

Q:

我国目前法律法规和其他部门规章对规范生成式人工智能的应用有哪些具体的规定?

A:

2023年8月15日发布的《生成式人工智能服务管理暂行办法》，是专门针对生成式人工智能所作出的规定。该暂行办法对提供者及使用者必须遵循社会公德和伦理道德提出了基本要求，其中包括必须坚持社会主义核心价值观，必须保护商业秘密、个人隐私等不受侵犯；同时进一步从服务规范和法律责任等具体方面对生成式人工智能服务进行了具体规范。

对于人工智能的发展和规范，我国早在2017年就颁布了《国务院关于印发新一代人工智能发展规划的通知》，将人工智能作为战略目标，分三步实现。同时，提出要建立人工智能安全监管和评估体系，加大对数据滥用、侵犯个人隐私、违背道德伦理等行为的惩戒力度。

2021年颁布施行的《关于加强互联网信息服务算法综合治理的指导意见》以及2022年施行的《互联网信息服务算法推荐管理规定》均明确了信息服务的规范要求，算法推荐服务提供者不得利用算法推荐服务侵犯他人合法权益。

除此之外还有《互联网信息服务深度合成管理规定》，主要规定深度合成服务提供者进行训练数据管理必须保障数据安全，同时对可能涉及国家安全、个人隐私等信息进行安全评估，采取技术或者人工方式对深度合成服务使用者的输入数据和合成结果进行审核。

Q:

生成式人工智能形成的作品是否具有独创性，能否受到著作权法的保护？

A:

在司法实务中，从裁判者的视角来看，“独创性”要求作品由作者独立完成，并体现出作者的个性化表达。因此，如果只是机械地按照一定的顺序、公式或结构完成作品，不同的人会得到相同的结果，通常会被认为不具有独创性。

对于生成式人工智能形成的作品是否具有独创性，不能一概而论。在一定的数据模型之下，使用者的提示词越详细越具体，则人工智能生成的作品结果与其他人呈现的区别就越大，也可以认为其作品更具备独创性的可能。

具有独创性的作品是应当受到保护的，但是由于目前我国著作权法中著作权主体不包括人工智能模型，因此人工智能模型本身是否能够成为著作权法保护的主体尚存在争议。在目前的司法实务中，通常是由生成式人工智能的使用者向法院请求著作权保护的方式进行维权。

5月15日，首例“AI视听作品侵权案”开庭。此前，在4月23日，北京互联网法院宣判了全国首例AI声音侵权案，原告因其声音被AI技术模仿并商业化使用而获得胜诉；1月，中国首例AI生成图片著作权侵权案已判决生效。该案经过五次审理，最终法院认定原告图片具备“独创性”，符合作品的定义，属于美术作品，受到著作权法保护。生成式人工智能的诞生为人类社会提供了极大便利，但其产生和应用过程也伴随着诸多风险，特别是侵权行为风险。为了更好地了解我国生成式人工智能风险防范立法规定和侵权类型，记者专访了北京大学国际知识产权研究中心研究员、北京云亭律师事务所主任唐青林。

用AI生成内容 注意规避这些风险

Q:

生成式人工智能的迅速发展可能引发哪些法律风险？

A:

生成式人工智能可能引发的法律风险既广泛又复杂，引发的侵权主要表现在几个方面：

对他人信息和隐私权的侵害。生成式人工智能服务提供者前期需要收集大量数据并进行数据训练形成模型，在数据收集的过程中，服务提供者可能未经他人同意即收集、存储、使用了他人的信息，而这些信息可能涉及高度敏感信息或者个人隐私，此时服务提供者可能已经侵犯他人个人信息或者隐私权。

对他人姓名权、肖像权、名誉权的侵害。例如最高法发布的侵犯人格权的典型案例中，有一起就是人工智能软件擅自使用自然人形象创设虚拟人物构成侵权的案例。其中被告作为人工智能的运营者，在软件中允许用户自行添加AI陪伴者，在未征得原告同意的情况下，软件使用了原告的姓名、肖像等信息，使其被标识为“AI陪伴者”，并将该角色开放给其他众多用户。最终北京互联网法院判决被告侵害了原告姓名权、肖像权。

生成式人工智能也比较容易引发知识产权侵权纠纷。比如数据中存在他人享有著作权的作品，而服务提供者未经他人授权即使用这些作品，则可能侵害他人的著作权。同理，这种人工智能生成作品的行为也可能基于使用他人的商标，在对使用生成图片的商品进行宣传时导致消费者产生某种混淆，也可能会构成对商标权的侵害。

Q:

日常生活中人们利用生成式人工智能娱乐和社交，应当注意哪些方面？

A:

首先，在使用人工智能时，我们需要注意个人信息安全和个人隐私的保护。根据个人信息保护法的规定，向他人提供个人信息或者对外公开个人信息的，均应取得用户个人的单独同意。因为用户使用生成式人工智能时需要向机器提供数据，例如提出问题或者给出提示词，而这些数据本身则可能被用于人工智能的训练。服务提供者数据训练的过程中极可能抓取个人未公开的信息，甚至是高度敏感的信息和个人隐私，并将这些抓取的信息作为训练数据，此时很容易造成对个人信息和隐私的泄露，给个人造成不必要的损失。

其次，尽量确保输入信息的真实性和准确度。因为人工智能产品尚不能辨别信息的真伪，其只能根据输入的信息和指令完成文本、图片、视频的输出生成，由此人工智能容易将虚假的、不正确的信息作为事实对待，从而导致生

成内容的错误或者让用户产生混淆。

再次，作为使用者也应当注重对他人人格尊严、个人信息以及其他合法权益的保护，如果在使用人工智能过程中发现生成式人工智能服务不符合法律、行政法规，使用者有权向有关主管部门投诉、举报。

最后，高度关注未成年人群体使用人工智能的限度问题。《生成式人工智能服务管理暂行办法》专门规定，服务提供者应当采取有效措施防范未成年人用户过度依赖或者沉迷生成式人工智能服务。未成年人过度使用生成式人工智能更容易引发未成年人个人信息和隐私被抓取和非法利用的问题，也会容易导致未成年人可能被人工智能所生成的错误信息误导，会对未成年人的思想、认知产生消极的影响。

据新华网