

SHIFT

SHIFT TECHNOLOGY INSURANCE PERSPECTIVES

**SONDERAUSGABE ZUR BETRUGSBEZOGENEN
BILDANALYSE**

Aus der Redaktion

Im Kampf gegen Versicherungsbetrug erweist sich Generative AI (GenAI) als zweiseitiges Schwert. Einerseits ist sie ein leistungsfähiges Instrument, mit dem Versicherer Informationen aus unterschiedlichen Versicherungsdokumenten präzise extrahieren und klassifizieren sowie die Ergebnisse analysieren können, um Anomalien und Unstimmigkeiten zu erkennen, die auf Betrug hindeuten. GenAI kann Ermittlern Zusammenfassungen von Dokumenten, Handlungsempfehlungen für die nächsten Schritte sowie automatisierte Berichte und Prüfpfade bereitstellen. Zudem lassen sich viele Aspekte der Betrugserkennung automatisieren: komplexe Betrugsnetzwerke ebenso wie opportunistische Täter, die einen Versicherer ins Visier nehmen, können identifiziert werden. Dabei analysiert GenAI strukturierte und unstrukturierte Daten zuverlässig – einschließlich Text, Bildern und Audio.

Andererseits erleichtert die Verbreitung frei verfügbarer GenAI-Werkzeuge Personen mit betrügerischer Absicht den Versuch, Versicherungsbetrug zu begehen. Mit einfachen Eingaben in frei zugängliche Online-Ressourcen kann ein Betrüger glaubwürdig aussehende Unfallberichte, Kostenvoranschläge und Rechnungen, Begleitdokumente und sogar Bilder erzeugen.

Die Branche steht vor einem klassischen „Gleiches mit Gleichem bekämpfen“-Szenario. Wie können Versicherer GenAI nutzen, um diejenigen, die ihnen wirtschaftlich schaden wollen, tatsächlich einen Schritt voraus zu sein? In dieser Ausgabe der Shift Insurance Perspectives betrachten wir genauer, wie Bildanalyse bereits heute Versicherern hilft, Deepfakes zu erkennen und Auszahlungen für unberechtigte Schadensfälle zu vermeiden..

Die aktuelle Situation

Ist es Betrug oder nicht? Für Versicherer, die mit einem verdächtigen Schadenfall konfrontiert sind, ist diese Entscheidung nur so belastbar wie die Verlässlichkeit der herangezogenen Quellen. Leider wird die Erstellung fingierter „Beweise“ für Schadensmeldungen in der Branche nicht nur erschreckend fotorealistisch, sondern für den Durchschnittsnutzer immer einfacher. Wir haben gesehen, dass in eingereichten Bildern zu Kfz-Schäden nachträglich Beulen, Dellen, Schrammen und Kratzer hinzugefügt wurden. Betrüger manipulieren Fotos, um Löcher in Wänden sowie Rauch- oder Wasserschäden darzustellen und so Sachschäden zu begründen. Es gibt Fälle, in denen Gegenstände wie Fernseher, Stereoanlagen und andere Elektronik in Fotos von Wohnräumen eingefügt werden, um Diebstahlsmeldungen zu untermauern – häufig zusammen mit gefälschten Polizeiberichten. Wir sind an einem Punkt angekommen, an dem sich menschliche Schadenbearbeiter und Ermittler beim Anblick solcher KI-generierten Bilder nicht mehr auf ihre eigenen Augen verlassen können.

Gleichzeitig müssen wir uns für diese kritischen Entscheidungen nicht mehr ausschließlich auf unsere menschlichen Sinne stützen. Was durch KI erzeugt werden kann, lässt sich durch KI auch erkennen. Wie sollten Versicherer den Einsatz von KI zur Bildanalyse als Teil des Schadenprozesses denken?

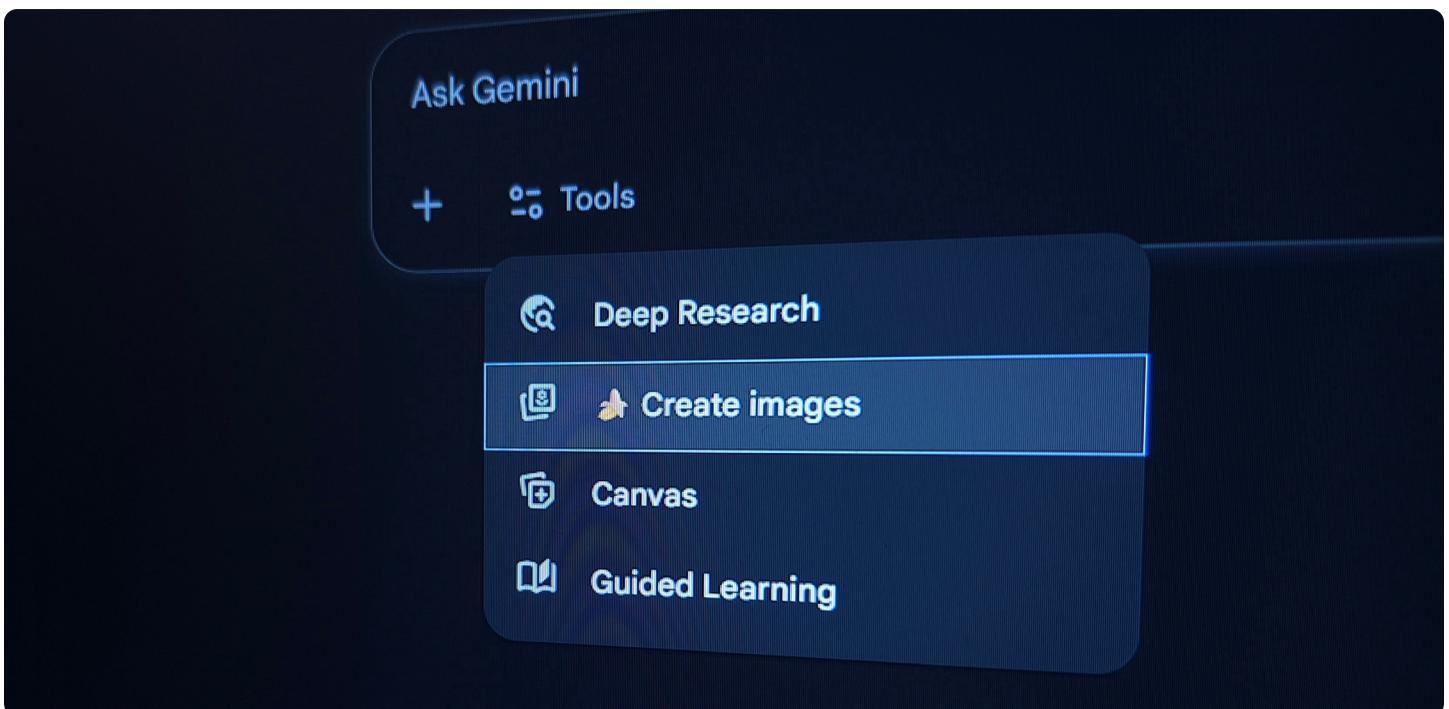


Analyse der zugrunde liegenden Metadaten

Der erste Schritt, um zu beurteilen, ob ein Bild legitim ist, besteht in der Analyse der zugehörigen Metadaten. Stimmt der Zeitstempel mit dem gemeldeten Zeitpunkt des Ereignisses überein? Weisen die Geolokationsdaten darauf hin, dass das Foto am Schadensort aufgenommen wurde – oder an einem weit entfernten Ort? Zeigen gerätebezogene Informationen, dass das Bild mit der im Schadenfall angegebenen Kamera/dem Smartphone/dem Tablet erzeugt wurde, oder deuten sie auf andere Unregelmäßigkeiten hin? Lassen die Metadaten erkennen, dass das Bild aus dem Internet heruntergeladen oder anderweitig manipuliert wurde? Jede dieser Konstellationen ist ein deutlicher Hinweis darauf, dass ein manipuliertes Bild zur Untermauerung einer betrügerischen Schadenmeldung verwendet wird.

Wenn Metadaten nicht verlässlich sind

In der Praxis der Bildanalyse gilt leider: Metadaten können nahezu genauso leicht manipuliert werden wie das Bild selbst. Bilder können in andere Formate konvertiert werden, wodurch die zugrunde liegenden Metadaten verfälscht oder verschleiert werden. Bilder können beschnitten werden oder ihre Metadaten können vollständig entfernt sein. Und obwohl solche Maßnahmen auf Betrugsabsichten hindeuten können, müssen sie es nicht. Versicherungsnehmer reichen mitunter Screenshots statt Originalfotos ein. Beschnittene Bilder werden mitunter genutzt, um Schäden hervorzuheben. Wie bei jedem Element einer Schadenakte, das Verdachtsmomente wecken kann, muss die Branche sorgsam vorgehen, um fehlerhafte Warnhinweise in der Bildanalyse zu vermeiden.



Nutzung der Bildanalyse im Kampf gegen Betrug

Es gibt mehrere Methoden, um unrechtmäßig verwendete Bilder zu identifizieren, die zur Begehung von Betrug eingesetzt werden. Die Error-Level-Analyse (ELA) ist eine etablierte Technik zur Überprüfung der Authentizität von fotografischen und anderen digitalen Beweismitteln. Obwohl nicht spezifisch GenAI-basiert, kann ELA helfen festzustellen, ob einzelne Elemente eines Fotos geschickt verändert oder zusammengesetzt wurden. Sie lässt sich auch nutzen, um zu prüfen, ob die Elemente eines digitalen Dokuments konsistent sind. Dieser Ansatz hat sich bewährt beim Erkennen von Deepfakes, beim Aufspüren neu hinzugefügter Bildelemente sowie beim Hervorheben von Informationen, die aus einem Dokument in ein anderes „kopiert und eingefügt“ wurden. Mit dem Aufkommen vollständig KI-generierter Bilder, wie sie heute von leicht verfügbaren Modellen erzeugt werden, ist ELA jedoch weniger wirksam geworden. Von LLMs generierte Bilder sind eine einzelne, kohärente Einheit; es gibt daher kein Zusammenfügen oder „Copy & Paste“, das aufgedeckt werden könnte.

Da breit verfügbare Modelle, etwa von OpenAI, Anthropic und anderen, inzwischen neben Text auch Bild-Inputs verarbeiten können, lassen sich GenAI-Werkzeuge einsetzen, um klar fiktive Inhalte, beispielsweise Zeichnungen, leichter auszusortieren. Auch wenn die Fähigkeit, eine KI-erzeugte Zeichnung im Kontext der Betrugserkennung zu identifizieren, zunächst unnötig erscheinen mag, ist eine Zeichnung ein echtes Bild. In einfachen Schadenfällen, insbesondere solchen, die für Dunkelverarbeitung (Straight-Through-Processing, STP) vorgesehen sind, könnte ohne diesen Check eine derartige Bilddatei ausreichen, um einen betrügerischen Anspruch erfolgreich zu machen. Diese Prüfung würde jedoch in der Regel neben einem komplexeren Bildanalysemodell laufen und eher wie eine administrative Vorsortierung von Beweismitteln fungieren.

Die Bildanalyse gewinnt zunehmend an Bedeutung, seit Transformers (eine Deep-Learning-Architektur) in größerem Umfang eingesetzt werden.

Die Bildanalyse gewinnt zunehmend an Bedeutung, seit Transformers (eine Deep-Learning-Architektur) in größerem Umfang eingesetzt werden. In den Shift-eigenen Tests liefern diese Modelle die besten und konsistentesten Ergebnisse, wenn es darum ging, KI-generierte Bilder zu erkennen. Transformers werden von den Modellen der generativen KI wie GPT abgeleitet und können in kleinere, spezialisierte Versionen modifiziert werden, die andere Eingabetypen – etwa Bilder – fokussiert berücksichtigen. Durch das Training von diesen Vision-Transformern auf „echte“ und „falsche“ Bilder wurden sehr überzeugende Resultate erzielt.

Fazit

Da moderne Smartphones standardmäßig automatisch irgendeine Form von KI einsetzen, um Fotos zu verbessern bzw. hochskalieren, ließe sich argumentieren, dass sämtliche zur Begründung eines Versicherungsanspruchs eingereichten Fotobelege als KI-generiert betrachtet werden könnten. Umso wichtiger ist es, zuverlässig zwischen Bildern zu unterscheiden, die den tatsächlichen Sachverhalt abbilden, und solchen, die eine falsche Darstellung stützen. Obwohl generative KI Betrügern neue Werkzeuge an die Hand gibt, stellt sie zugleich den Versicherern die Mittel bereit, wirksam gegenzuhalten.

SHIFT

Über Shift Technology

Shift Technology ist die führende KI-Plattform für Versicherungen. Shift kombiniert generative, agentenbasierte und prädiktive KI, um das Underwriting, die Schadensabwicklung sowie die Betrugserkennung und das Risikomanagement zu transformieren und so die betriebliche Effizienz und das Kundenerlebnis zu verbessern und konkret betriebswirtschaftliche Ergebnisse zu erzielen. Shift wird von weltweit führenden Versicherern anerkannt und eingesetzt und liefert KI genau dann, wenn sie am dringendsten benötigt wird – in großem Umfang und mit nachgewiesenen Ergebnissen.

Erfahren Sie mehr unter www.shift-technology.com/de.