

Künstliche Stimmen, echte Gefahr

24.06.2026, 13:49 | IT, New Media & Software

Pressemitteilung von: *idw - Informationsdienst Wissenschaft*



„Hallo Oma, hier ist Florian, ich brauche dringend Geld.“ Mithilfe Künstlicher Intelligenz können Kriminelle inzwischen Stimmen von bekannten oder vertrauten Personen täuschend echt nachahmen – automatisiert, individuell und massenhaft. Dazu kommt, dass Cyberkriminelle Informationen aus sozialen Netzwerken oder Datenlecks abgreifen, um mit persönlichen Details gezielter vorzugehen. So gewinnen sie schnell Vertrauen. Im Forschungsprojekt „AntiScam“ (deutsch: Anti-Betrug) der Universität Siegen untersuchen Wissenschaftler*innen sogenannte „Conversational Scams“, also digitale Betrugsversuche, die wie echte Unterhaltungen wirken. Besonders gefährlich ist dabei „Voice Phishing“: Betrüger geben sich am Telefon etwa als Familienmitglied, Bank oder Behörde aus. Das Bundesministerium für Bildung und Forschung (BMBF) fördert das Vorhaben. Insgesamt steht eine Projektsumme von 1,4 Millionen Euro zur Verfügung. Die Universität Siegen koordiniert das Verbundprojekt. Kooperationspartner sind die Hochschule Bonn-Rhein-Sieg und das Siegener Unternehmen open.INC.

„Die Technik entwickelt sich rasant. Viele Menschen merken inzwischen gar nicht mehr, ob sie mit einem echten Menschen oder mit einer KI sprechen“, sagt Dr. Md Shajalal, Projektverantwortlicher von AntiScam. Das haben die Wissenschaftler*innen bereits herausgefunden. „Oft verlässt man sich auf sein Bauchgefühl und liegt dabei falsch. Genau deshalb brauchen wir neue Schutzmechanismen und mehr Aufklärung.“ Oft werde bei Betrugsanrufen mit Angst und Druck gearbeitet, berichtet Shajalal. „Durch den Stress fällt es vielen Betroffenen schwer, Manipulation rechtzeitig zu erkennen.“ Das seit langem bekannte Phishing – Betrugsversuche etwa per E-Mail – ist bereits gut erforscht. Voice-Phishing bisher nicht. Das Projekt untersucht, wie sich solche Angriffe erkennen und verhindern lassen. Die Forschenden entwickeln dafür Systeme, die künstlich erzeugte Stimmen identifizieren können.

AntiScam verfolgt dabei einen interdisziplinären, sozio-technischen Forschungsansatz, der Methoden des Maschinellen Lernens mit Methoden der Mensch-Computer-Interaktion verbindet. Dafür entsteht zunächst eine systematische Bedrohungskarte KI-gestützter Gesprächsbetrugsformen, die typische Angriffsmuster, Manipulationsstrategien und potenzielle Gegenmaßnahmen erfasst. Auf dieser Grundlage entwickeln die Forschenden technische Erkennungssysteme sowie Strategien zur digitalen Selbstverteidigung.

Im Zentrum der technischen Forschung steht ein KI-Modell, das echte von künstlich erzeugten Stimmen unterscheiden kann. Trainiert wird das System mit einer eigens aufgebauten Datenbank aus mehreren tausend Sprachaufnahmen menschlicher Sprecher*innen sowie KI-generierter Stimmen. Die Analyse konzentriert sich auf akustische Merkmale, die vom menschlichen Ohr kaum wahrgenommen werden können, wie etwa ungewöhnliche Gleichförmigkeit, fehlende Atemgeräusche oder unnatürliche Modulationsmuster.

Besonderen Wert legen die Forschenden dabei auf den Einsatz sogenannter erklärbarer KI (Explainable AI, XAI): Das

System soll nicht nur verdächtige Stimmen erkennen, sondern nachvollziehbar machen, warum eine Stimme als potenziell manipuliert eingestuft wurde. Nutzerinnen und Nutzer sollen Warnungen dadurch besser verstehen und ihre eigene Fähigkeit stärken, betrügerische Kommunikationssituationen künftig selbst zu erkennen.

Wenn Vertrauen zur Falle wird

Gleichzeitig sollen Nutzerinnen und Nutzer verstehen, woran sich Manipulation erkennen lässt. Deshalb setzt AntiScam nicht nur auf Technik, sondern auch auf Aufklärung. Geplant sind unter anderem interaktive Lernangebote und ein Lernspiel, das Verbraucherinnen und Verbraucher für typische Betrugsstrategien sensibilisiert. Verbraucher*innen sollen typische Betrugsstrategien dort realitätsnah kennenlernen und trainieren können. „Kriminelle setzen immer stärker auf psychologische Manipulation“, sagt Prof. Dr. Gunnar Stevens von der Professur für Wirtschaftsinformatik / Datenschutz und IT-Sicherheit. „Deshalb reicht klassische IT-Sicherheit allein nicht mehr aus. Wir müssen verstehen, wie Menschen Entscheidungen treffen und wie sich Vertrauen digital manipulieren lässt.“

Parallel prüfen die Forschenden rechtliche Fragen rund um KI-Betrug und erarbeiten Empfehlungen für besseren Verbraucherschutz und neue Regeln im Umgang mit Deepfakes und digitaler Manipulation. Den rechtlichen Part übernimmt Prof. Dr. Maximilian Becker von der Professur für Bürgerliches Recht und Wirtschaftsrecht, insbesondere Immaterialgüterrecht sowie Medienrecht.

Die bisherigen Forschungsergebnisse wurden bereits auf internationalen Fachkonferenzen vorgestellt, unter anderem in Deutschland, Japan und Kanada. Weitere wissenschaftliche Publikationen und Konferenzbeiträge sind im Projektverlauf geplant. Sowohl die entwickelten Softwarelösungen als auch die wissenschaftlichen Erkenntnisse sollen darüber hinaus offen zugänglich gemacht werden, um Forschung, Wirtschaft und Gesellschaft langfristig bei der Bekämpfung KI-gestützter Betrugsformen zu unterstützen. Die Software soll Ende 2027 verfügbar sein. Die Projektergebnisse bieten Potenziale für die Entwicklung neuer Sicherheitsprodukte und neue Marktchancen für Unternehmen im Bereich der Cybersicherheit.

wissenschaftliche Ansprechpartner:

Prof. Dr. Gunnar Stevens

Wirtschaftsinformatik / Datenschutz und IT-Sicherheit, Universität Siegen

gunnar.stevens@uni-siegen.de

Universität Siegen

NoraRatmann (Mitarbeiter in der Presse- und Öffentlichkeitsarbeit)

0271 740-4836

ratmann@presse.uni-siegen.de

News-ID: 1315942 • Views: 188 (Stand: 12.07.2026)

Link zur Pressemitteilung:

<https://www.openpr.de/news/1315942/Kuenstliche-Stimmen-echte-Gefahr-idw.html>