

24.04.2026

Die Industrialisierung des Betrugs: Wie KI die Finanzkriminalität revolutioniert

Die Zeiten von offensichtlich fehlerhaften Phishing-Mails und plumpen Betrugsversuchen sind endgültig vorbei. Im Jahr 2026 erleben wir eine echte Zäsur: Künstliche Intelligenz (KI) hat Cyberkriminalität von mühsamer, manueller Arbeit in eine hochautomatisierte, global skalierende Industrie verwandelt.

Was früher einzelne Täter mit begrenzter Reichweite waren, sind heute organisierte Netzwerke mit spezialisierten KI-Tools, datengetriebenen Strategien und beinahe unbegrenzter Skalierbarkeit. Für Nutzer von E-Banking und digitalen Finanzdiensten bedeutet das: Herkömmliche Warnsignale, wie Rechtschreibfehler, unpersönliche Anrede oder merkwürdige Absender, verlieren zunehmend ihre Aussagekraft.

Während Finanzinstitute ihre technischen Schutzmechanismen stetig verbessern, verlagern Angreifer ihren Fokus konsequent auf die grösste Schwachstelle im System: den Menschen. KI fungiert dabei als Reichweitenverstärker, Präzisionswerkzeug und Täuschungsgenie zugleich.

Deepfakes und Voice Cloning: Die Identität als Beute

Die wohl gefährlichste Entwicklung ist der sogenannte Deepfake – die Fähigkeit, Stimmen und Videos täuschend echt zu imitieren. Während früher umfangreiches Audiomaterial notwendig war, reichen heute bereits wenige Sekunden aus Social Media oder Sprachnachrichten aus, um einen überzeugenden digitalen Zwilling zu erzeugen.

- **Der «Chef-Betrug» (CEO-Fraud) 2.0:**

In einer scheinbar legitimen Videokonferenz fordert ein «Geschäftsführer» eine dringende Überweisung. Stimme, Gesicht und Verhalten stimmen, nur diese Person nimmt gerade nicht an der Videokonferenz teil.

- **Der Schockanruf:**

Opfer erhalten Anrufe von angeblichen Familienmitgliedern in Not. KI simuliert nicht nur die Stimme, sondern erzeugt ein komplettes akustisches Szenario inklusive Stress, Emotion und Umgebung.

Massen-Personalisierung durch LLMs

Ein Large Language Model (LLM) ist ein KI-System, das durch das Training mit gigantischen Datenmengen statistische Muster in der Sprache gelernt hat, um eigenständig logisch klingende Texte zu generieren und Aufgaben zu lösen. Dazu gehören ChatGPT, Gemini oder Claude. Solche LLMs ermöglichen eine völlig neue Qualität von Betrug, die sich durch Individualisierung, Kontextsensitivität und Echtzeitfähigkeit auszeichnet. Anstatt auf generische Massenmails zu setzen, nutzen Angreifer KI, um täuschend echte und hochgradig personalisierte Szenarien zu erzeugen. Dabei analysieren die Systeme öffentlich zugängliche Informationen wie Social-Media-Profile, Firmenwebsites oder Pressemitteilungen und verknüpfen diese zu glaubwürdigen Geschichten. So entstehen beispielsweise Zahlungsaufforderungen oder Projektbezüge, die real existieren und daher kaum Verdacht erregen.

Hinzu kommt, dass Sprachbarrieren praktisch verschwunden sind. Moderne KI kann Inhalte nicht nur korrekt übersetzen, sondern auch kulturell und sprachlich anpassen. Betrüger treten dadurch in fehlerfreiem Hochdeutsch,

Schweizerdeutsch oder sogar in regionalen Dialekten auf und wirken dadurch deutlich authentischer als früher.

Ein weiterer entscheidender Unterschied liegt in der Interaktivität: Während klassische Phishing-Mails statisch waren, reagieren KI-gestützte Systeme heute dynamisch auf Rückfragen. Sie führen überzeugende Dialoge, liefern plausible Erklärungen, gehen auf Einwände ein und passen ihre Argumentation in Echtzeit an die Reaktionen des Opfers an. Dadurch entsteht eine Form der Kommunikation, die kaum noch von echten menschlichen Interaktionen zu unterscheiden ist und das Risiko erfolgreicher Täuschung erheblich erhöht.

Automatisierung und Skalierung: Betrug als Geschäftsmodell

Ein zentraler Unterschied zur Vergangenheit ist die Industrialisierung:

- **Crime-as-a-Service:**

Kriminelle bieten fertige KI-Tools, Deepfake-Services oder komplette Betrugskampagnen im Darknet an – inklusive Support und Updates.

- **A/B-Testing von Betrug:**

Angriffe werden wie Marketingkampagnen optimiert. Verschiedene Varianten werden getestet, Erfolgsquoten analysiert und Strategien kontinuierlich angepasst.

- **24/7-Betriebsmodelle:**

KI-Systeme arbeiten rund um die Uhr, ohne Pausen oder Ermüdung und erhöhen so die Effizienz drastisch.

Die psychologische Falle: Social Engineering in Perfektion

Trotz aller technologischen Fortschritte bleibt der Kern von Betrugsversuchen unverändert: die gezielte Manipulation menschlichen Verhaltens. Künstliche Intelligenz verstärkt dabei bekannte psychologische Hebel und setzt sie deutlich präziser ein als je zuvor. Besonders häufig wird ein Gefühl der Dringlichkeit erzeugt, indem Opfer glauben, sofort handeln zu müssen, um einen Schaden abzuwenden. Gleichzeitig wird Autorität inszeniert, etwa durch vermeintliche Anweisungen von Vorgesetzten, Banken oder Behörden, die kaum hinterfragt werden. Auch das Prinzip der Knappheit kommt gezielt zum Einsatz, indem eine angeblich einmalige Gelegenheit oder ein begrenztes Zeitfenster suggeriert wird. Ergänzt wird dies durch den Aufbau von Vertrauen, beispielsweise durch das Vortäuschen bekannter Personen oder persönlicher Beziehungen.

Mögliche Schutzstrategien

Technische Lösungen allein reichen nicht aus. Entscheidend ist eine geschärfte «Human Firewall» – also das Sicherheitsbewusstsein jedes Einzelnen. KI kann auch für die Detektion verwendet werden, doch diese Verfahren sind noch nicht zuverlässig.

- Zwei-Faktor-Authentifizierung (2FA), biometrische Absicherung und Transaktionsbestätigungen bieten wichtige Schutzebenen und sollten verwendet werden.
- Bei ungewöhnlichen Zahlungsaufforderungen immer einen zweiten Kanal nutzen. Rückruf nur über bekannte Nummern oder alternative Kommunikationswege.
- Zeitdruck ist ein klassisches Manipulationsmittel. Seriöse Institutionen setzen keine unrealistischen Fristen.
- Bei Unsicherheit «Vier-Augen-Prinzip» anwenden.
- Weniger öffentlich verfügbare persönliche Daten (in sozialen Medien) reduzieren die Angriffsfläche deutlich.