 年 8 月 2 日后，Article 50 透明度义务正式生效，覆盖聊天机器人、深度伪造内容、合成媒体等所有在欧盟市场运行的 AI 系统，违规者面临最高 3500 万欧元或全球年营业额 7% 的罚款。对于已上市模型，过渡期至 2027 年 8 月 2 日。这一合规时间表对中国企业形成了紧迫的倒计时压力。

更值得关注的是，欧盟通过《人工智能法案》第 53(1)(c) 条和 DSM 指令第 4(3) 条，实际上将全球 AI 训练数据的版权合规标准与欧盟法绑定。即使模型训练行为发生在中国或美国，只要最终产品投放欧盟市场，就必须满足欧盟的 opt-out 合规要求。这种市场接入条件逻辑，使得中国企业在单一法域的合规不足可能演变为全球市场的系统性风险。

(二) 全链条合规运营成本持续攀升

欧盟 AI 知识产权合规要求中国企业建立覆盖数据全生命周期的治理体系，这将显著增加运营成本。具体而言：

第一，训练数据溯源与清洗成本。企业须对海量训练数据进行来源合法性审查，识别并移除可能触发 opt-out 的版权内容，建立 robots.txt 等机器可读协议的自动识别与遵守系统。对于已训练完成的模型，若发现训练数据存在合规瑕疵，

可能面临重新训练或微调的高昂成本。2026 年欧盟 Article 50 的生效进一步要求所有在欧盟市场运行的 AI 系统建立机器可读标记机制，技术改造成本不菲。

第二，透明度披露成本。根据 AI 办公室模板，企业须详细披露训练数据的模态、规模、语言、主要数据来源、爬取域名排名、版权合规措施等信息。2026 年欧洲议会决议进一步要求强化透明度规则，预示着未来可能面临更细化的披露要求。这不仅涉及技术系统的改造，还可能暴露企业的数据策略和竞争优势，在商业秘密保护与合规披露之间形成张力。

第三，法律与合规团队建设成本。企业须在欧盟或熟悉欧盟法的司法管辖区配备专业合规人员，建立与 AI 办公室、成员国主管机关的沟通渠道，及时响应权利持有人的投诉和监管机关的信息请求。2026 年 8 月 Article 50 生效后，即使是中小企业使用现成 AI 工具，亦须配备合规人员处理披露义务。

第四，多法域合规协调成本。由于美国、日本、英国、中国等法域对 AI 训练数据的版权规则存在显著差异，企业可能面临合规冲突困境：同一训练行为在美国可能构成合理使用（但 2026 年 Suno 案显示音乐领域合理使用抗辩被驳回），在欧盟可能因 opt-out 而需授权，在日本可能适用非享受目的例外（但 2026 年拟建立补偿框架），在中国则须确保数据来源合法。协调这些冲突性要求需要复杂的法律架构设计。

（三）训练数据使用存在多重法律侵权风险

中国 AI 企业在训练数据方面面临三重法律风险：

第一，欧盟 TDM 例外的适用不确定性。虽然德国 LAION 案确认 AI 训练数据集制作可纳入 TDM 例外，但该判决针对的是非营利科研组织，且强调机器可读 opt-out 的有效性。对于商业 AI 企业，TDM 例外的适用边界仍不清晰。特别是，若企业从第三方数据中介获取训练数据，而原始权利人已行使 opt-out，则企业可能无法援引 TDM 例外，面临直接侵权风险。2026 年欧洲议会决议要求强化 opt-out 机制有效性，预示着商业 TDM 的适用空间可能进一步收窄。

第二，法国推定侵权机制的潜在冲击。若法国推定利用法案最终通过，将从根本上改变 AI 版权纠纷的举证责任分配。在现行法下，权利人须证明 AI 企业实际使用了其作品；而在新机制下，权利人仅需证明 AI 输出与原作高度相似即可推定侵

权，由 AI 企业证明其训练数据未包含该作品。鉴于生成式 AI 的黑箱特性，企业自证清白的技术难度极高，可能面临大规模的版权索赔。该法案还将适用于其生效之日尚未审结的诉讼程序，具有溯及力。

第三，美国训练数据诉讼的示范效应。2026 年 Suno 案法官驳回合理使用抗辩，Anthropic 15 亿美元和解，以及 New York Times 诉 OpenAI 案的持续进展，表明美国法院对 AI 训练数据的合理使用认定日趋严格。特别是，法官区分了合法获取与盗版获取：合法获取的书籍用于训练可能构成合理使用，但从 LibGen 等盗版网站下载则明确被排除。这一区分对中国企业具有重要警示意义：训练数据来源的合法性正成为全球司法共识，而非仅仅是欧盟的特殊要求。

第四，中国国内法与欧盟法的衔接风险。中国《生成式人工智能服务管理暂行办法》第 7 条要求使用具有合法来源的数据，涉及知识产权的，不得侵害他人依法享有的知识产权。然而，中国《著作权法》第 24 条规定的 13 种合理使用情形并未明确涵盖 AI 模型训练，导致中国企业在满足欧盟 opt-out 要求的同时，可能在中国法下仍面临先授权后使用的合规压力。这种内外夹击的合规困境，对依赖海量互联网数据训练的企业尤为严峻。

（四）全球知识产权布局面临全新挑战

欧盟 AI 知识产权治理框架对中国企业的知识产权布局策略产生深刻影响：

在专利领域，欧盟及欧洲专利公约（EPC）第 52 条将计算机程序本身排除在可专利客体之外，要求发明须具备技术性。AI 算法、模型架构等核心创新在欧盟获得专利保护的门槛较高。德国联邦法院已明确 AI 不能作为发明人，英国最高法院在 Thaler 案中也确认了同一立场。2026 年 3 月美国最高法院拒绝审理 Thaler 案，进一步巩固了全球主要法域对 AI 发明人资格的否定立场。这意味着中国企业的 AI 辅助发明必须由自然人作为发明人，且须清晰界定人类与 AI 的贡献边界，否则可能面临专利无效风险。

在版权领域，欧盟对纯 AI 生成内容持否定保护态度，美国 2026 年 3 月最高法院裁决确认同一立场，日本和中国虽承认人类实质性贡献可获保护，但纯 AI 自主输出仍不受保护。中国企业在欧盟市场投放的 AI 生成内容（如营销文案、设计图案、音乐等）可能无法获得版权保护，导致其易被竞争对手复制而缺乏救济手段。

同时，若 AI 生成内容被认定为与受版权保护作品实质性相似，企业还可能面临输出阶段的侵权指控。2026 年美国 Suno 案和法国推定法案均显示，输出阶段的侵权风险正急剧上升。

在数字复制品领域，2026 年全球出现立法竞赛态势。美国 NO FAKES Act 已获参议院司法委员会全票通过，拟创设联邦数字复制品权利；英国计划 2026 年夏季就数字复制品权开展咨询；日本拟修订《不正当竞争防止法》规制声音模仿。中国企业若产品涉及 AI 数字人、虚拟主播、声音克隆等功能，须同时跟踪美、英、日三国的立法动态，提前布局合规方案。

在商业秘密领域，欧盟《人工智能法案》的透明度要求与商业秘密保护之间存在内在张力。训练数据摘要模板要求披露数据来源、处理方法和版权合规措施，可能迫使企业公开其数据策略和技术细节。虽然模板允许基于商业秘密理由进行有限删减，但充分详细的底线要求仍可能侵蚀企业的竞争优势。2026 年欧洲议会决议要求进一步强化透明度，这一张力可能加剧。

在设计领域，欧盟设计法新规拓宽了保护范围，但 AI 生成设计的作者身份问题尚未解决。中国企业若将 AI 生成设计作为核心资产，在欧盟申请设计保护时可能面临权属挑战和无效宣告风险。

（五）国际 AI 规则制定话语权有待提升

欧盟 AI 知识产权治理框架的布鲁塞尔效应正在重塑全球规则走向。通过设定高标准的市场准入条件，欧盟实际上将其版权价值观（如强保护、opt-out 机制、透明度义务）向全球输出。中国企业在被动接受这些规则的同时，也面临话语权缺失的困境：当前全球 AI 知识产权规则主要由欧美日等发达经济体主导，中国虽然拥有庞大的 AI 产业规模和数据资源，但在国际规则制定中的影响力与产业实力不匹配。

这种话语权失衡可能导致规则错配：欧盟规则更多反映其创意产业（如出版、音乐、影视）的利益诉求，强调权利保留和透明度；美国规则更多反映其科技巨头的利益，强调合理使用和市场自由（但 2026 年 Suno 案显示这一立场正在收紧）；日本规则则明确服务于其 AI 友好国战略（但 2026 年补偿框架的提出显示其正向中

间立场调整)。中国作为兼具强大 AI 产业和内容产业的经济体,需要在全球规则博弈中寻找平衡,但目前的被动合规模式难以有效维护本国产业利益。

六、我国 AI 企业跨境合规应对工作举措

(一) 搭建适配全球业务的合规管理体系

中国 AI 企业必须从根本上转变技术先行、合规补救的被动模式,将合规要求系统性嵌入技术研发、产品设计与商业运营的全过程,建立设计即合规 的组织文化。

第一,设立首席 AI 合规官 或类似职位,统筹全球 AI 合规事务,直接向董事会或 CEO 汇报。该职位应具备法律、技术和跨文化沟通能力的复合背景,负责跟踪欧盟、美国、日本、英国、中国等主要法域的立法动态,评估合规风险,制定应对策略。鉴于 2026 年欧盟 Article 50、美国 NO FAKES Act、日本补偿框架等关键立法节点密集到来,首席 AI 合规官须建立立法雷达系统,对主要法域的法案进展、司法裁判、监管指南进行实时监测和预警。

第二,建立双轨合规委员会,由法务、技术、产品、数据、安全等部门组成,定期审议 AI 产品的合规状态。针对欧盟市场,设立专门的欧盟合规工作组,负责与 AI 办公室、成员国主管机关、行业协会的沟通协调,确保信息对称和响应及时。针对美国市场,建立诉讼预警小组,跟踪 Suno、Anthropic、New York Times 等标志性案件的进展,评估对中国企业商业模式的映射影响。

第三,实施合规前置流程,在产品立项阶段即进行合规评估,将训练数据合法性、版权政策、透明度义务、高风险系统认定、数字复制品合规等作为产品上线的硬性门槛。对于已上市产品,建立定期合规审计机制,至少每六个月审查一次训练数据摘要和版权政策的更新情况。2026 年 8 月欧盟 Article 50 生效前,须完成所有面向欧盟用户的 AI 系统的透明度披露改造。

第四,建立外部顾问网络,聘请熟悉欧盟 AI 法案、DSM 指令、GDPR、美国版权法、日本版权法、英国 CDPA 等法规的当地律师事务所和咨询公司,特别是在布鲁塞尔、柏林、巴黎、华盛顿、东京、伦敦等监管中心设立常驻顾问,确保第一时间获取监管指南和执法动态。

（二）落实训练数据全生命周期合规管控

训练数据合规是 AI 知识产权治理的核心战场。中国企业须建立覆盖数据从收集到销毁全生命周期的闭环管理机制：

第一，数据来源分类与合法性审查。将训练数据来源分为公共领域数据、商业授权数据、用户生成数据、合成数据、网络爬取数据等类别，分别建立合法性审查标准。对于网络爬取数据，部署自动化的 robots.txt 和 TDM 保留协议识别系统，确保在爬取前即判断目标网站是否行使了 opt-out 权利。对于从第三方数据中介获取的数据，建立供应商尽职调查制度，要求提供数据来源合法性证明、权利清理文件和不侵权承诺。2026 年美国 Suno 案和 Anthropic 案表明，从盗版网站获取数据即使在美国也可能被排除在合理使用之外，因此须建立盗版数据源黑名单，将 LibGen、Pirate Library Mirror 等已知盗版网站永久排除在爬取范围之外。

第二，数据清洗与权利保留。建立 opt-out 数据库，持续收集和更新全球主要权利人、出版社、图片社、音乐厂牌等的机器可读 opt-out 信号，在数据预处理阶段自动过滤或标记受保留权利约束的内容。对于无法确定版权状态的数据，采取保守策略，宁可剔除也不冒险使用。建立数据清洗日志，记录每一步处理操作，以备监管审查和诉讼抗辩。2026 年法国推定法案若通过，企业须保存完整的训练数据溯源记录，以在诉讼中自证未使用特定作品。

第三，合成数据与内部数据替代。加大合成数据的研发投入，利用自研模型生成训练数据，从根本上规避第三方版权风险。同时，充分利用企业自有数据资产（如用户授权数据、内部文档、产品数据等），通过数据增强、迁移学习等技术降低对第三方公开数据的依赖。2026 年日本 APPI 修正案为去标识化个人数据的 AI 训练使用创设了合法路径，中国企业可考虑在日本设立数据中心，利用该例外进行模型训练。

第四，数据跨境传输合规。若训练数据涉及个人信息，须遵守欧盟 GDPR 和中国《个人信息保护法》的双重约束，通过标准合同条款（SCCs）、加密传输、匿名化处理等措施确保跨境传输合规。建立数据本地化策略，对于敏感训练数据，考虑在欧盟境内设立数据中心，避免跨境传输的合规风险。2026 年日本 APPI 修正案虽

放松了去标识化数据的跨境限制，但仍要求传输风险评估和合同保障，不可掉以轻心。

第五，数据保留与销毁。根据处理目的和法律要求为不同类别数据设定严格的保留期限，在期限届满或目的达成后安全、彻底地删除或匿名化相关数据。特别对于可能触发版权争议的数据，建立争议数据快速销毁机制，一旦收到权利人投诉或监管调查通知，立即启动数据隔离和销毁程序。

（三）统筹开展全球知识产权前瞻布局

面对欧盟及全球 AI 知识产权规则的不确定性，中国企业须采取更加积极、前瞻的知识产权布局策略：

第一，专利布局策略。针对欧盟市场，调整专利申请策略，避免将 AI 算法本身作为专利客体，而是将其与具体技术领域（如医疗诊断、工业控制、自动驾驶等）结合，强调其技术效果和技术贡献，以满足 EPC 第 52 条的要求。充分利用《专利合作条约》（PCT）进行国际申请，在优先权期限内评估各目标市场的专利可授权性，动态调整申请策略。对于 AI 辅助发明，确保发明人身份文件清晰记载自然人发明人的实质性贡献，避免因发明人资格问题导致专利无效。

第二，版权与商业秘密双轨保护。对于 AI 生成内容，在欧盟和美国市场不依赖版权保护，而是转向商业秘密保护（如通过技术措施限制访问、签订保密协议、分段披露等），同时在中国等承认 AI 生成内容版权的法域积极登记版权。对于核心算法、模型架构、训练方法等，优先采用商业秘密保护，避免专利申请导致的公开披露。建立商业秘密分级管理制度，对核心资产实施最严格的访问控制和加密保护。2026 年上海法院判决明确简单提示词不受版权保护，提示企业须在创作过程中详细记录人类的创造性贡献（如提示词迭代日志、参数调整记录、后期编辑痕迹等），以备权属争议时举证。

第三，商标与品牌保护。通过马德里体系进行国际商标注册，覆盖欧盟、美国、日本、英国等主要市场，避免市场未动、商标先失。对于 AI 产品的品牌名称、LOGO、界面设计等，提前进行商标检索和注册，防范抢注风险。对于 AI 生成的品牌内容（如广告文案、包装设计），在欧盟和美国市场须注意其可能无法获得版权保护，应通过商标注册和外观设计申请构建替代保护体系。

第四，设计权布局。针对欧盟设计法新规，及时将 AI 生成设计在欧盟知识产权局（EUIPO）注册为欧盟设计，确保在拓宽后的设计定义下获得保护。同时，保留完整的设计创作过程文档，证明人类设计师在 AI 生成过程中的实质性指导和选择，以应对可能的权属挑战。

第五，数字复制品合规前置。鉴于 2026 年美国 NO FAKES Act、英国数字复制品咨询、日本声音模仿立法均指向数字人格权保护，中国企业若产品涉及 AI 数字人、虚拟主播、声音克隆、深度伪造等功能，须立即开展合规评估：建立数字人格权数据库，记录所有使用真实人物声音、形象、肖像的授权情况；在产品的设计阶段嵌入授权验证模块，确保未获授权的真实人物不会被用于生成数字复制品；建立快速下架机制，响应权利人的删除请求。

第六，开源合规管理。建立开源软件审批流程，对训练框架、模型库、数据处理工具等使用的开源组件进行许可扫描和溯源审查，避免 GPL、AGPL 等传染性许可对商业产品的污染。对于基于开源模型微调或改进的产品，清晰界定原模型与改进部分的知识产权边界，避免权属纠纷。

（四）完善合规技术手段，落实信息披露要求

欧盟《人工智能法案》的透明度要求迫使企业从黑箱运作转向可解释合规，这既是挑战，也是建立信任的机会。

第一，训练数据摘要系统。按照 AI 办公室模板要求，建立自动化的训练数据摘要生成系统，确保能够实时或准实时地提供模型/提供者元数据、数据来源清单、处理措施等信息。对于多模态模型，分别统计文本、图像、视频、音频等模态的数据规模和来源分布。建立版本控制机制，在模型更新或再训练时自动触发摘要更新，确保披露信息的时效性。2026 年欧洲议会决议要求进一步强化透明度规则，企业应建立透明度升级预案，确保能够快速响应更严格的披露要求。

第二，版权合规技术措施。开发或采购能够自动识别和遵守机器可读 opt-out 信号的技术工具，包括 robots.txt 解析器、TDM 保留协议识别器、元数据标签扫描器等。在模型训练阶段，部署版权防火墙，自动拦截和过滤受权利保留约束的内容。在模型推理阶段，部署输出内容过滤系统，通过关键词过滤、相似度检测、模

型分类器等技术手段，降低生成内容侵犯版权的风险。2026 年法国推定法案若通过，输出过滤系统将成为企业自证清白的关键技术证据。

第三，可解释 AI 与审计追踪。建立模型决策的可解释性机制，记录训练数据的使用路径、模型参数的调整历史、生成内容的来源依据等，确保在监管审查或诉讼中能够提供完整的审计追踪。对于高风险 AI 系统，按照《人工智能法案》要求建立日志记录系统，保存输入数据、输出结果、人工监督记录等，保存期限不少于系统生命周期。2026 年在 New York Times 诉 OpenAI 案的证据开示程序中，据报道涉及大量用户对话记录的调取，但具体数量以法院公开记录为准。这一案例表明，完整的日志记录已成为诉讼中的关键证据类型。

第四，用户告知与标识。对于 AI 生成内容，按照欧盟《人工智能法案》Article 50 和中国《互联网信息服务深度合成管理规定》等要求，建立清晰、显眼的 AI 生成标识机制，确保用户能够明确知晓其正在与 AI 交互或接收 AI 生成内容。2026 年 8 月 2 日后，欧盟要求聊天机器人在首次交互时即披露其 AI 身份，深度伪造内容须清晰标注人工生成，机器可读标记须嵌入合成音频、图像、视频或文本的元数据中。企业须在产品界面、用户旅程的适当节点实施披露，而非将告知信息 buried 在法律声明中。

第五，投诉响应机制。设立专门的权利持有人投诉联络点，建立标准化的投诉受理、调查、处理和反馈流程。对于版权侵权投诉，承诺在合理期限内（如 14 个工作日）作出初步回应，在调查期间可采取临时性措施（如暂停相关模型功能、移除涉嫌侵权内容等），以展示合规诚意并降低监管处罚风险。2026 年美国 NO FAKES Act 拟引入通知—保持下架机制，要求平台在收到有效通知后不仅删除内容，还须防止相同内容重新出现，这意味着企业须建立内容指纹识别和监控系统。

（五）主动参与国际 AI 知识产权规则制定

面对全球 AI 知识产权规则的竞争与博弈，中国企业不能仅做被动的合规接受者，而应积极参与国际规则制定，提升话语权和影响力。

第一，深度参与 WIPO 对话。利用 WIPO AI 与知识产权政策对话平台，积极提交立场文件、参与技术会议、加入专家工作组，就 AI 发明人资格、AI 生成内容保护、TDM 例外国际协调等议题表达中国产业界诉求。推动 WIPO 制定 AI 知识产权国

际指南或示范法，为各国立法提供参考，减少规则碎片化。2026 年全球规则碎片化趋势加剧，WIPO 的协调作用愈发重要。

第二，借助一带一路与 RCEP 框架。在一带一路沿线国家和《区域全面经济伙伴关系协定》（RCEP）成员国中，推动建立兼顾 AI 产业发展与版权保护的平衡性规则。利用中国与东盟、中东、非洲等地区的数字合作协议，输出中国 AI 治理经验，构建与欧盟、美国并行的第三极规则影响力。2026 年日本通过 RCEP 框架内的双边 IP 行动计划与欧盟协调，中国企业可借鉴这一模式，在 RCEP 框架内推动 AI 知识产权规则的互认与协调。

第三，加强中欧监管对话。通过中欧数字高层对话、中欧 AI 工作组等机制，就《人工智能法案》的实施细节、训练数据摘要模板的标准解释、opt-out 技术标准互认、Article 50 透明度义务的实施指南等议题与欧盟委员会、AI 办公室开展深入沟通。争取在欧盟次级立法（如实施指南、标准）制定过程中反映中国企业的合理关切，避免规则过度偏向欧盟本土产业。2026 年欧洲议会决议要求立即立法行动，中国企业应抓住这一窗口期，通过行业协会和商会组织集体发声。

第四，推动国内立法完善。积极向国家网信办、国家版权局、国家知识产权局等部门建言献策，推动完善《著作权法》中关于 AI 训练的合理使用条款，明确 TDM 例外的适用条件和边界，为中国企业提供与欧盟、日本对等的国内法基础。同时，推动建立 AI 生成内容的版权登记和保护制度，在保护人类作者权益的前提下，为 AI 辅助创作提供适度的法律保障。2026 年中国法院已就 AI 生成内容保护形成光谱式裁判标准，立法机关应及时回应，将司法经验上升为法律规则。

第五，构建产业联盟与标准体系。联合国内 AI 企业、内容企业、科研机构，成立中国 AI 知识产权联盟，共同制定行业自律准则、技术标准（如机器可读 opt-out 的中文标准、训练数据溯源格式、AI 生成内容水印协议等）、最佳实践指南。通过团体标准、行业标准乃至国家标准的制定，提升中国在全球 AI 知识产权治理中的规则贡献度。2026 年欧盟 AI 办公室已就训练数据摘要模板建立标准，中国应同步建立中文标准，确保中国企业的合规文档能够与国际标准互认。

第六，跟踪并应对数字复制品全球立法。建立数字人格权立法跟踪系统，实时监测美国 NO FAKES Act、英国数字复制品咨询、日本声音模仿立法、欧盟 GDPR 个

人数据保护等法域的动态，评估对中国企业产品的影响。在产品的设计阶段即嵌入全球数字人格权合规模块，确保产品能够根据不同法域的要求自动调整数字复制品的生成和使用规则。

七、结论与发展趋势展望

全球人工智能知识产权治理正处于从各自为政向规则竞争的关键转折期。2026年以来，这一转折明显加速：欧盟从法案框架搭建进入实质性实施阶段，Article 50 透明度义务的生效标志着全球首个强制性 AI 内容标识制度落地；美国通过最高法院裁决确立了纯 AI 生成内容的不可版权性，同时通过 NO FAKES Act 推进数字复制品联邦立法；日本从无条件 AI 友好转向有条件权利补偿，首部 AI 专门立法生效并启动版权补偿框架建设；英国维持版权现状但积极布局数字复制品权；中国则通过活跃的司法实践形成了独特的光谱式保护标准。

欧盟以其《人工智能法案》和《数字单一市场版权指令》为核心，构建了全球最为严格、最为系统的 AI 知识产权监管框架，通过域外效力、分阶段实施、高额处罚等机制，深刻影响着全球 AI 企业的合规行为和市场策略。德国 LAION 案、法国推定法案等司法和立法动态，进一步加剧了合规环境的不确定性。美国 2026 年 Suno 案和 Anthropic 15 亿美元和解表明，训练数据来源合法性已成为全球司法共识，而非仅仅是欧盟的特殊要求。

对中国企业而言，欧盟 AI 知识产权治理框架既是市场准入的高门槛，也是转型升级的催化剂。短期内，合规成本上升、训练数据受限、法律风险增加将构成显著挑战；长期来看，建立全球领先的合规能力、参与国际规则制定、构建自主知识产权体系，将成为企业核心竞争力的重要组成部分。2026 年全球立法竞赛的加剧，特别是数字复制品权的兴起，要求中国企业必须具备同时跟踪多法域立法动态、快速调整产品合规策略的能力。

未来三至五年，全球 AI 知识产权治理将呈现以下趋势：

第一，规则趋同与分化并存。欧盟规则将通过布鲁塞尔效应持续向全球输出，但美国、日本等法域不会完全照搬，而是在核心原则（如透明度、opt-out）上趋同，在具体标准（如 TDM 例外范围、举证责任、数字复制品权）上保持分化。中国

企业须建立模块化合规体系，在统一治理框架下灵活适配不同法域的具体要求。2026 年日本补偿框架的提出和英国数字复制品权的酝酿，预示着中间道路选项正在增多。

第二，技术标准与法律规则深度融合。机器可读 opt-out、训练数据溯源、AI 生成内容标识、数字复制品检测等要求，将推动技术标准与法律规则的一体化发展。掌握相关技术标准制定权的企业，将在规则竞争中占据先机。中国企业应加大在 opt-out 技术标准、数据溯源协议、AI 水印技术、数字人格权检测等领域的研发投入和标准布局。2026 年欧盟 AI 办公室已就训练数据摘要和透明度标记建立标准模板，中国应同步建立中文标准体系。

第三，从事后救济到事前预防。法国推定法案、美国 NO FAKES Act 的通知—保持下架机制、日本补偿框架等立法动向表明，AI 知识产权治理正从依赖权利人事后维权，转向要求企业事前自证合规。这意味着企业须建立更加完善的内部合规体系和文档留存机制，将合规证明能力作为核心管理能力。2026 年 New York Times 诉 OpenAI 案中法官要求提供 2000 万条对话日志，进一步印证了文档即合规的新范式。

第四，数据主权与知识产权交织。随着各国数据主权意识的增强，训练数据的跨境流动将面临更多限制，数据本地化要求可能与知识产权规则产生复杂互动。日本 2026 年 APPI 修正案、欧盟 GDPR 跨境传输规则、中国《个人信息保护法》出境安全评估等制度，正在构建一个数据主权—知识产权双重治理框架。中国企业须在全球数据治理与知识产权治理的双重框架下，重新设计数据战略和合规架构。

第五，数字复制品权成为新战场。2026 年，美国、英国、日本、欧盟均在数字复制品或人格权领域加速立法，标志着 AI 知识产权治理从内容保护向身份保护扩展。AI 数字人、虚拟主播、声音克隆、深度伪造等产品形态，将成为未来合规争议的焦点。中国企业须将数字人格权合规纳入产品设计的核心考量，建立全球数字复制品合规体系。

总之，人工智能知识产权合规已从可选项变为必选项，从成本中心变为战略资产。2026 年全球立法和司法的密集动态，进一步压缩了企业的合规窗口期。中国企业唯有以欧盟治理框架为镜鉴，同时跟踪美、日、英、中等法域的最新动向，立

足全球视野，构建主动、系统、前瞻的合规管理体系，方能在全全球 AI 产业的激烈竞争中行稳致远，实现技术创新与规则引领的协同发展。