



新华社北京2月12日电（新华社“新华视点”记者 邵鲁文 温竞华）亲耳听到的就是真的吗？未必。网络平台上，AI声音随处可见。

从“张文宏医生”推销蛋白棒视频“走红”网络，后被本人“打假”，到多位配音演员称声音被AI“偷走”，公开维权……“新华视点”记者调查发现，随着人工智能技术和语音大模型应用的发展，AI合成声音App大量出现，最快只需十几秒便可“克隆”出来。与此同时，AI声音滥用现象愈发突出，不时引发争议。

你的声音被谁“偷”走了？

——AI声音滥用现象调查

AI声音滥用不时发生

记者在某短视频平台以“AI克隆声音”为关键词检索发现，明星翻唱、新闻播报、吐槽点评等大量视频涉及AI声音，有些甚至出现不雅词汇，不少视频点赞和评论量过千。

而AI声音滥用事件也不时有发生，引发关注和讨论。

在社交平台上，通过AI模仿恶搞各领域名人的音视频不在少数。此前，短视频平台涌现了大量AI模仿某知名企业家声音吐槽堵车、调休、游戏等热门话题的视频，个别视频甚至还有脏话出现，一度登上热搜。该企业随后发现视频回应称：“相关事件的确让自己挺困扰，也挺不舒服，希望大家都不要‘玩’了。”

一些商家在短视频平台带货时，通过AI模仿声音技术将主播“变”为知名女明星、知名医生，销售服装、保健品等相关产品，对消费者造成了严重误导。国家传染病医学中心主任、复旦大学附属华山医院感染科主任张文宏接受媒体采访时表示，通过语音合成来模仿他的声音进行直播带货，这样的账号“不止一个，且一直在变”，他多次向平台投诉但屡禁不绝。

胖东来创始人于东来的声音也曾

频遭AI模仿，相关平台出现了大量与事实不符的合成视频。胖东来商贸集团为此发布声明称，多个账号未经授权擅自利用AI技术手段生成于东来的声音，加入误导性文案，对公众造成误导和混淆。

记者了解到，有不法分子通过“AI换声”仿冒一位老人的孙子，以“打人须赔偿，否则要坐牢”为由，诈骗老人2万元。类似的诈骗案件在全国已发生多起，有的诈骗金额达到上百万元。

中国社会科学院大学互联网法治研究中心主任刘晓春表示，在未经过授权、未进行标注的情况下，用他人声音制作AI语音产品，尤其是“借用”公众人物的声音，很容易引起误解，不仅会侵害个人信息安全，还可能扰乱网络空间生态和秩序。

声音是如何被“偷”走的？

AI如何生成以假乱真的声音？受访专家介绍，AI能够“克隆”声音，主要是依靠深度学习算法，即短时间内从采集的声音样本中提取关键特征，包括频率、音色、声调、语速、情感等，将这些特征记录为数学模型，再通过算法合成。

中国科学院自动化研究所模式识别实验室工程师牛少东说，随着算法越来越先进，在高性能设备和高精度模型的加持下，AI生成的语音内

容从两年前的“一眼假”升级到如今的“真假难辨”。

大四学生耿孝存最近经常在网络音乐播放器中收听几首翻唱歌曲，他一直以为这些歌曲由某知名女歌手翻唱，后来才得知其全部是AI合成的。“声音逼真到我从来没怀疑过。”耿孝存说。

AI声音在最近一两年时间内变得格外“流行”。清华大学新闻与传播学院教授沈阳说，人工智能技术的普及，让AI模拟声音的门槛大幅降低。通过一些开源软件和平台，没有专业知识的普通用户也能操作。

大量App能够进行AI合成声音，最快只需十几秒。记者在应用商店搜索发现，相关App有数十款，下载量最高超千万次。

记者联系了一款App的客服人员，对方表示，花198元就能解锁付费会员，对着镜头说几遍“12345”，AI就会根据声音生成各类内容的出镜口播视频。记者操作后发现，通过这款软件生成的名人声音，基本可以以假乱真，且录入名人声音不需要提供任何授权证明。

业内人士告诉记者，AI模拟人声在互联网“流行”，有追逐流量和变现的目的。通过“克隆”名人声音制作的恶搞、猎奇类视频，在相关平台播放和点赞量均不低，有的甚至还被推上热搜。发布者也相应获得流量曝光、粉丝增长、广告收入等播放收益。

此外，“偷”人声音也有不法利益驱动。国家金融监管总局2024年7月发布的《关于防范新型电信网络诈骗的风险提示》中提到，不法分子可能对明星、专家、执法人员等音视频进行人工合成，假借其身份传播虚假信息，从而实现诈骗目的。

多措并举强化治理

用AI生成他人声音，是否违法违规？多位受访专家表示，个人声音中包含的声纹信息具备可识别性，能以电子方式记录，能关联到唯一自然人，是生物识别信息，属于个人信息保护法规定的敏感个人信息之一。

2024年4月，北京互联网法院宣判全国首例“AI声音侵权案”，明确认定在具备可识别性的前提下，自然人声音权益的保护范围可及于AI生成声音。该法院法官认为，未经权利人许可，擅自使用或许可他人使用录音制品中的声音构成侵权。

近年来，有关主管部门出台《人工智能生成合成内容标识办法（征求意见稿）》《互联网信息服务深度合成管理规定》《生成式人工智能服务管理暂行办法》等规定，一定程度上给AI技术使用

划定了红线。

沈阳等专家认为，关于人工智能应用产生的造谣侵权、刑事犯罪、道德伦理等问题，建议有关部门细化完善相关规则，通过典型案例、司法解释等方式给予更为明确的规范指引，厘清法律法规边界。

中国科学院科技战略咨询研究院院长潘教峰认为，需进一步强化人工智能伦理规制，超前部署人工智能风险研究，提前预判人工智能技术应用可能带来的社会影响。

2024年12月，广电总局网络视听司发布《管理提示（AI魔改）》，要求严格落实生成式人工智能内容审核要求，对在平台上使用、传播的各类相关技术产品严格准入和监看，对AI生成内容做出显著提示。

多位专家表示，各类社交网络、短视频平台要强化主动监管意识，及时发现、处理可能涉及侵权的AI生成作品；相关部门应继续加大对利用AI技术进行诈骗等违法犯罪行为的打击力度，形成更加完善的常态化治理机制。

牛少东说，在AI时代，个人也要更加注意保护自己的生物特征信息，增强法律意识，抵制他人侵权等行为。

