

Une sommité mondiale lance un avertissement: il n'est plus à même de repérer les vidéos générées par l'IA

par Frank Bergman*



Frank Bergman. (Photo <https://slaynews.com>)

L'un des plus grands experts mondiaux en criminalistique numérique tire la sonnette d'alarme après avoir admis que l'intelligence artificielle (IA) est désormais si performante que même les spécialistes les plus chevronnés ne sont plus en mesure de distinguer de manière

fiable ce qui est authentique de ce qui est falsifié.

Cet avertissement intervient à un moment où la technologie du «hypertrucage» ou «deepfake» connaît une croissance explosive à l'échelle mondiale et favorise la fraude, la manipulation politique, l'usurpation d'identité ainsi que des escroqueries sophistiquées, dont les experts craignent qu'elles ne dépassent bientôt la capacité de la société à distinguer la réalité de la fiction.

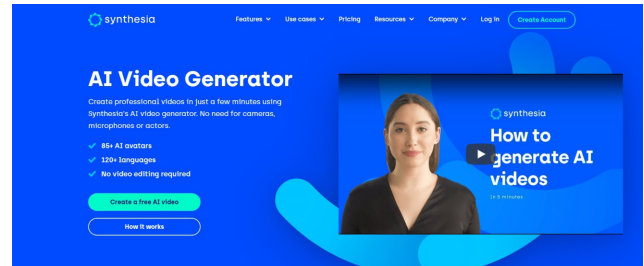
Même les experts se laissent tromper

Selon un article du «New York Times», le célèbre expert en criminalistique numérique *Hany Farid* affirme que l'intelligence artificielle a évolué si rapidement que bon nombre des méthodes traditionnelles d'identification des images et vidéos falsifiées ne sont plus efficaces.¹

Farid, professeur à la UC *Berkeley School of Information* et conseiller scientifique principal au Center for *Long-Term Cybersecurity*, se consacre depuis des décennies à la détection des manipulations numériques et des supercheries en ligne.

Il affirme aujourd'hui que la technologie progresse plus vite que les humains ne peuvent suivre. «Je ne crois plus en rien», a admis Farid.

* Frank Bergman est un journaliste spécialisé dans la politique et l'économie qui vit sur la côte Est des Etats-Unis. Outre ses reportages d'actualité, M. Bergman réalise également des entretiens avec des chercheurs et des experts dans divers domaines, et mène des enquêtes sur des personnalités et des organisations influentes du monde sociopolitique. <https://slaynews.com/author/frank-bergman/>



(Photo <https://filmora.wondershare.de>)

«Pour chaque image que je vois, je trace mentalement des lignes d'ombre et je calcule la géométrie pour déterminer ce que je suis réellement en train de regarder.»

«C'est fini», a-t-il averti. «D'ici un ou deux ans, tout notre système visuel sera complètement inutile.» Farid a expliqué que les systèmes d'IA sont désormais remarquablement habiles à reproduire des détails qui, auparavant, permettaient de démasquer les contenus falsifiés – notamment des ombres réalistes, des expressions faciales, l'éclairage ambiant, des sons, des plantes et même des mouvements corporels.

Il en résulte un déferlement croissant de contenus synthétiques qui deviennent de plus en plus impossibles à distinguer de la réalité.

L'explosion des hypertrucages (deepfake) alimente la hausse mondiale de la fraude

Cet avertissement intervient à un moment où la technologie des hypertrucages est utilisée comme une arme à une échelle sans précédent. Selon la société de cybersécurité *DeepStrike*, le nombre de fichiers de hypertrucage circulant en ligne est passé d'environ 500 000 en 2023 à plus de 8 millions en 2025.²

Au cours de la même période, les tentatives de fraude liées à des contenus générés par l'IA auraient augmenté de plus de 3000% à l'échelle mondiale.

Les experts mettent en garde contre le fait que les criminels utilisent désormais l'intelligence artificielle pour se faire passer pour des membres de la famille, des dirigeants d'entreprise, des responsables politiques, des célébrités et des fonctionnaires. Dans de nombreux cas, les victimes

ne sont pas en mesure de distinguer les voix, vidéos ou photos falsifiées des authentiques.

On reproche aux «Big Tech» d'ignorer les risques

Farid a adressé ses critiques les plus virulentes aux entreprises de la Silicon Valley, engagées dans une course au développement de systèmes d'IA toujours plus performants. «Je ne supporte plus cet endroit», a-t-il déclaré. «Ces géants de la tech sont prêts à tout raser tant qu'ils en tirent des bénéfices.»

«Ils ne s'intéressent à rien qui pourrait les ralentir.» Farid a acquis une vaste expérience tant dans le secteur public que dans le secteur privé.

Il a contribué au développement de la technologie PhotoDNA de *Microsoft*,³ utilisée pour identifier les contenus liés à la pédopornographie, et a collaboré avec l'armée américaine sur des systèmes de détection des vidéos d'hypertrucage.

Mais malgré des années passées à lutter contre la désinformation numérique, il affirme aujourd'hui que le rythme du développement de l'IA est devenu profondément inquiétant.

L'IA contre l'IA

Face à l'inefficacité croissante de la détection humaine, les experts se tournent vers l'intelligence artificielle elle-même comme moyen de défense potentiel. Des outils tels que *Deepfake-O-Meter*⁴ et la société de Farid, *GetReal Security*,⁵ utilisent des systèmes d'IA pour analyser des images, des vidéos et des enregistrements audio à la recherche de signes de manipulation.

Parallèlement, les entreprises technologiques et les chercheurs militent en faveur de nouvelles normes d'authentification conçues pour vérifier l'origine des contenus numériques. Une initiative, connue sous le nom de *Content Credentials*,⁶ permettrait aux images d'être accompagnées d'un enregistrement numérique indiquant comment elles ont été créées et si elles ont été modifiées.

Les partisans de cette initiative affirment que de tels systèmes pourraient devenir indispen-

sables à mesure que les contenus générés par l'IA inondent Internet.

Des escrocs exploitent cette technologie

Les experts avertissent que les «hypertrucages» sont déjà utilisés dans un nombre croissant d'escroqueries ciblant les Américains lambda. Selon la société de détection par IA *Copyleaks*,⁷ les menaces émergentes comprennent de faux témoignages de célébrités, des escroqueries sentimentales utilisant des identités générées par l'IA, des appels d'urgence frauduleux imitant la voix de proches, des séquences de vidéosurveillance fabriquées de toutes pièces et de fausses vidéos politiques destinées à manipuler les électeurs.

A mesure que cette technologie devient moins coûteuse et plus accessible, les chercheurs craignent que la menace ne fasse que s'intensifier.

L'avertissement de Farid souligne une inquiétude croissante parmi les professionnels de la cybersécurité: la société pourrait entrer dans une ère où «voir, ce n'est plus croire». Pendant des décennies, les photographies et les vidéos ont été considérées comme parmi les preuves les plus solides qui soient.

Aujourd'hui, même l'un des plus grands experts mondiaux affirme que cette confiance s'effrite rapidement.

Source: <https://slaynews.com/world-leading-expert-warns-no-longer-detect-ai-videos>, 20 juin 2026

(Traduction «Point de vue Suisse»)

¹ <https://www.nytimes.com/2026/06/14/us/ai-deepfake-hany-farid.html>, 14 juin 2026

² <https://deepstrike.io/blog/deepfake-statistics-2025>, 19 avril 2026

³ <https://www.ischool.berkeley.edu/news/2020/hany-farid-leads-microsoft-project-detect-child-sex-predators>, 9 janvier 2020

⁴ <https://zinc.cse.buffalo.edu/ubmdfl/deep-o-meter/>

⁵ <https://www.getrealsecurity.com/teams/labs/pr-hany-farid>

⁶ <https://contentcredentials.org/>

⁷ <https://copyleaks.com/blog/5-most-common-types-of-video-deepfake-scams>