

相反判决中的真谛(一): 人工智能训练数据的非侵权要求

作者: 杨迅

在人工智能飞速发展的当下, 训练数据的合法性问题成为了各界关注焦点。训练数据是否需要获得著作权人授权, 在国际上一直存在争议, 美国 Thomson Reuters 诉 Ross Intelligence 案件堪称里程碑。而在中国, 虽暂无针对训练数据侵犯著作权的生效判决, 但出现了两起涉及人格权的声音数据训练人工智能的案件, 且判决结果截然相反, 引发了诸多思考。

一. 训练数据非侵权性问题的提出

在当今数字化时代, 人工智能技术正以前所未有的速度发展, 深刻地改变着我们的生活方式、工作模式以及社会的方方面面。在这其中, 机器学习作为人工智能的核心领域之一, 发挥着至关重要的作用。而机器学习的关键要素之一, 就是海量的数据。无论是图像识别技术, 能够精准地辨别照片中的物体、场景和人物; 还是语音识别系统, 能够将人类的语音准确地转换为文字; 抑或是自然语言处理, 让计算机能够理解、生成并有效地处理人类语言, 这些人工智能的应用和进步都离不开大量的数据来进行模型训练。只有通过不断地用海量数据对模型进行训练, 人工智能系统才能逐步具备相应的智能和能力, 从而为各种复杂的应用场景提供高效、精准的服务, 满足人们在不同领域的需求。

随着人工智能技术的广泛应用和快速发展, 《生成式人工智能服务管理暂行办法》应运而生, 其第七条明确提出了训练数据合法性的原则要求, 为人工智能的健康发展划定了基本的底线和框架。然而, 尽管有了这一原则性的规定, 但在具体的实践操作中, 对于如何评判训练数据是否构成侵权, 目前法律层面尚未有更加明确、细致且具有可操作性的规定。这就导致在实际的数据使用过程中, 存在一定的模糊地带, 使得企业和开发者在利用数据进行模型训练时面临诸多不确定性和潜在的法律风险。

.....
如您需要了解我们的出版物,
请联系:

Publication@llinkslaw.com

在我国,《著作权法》采用列举制的方式对“合理使用”的情形进行了规定,这一规定在传统的著作权保护领域发挥了重要作用,为个人学习研究、评论等合理使用行为提供了一定的法律依据。然而,随着人工智能技术的兴起,模型训练目的的数据使用需求与日俱增,这种使用方式往往需要对大量的受版权保护的作品进行使用,而这种大规模、系统性的使用与传统的合理使用场景存在着显著的差异。现有的合理使用规定难以直接适用于模型训练目的的数据使用,这就使得在判定其合法性时面临着诸多困难和挑战。

同样地,我国《民法典》中对于肖像权和类似肖像的人格权的规定,也面临着类似的困境。在涉及声音数据训练等新兴的人工智能应用场景下,如何在保护个人声音权益等人格权的同时,又不阻碍人工智能技术的正常发展,成为了一个亟待解决的问题。

若对训练数据施加过于严格的要求,可能会对人工智能技术的发展造成严重的阻碍。过于严苛的数据获取限制,将使企业在数据收集和使用上面临巨大的困难,不仅增加了研发成本和时间成本,还可能导致企业在市场竞争中处于劣势地位,进而影响整个行业的发展速度和竞争力。因此,如何在保护各方权益的基础上,合理地规范人工智能训练数据的使用,为技术发展创造良好的法律环境,是一个需要社会各界共同探讨和解决的重要课题。

二. 有关训练数据的司法案例

当下,人工智能的训练数据合法性问题在全球范围内都备受关注,不同国家和地区陆续出现了一系列相关案件。在美国,有 Digital News Media 诉 OpenAI,《纽约时报》诉 OpenAI 和微软,而 Thomson Reuters 诉 Ross Intelligence 案件则是首例对 AI 训练数据版权争议作出明确裁定的案件,具有重要的里程碑意义,其对 AI 模型训练中使用受版权保护材料的边界进行了初步界定,引发了行业对 AI 训练数据合规性的深度思考。

而在中国,虽然暂无直接针对训练数据侵犯著作权的生效判决,但出现了两起较为典型的声音数据训练人工智能构成人格权纠纷的案件,即殷某桢诉北京某智能科技公司等人格权纠纷案以及某配音员诉成都声入人心网络科技有限公司声音保护纠纷案,这两起案件的判决结果截然相反,为我国在人工智能训练数据的法律规制方面提供了有益的借鉴和参考。

(一) 美国 Thomson Reuters 诉 Ross Intelligence 案件

Thomson Reuters 运营的 Westlaw 数据库收录了几乎美国所有法院判决,并由编辑团队为每个判决要点撰写主题提要,配以相应的关键编码系统。Ross Intelligence 作为一家法律人工智能初创企业,欲开发一款 AI 法律检索工具,它与第三方咨询公司 LegalEase Solutions 合作,使用 LegalEase 提供的大量法律问题摘要和相关案例集合——“Bulk Memos”供 AI 训练使用,而這些“Bulk Memos”在内容和结构上与原 Westlaw 数据库的内容高度相似。Thomson Reuters 发现后起诉 Ross 侵犯其主题提要 and 关键编码系统的著作权。

2025 年 2 月 11 日, 美国特拉华州联邦地区法院对 Thomson Reuters 诉 Ross Intelligence 案作出部分简易判决。法院在判断合理使用时综合考虑了四要素: 第一要素是使用的目的和性质, 包括使用是否具有商业性及是否具有变革性。法官参考美国联邦最高法院于 2023 年在 *Andy Warhol Foundation for the Visual Arts 诉 Goldsmith* 一案中对合理使用要素一的分析标准, 认为 Ross 对案头批注的使用目的与 Thomson Reuters 并无本质不同, 是专门用于开发与 Westlaw 直接竞争的同类法律检索工具, 缺乏“转换性目的”, 且具有商业性质, 因此要素一更有利于 Thomson Reuters。第二要素是受保护作品的性质, 法院认为 Westlaw 的主题提要虽满足最低限度的创造性标准, 但独创性较低; 并且关键编码系统本质上是一个事实汇编, 创造性有限, 因此该要素有利于 Ross。但该要素在本案中权重相对较轻。第三要素是使用的数量和实质性, Ross 通过 LegalEase 实际复制了数千条主题提要用于训练, 属于大规模全文复制, 但由于 Ross 的最终产品给用户看的并不是 Westlaw 的提要, 而是相关联的司法意见全文, 且这些法院判决本身不受著作权保护, 所以第三要素总体上对 Ross 有利。第四要素是对原作品市场或价值的影响, 法院认为 Ross 使用原告内容训练 AI 的主要目的就是希望其法律检索工具能替代 Westlaw, 对原告的市场利益构成了潜在威胁, 该要素明显支持原告。最终, 法院综合认定 Ross 未经许可使用 Westlaw 编辑内容训练 AI 不属于合理使用, 构成对 Thomson Reuters 著作权的侵犯。

(二) 中国两起声音数据训练案件

(2023)京 0491 民初 12142 号案件

在该案中, 殷某桢系配音演员, 发现他人利用其配音制作的作品在多个知名 APP 广泛流传。经溯源, 作品中的声音来自被告一北京某智能科技公司运营平台的文本转语音产品。法院查明, 殷某桢曾受被告二某文化公司委托录制录音制品, 该公司将录音数据提供给被告三某软件公司, 后者未经殷某桢同意, 采用人工智能技术处理其录音, 生成文本转语音产品并出售, 被告一北京某智能科技公司直接调取并在其平台使用该产品。

北京互联网法院依据《中华人民共和国民法典》第一千零二十三条第二款规定: “对自然人声音的保护, 参照适用肖像权保护的有关规定。” 认定殷某桢声音权益的保护范围包括经人工智能处理的声音。本案中, 经当庭勘验, 案涉 AI 声音与殷某桢的音色、语调、发音风格等具有高度一致性, 通过该声音能够引起一般人产生与殷某桢有关的思想或感情活动, 能够将该声音联系到殷某桢本人, 进而识别出殷某桢的主体身份, 具备可识别性。而被告二某文化公司对录音制品虽享有著作权等权利, 但不包括授权他人对殷某桢声音进行 AI 化使用的权利, 且《数据授权书》并非殷某桢本人签署, 故被告二与被告三某软件公司签订协议, 在未经殷某桢本人知情同意的情况下, 授权被告三某软件公司 AI 化使用原告声音的行为无合法权利来源, 构成对殷某桢声音权益的侵犯。

(2024)川 7101 民初 8969 号案件

在该案中，原告是一名资深配音员，其配音网名有一定知名度，曾为多个知名节目和广告等进行配音。2023 年 3 月份，原告在网络上发现被告成都声入人心网络科技有限公司开发的应用程序“超级配音 APP”采用人工智能技术，未经原告本人允许克隆了原告的声音，以“孟师”的角色提供文字转语音服务获利，该音频已被众多第三方配音公司所广泛传播使用。

成都铁路运输第一法院认为，认定声音权益是否受侵害，首先需判断声音是否具有可识别性。本案中，原告提交的证据虽能证明其通过配音能产生经济收益，但原告声音本身多变，其上传至网络的录音作品也反映出根据对应的作品或是文案不同，发声方式、语调等存在差异性，与原告本人身份之间的联系不够密切，通过其声音识别其本人存在一定难度。针对被诉侵权的软件中的“孟师”的声音，法院认为在未见到附有原告肖像的视频情形下，通过聆听“孟师”的多个声音文件并不能使一般大众将该声音与原告产生联系，也不能达到通过该声音即可识别出声音的主体为原告的程度。同时，在案证据无法证明被告存在侵害原告声音权益的行为。原告虽主张被告的配音软件通过使用 AI 技术手段合成其声音而生成了“孟师”的声音，但从人耳聆听还是录音比对来看，均不能认定两声音之间存在高度相似性，原告亦未提交基础证据证明被告存在侵权行为，故法院驳回了原告的全部诉讼请求。

（三）两起声音案件差异中的规律

这两起案件虽然结果不同，对声音相似性的判断也不尽一致，但是法院在审理时都将重点聚焦于人工智能训练后的成果是否能够识别到具体个人，即原始声音数据提供者，从而与其相竞争。在殷某桢案中，由于案涉 AI 声音与殷某桢本人声音高度一致，具备可识别性，容易引起混淆，进而构成侵权；而在成都案中，被诉侵权声音与原告声音不具备可识别性，无法建立直接联系，法院认定不构成侵权。这表明，法院在审理此类案件时，关键在于判断训练成果是否会与原权利人形成竞争关系，导致原权利人的利益受损。

上述判断与美国 Thomson Reuters 诉 Ross Intelligence 案件异曲同工。可见，在人工智能训练数据的侵权判断中，不仅要关注数据来源的合法性，更要重点考量训练成果是否与原数据来源形成替代或竞争关系。

三. 两起声音训练数据案件的启示

从这两起声音案件以及美国 Thomson Reuters 诉 Ross Intelligence 案件来看，尽管诉因都是围绕训练数据是否需要获得原权利人授权，但最终的判决重心都集中于训练成果是否与原权利人混淆，构成替代和竞争。这一共通点为人工智能训练数据的合法性判断提供了一种新的视角和思路，即不仅要关注数据的来源是否合规，更要考量训练成果的应用是否会侵犯原权利人的合法权益，与之形成不正当竞争关系。

这两起声音案件为更广阔的数据获取用于模型训练提供了重要指引。对于企业或其他主体而言，在利用他人数据进行人工智能模型训练时，应重点关注训练模型的使用方式，避免与数据来源形成直接竞争。在数据收集阶段，应尽量获取合法授权，确保数据来源的正当性；在模型训练和应用过程中，

要充分考虑训练成果的特点和用途，通过技术手段和合理设计，降低与原始数据提供者产生混淆和竞争的可能性，从而在一定程度上降低侵权风险，促进人工智能技术的健康、合法发展。

如您希望就相关问题进一步交流，请联系：



杨 迅
+86 21 3135 8799
xun.yang@llinkslaw.com

如您希望就其他问题进一步交流或有其他业务咨询需求，请随时与我们联系: master@llinkslaw.com

上海

上海市银城中路 68 号
时代金融中心 19 楼
T: +86 21 3135 8666
F: +86 21 3135 8600

北京

北京市朝阳区光华东里 8 号
中海广场中楼 30 层
T: +86 10 5081 3888
F: +86 10 5081 3866

深圳

深圳市南山区科苑南路 2666 号
中国华润大厦 18 楼
T: +86 755 3391 7666
F: +86 755 3391 7668

香港

香港中环遮打道 18 号
历山大厦 32 楼 3201 室
T: +852 2592 1978
F: +852 2868 0883

伦敦

1/F, 3 More London Riverside
London SE1 2RE
T: +44 (0)20 3283 4337
D: +44 (0)20 3283 4323



www.llinkslaw.com



Wechat: LlinksLaw

本土化资源 国际化视野

免责声明：

本出版物仅供一般性参考，并无意提供任何法律或其他建议。我们明示不对任何依赖本出版物的任何内容而采取或不采取行动所导致的后果承担责任。我们保留所有对本出版物的权利。

© 通力律师事务所 2025