

CENTRO DE AVALIAÇÃO DE SUFICIÊNCIA EM LÍNGUAS ESTRANGEIRAS
EDITAL 04/2025

LÍNGUA INGLESA

TEXTO:

**GOOGLE'S NEW AI TOOL GENERATES CONVINCING DEEPPAKES OF RIOTS, CONFLICT, AND ELECTION
FRAUD**

by Andrew R. Chow and Billy Perrigo

Updated: Jun 3, 2025

TIME MAGAZINE

Google's recently launched AI video tool can generate realistic clips that contain misleading or inflammatory information about news events, according to a TIME analysis and several tech watchdogs.

TIME was able to use Veo 3 to create realistic videos, including a Pakistani crowd setting fire to a Hindu temple; Chinese researchers handling a bat in a wet lab; an election worker shredding ballots; and Palestinians gratefully accepting U.S. aid in Gaza. While each of these videos contained some noticeable inaccuracies, several experts told TIME that if shared on social media with a misleading caption in the heat of a breaking news event, these videos could conceivably fuel social unrest or violence.

While text-to-video generators have existed for several years, Veo 3 marks a significant jump forward, creating AI clips that are nearly indistinguishable from real ones. Unlike the outputs of previous video generators like OpenAI's Sora, Veo 3 videos can include dialogue, soundtracks and sound effects. They largely follow the rules of physics, and lack the telltale flaws of past AI-generated imagery.

Users have had a field day with the tool, creating short films about plastic babies, pharma ads, and man-on-the-street interviews.

But experts worry that tools like Veo 3 will have a much more dangerous effect: turbocharging the spread of misinformation and propaganda and making it even harder to tell fiction from reality. Social media is already flooded with AI-generated content about politicians. In the first week of Veo 3's release, online users posted fake news segments in multiple languages, including an anchor announcing the death of J.K. Rowling and of fake political news conferences.

"The risks from deepfakes and synthetic media have been well known and obvious for years, and the fact the tech industry can't even protect against such well-understood, obvious risks is a clear warning sign that they are not responsible enough to handle even more dangerous, uncontrolled AI and AGI," says Connor Leahy, the CEO of Conjecture, an AI safety company. "The fact that such blatant irresponsible behavior remains completely unregulated and unpunished will have predictably terrible consequences for innocent people around the globe."

Days after Veo 3's release, a car plowed through a crowd in Liverpool, England, injuring more than 70 people. Police swiftly clarified that the driver was white, to pre-empt racist speculation of migrant involvement. (Last summer, false reports that a knife attacker was an undocumented Muslim migrant sparked riots in several cities.) Days later, Veo 3 obligingly generated a video of a similar scene, showing police surrounding a car that

had just crashed—and a Black driver exiting the vehicle. TIME generated the video with the following prompt: “A video of a stationary car surrounded by police in Liverpool, surrounded by trash. Aftermath of a car crash. There are people running away from the car. A man with brown skin is the driver, who slowly exits the car as police arrive- he is arrested. The video is shot from above - the window of a building. There are screams in the background.”

After TIME contacted Google about these videos, the company said it would begin adding a visible watermark to videos generated with Veo 3. The watermark now appears on videos generated by the tool. However, it is very small and could easily be cropped out with video-editing software. In a statement, a Google spokesperson said: “Veo 3 has proved hugely popular since its launch. We’re committed to developing AI responsibly and we have clear policies to protect users from harm and governing the use of our AI tools.”

Videos generated by Veo 3 have always contained an invisible watermark known as SynthID, the spokesperson said. Google is currently working on a tool called SynthID Detector that would allow anyone to upload a video to check whether it contains such a watermark, the spokesperson added. However, this tool is not yet publicly available.

Veo 3 is available for \$249 a month to Google AI Ultra subscribers in countries including the United States and United Kingdom. There were plenty of prompts that Veo 3 *did* block TIME from creating, especially related to migrants or violence. When TIME asked the model to create footage of a fictional hurricane, it wrote that such a video went against its safety guidelines, and “could be misinterpreted as real and cause unnecessary panic or confusion.” The model generally refused to generate videos of recognizable public figures, including President Trump and Elon Musk. It refused to create a video of Anthony Fauci saying that COVID was a hoax perpetrated by the U.S. government.

Veo’s website states that it blocks “harmful requests and results.” The model’s documentation says it underwent pre-release red-teaming, in which testers attempted to elicit harmful outputs from the tool. Additional safeguards were then put in place, including filters on its outputs.

A technical paper released by Google alongside Veo 3 downplays the misinformation risks that the model might pose. Veo 3 is bad at creating text and is “generally prone to small hallucinations that mark videos as clearly fake,” it says. “Second, Veo 3 has a bias for generating cinematic footage, with frequent camera cuts and dramatic camera angles – making it difficult to generate realistic coercive videos, which would be of a lower production quality.”

However, minimal prompting did lead to the creation of provocative videos. One showed a man wearing an LGBT rainbow badge pulling envelopes out of a ballot box and feeding them into a paper shredder. (Veo 3 titled the file “Election Fraud Video.”) Other videos generated in response to prompts by TIME included a dirty factory filled with workers scooping infant formula with their bare hands; an e-bike bursting into flames on a New York City street; and Houthi rebels angrily seizing an American flag.

Some users have been able to take misleading videos even further. Internet researcher Henk van Ess created a fabricated political scandal using Veo 3 by editing together short video clips into a fake newsreel that suggested a small-town school would be replaced by a yacht manufacturer. “If I can create one convincing fake

story in 28 minutes, imagine what dedicated bad actors can produce,” he wrote on Substack. “We’re talking about the potential for dozens of fabricated scandals per day.”

Companies need to be creating mechanisms to distinguish between authentic and synthetic imagery right now,” says Margaret Mitchell, chief AI ethics scientist at Hugging Face. “The benefits of this kind of power—being able to generate realistic life scenes—might include making it possible for people to make their own movies, or to help people via role-playing through stressful situations,” she says. “The potential risks include making it super easy to create intense propaganda that manipulatively enrages masses of people or confirms their biases so as to further propagate discrimination—and bloodshed.”

In the past, there were surefire ways of telling that a video was AI-generated—perhaps a person might have six fingers, or their face might transform between the beginning of the video and the end. But as models improve, those signs are becoming increasingly rare. For now, Veo 3 will only generate clips up to eight seconds long, meaning that if a video contains shots that linger for longer, it’s a sign it could be genuine. But this limitation is not likely to last for long.

Eroding trust online

Cybersecurity experts warn that advanced AI video tools will allow attackers to impersonate executives, vendors or employees at scale, convincing victims to relinquish important data. Nina Brown, a Syracuse University professor who specializes in the intersection of media law and technology, says that while there are other large potential harms—including election interference and the spread of nonconsensual sexually explicit imagery—arguably most concerning is the erosion of collective online trust. “There are smaller harms that cumulatively have this effect of, ‘can anybody trust what they see?’” she says. “That’s the biggest danger.” Already, accusations that real videos are AI-generated have gone viral online. One post on X, which received 2.4 million views, accused a Daily Wire journalist of sharing an AI-generated video of an aid distribution site in Gaza. A journalist at the BBC later confirmed that the video was authentic. Conversely, an AI-generated video of an “emotional support kangaroo” trying to board an airplane went viral and was widely accepted as real by social media users.

Veo 3 and other advanced deepfake tools will also likely spur novel legal clashes. Issues around copyright have flared up, with AI labs including Google being sued by artists for allegedly training on their copyrighted content without authorization. (DeepMind told TechCrunch that Google models like Veo “may” be trained on YouTube material.) Celebrities who are subjected to hyper-realistic deepfakes have some legal protections thanks to “right of publicity” statutes, but those vary drastically from state to state. In April, Congress passed the Take it Down Act, which criminalizes non-consensual deepfake porn and requires platforms to take down such material.

Industry watchdogs argue that additional regulation is necessary to mitigate the spread of deepfake misinformation. “Existing technical safeguards implemented by technology companies such as ‘safety classifiers’ are proving insufficient to stop harmful images and videos from being generated,” says Julia Smakman, a researcher at the Ada Lovelace Institute. “As of now, the only way to effectively prevent deepfake videos from

being used to spread misinformation online is to restrict access to models that can generate them, and to pass laws that require those models to meet safety requirements that meaningfully prevent misuse.”

Available at: <<https://time.com/7290050/veo-3-google-misinformation-deepfake/>>

QUESTÃO 01 (1,0)

Qual é os significados das palavras “misleading” e “watchdogs” no parágrafo 1?

- (A) Que lidera, moderadores.
- (B) Que mostra, policiais.
- (C) **Que engana, vigilantes.**
- (D) Que desaponta, espiões.

QUESTÃO 02 (1,0)

Escolha a alternativa que melhor expressa a ideia principal da frase, "While each of these videos contained some noticeable inaccuracies, several experts told TIME that if shared on social media with a misleading caption in the heat of a breaking news event, these videos could conceivably fuel social unrest or violence.", no parágrafo 2.

- (A) **Os especialistas disseram que poderiam provocar distúrbios sociais ou violência se mostrados em boletins de notícias de última hora.**
- (B) Os vídeos, por serem perfeitamente fiéis aos fatos, protegem a população de boatos e pânico, mesmo quando mostrados em boletins de notícias.
- (C) Compartilhar vídeos nas redes sociais sempre evita a disseminação de notícias falsas e reduz a violência causados por razões sociais.
- (D) Os especialistas afirmaram que os vídeos eram 100% precisos e não causariam nenhum impacto nas redes sociais.

QUESTÃO 03 (1,0)

Qual o argumento principal de Connor Leahy sobre AI e IAG?

- (A) Leahy afirma que produções de IA e IAG são inofensivas e que a regulação excessiva seria a principal causa de danos globais.
- (B) Leahy sugere que apenas governos, não empresas, deveriam desenvolver IA e IAG sem supervisão de vigilantes públicos.
- (C) **Leahy argumenta que a incapacidade da indústria de proteger contra riscos de mídia sintética indica falta de responsabilidade para lidar com IA mais perigosa.**
- (D) Leahy defende que, apesar dos riscos conhecidos de *deepfakes*, a indústria de tecnologia tem gerido bem essas ameaças e não precisa de mais regulação.

QUESTÃO 04 (1,0)

Por que a polícia esclareceu rapidamente que o motorista era de raça branca? Cite duas razões. Parágrafo 7.

- (A) Para evitar especulação sobre o motorista e porque a Inglaterra é um país com problemas de tumultos raciais constantes.

- (B) Para esclarecer tudo sobre o acidente e porque em 2024 o racismo na Inglaterra cresceu e causou violência em Liverpool.
- (C) Para esclarecer o envolvimento da polícia e porque no ano passado houve problemas com uma notícia falsa em Liverpool.
- (D) Para evitar especulação racista sobre o envolvimento migrante e porque em 2024 uma notícia falsa provocou tumultos.

QUESTÃO 05 (1,0)

Qual a explicação para o uso de “did block” na frase, “There were plenty of prompts that Veo 3 *did* block TIME from creating”? Parágrafo 10.

- (A) Para indicar uma ação habitual no passado.
- (B) Para enfatizar que Veo 3 realmente bloqueou TIME.
- (C) Para expressar uma possibilidade no passado.
- (D) Para sugerir uma ação condicional que aconteceu por acaso.

QUESTÃO 06 (1,0)

De acordo com o artigo técnico divulgado pelo Google sobre o Veo 3, qual das seguintes afirmações é verdadeira?

- (A) O Veo 3 é propenso a pequenas alucinações que fazem com que os vídeos se mostrem falsos.
- (B) O artigo afirma que o Veo 3 não tem distorções ou limitações quando se trata de geração de vídeo.
- (C) O modelo é altamente eficaz na geração de vídeos coercitivos com alta qualidade de produção.
- (D) O Veo 3 se destaca na criação de vídeos realistas com ângulos de câmera dramáticos.

QUESTÃO 07 (1,0)

Qual a tradução correta dos seguintes sintagmas nominais: “na LGBT Rainbow badge”, “Election Fraud Video”, “their bare hands”, “a New York City Street”?

- (A) “um distintivo LGBT”, “Eleição Fraude Vídeo”, “mãos nuas”, “na rua Nova York”.
- (B) “um arco-íris LGBT”, “Um vídeo com fraude”, “com suas mãos”, “em Nova York na rua”.
- (C) “um distintivo de arco-íris LGBT”, “Vídeo de Fraude Eleitoral”, “suas próprias mãos”, “uma rua em Nova York”.
- (D) “um distintivo de arco-íris”, “um Vídeo Eleitoral Fraudulento”, “suas próprias mãos”, “uma rua em Nova York”.

QUESTÃO 08 (1,0)

O que levará à erosão da confiança coletiva online segundo Nina Brown?

- (A) A existência de vídeos representando executivos e funcionários de empresas.
- (B) O fato de que ninguém saberá reconhecer se um vídeo é verídico ou falso.
- (C) O fato de que Veo 3 possibilita a criação de vídeos com conteúdo falso.
- (D) A possibilidade do uso de vídeos falsos para interferir em eleições.

QUESTÃO 09 (1,0)

Quais são algumas questões que gerarão conflitos legais por causa de vídeos *deepfake*? Escolha a opção **incorreta**.

- (A) Haverá mais ações legais sobre direitos autorais por parte de artistas.
- (B) Celebridades serão cada vez mais sujeitos a vídeos *deepfakes* super-realistas.
- (C) As vítimas de vídeos pornográficos *deepfake* terão que exigir a retirada do vídeo.
- (D) Youtube terá que exigir que DeepMind não use os vídeos da plataforma para *deepfakes*.**

QUESTÃO 10 (1,0)

Segundo Julia Smakman quais são as maneiras de controlar a produção de vídeos *deepfake*? Escolha a opção **correta**.

- (A) Aprovar leis que exijam que os modelos IA atendam a requisitos de segurança que previnem toda a produção de *deepfakes*.
- (B) Restringir o acesso a modelos que geram *deepfakes* e aprovar leis que exijam que esses modelos atendam a requisitos de segurança.**
- (C) Restringir de forma contundente o acesso a todos os modelos que geram *deepfakes* e assim controlar o problema na fonte.
- (D) Aprovar leis que restringem o acesso a modelos *deepfakes* e que punem aqueles que compartilham este tipo de vídeo.