

Deepfakes are everywhere but it is the harm not the tech that needs fixing



indiatoday.in/opinion/story/deepfakes-videos-technology-ban-on-genai-clips-expert-explains-2749018-2025-07-01

Nikhil Narendran

July 1, 2025



As deepfakes blur the line between real and fake, concerns around trust and authenticity grow stronger in the digital age.

What do Elon Musk, Taylor Swift, Sachin Tendulkar, and Asaduddin Owaisi have in common? They have all been targeted by deep fake videos. Likelihood with celebrities across has been exploited for deep fakes promoting crypto schemes, financial scams, and fake nudes.

If you are reasonably active on social media, you might have already spotted AI-generated influencers that confuse you as to whether you are interacting with a real person or an AI bot. This phenomenon is so prevalent that we have started losing touch with what's real.

Seeing is not believing

Generative AI (GenAI) has made it easy for anyone with an internet connection to create content that looks and sounds indistinguishably real. It is leading us to question our age-old wisdom that “seeing is believing.”

Gen-AI is altering our relationship with our visual and auditory senses. In the last century, we considered the camera the arbiter of truth and relied on this medium to tell us the truth.

However, we often forget that the same camera was also used in the last century to capture genuine moments of life, as well as for cinema and propaganda. Iconic images of Soviets capturing Berlin and the Toppling of Saddam Hussein's statue were all carefully staged images. Most social media content is also carefully planned, involving make-believe image building.

So, seeing was not believing even before the Deep Fakes.

Spoofing of Senses

In fact, all our senses have been manipulated, and we have accepted this truth. We are constantly encouraged to laugh during Sitcoms using canned laughter, and most music, including vocals, has been heavily processed.

Synthetic fabrics such as Nylon, Rayon, and PU leather have long hijacked our sense of touch. 99% of the vanilla we consume is not from orchid pods but from synthetic vanilla made from petrochemicals. Synthetic meat, Dalda, ground oil, etc., are other synthetic oils that have taken over our sense of taste.

The fresh bread smell that you get from bakeries is not just from the freshly baked bread but also from maltol, diacetyl and yeast esters released by their ventilation systems.

So, the visual sense is not the only one being taken over by synthetic items. It is the latest one to be altered by technology.

How do we regulate Deep Fakes?

Whenever a new technology comes up, whether it was automobiles, the internet, or electricity, human civilisations have always grappled with this question. Deep fakes using Gen AI are no different.

Should we then ban deep fakes altogether? Do we need to regulate tech or harms?

A knife can be used in the kitchen, so it can also be used to harm a person. We do not regulate knives in the Kitchen, but have outlawed using them to harm.

We have seen how our camera has been used in the past, both for good and bad. It has been used for spreading knowledge and information at the same time, for hatred and falsehood. With deep fakes, only the mode of origination has changed. Instead of a camera being used to commit harm, now an algorithm has been used to commit the same harm.

Therefore, just like we didn't ban the camera for all it's good, we need not ban the Gen AI video or audio generation capabilities. We just need to continue regulating the harmful uses of this technology.

Are we concerned about an AI-generated adorable cat video? If his family consents, do we have any issues with a virtual replica of an erstwhile Bollywood actor acting in a new film?

AI-generated videos need not be regulated altogether. We just need to regulate its bad uses, i.e., the generation of Gen AI content that harms a person or their rights.

For instance, a deep fake that represents a celebrity in a scam or a pornographic video needs to be regulated, as should a fake video in which a politician makes a communally sensitive speech.

Do we need new laws for this?

Interestingly, we may not need new laws to regulate deep fakes. Since we never regulated the camera but only its harmful uses, we can apply the same logic here.

Our Bhartiya Nyaya Sanhita and Information Technology Act are equipped enough to deal with communally sensitive content, obscene content, child abuse materials, and defamatory content.

Victims can approach the police station if they are affected by deepfakes or other harms caused by Gen AI. We have seen cases where the police are unaware of how certain provisions under the Bhartiya Nyaya Sanhita or the Information Technology Act apply to Gen AI content as well. This emphasizes the need for education and greater awareness.

Building Trust

While we may not immediately enact new laws, there is indeed a need for transparency regarding synthetic visual and audio content (i.e., AI-generated content). Gen AI has arguably empowered both good and bad uses of the technology. The technology itself has amplified the problem of Deep Fakes as the tools are becoming increasingly accessible.

The government has already proposed that synthetic content should be obligatorily flagged to differentiate between genuine and AI-produced material. Many respected AI companies have adopted responsible AI principles, incorporating transparency into their outputs. These companies use a watermark in their outputs, helping users to determine whether the content is generated by AI. However, this practice is not universally adopted, leading to confusion among the public about whether content is AI-generated or real. The effectiveness of voluntary labelling remains limited, especially since malicious actors are unlikely to comply.

To make progress, we need systems that automatically watermark or label AI-generated content at the point of creation—ideally at the model or platform level. Such watermarking mechanisms must be tamper-resistant and interoperable across platforms. Simultaneously, the government is rightly advocating for technical capabilities that enable

tracing the originators of synthetic content. This could include embedded metadata, secure logs, or authenticated signatures that help identify the device, model, or user that produced the deep fake.

Once such measures become mandatory and widely adopted, enforcement agencies will be better equipped to trace and act against those misusing Gen AI tools to create deep fakes.

Just as the camera was once feared for capturing the soul and invading privacy, GenAI is often perceived as a threat to our senses and reality. But over time, society learnt to distinguish the tool from its misuse. We regulated invasions of privacy and national security, not the act of photography itself. A similar approach can help us manage deep fakes by not banning the tool but addressing the harms arising from deep fakes.

(Nikhil Narendran is a Partner in Trilegal's Bengaluru office and part of the TMT practice of the firm. He is a subject matter expert in the technology, media, and telecom communication space)