
756/A(E) XXVIII. GP

Eingebracht am 25.02.2026

Dieser Text wurde elektronisch übermittelt. Abweichungen vom Original sind möglich.

ENTSCHLIESSUNGSANTRAG

der Abgeordneten Süleyman Zorba, Meri Disoski, Freundinnen und Freunde

betreffend Schutzmechanismen der KI-Verordnung dürfen nicht per Digital-Omnibus ausgehebelt werden

BEGRÜNDUNG

Nudify-Apps und KI-Deepfakes sorgen weltweit für Skandale: Aus normalen Fotos werden damit täuschend echt wirkende Nackt- oder Pornobilder erzeugt, und das ohne Einwilligung der Betroffenen.

Am 25.12.2025 lancierte Elon Musk bei seinem KI-Chatbot Grok einen „spicy mode“. In den folgenden Tagen entlud sich eine Lawine an sexualisierten KI-Deepfakes über X. Fotos von Frauen, aber auch von Mädchen, wurden millionenfach für die Erstellung von Porno-Deepfakes missbraucht. Prominente Frauen, wie Taylor Swift wurden ebenso Opfer wie Politikerinnen oder Frauen, deren Fotos für Revenge Porn missbraucht wurden. Die Flut an Deepfake Pornos machte aber auch vor Kindern nicht halt – so wurden nach Erhebungen des Center for Countering Digital Hate (CCDH) mehr als 23.000 pornografische Deepfakes von Kindern mit Grok erstellt.¹

Während Elon Musk im Bikini als vermeintlicher Zensur-Kämpfer posierte, warnen Regierungen und NGOs vor zunehmender sexueller Gewalt, befeuert durch KI. Das volle Ausmaß des Grok-Skandals zeichnet sich nun in ersten Erhebungen ab: 85% aller Grok-Bilderstellungsanfragen waren explizit sexualisiert. In weniger als zwei Wochen wurden etwa 3 Millionen sexualisierte Deepfakes, primär von Frauen erstellt. Am Höhepunkt des Skandals am 5.1.2026 wurden bis zu 6.700 sexualisierte Deepfakes pro Stunde (!) erstellt.²

¹ <https://counterhate.com/research/grok-floods-x-with-sexualized-images/#about>
<https://www.heise.de/news/Analysen-Zwei-bis-drei-Millionen-sexualisierte-Deepfakes-von-Grok-generiert-11150776.html>

² <https://www.aicerts.ai/news/global-fallout-from-grok-deepfake-scandal/>

Skandale wie dieser offenbaren nicht nur gesellschaftliche Abgründe sondern zeigen auch, wie (generative) KI millionenfach für Missbrauch und strafbare Handlungen genutzt werden kann. Und sie zeigen, wie wesentlich es ist, dass KI-Anbieter ihre Dienste mit funktionierenden Guardrails – digitalen „Leitplanken“ – ausstatten, die rechtswidrige Outputs verhindern. Bei Grok fehlten hinreichende Guardrails bzw waren sie – wie auch schon die von xAI gewählte Bezeichnung als „spicy mode“ nahelegt – absichtlich lasch ausgestaltet.³

Wer glaubt, dass Grok ein Einzelfall ist, liegt leider falsch. Untersuchungen zeigen, dass große App-Stores (zB Apple oder Google) und Plattformen wie Metas Instagram oder Facebook zeitweise Dutzende „nudify“-Apps sowie hunderte entsprechende Werbeanzeigen zugelassen haben, obwohl die Inhalte klar auf nicht-einvernehmliche Deepfakes zielten.⁴ So wurde explizit mit „Lade ein Foto hoch“ und „sieh jeden nackt“ geworben.

Aber auch abseits der Schmutzdecke zeigt sich zunehmend die dunkle Seite von KI-Anwendungen, wenn KI-Chatbots vulnerable Menschen psychologisch schädigen: So haben etwa Chatbots wie ChatGPT oder Character.AI bei Jugendlichen Suizidgedanken nicht nur validiert, sondern auch noch gefördert, inklusive Ratschlägen zu Suizid-Methoden und Abschiedsbriefen.⁵ Gleichzeitig pushen KI-Image-Generatoren und Filter-Apps bei Teenagern ein toxisches Idealbild, lösen Essstörungen und Selbsthass aus. Studien zeigen, dass 70% der Mädchen nach KI-generierten „perfekten“ Selfies ihr echtes Äußeres ablehnen.

All das zeigt: Ohne robuste Guardrails und verlässliche Regeln wird KI zum unsichtbaren gesellschaftspolitischen Gift. Angesichts der aktuellen Befunde zum Risikopotenzial von KI, insbesondere von generativer KI ist eine Aufweichung der KI-Verordnung, die einen verlässlichen Rahmen für diese technologische Revolution schafft, durch einen Digital-Omnibus völlig verfehlt.

Die KI-Verordnung wurde erst 2024 verabschiedet und wird – zB im Hinblick auf KI-Kennzeichnung oder Regeln für Hochrisiko-KI – erst am 2.8.2026, mit 2-jähriger Legisvakanz in Kraft treten. Dieses ausdifferenziert gestaltete, mit Experten in viel aufwändiger Detailarbeit entwickelte Instrument zum Schutz von Grundrechten, Demokratie und Bevölkerung sollte nicht mit einem „Digital-Omnibus“⁶ zugunsten von Big Tech Konzernen ausgehebelt werden. Klar ist, dass dieser Digital-Omnibus auf die

³ <https://www.diepresse.com/20572034/ki-als-waffe-gegen-frauen-zeit-fuer-strenge-regeln>

⁴ <https://www.cbsnews.com/news/meta-instagram-facebook-ads-nudify-deepfake-ai-tools-cbs-news-investigation/>

⁵ <https://www.heise.de/news/Suizid-nach-KI-Chats-Google-und-Character-ai-wenden-womoeglich-Urteile-ab-11135053.html> <https://www.handelsblatt.com/technik/ki/ki-chats-selbsttoetung-nach-ki-chat-wie-gefaehrlich-ist-chatgpt/100161344.html> <https://www.forschung-und-lehre.de/forschung/chatgpt-und-co-geben-gefaehrdende-inhalte-weiter-7266>

⁶ <https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=CELEX:52025PC0836>

Lobby-Wunschliste von Tech-Konzernen wie Meta, Microsoft und Google zurück geht.⁷ Mit digitaler Souveränität ist so ein Einknicken vor Big Tech nicht vereinbar.

Die unterfertigenden Abgeordneten stellen daher folgenden

ENTSCHLIESSUNGSANTRAG

Der Nationalrat wolle beschließen:

„Die Bundesregierung wird aufgefordert, auf europäischer Ebene auf eine Beibehaltung des hohen Schutzniveaus der KI-Verordnung hinzuwirken und Vorschlägen, die Digitalgesetze aufweichen, nicht zuzustimmen. Weiters wird die österreichische Bundesregierung aufgefordert, endlich die in der KI-Verordnung gesetzlich vorgesehene KI-Behörde einzurichten.“

In formeller Hinsicht wird die Zuweisung an den Ausschuss für Wissenschaft, Forschung und Digitalisierung vorgeschlagen.

⁷ <https://www.derstandard.at/story/3000000304601/wie-tech-unternehmen-der-eu-kommission-die-aufweichung-der-digitalgesetze-diktierten>