

KI ZUR VERHINDERUNG VON IDENTITÄTSBETRUG

VON DER KUNDENIDENTIFIKATION ZUR PRÄVENTION VON VERBRAUCHERBETRUG

HRSG.: WILHELM BAUER | OLIVER RIEDEL | THOMAS RENNER | MATTHIAS PEISSNER





Christian H. Schunck, Rachelle Sellung, Heiko Roßnagel

KI ZUR VERHINDERUNG VON IDENTITÄTSBETRUG

Von der Kundenidentifikation zur Prävention von Verbraucherbetrug

Herausgeber

Wilhelm Bauer, Oliver Riedel, Thomas Renner, Matthias Peissner

VORWORT

Künstliche Intelligenz (KI) ist eine der zentralen Technologien für die Zukunft. Ihre Einführung und der Einsatz fordern Unternehmen im besonderen Maß heraus. Es gilt, das Potenzial zu erkennen und dieses wirtschaftlich nutzbar zu machen. Lassen Sie sich dabei durch Europas größte Forschungsk Kooperation auf dem Gebiet der KI, Cyber Valley, begleiten.

Mit dem KI-Fortschrittszentrum von Fraunhofer IAO und Fraunhofer IPA unterstützen wir Unternehmen dabei, das Potenzial von KI nutzbringend einzusetzen. An der Schnittstelle zwischen anwendungsorientierter Wirtschaft und exzellenter Forschung des Cyber-Valley-Konsortiums entwickeln wir innovative KI-Anwendungen für die Praxis und treiben damit die Kommerzialisierung von KI voran. Erklärtes Ziel ist dabei, menschenzentrierte KI-Lösungen zu entwickeln. Denn nur wenn Menschen mit einer neuen Technologie intuitiv interagieren und vertrauensvoll zusammenarbeiten, kann ihr Potenzial optimal ausgeschöpft werden.

Die Studienreihe »Lernende Systeme« des KI-Fortschrittszentrums gibt Einblick in die Potenziale und die praktischen Einsatzmöglichkeiten von KI. Dabei werden übergreifende Themen wie Zuverlässigkeit, Erklärbarkeit (xAI), cloudbasierte Plattformen, Technologien und Einführungsstrategien diskutiert. Zudem werden einzelne Anwendungsbereiche in der Wissensarbeit, Bauwirtschaft, Produktion und dem Kundenservice im Detail beleuchtet.



Die vorliegende Studie betrachtet den Einsatz von KI in der Kundenidentifikation, der Transaktionsbewertung und der Verhinderung von Identitäts- und Verbraucherbetrug. Immer mehr Endkund*innen und Vertriebskanäle nutzen das Internet, und KI leistet schon heute einen entscheidenden Beitrag zu skalierbaren Lösungen in der Betrugsprävention. Unternehmen profitieren so von einem signifikanten Zuwachs an neuem, sicherem Geschäft.

Wir wünschen Ihnen eine spannende Lektüre, und freuen uns, wenn wir in Zukunft auch Sie mit unserer Expertise auf Ihrem Weg zur menschenzentrierten KI unterstützen dürfen.

Wilhelm Bauer, Oliver Riedel, Thomas Renner, Matthias Peissner
Fraunhofer-Institut für Arbeitswirtschaft und Organisation IAO

INHALT

1	Management Summary	6
2	Einleitung	8
3	Vorgehen der Studie	10
4	Studienergebnisse	12
4.1	Vorstellung der verschiedenen Verfahren	13
4.1.1	Fernidentifikationen von natürlichen Personen	13
4.1.2	Geräteerkennung	14
4.1.3	Lösungen zum Datenabgleich mit gepoolten Daten/großen Datenbanken	15
4.1.4	Behavioral Biometrics	16
4.1.5	Kundenindividuelle Lösungen	17
4.2	Betrugserkennung und Identitätsmanagement	17
4.3	Aktueller Stand	18
4.3.1	Haupteinsatzfelder für die KI	18
4.3.2	Kurze Übersicht der verwendeten Methoden	20
4.3.3	Datengrundlage und -bereitstellung	23
4.3.4	Qualitätssicherung	25
4.3.5	Kundeninteraktion mit der Lösung und Umgang mit Verdachtsfällen	26
4.3.6	Vorteile durch den KI-Einsatz	27
4.4	KI – Potenzial und Herausforderungen	28
4.4.1	Zukünftiges Potenzial	28
4.4.2	Die größten Herausforderungen	29
4.4.3	Verbesserungspotenzial	32
5	Zusammenfassung	34

6	Steckbriefe der Unternehmen	35
6.1	CRIF Bürgel GmbH	36
6.2	d-fine GmbH	38
6.3	NECT GmbH	40
6.4	RISK.IDENT GmbH	42
6.5	SCHUFA AG	44
6.6	WidasConcepts GmbH	46
	Literatur	48
	KI-Fortschrittszentrum	50
	Fraunhofer-Gesellschaft	51
	Autorinnen und Autoren	56

ABBILDUNGEN

<i>Abbildung 1: Vorstellung der verschiedenen Verfahren.</i>	<i>14</i>
<i>Abbildung 2: Verfahren der Geräteerkennung.</i>	<i>15</i>
<i>Abbildung 3: Hauptaufgaben der KI in der Betrugserkennung.</i>	<i>19</i>
<i>Abbildung 4: Ergebnisse bezüglich des Aufwands für die Bereitstellung der Daten von den Kunden.</i>	<i>24</i>
<i>Abbildung 5: KI-Potenzial und Herausforderungen.</i>	<i>32</i>
<i>Abbildung 6: Verbesserungspotenzial.</i>	<i>33</i>

1 MANAGEMENT SUMMARY

Identitätsbetrug ist in Deutschland und weltweit zu einer ernst zu nehmenden Bedrohung geworden. Die Täter*innen agieren höchst professionell, und die Schäden sind für Bürger*innen so wie für den Handel erheblich. Die Kosten von Betrugsfällen mit Verbraucher- bzw. Endkundenbezug betragen für Unternehmen im E-Commerce [1] schätzungsweise ca. 2 Prozent des Umsatzes. Auf Deutschland bezogen, entspricht dies einem Schadensvolumen von mehreren Milliarden Euro pro Jahr.

Betrugsabwehrsysteme in Unternehmen basieren oft weitgehend noch auf manueller Überprüfung, der Aufholbedarf ist dementsprechend hoch und wird durch die gegenwärtige Coronakrise für viele Unternehmen noch verschärft, da sie zu einer erheblichen Verlagerung von Geschäftsbeziehungen in den Cyberbereich führt.

In dieser Studie analysieren wir »good practices« für die schnelle, automatisierte und KI-basierte Prüfung von Kundenidentitäten und Onlinetransaktionen. Das Ziel ist, skalierbare Verfahren und Lösungen zu identifizieren und zu analysieren und ihre Einsatzmöglichkeiten aufzuzeigen.

Die Studie gibt Einblicke in die am Markt verfügbaren Lösungsansätze; die Haupteinsatzfelder von KI, die Datengrundlage und die verwendeten Modelle; Verfahren zur Qualitätssicherung; die Interaktion von Kund*innen mit den am Markt verfügbaren Lösungen; die Vorteile des Einsatzes und das zukünftige Potenzial von KI; die an der Studie beteiligten Unternehmen werden in Form von Steckbriefen vorgestellt.

Wesentliche Studienergebnisse sind:

- Eine effektive Prävention von Identitäts- und Verbraucherbetrug ohne den unterstützenden Einsatz von KI ist kaum noch denkbar.
- Die auf dem Markt verfügbaren Lösungen spiegeln die Vielfalt der Betrugsmuster wieder und werden oft in Kombination eingesetzt.
- Auch der Einsatz von KI ist sehr vielfältig und reicht von Text- über Bild- bis zur Anomalieerkennung und der Optimierung von Regelsystemen.
- Unternehmen, die entsprechende Produkte einsetzen möchten, sollten darauf vorbereitet sein, dass kontinuierlich an der Qualitätssicherung gearbeitet werden muss. Nur so lässt sich das Verbesserungspotenzial durch den KI-Einsatz voll ausschöpfen.
- Durch den Einsatz von KI kann man eine Erhöhung der Trennschärfe in der Beurteilung von Verdachtsfällen und damit einen signifikanten Anstieg an neuem, sicherem Geschäft erwarten.

2 EINLEITUNG

Durch Betrug mit Verbraucher- bzw. Endkundenbezug entstehen der deutschen Wirtschaft vor allem in den Bereichen Einzelhandel, E-Commerce und Finanzdienstleistungen Milliarden-schäden pro Jahr. Dem Betrug, unter anderem mit gefälschten oder erfundenen Identitäten sowie unautorisierten oder betrügerischen Transaktionen, muss daher zunehmend mit innovativen und automatisierten Verfahren begegnet werden.

»Zu Beginn der meisten Cybercrime-Straftaten steht der Diebstahl einer digitalen Identität.« [2]

In dieser Studie analysieren wir »good practices« für die schnelle, automatisierte und KI-basierte Prüfung der Identitäten von Kund*innen und Onlinetransaktionen. Das Ziel ist die Identifikation und Analyse von skalierbaren Verfahren zur Identitäts- und Transaktionsprüfung.

Dabei gilt es zu beachten, dass das Thema »Betrug« im Allgemeinen, aber insbesondere auch im Bereich der Identitäts- und Transaktionsprüfung sehr komplex ist, und oft eindeutige Definitionen fehlen. So können z. B. Identitäten natürlichen Personen, aber auch Unternehmen (d. h. juristischen Personen) zugerechnet werden, und Identitätsbetrug kann mit erfundenen Identitäten, aber auch mit gestohlenen Identitäten erfolgen. Für jeden dieser – und vieler weiterer – Fälle werden andere Ansätze sowohl für die Identifikation von potenziellen Betrugsfällen als auch im Umgang mit Verdachtsfällen benötigt.

Zudem muss Betrug erst einmal richtig als solcher erkannt werden: Nicht jeder Verdachtsfall ist ein Betrugsfall. Immer gilt: Betrug liegt vor, wenn eine Person mit betrügerischer Intention ans Werk geht. Dies ist ein fundamentaler Unterschied zum Bereich der – nicht betrügerischen – Kreditausfälle bzw. Bonität. Hier haben Endkund*innen die Intention zu zahlen, aber sie überschätzen ihre finanziellen Möglichkeiten oder können aufgrund von unvorhergesehenen Ereignissen wie Scheidung, Arbeitslosigkeit oder Berufsunfähigkeit, ihren Zahlungsverpflichtungen nicht nachkommen.

Für ein Unternehmen ist in diesem Fall die Fragestellung eindeutig: »Können meine Kund*innen so viel zahlen, dass ich Gewinn mache?« Im Gegensatz dazu versuchen im Betrugsfall die Täter*innen aktiv die Entdeckung des Betrugsversuchs als Verdachtsfall zu verhindern, und es ist oft nicht leicht zu entscheiden, ob es sich bei einem Verdachtsfall wirklich um Betrug handelt. Zum Beispiel kann Betrug auch erst bei der Rücksendung von bestellter Ware erfolgen

(z. B. echte Ware wird gegen gefälschte Produkte getauscht, es werden aus der Ware Ersatzteile entnommen). Die Antwort auf die Frage »Was soll verhindert werden?« ist im Betrugsfall daher wesentlich schwerer zu beantworten.

Letztlich ist Betrugsprävention zudem in der Regel ein nachrangiger Prozess, Priorität hat das Kerngeschäft bzw. der Verkauf. Das impliziert, dass die Daten, die zur Betrugsprävention genutzt werden können, meistens diejenigen Daten sind, die im Verkaufsprozess ohnehin anfallen bzw. die automatisch, ohne Mitwirkung von Endkund*innen, abgegriffen werden können. Selten wird der Verkaufsprozess zu Zwecken der Betrugsprävention verändert.

In diesem Umfeld bietet der Einsatz von Künstlicher Intelligenz (KI) ein enormes Potenzial, um Betrug zu verhindern und Unternehmen einen möglichst hohen Anteil an sicherem (Neu)geschäft zu ermöglichen.

So vielfältig, wie sich der Bereich des Onlinebetrugs darstellt, so vielfältig sind auch die Lösungen, die Unternehmen entwickeln, um andere Unternehmen bei der Betrugserkennung zu unterstützen. In dieser Studie gehen wir auf einige dieser Ansätze ein, die unter Einbeziehung von KI von Unternehmen wie Auskunftsteilen, Beratungsunternehmen, Betreiberfirmen von Portalen für das Identitätsmanagement, Diensten zur Videoverifikation von Konsumierenden und auf Betrugserkennung spezialisierten Dienstleistern angeboten werden.

Dabei wird insbesondere

- untersucht, wie die Lösungen angeboten und zur Verfügung gestellt werden;
- die Integration von KI und Daten und der gegenwärtige Stand des Einsatzes von KI in den Lösungen vorgestellt;
- das Potenzial und die Herausforderungen für den weitergehenden Einsatz von KI für die Betrugserkennung in der Zukunft betrachtet.

Schließlich werden die Unternehmen, die sich an der Studie beteiligt haben, anhand eines kurzen Steckbriefs vorgestellt. Diese Übersicht erhebt keinesfalls einen Anspruch auf Vollständigkeit, sondern soll in erster Linie einen Überblick über die auf dem Markt verfügbaren Lösungen geben.

3 VORGEHEN DER STUDIE

Die Studie wurde durch eine eingehende Analyse der bestehenden Literatur zum Thema KI-Einsatz zur Kundenidentifikation und Verhinderung von Identitäts- und Verbraucherbetrug vorbereitet. Danach wurden die unterschiedlichen Lösungsansätze, die auf dem Markt zur Verfügung stehen, identifiziert und ein Fragebogen für semistrukturierte Interviews mit Anbietern dieser Lösungen vorbereitet. Wir haben dann verschiedene Anbieter relevanter Lösungen kontaktiert und sechs Expert*innen auf Grundlage des Fragebogens interviewt. Basierend auf den Vorarbeiten und den Interviews wurde die vorliegende qualitative Studie erstellt. Im Folgenden geben wir weitere Informationen zur Vorgehensweise.

Inhalt und Aufbau des Fragebogens

Der Fragebogen wurde so erstellt, dass er als Basis für ein 30- bis 45-minütiges, semistrukturiertes Interview dient. Es ergaben sich insgesamt 24 Fragen in vier Hauptkategorien. Die erste Kategorie konzentrierte sich auf einen Einblick in das von der interviewten Person vertretene Unternehmen, wobei Fragen zu den entwickelten Produkten gestellt wurden. Die zweite Kategorie konzentrierte sich auf den aktuellen Stand der entwickelten Lösungen in Bezug auf den Einsatz von KI. Die dritte Kategorie konzentrierte sich auf das Potenzial von KI und mögliche Hindernisse für den Einsatz von KI in der Zukunft. Die letzte Kategorie beinhaltet demografische Fragen und schließt den Fragebogen ab.

Umfang der Befragung

Ziel war es, eine vielfältige Auswahl an interviewten Personen von Lösungsanbietern aller Größen und mit unterschiedlichen Schwerpunkten zu gewinnen, um einen umfassenden Überblick über die Herausforderungen und bestehenden Lösungsansätze für die Verhinderung von Identitätsbetrug und KI zu erhalten. Dazu kontaktierten wir verschiedene Unternehmen mit relevanten Lösungen. Jede für das Interview ausgewählte Person durchlief einen anfänglichen Screeningprozess, um sicherzustellen, dass ihre Expertise und die Lösungen zu den Anforderungen passen, die wir für die Beantwortung des Fragebogens benötigten.

Ablauf und Auswertung der Interviews

Die Interviews fanden als Online-Videointerviews statt, wobei jeweils zwei interviewende Personen und eine befragte Person anwesend waren. Mit Einverständnis der befragten Personen wurden die Interviews aufgezeichnet und transkribiert. Das transkribierte semistrukturierte Interviewformular wurde dann den befragten Personen zur Überprüfung vorgelegt. Die Expert*innen durften Änderungen und Klarstellungen vornehmen und haben dann die Transkription akzeptiert. Danach wurden die Interviews für diese Studie ausgewertet und die Aufnahmen der Interviews gelöscht.

4 STUDIENERGEBNISSE

Um den Einsatz von KI zur Kundenidentifikation und Verhinderung von Identitäts- und Verbraucherbetrug richtig verstehen und beurteilen zu können, ist es wichtig, sich zuerst ein paar grundlegende Charakteristiken von Betrug zu vergegenwärtigen: Betrug impliziert immer Vorsatz. Betrugsmuster sind sehr vielfältig, Betrüger*innen zeigen ein kreatives, sich schnell veränderndes Verhalten und versuchen mit immer neuen Ansätzen erfolgreich zu sein. Es mangelt in der Regel an allgemein anerkannten Klassifikationen für Betrugsmuster.

Dies hat wichtige Konsequenzen:

1) Verdachtsfälle sind nicht immer Betrugsversuche, und als unverdächtig eingestufte Vorkommnisse können unerkannte Betrugsversuche sein. Dies erschwert eine präzise Zielbestimmung und damit die Definition von Zielparametern, um abweichende Ergebnisse erkennen zu können.

2) Es ist sehr schwierig, von der Vergangenheit in die Zukunft zu prognostizieren, oft muss betrügerisches Verhalten, wenn es eindeutig als solches erkannt wurde, möglichst schnell »nachgelernt« werden.

Für den Einsatz von KI heißt dies: Es ist oft nicht einfach, Betrugsfälle richtig zu klassifizieren, und bei geringen Fallzahlen mit großen Schäden kann man nicht immer warten, bis neue Fallmuster von der KI erkannt werden. Andererseits kann es sinnvoll sein, bestimmte Fallmuster, die nur kurzfristig bestehen (z. B. wenn sich in einer Stadt Betrugsversuche mit Ausweisen aus einem bestimmten Land kurzfristig häufen), schnell wieder »auszusortieren«.

Im Folgenden werden die Konsequenzen für den Einsatz von KI im Detail diskutiert. Aber schon hier kann erwähnt werden, dass aus diesen Gründen KI oft im Zusammenhang mit Regelsystemen, z. B. zur Regeloptimierung, eingesetzt wird. Regeln lassen sich so schnell verbessern, aber auch wieder ausschalten, ohne dass das gesamte System rekaliert werden muss.

4.1 Vorstellung der verschiedenen Verfahren

Für diese Studie konnten wir mit Anbietern verschiedener Lösungen sprechen. Deren Lösungsansätze nutzen die folgenden Verfahren:

- Fernidentifikationen von natürlichen Personen
- Geräteerkennung
- Lösungen zum Datenabgleich mit gepoolten Daten/großen Datenbanken
- Behavioral Biometrics
- Kundenindividuelle Softwareentwicklung

Oft werden von Anbietern verschiedene Verfahren miteinander kombiniert. Im Folgenden stellen wir diese Verfahren kurz vor:

4.1.1 Fernidentifikationen von natürlichen Personen

Bei der Fernidentifikation von natürlichen Personen geht es im Wesentlichen um die Erfüllung von regulatorischen »know your customer (KYC)«-Anforderungen in stark regulierten Branchen, z. B. im Banken- und Versicherungsbereich, unter anderem zur Verhinderung von Geldwäsche. Kund*innen sollen durch die Vorlage und den Abgleich mit amtlichen Dokumenten wie dem Personalausweis ihre Identität eindeutig nachweisen. Im Offlinebereich wäre hier ein bekanntes Standardverfahren das von der deutschen Post entwickelte Postident®-Verfahren, bei dem sich Kund*innen in einer Poststelle mithilfe eines Reisepasses oder Personalausweises persönlich ausweisen. Zudem werden weitere Daten auf der Vorder- und Rückseite des Ausweisdokuments erfasst und mit dem von den Nutzer*innen angegebenen Daten abgeglichen wie z. B. Adresse, Geburtsdatum und Unterschrift.

Im Rahmen der Online-Fernidentifikation wird die persönliche Vorsprache mithilfe eines video-basierten Verfahrens, bei dem Endkund*innen sich und ihr Ausweisdokument in einem Videostream vorstellen, ersetzt.

Der Anbieter des Verfahrens zur Fernidentifikation erhält den Auftrag von seinem Kunden. Dieser Kunde leitet die relevanten Daten der Endkund*innen oder Konsumierenden an den Anbieter weiter. Die Endkund*innen laden dann eine App des Anbieters auf ihre Smartphones und nutzen die Funktionen der App und auch den Videostream für die Fernidentifikation.

Für die Betrugserkennung in diesem Kontext ist es wesentlich, die Echtheit des Ausweisdokuments zu prüfen und die Übereinstimmung des Fotos im gezeigten Ausweisdokument mit der partizipierenden Person im Videostream zu verifizieren (Gesichtserkennung und -abgleich).

Für die meisten gegenwärtigen Anwendungen ist die Akzeptanz des Verfahrens durch die zuständigen Regulierungsbehörden zentral. Eine Ausweitung des Verfahrens zur Identifikation von Kund*innen auch für andere Anwendungsfälle (z. B. Online-Partnerschaftsvermittlungsportale) wäre denkbar. Allerdings müssen dafür die Kosten einer Fernidentifikation deutlich sinken, während sich die Zuverlässigkeit nicht wesentlich ändern sollte.

Einsatzfelder für KI sind die Gesichtserkennung und der Gesichtsabgleich, die Erkennung, ob es sich bei dem Videobild der Person um das Bild einer echten, lebendigen Person handelt und die Echtheitsprüfung der Ausweisdokumente einschließlich ihrer Sicherheitsmerkmale.



Abbildung 1: Vorstellung der verschiedenen Verfahren.

4.1.2 Geräteerkennung

Die Geräteerkennung ist oft ein wichtiger Bestandteil von Systemen zur Betrugserkennung. Das Ziel ist es, Endgeräte wie Tablets, Smartphones oder Computer eindeutig zu identifizieren und so Geräte wiederzuerkennen, die in der Vergangenheit in Zusammenhang mit Betrugsfällen eingesetzt wurden. Dies kann ein wichtiges Warnsignal für einen möglichen Betrugsversuch sein. In diesem Zusammenhang können die folgenden Verfahren eingesetzt werden (siehe Abbildung 2).



Es werden Cookies gesetzt, die als ein künstlicher »Fingerabdruck« agieren und so eine Wiedererkennung ermöglichen. Viele Browserhersteller versuchen dies allerdings zu unterbinden und auch Betrüger *innen versuchen solche Cookies zu entfernen.



Es werden eindeutige Daten wie IP- und MAC-Adressen ausgelesen und gespeichert. Ein solches Vorgehen ist datenschutzrechtlich oft problematisch, da man dazu tief in die Geräte eingreifen muss bzw. auch dynamische IP-Adressen im Sinne der DSGVO als personenbezogene Daten einzustufen sind. Daher wird auf die Erhebung solcher Daten in der Regel verzichtet.



Es werden Browsereinstellungen und andere Daten, die über den Browser zugänglich sind, genutzt und gemessen. Dies kann mehr als hundert Datenpunkte ergeben, aus denen sich eindeutige Erkennungsmerkmale ableiten lassen. Einige dieser Merkmale können auch direkt Hinweise auf ein kompromittiertes Gerät bzw. auf Muster liefern, die auf ein betrügerisches Verhalten hinweisen. Bei diesen Merkmalen handelt es sich nicht um personenbeziehbare Daten, und eine Verarbeitung ist datenschutzrechtlich daher in der Regel unbedenklich.



Es ist möglich, einem Endgerät kleinere Aufgaben zu geben, z. B. Berechnungen mit Pixeln, und dann zu beobachten, wie das Gerät eine solche Aufgabe löst. Damit lässt sich beispielsweise prüfen, ob ein Gerät, das vorgibt, ein iPhone® zu sein, auch wirklich ein iPhone® ist.

Abbildung 2: Verfahren der Geräteerkennung.

Einsatzfelder für die KI liegen in der Erkennung von Gerätemerkmalen und der Erkennung von Unplausibilitäten und Mustern, die auf kompromittierte Geräte hinweisen.

4.1.3 Lösungen zum Datenabgleich mit gepoolten Daten/großen Datenbanken

Bei diesem Lösungsansatz werden die Informationen, die zu einem Vorgang (Verbraucher*in, Bestellung) erhoben werden, mit Daten in proprietären Datenpools abgeglichen und/oder in einem großen Datenpool – oft unternehmensübergreifend – vernetzt. Dabei wird nach Verbindungen zu bekannten Betrugs- oder Verdachtsfällen gesucht, indem Muster und Unplausibilitäten identifiziert werden. Lösungsanbieter ermöglichen ihren Kunden zusätzlich auch oft einen gezielten Informationsaustausch zu Betrugsfällen und -versuchen. Dies bietet sich besonders

bei den Sachverhalten an, in denen Betrugsversuche ähnliche Charakteristiken für die meisten Unternehmen einer Branche haben. Dies ist vor allem bei Betrugsphänomenen im E-Commerce und Digital Business (digitale Güter) der Fall.

»Es ist oft sinnvoll, Daten für die Analysen subscriptionsübergreifend zusammenzuführen.«

Thomas Widmann, Gründer und Geschäftsführer, WidasConcepts GmbH

Im Wesentlichen können hier zwei Einsatzbereiche unterschieden werden:

Kundenidentitäten (KYC – know your customer):

Das Ziel ist, Kundenidentitäten eindeutig festzustellen und dabei gefälschte bzw. gestohlene Identitäten zu erkennen. Dabei müssen die Systeme in der Lage sein, auch mit unvollständigen oder fehlerhaften Daten (z. B. aufgrund von Tippfehlern) zu arbeiten. Im E-Commerce ist oft eine zuverlässige Zuordnung zwischen Adresse und Endkund*in wichtig. Zudem muss gegebenenfalls zwischen natürlichen Personen (B2C) und juristischen Personen (B2B) als Endkund*innen unterschieden werden. Unternehmen sind in der Regel eindeutig zu identifizieren, aber um regulatorischen Anforderungen gerecht zu werden, muss oft die wirtschaftlich berechnete Person ermittelt werden.

Transaktionen (KYT – know your transaction):

Für Transaktionen werden prädiktive Lösungen eingesetzt, die basierend auf Transaktionsmustern Abhängigkeiten identifizieren und Transaktionen als auffällig bzw. unauffällig bewerten. Dabei werden transaktionale Zusammenhänge und Transaktionshistorien über längere Zeiträume hinweg verglichen und ausgewertet.

Wie schon weiter oben erwähnt, wird KI in jedem dieser Bereiche – jeweils auf unterschiedliche Weise – zur Mustererkennung und Regeloptimierung eingesetzt.

4.1.4 Behavioral Biometrics

In diesem Bereich werden Technologien eingesetzt, die beobachten und auswerten, wie Endkund*innen sich auf digitalen Plattformen und in »sales channels« verhalten (z. B. Schreibgeschwindigkeit, Klickraten), um Auffälligkeiten zu erkennen [3], [4]. Viele diese Technologien werden schon im Marketing und zur Optimierung von Webseiten eingesetzt. Allerdings ist hier das vorrangige Ziel, das Verhalten von Betrüger*innen vom Verhalten »normaler« Nutzer*innen zu unterscheiden.

4.1.5 Kundenindividuelle Lösungen

Oft stehen interessierten Unternehmen selbst viele Daten zur Verfügung, die zur Betrugserkennung notwendig wären. Hier liegen dann die Herausforderungen an drei wesentlichen Stellen: Erstens müssen die Daten aus verschiedenen Silos zusammengeführt werden, zweitens müssen die Daten auf spezielle Anforderungen hin ausgewertet und analysiert werden, und drittens muss eine geeignete Entscheidungslogik entwickelt werden. Hierbei bieten spezialisierte Beratungsunternehmen ihre Dienstleistungen zur Entwicklung von hochgradig zugeschnittenen Modellen und Lösungen an. Für solche individuellen Lösungen gibt es insbesondere in der Finanzwirtschaft großen Bedarf, z. B. zur Verhinderung von Geldwäsche und Terrorismusfinanzierung und für das Compliance-Monitoring (z. B. Befolgung von Compliance-Richtlinien durch Mitarbeiter*innen).

Die finale Bewertung müssen die Kunden treffen

Abschließend ist noch festzuhalten, dass in den meisten Fällen die Kunden der Lösungsanbieter die Entscheidung treffen müssen, wie mit einem verdächtigen Vorgang umzugehen ist. Es werden aus den hier genannten Lösungsansätzen Indikationen, Scores bzw. Wahrscheinlichkeiten abgeleitet und die Kund*innen müssen diese Anhaltspunkte und die weiteren damit zur Verfügung gestellten Informationen in ihre internen Prozesse integrieren, diese dann mit einer eigenen Gesamtentscheidungslogik final bewerten, und dann eine geeignete Reaktion auslösen.

4.2 Betrugserkennung und Identitätsmanagement

Im Zusammenhang mit dem Identitätsmanagement sollte noch zwischen der Verifikation von Kundenidentitäten und dem sogenannten Identitäts- und Accessmanagement unterschieden werden.

Beim Ersteren geht es um die Verifikation der Identität von natürlichen bzw. juristischen Personen.

Beim Letzteren geht es in erster Linie um das technische Management der Nutzerkonten von Endkund*innen und die Gewährung eines Zugangs zu Ressourcen eines Unternehmens. Der Grad einer notwendigen Verifikation der Identität der Nutzer*innen kann je nach Anwendungsfall stark variieren. Aber auch hier sollte unabhängig vom Grad der Verifikation von Nutzeridentitäten die betrügerische Erstellung »falscher« Nutzerkonten bzw. der betrügerische Zugriff auf Ressourcen und APIs eines Unternehmens verhindert werden. In diesen Bereich fällt damit im weiteren Sinne auch die Erkennung von Botnets, für die sich ebenfalls der Einsatz von KI anbietet.

Identitäts- und Accessmanagementlösungen von Drittanbietern geben Unternehmen im E-Commerce die Möglichkeit, sich auf ihr Kerngeschäft zu konzentrieren und die Identifikation und das Management der Konten ihrer Endkund*innen outzusourcen. Drittanbieter haben dann die Aufgabe, die Identitäten der Endkund*innen hinreichend zu überprüfen und Betrugsversuche zu unterbinden. Dabei wird im Wesentlichen auf die schon oben genannten Verfahren zurückgegriffen.

Bei den an dieser Studie beteiligten Unternehmen konzentrieren sich die meisten Anbieter auf die Überprüfung der Kundenidentitäten und nicht auf das Identitäts- und Accessmanagement. Ein Unternehmen bietet Lösungen in beiden Bereichen an.

Die beste Betrugsabwehr ist im Zusammenhang mit einer eindeutigen Identifikation der Nutzer*innen möglich. Viele Unternehmen, die ihre Betrugserkennung online automatisieren möchten, werden wahrscheinlich einen Mix von Lösungen brauchen.

4.3 Aktueller Stand

In diesem Abschnitt gehen wir auf mehrere wesentliche Aspekte ein: die Einsatzfelder von KI in der Betrugserkennung, die verwendeten KI-Methoden und Verfahren, die Datengrundlage und -bereitstellung, die Qualitätssicherung, die Art der Interaktion der Kunden mit den Lösungen der Anbieter und schließlich die Vorteile, die durch den Einsatz von KI erzielt werden.

4.3.1 Haupteinsatzfelder für die KI

Die Art und Weise des Einsatzes von KI hängt von der Art der genutzten Lösung ab. Im Wesentlichen werden von unseren Interviewpartner*innen die folgenden Verfahren eingesetzt:

»KI wird in allen Bereichen der Betrugsprävention eingesetzt. Die Trennschärfe der Prognosen wird so erhöht. Je dynamischer die Phänomene und das Verhalten, umso relevanter wird die KI. Überall dort, wo eine Vielzahl von Daten verarbeitet werden muss und bei sich schnell bewegendem Phänomenen ist der Wert von KI hoch.«

Dr. Frank Schlein, CEO, CRIF Bürgel GmbH

Verfahren zur Optimierung von Regelsystemen

Regelsysteme werden oft vorteilhaft zur Betrugserkennung eingesetzt, da man so immer einen präzisen Überblick über die aktiven Regeln hat und Regeln gegebenenfalls einfach und ohne Rekalisierung des Systems »abschalten« kann. Betrugsereien bauen oft opportunistisch auf kurzfristig verfügbaren Hilfsmitteln oder Schwachpunkten auf und sind daher zeitlich begrenzt. Es ist daher sinnvoll, wenn man Regeln leicht wieder aussetzen kann.

Ein Problem ist, dass oft eindeutige Definitionen für die verschiedenen Betrugsmuster fehlen, was die Definition von geeigneten Labels und Zielparametern erschwert – mit negativen Auswirkungen auf die Trennschärfe.

Insbesondere unterstützt KI die:

1. Identifizierung von trennscharfen Parametern und Verarbeitung dieser Parameter in Regelsystemen
2. Optimierung von Regelsystemen, d. h. die Gewichtung von Regeln abhängig vom Einsatzbereich

Abbildung 3: Hauptaufgaben der KI in der Betrugserkennung.

Modellbasierte Ansätze zur Identifikation von neuen Betrugsmustern

Diese Verfahren werden meistens bei kundenindividuellen Lösungen eingesetzt und oft mit regelbasierten Ansätzen kombiniert. In der Praxis sind diese Verfahren oft nicht leicht zu realisieren. Der Einsatz von selbstlernenden Erkennungsverfahren ist insbesondere auch immer dann problematisch, wenn die verfügbaren Datenmengen nicht groß genug sind, wenn Betrüger*innen die Möglichkeit haben, weitgehend verfälschte Daten einzugeben und wenn man nicht warten kann, bis man genügend Fälle zur Klassifikation hat. Daher bewähren sich solche Verfahren insbesondere bei der Findung von neuen Betrugsmustern im Transaktionsbereich mit langen Transaktionshistorien und weniger im Antragsbereich.

Bildverarbeitende Ansätze

Zur Beurteilung der Echtheit von Hologrammen und anderen Sicherheitsmerkmalen auf Ausweisdokumenten und zur Liveness detection, d. h. zur Beurteilung, ob es sich bei der Person, die ihr Dokument vorweist, um eine lebende Person handelt.

Weitere Ansätze

KI wird zudem in vielen weiteren Bereichen wie der automatisierten Textverarbeitung (z. B. zur Analyse von Texten im Zusammenhang von Transaktionen), der Erkennung unerlaubter Zugriffe auf APIs (z. B. unter Nutzung falscher Identitäten) und zur Erkennung von Botnets verwendet.

4.3.2 Kurze Übersicht der verwendeten Methoden

Es gibt eine Vielzahl von verschiedenen Algorithmen und Lernmethoden, die für die Datenanalyse, insbesondere für die Erkennung von Anomalien, verwendet oder angewendet wurden. Da es eine umfangreiche Anzahl von Publikationen zum Thema Maschinelles Lernen gibt, waren die folgenden Publikationen nützlich, um die Methoden des Maschinellen Lernens und ihre Anwendungen im Allgemeinen und auf Identitätsbetrug zu überprüfen; [5]–[11] Die für die Betrugserkennung wichtigen maschinellen Lernmethoden werden typischerweise in Supervised Learning (überwachtes Lernen), Unsupervised Learning (unüberwachtes Lernen) und Lernmethoden mit Artificial Neural Network (künstliches neuronales Netzwerk) unterteilt. Unsere Interviewpartner erwähnten auch den Einsatz von häufig verwendeten überwachten Methoden, wie z.B. Decision Trees. Andere Methoden, die von den befragten Personen häufig angewendet wurden, waren Random Forest, Gradient-Tree-Boosting, Ensemble-Methoden und Neural Networks, auf die wir im Folgenden näher eingehen.

»Wir unterteilen unsere Auswertung in Identitätsanalyse und Transaktionsanalyse, da sich diese konzeptionell unterscheiden und mit unterschiedlicher Methodik umgesetzt werden. Ziel der Identitätsanalyse ist eine möglichst instantane Identifikation von geklauten, gefälschten oder verfälschten Identitäten direkt bei der Registrierung/Anmeldung bzw. Einführung eines neuen Nutzers in das betroffene System. Die Modelle arbeiten hier mit dem zum Einführungszeitpunkt verfügbaren statischen Daten. Die für die Transaktionsanalyse angewandten Modelle unterscheiden sich insofern, dass hier mitunter auf eine bekannte Transaktionshistorie zur Anomalie-detektion zurückgegriffen werden kann. Das Problem hat dann entsprechend auch eine zeitliche Struktur.«

Todor Dobrikov, Partner, d-fine GmbH

Anwendungen von Supervised-Learning-Methoden

Beim Supervised Learning wird die KI mit Datensätzen von Eingabe-Ausgabe-Paaren trainiert, damit die KI die Abbildung neuer Eingabedaten auf eine erwartete Ausgabe oder Klasse lernen kann. Der Trainingsdatensatz ist bereits »gelabelt«, d.h. jeder Eingabe ist eine Klasse zugewiesen. Die Genauigkeit des Mappings wird mit einer Verlustfunktion gemessen und so lange angepasst, bis der Fehlerwert minimiert ist. Es gibt zwei Arten von Problemen, auf die Methoden des Supervised Learning hauptsächlich angewendet werden: Klassifizierung und Regression.

Bei der Klassifikation werden Algorithmen verwendet, um Testdaten bestimmten Kategorien zuzuordnen, d.h. um die spezifischen Entitäten in einem Datensatz zu erkennen und dann Schlussfolgerungen zu ziehen, wie sie »gelabelt« werden sollten. Es existieren zahlreiche Algo-

rithmen, die diese Methode unterstützen und die von den Teilnehmenden unserer Studie verwendet wurden, wie z. B. Support Vector Machines (SVM), Decision Trees, k-Nearest Neighbor, Linear Classifiers und Random Forest.

Das Regressionsproblem betrachtet die Beziehung zwischen abhängigen und unabhängigen Variablen. Die Supervised-Learning-Algorithmen, die für dieses Problem verwendet werden, sind in der Regel lineare, logistische und polynomiale Regressionen.

Eine weitere gängige Methode, die beim Supervised Learning eingesetzt wird, ist die »Gradient Boosting Method«, eine Technik, die sowohl bei Klassifikations- als auch bei Regressionsproblemen angewendet wird. Einige der interviewten Personen erwähnten, dass sie Decision Trees anwenden, um vorherzusagen, ob eine Transaktion normal oder betrügerisch ist, indem sie Klassifikatoren verwenden, um normale Transaktionen von betrügerischen zu unterscheiden. Die Stärke der Verwendung von Decision Trees, Forest oder Classification And Regression Trees (CART) liegt darin, dass sie einfach zu implementieren und zu verstehen sind [8]. Außerdem erfordert das Training und ihr Betrieb eine relativ bescheidene Rechenleistung, was für den Echtzeitbetrieb ein Vorteil ist. Als Einschränkung wurde identifiziert [8], dass diese Methode ein ständiges Retraining erfordern kann und dass die Optimierungen, die während der Erstimplementierung vorgenommen werden, eine hohe Rechenleistung erfordern.

In Bezug auf die Techniken für Regressionsprobleme erläuterte ein Unternehmen, wie es Prognosemodelle und Regressionsmodelle für sequenzielle Daten auf Transaktionshistorien anwendet. So werden die Modelle in Bezug auf sequenzielle Daten, z. B. die Transaktionshistorie, trainiert. Darüber hinaus wurden auch kundenspezifische Modellierungstechniken zur Anomalieerkennung eingesetzt. Diese Techniken sind besonders gut einsetzbar, wenn ein Kunde eine längere Transaktionshistorie zur Verfügung stellen kann.

In Anbetracht der Techniken, die sowohl für Regressions- als auch für Klassifikationsprobleme verwendet werden, waren Ensemble-Methoden eine häufige Antwort unserer Befragten für Supervised Methods, die für die Betrugserkennung verwendet werden, insbesondere für die Mustererkennung. Ensemble-Methoden kombinieren verschiedene Modelle, um ein optimiertes Vorhersagemodell zu erstellen. Zu den Ensemble-Methoden, die in unserer Studie am häufigsten genannt wurden, gehörten die Random-Forest-Techniken. Eine weitere Methode, die von unseren Interviewpartner*innen mehrfach genannt wurde, war die Gradient-Tree-Boosting-Methode. Zum Beispiel erwähnte d-fine, dass die Gradient-Tree-Boosting-Methode auf eindimensionale Datensätze angewandt wurde (z. B. bei der Ersteinführung des Datensatzes einer Kundin oder eines Kunden in ein System). Andere Unternehmen, die den Einsatz dieser Methode erwähnten, waren SCHUFA und CRIF Bürgel.

Anwendung von Unsupervised-Learning-Methoden

Unsupervised Learning verwendet nicht gelabelte Datensätze, wobei das Ziel darin besteht, Muster zu entdecken, die bei der Lösung von Clustern oder Assoziationsproblemen helfen. Die Algorithmen, die hier verwendet werden, sind hierarchische, k-means- und Gaußsche Mix-Modelle. Der Hauptvorteil von Unsupervised-Learning-Methoden ist, dass es nicht notwendig ist, »gelabelte« Daten bereitzustellen, um Muster mit einer multivariaten Verteilung der Eingabe erkennen zu können. So ist es möglich, Transaktionen in Kategorien wie betrügerisch und nicht betrügerisch (d-fine) zu klassifizieren. Ein weiterer wichtiger Faktor, den es zu berücksichtigen gilt, ist die nicht spezifische oder segmentspezifische Modellierung. d-fine hob hervor, dass globale Modelle trainiert werden, um zu erkennen, welche Transaktionen von Endkund*innen normal oder eine Anomalie in Bezug auf die Transaktionshistorie ähnlicher Endkund*innen sind. Hier ist das Clustern von Endkund*innen in Gruppen über ein Cluster Model besonders wichtig, vor allem im Hinblick auf die ersten Transaktionen von neuen Endkund*innen.

Neural Network Methods

Eine weitere wichtige Gruppe von Methoden, die in den Interviews häufig genannt wurde, sind die Neural-Network-Methoden. Neural-Network-Methoden werden als eine Teilmenge des Machine Learning betrachtet und gehören zu den Deep-Learning-Algorithmen. Neural-Network-Methoden können auch Supervised-, Unsupervised-, und Reinforcement-Methoden sein [6]. In einem Betrugskontext funktioniert die Anwendung eines künstlichen Neural-Network-Algorithmus im Rahmen einer überwachten Klassifizierung gut [7], da er Merkmale rechtzeitig erkennen und Vorhersagen treffen kann. Darüber hinaus haben die Neuronal Networks in der Betrugserkennung ihre Stärken in Bezug auf ihre Nachhaltigkeit bei anderen nichtalgorithmischen, binären Klassifikationsproblemen [8]. Eine Einschränkung jedoch bedeutet, dass für die Anwendung dieser Methode eine hohe Rechenleistung für das Training und den Betrieb erforderlich ist. Dies kann einen Einsatz dieser Methode behindern und insbesondere für Echtzeitfunktionen ungeeignet machen. Darüber hinaus erfordert diese Methode ständige Rekalibrierungen und Trainings, um die Anpassung an neue Betrugsmethoden zu gewährleisten [8].

Der Einsatz von Neural-Network-Algorithmen in Betrugserkennungslösungen wurde von vielen interviewten Personen hervorhoben. In einem Interview wurde die Anwendung von Neural Networks beim Unsupervised Learning zur Erkennung von Anomalien erwähnt. Andere Befragte erwähnen den Einsatz von Recurrent Neural Networks für die textbasierte Analyse.

4.3.3 Datengrundlage und -bereitstellung

Die verwendeten Daten sind abhängig vom Anwendungsfall und verwandten Verfahren. Oft werden Daten von den Lösungsanbietern für Analysen auch subscriptionsübergreifend zusammengeführt.

Fernidentifikationen von natürlichen Personen

Hier sind die relevanten Daten die vom Unternehmenskunden übermittelten zu überprüfenden Endkundendaten wie Name, Geburtsdatum und Adresse sowie die Bilddaten, die der Video-stream zur Fernidentifikation liefert. Insbesondere die Daten des Videostreams werden auch mithilfe von KI analysiert.

Als Trainingsdaten werde u. a. gefälschte Ausweise, die von den Landeskriminalämtern zur Verfügung gestellt werden, genutzt. Zudem erfolgt ein Abgleich mit Datenbanken, in denen Ausweisdaten, die betrügerisch genutzt worden sind, gelistet werden.

Zudem werden Verfahren zur Geräteerkennung eingesetzt (siehe unten). Insbesondere wird analysiert, ob das benutzte Gerät schon für Betrugsversuche eingesetzt war und ob mit dem Gerät auffällig viele Fernidentifikationen in einem sehr kurzen Zeitraum durchgeführt wurden.

Geräteerkennung

Auf die für die Geräteerkennung notwendigen Daten wurde schon im Abschnitt 4.1.2 eingegangen. Eine der meistgenutzten Datenquellen ist der Browser und insbesondere die Daten, die über den Browser zugänglich sind. Auch IP-Adressen und andere Daten, die Rückschlüsse auf den Aufenthaltsort des Geräts erlauben, werden in diesem Zusammenhang genutzt.

Hinzu kommen Validierungsdaten, mit denen man überprüft, ob sich ein Gerät, dass sich z. B. als iPhone® ausgibt, bei einem Test auch wie ein iPhone® verhält.

Bei der Geräteerkennung ist zudem die Größe und Aktualität der Datenbank wichtige Kriterien, die der Lösungsanbieter zum Abgleich der Geräteerkennung mit Gerätezeichnungen, mit denen Betrugsversuche unternommen wurden, zur Verfügung hat.

Lösungen zum Datenabgleich mit gepoolten Daten/großen Datenbanken

Bei diesen Verfahren werden Antrags- bzw. Transaktionsdaten gesammelt und verarbeitet. Um Antragsdaten zu validieren und zu überprüfen, werden in der Regel alle verfügbaren Daten genutzt, u. a. Name, Adressinformationen von Privatpersonen oder Unternehmen, Geburtsdaten, E-Mail-Adressen, Telefonnummern, Ortsangaben, Geolokationsdaten z. B. von Häusern und Gebäuden, aber auch Produktdaten, der Warenkorb oder das Registrierungsdatum.

Die Ressourcen von Betrüger*innen sind begrenzt, daher werden beispielsweise Telefonnummern öfter wiederverwendet. Dies macht die Integration dieser personenbezogenen Daten in den Betrugserkennungsprozess notwendig.

Betrüger*innen manipulieren Daten oft inkremental, sodass sie bei einem Abgleich mit Auskunftsteilen als unbekannte neue Kund*innen erscheinen, über die keine positiven oder negativen Informationen vorliegen. Daher sind Änderungshistorien der Daten von Kund*innen und vor allem Adressdaten sehr wichtig und nützlich. Teilweise werden diese Änderungen erst in den gepoolten Daten der Lösungsanbietenden über die Zeit verfolgbar und sichtbar.

Behavioral Biometrics

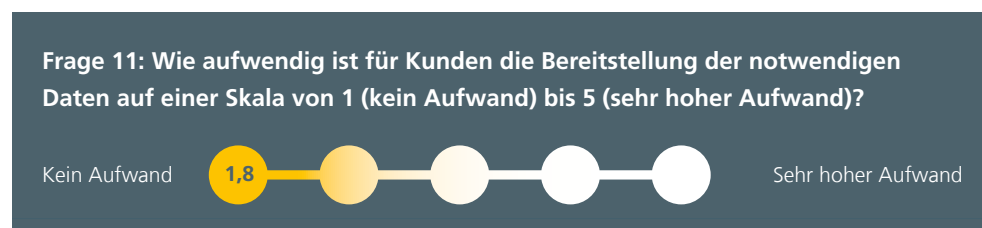
Hier werden Hintergrundinformationen und Informationen zum Verhalten von Nutzer*innen, z. B. in digitalen Kanälen für Einkaufs- und Vertragsabschlüsse gewonnen und ausgewertet. Dazu gehören Marketingdaten wie Absprungpunkte, Klickraten oder typische Verweildauern. Diese Daten werden dann mit Nutzernamen, Gerätemerkmalen, E-Mail-Adressen oder auch Telefonnummern assoziiert und über längere Zeiträume hin ausgewertet.

Kundenindividuelle Softwareentwicklung

Bei der kundenindividuellen Softwareentwicklung liegt ein besonderes Interesse auf der Transaktionsbewertung. Dazu werden insbesondere Transaktionsbeschreibungen (Natural Language Processing), Transaktionsarten, -volumen, -höhe, -frequenzen, und -historien berücksichtigt. Zur Erkennung von verdächtigen Mustern und Anomalien sind vor allem Zeitverläufe sehr wichtig.

Der Aufwand für die Bereitstellung der Daten wird von fast allen Anbietern – zumindest nach der Erstimplementierung – im Regelbetrieb automatisiert und damit als gering angesehen. Meist werden dafür APIs eingesetzt. Einige Unternehmen bieten eine reine Cloud-Lösung (oft auf eigenen Servern gehostet) an, andere bieten sowohl Cloud-, als auch On-premise- oder Hybridlösungen je nach Kundenwunsch an. Wenn Daten kundenübergreifend gepoolt werden, ist eine reine On-premise-Lösung natürlich nicht umzusetzen. Kundenindividuelle Lösungen werden meist als On-premise-Lösung entwickelt.

Abbildung 4: Ergebnisse bezüglich des Aufwands für die Bereitstellung der Daten von den Kunden.



Bei Erstimplementierungen und -analysen ist der Aufwand abhängig von den verwendeten Daten und der IT-Infrastruktur der Unternehmenskunden. Da Betrugsmethoden sich immer weiterentwickeln und verändern, ändern sich allerdings auch die Anforderungen an Datenstruktur und Qualität kontinuierlich. Daher ist in der Regel eine andauernde fachliche Auseinandersetzung aller Beteiligten mit dem System und den Datenanforderungen erforderlich. Auf diesen Aspekt gehen wir im nächsten Abschnitt zur Qualitätskontrolle noch einmal gesondert ein.

4.3.4 Qualitätssicherung

Da – abgesehen von einem Grundrauschen – Betrug organisiert stattfindet und sich Betrugsmuster fortlaufend ändern und oft komplex und nicht einfach zu klassifizieren sind, ist Betrugsprävention ein fortlaufender Prozess. Unternehmenskunden müssen sich hier in der Regel aktiv mit einbringen, wenn sie die bestmöglichen Ergebnisse erzielen möchten.

Wichtig ist es dabei, immer im Auge zu behalten, dass ein Verdachtsfall nicht notwendigerweise ein Betrugsfall ist: Endkund*innen sollte nicht ungerechtfertigterweise ein Betrugsversuch unterstellt werden.

Daher müssen insbesondere unsichere Fälle intensiv analysiert werden. Bei der SCHUFA schauen z. B. mehrere Expert*innen mit unterschiedlichen Expertisen und Hintergründen gleichzeitig über unsichere Fälle, um das menschliche Bias so weit wie möglich zu reduzieren. Dann wird eine gewichtete Mehrheitsentscheidung getroffen, und daraus werden dann Labels generiert. So lassen sich unsichere Fälle weiter adjustieren. Dieses Vorgehen basiert auf einem Verfahren des Active Learnings, das als Qualitätssicherungsprozess kontinuierlich mitläuft.

Auch bei allen weiteren Anbietern wird Qualitätssicherung – oft in enger Abstimmung mit den Kunden – großgeschrieben. Dazu gehören: Bias-Reduktion, kontinuierliches Sampling von unsicheren Fällen, manuelle Überprüfung, Rückführung des Feedbacks von Expert*innen in das System, und ein enger und regelmäßiger, teilweise wöchentlicher Austausch zur weiteren Verbesserung und Optimierung der Prognosen mit den Kunden.

Eine regelmäßige Kalibrierung und Suche nach neuen Regeln des Systems wird als sehr sinnvoll erachtet – wie und wie oft, hängt von den Angriffsvektoren und Anwendungsfällen ab. Nach Einschätzung von RISK IDENT müssen hierbei vor allem folgende Aspekte berücksichtigt werden:

- Veralterte Angriffsmuster aussortieren
- Neue Angriffsmuster erkennen und integrieren
- Neue Datenquellen und Erfahrungswerte integrieren.

Teilweise können diese Prozesse durch MLOps (automatisiertes Training, Setzen von Triggern) vereinfacht werden. Eine fachliche Validierung muss trotzdem immer vorgenommen werden, im Bankensektor gibt es auch regulatorische Anforderungen, die in diesem Zusammenhang zu berücksichtigen sind.

Insbesondere erstreckt sich die fachliche Validierung auf neue Regeln, die automatisch generiert werden. Diese können nicht ohne fachliche Prüfung übernommen und eingesetzt werden. Dabei geht es nicht nur um den Trennschärfezugewinn bei der Entscheidung, sondern die Unternehmen müssen vor dem Einsatz neuer Regeln immer auch ihre Policies, Risikofreudigkeit und die rechtlichen, prozessualen und technischen Rahmenbedingungen berücksichtigen.

Während der technische Wartungsaufwand für Lösungen zur Betrugserkennung nach der Erstimplementation als sehr gering eingeschätzt wird, wird der mit der notwendigen Qualitätssicherung verbundene fachliche Wartungsaufwand als deutlich höher bewertet.

4.3.5 Kundeninteraktion mit der Lösung und Umgang mit Verdachtsfällen

Bei den meisten Lösungsanbietern liegt die endgültige Entscheidung, wie mit einem Betrugsverdachtsfall umzugehen ist, bei den Kunden. Die Ausnahme gibt es bei den Lösungen, bei denen diese ihr gesamtes Management der Endkund*innen online an den Lösungsanbieter outsourcen. Hier bekommen nur die Endkund*innen gegebenenfalls eine Information.

In der Praxis heißt dies, dass den Unternehmen automatisch – in der Regel in Echtzeit über Schnittstellen/APIs – Indikationswerte und Scores, die eine Einschätzung des Betrugsrisikos widerspiegeln, zur Verfügung gestellt werden. Sie integrieren dann diese Informationen in ihre Prozesse und analytischen Verfahren. Basierend auf den Scores, der fachlichen Beratung bei der Implementierung und unter Abwägung ihrer Policies müssen sie dann die Verdachtsfälle selbstständig prüfen und unabhängig eine Entscheidung treffen. Der Automatisierungsgrad ist auch hier in der Regel sehr hoch, allerdings kann in Einzelfällen eine manuelle Prüfung (»manual order review«) notwendig werden.

Wenn eine manuelle Prüfung notwendig ist, werden den Kunden oft weitere Hintergrundinformationen über Dashboards zugänglich gemacht. Dazu gehören dann z. B. Hinweise auf die getroffenen Regeln, die zur Aussteuerung geführt haben.

Ein weiterer wichtiger Aspekt in diesem Zusammenhang ist die Information der Endkund*innen. Im E-Commerce erfahren Endkund*innen oft nur, ob sie in Vorkasse treten müssen oder auf Rechnung bestellen können. Bei Verdachtsfällen von Identitätsdiebstahl, d. h. wenn die Identität von Endkund*innen kompromittiert wurde, sollte aber auch eine Information der Betroffenen erfolgen, wenn nötig über andere Kommunikationswege (z. B. per Telefon statt Onlinekontakt).

In manchen Situationen lösen Endkund*innen einen Betrugsverdacht auch unbeabsichtigt aus. NECT berichtet, dass in einigen Fällen Endkund*innen ohne Betrugsabsicht Ausweiskopien auf Papier für die Fernidentifikation vorweisen. In diesen Fällen erfolgt eine Differenzierung manuell, die Endkund*innen werden entsprechend informiert und gebeten, die Fernidentifikation mit dem Ausweisdokument im Original noch einmal zu durchlaufen.

4.3.6 Vorteile durch den KI-Einsatz

Alle befragten Unternehmen bestätigen, dass der Einsatz von KI die Trennschärfe der Prognosen signifikant erhöht. Für deren Unternehmenskunden ist die »figure of merit« in der Regel, wie viel sicheres (Neu)geschäft durch den Einsatz von KI zusätzlich ermöglicht wird. Die KI trägt damit im Wesentlichen zu einer Reduktion der »false positives« (unbegründete Verdachtsfälle) bei.

Die KI funktioniert immer dann gut, je dynamischer die Phänomene und das Verhalten sind, und wo eine Vielzahl von Daten verarbeitet werden muss und kann. Die Performance von KI ist daher im Transaktionsbereich oft besser als im Antragsbereich. Bei Unternehmen mit hohen Antragsaufkommen (> 1 000 000 p. a.) und entsprechenden Anzahlen an Betrugsversuchen erhöht sich der Nutzen von KI dementsprechend deutlich.

Eine quantitative Einschätzung der Verbesserung ist nur schwer möglich und hängt von Angriffsszenarien bzw. -mustern einerseits und der verwendbaren Datengrundlage, Prozessen und existierenden Regeln bei Kund*innen andererseits ab. Zudem benötigt man genügend bestätigte Betrugsfälle für eine Analyse. Nichtsdestotrotz entnehmen wir den Gesprächen, dass sich durch KI eine Verbesserung der Trennschärfe im Antragsbereich im einstelligen, in Einzelfällen im niedrigen zweistelligen Prozentbereich erwarten lässt. Im Transaktionsbereich sind die Zahlen in der Regel höher.

Im Vergleich von generischen mit individuellen KI-Lösungen für Unternehmen ist der KI-Impact in beiden Fällen signifikant, aber bei der individuellen Modellierung für sie in der Regel noch stärker ausgeprägt. RISK IDENT weist aber darauf hin, dass eine überraschend gute Verbesserung bei den Unternehmen oft dafür spricht, dass man ein einzelnes spezifisches Angriffsmuster sehr gut erkannt hat. Dies könnte man aber auch durch eine einzelne entsprechende Regel – ganz ohne KI – abbilden. Wenn KI zur Regeloptimierung und Verbesserung der Gewichtung benutzt wird, erhöht sich die Trennschärfe deutlich, aber nicht so drastisch.

Bei der Fernidentifikation ist die Situation etwas anders. Hier ist die »figure of merit« der Vergleich mit den Fähigkeiten ausgebildeter Grenzbeamten*innen. Auch hier sind quantitative Angaben schwierig.

»Beim Gesichtsabgleich hat die KI eine geringere Fehlerquote als ein Grenzbeamter.«

Benny Bennet Jürgens, CEO, Founder, NECT GmbH

Dies gilt auch für die Überprüfung der Sicherheitsmerkmale, und »liveliness detection« (beim Video-Ident). Schwierig für die KI ist noch die Kontextwahrnehmung (das »Bauchgefühl« bzw. die Beurteilung der Schlüssigkeit der Gesamtsituation). Allerdings wird in den entsprechenden Studien [8] die KI mit der Performance von motivierten, ausgeruhten Expert*innen verglichen. Diese Studiensituation sei aber deutlich vom Einsatz im täglichen Betrieb zu unterscheiden.

4.4 KI – Potenzial und Herausforderungen

In diesem Abschnitt werden die Studienergebnisse in Bezug auf das Potenzial und die Herausforderungen, die der Einsatz von KI in der Betrugserkennung mit sich bringen wird, vorgestellt. Zunächst wird erläutert, wie die Befragten das zukünftige Potenzial von KI in der Betrugserkennung einschätzen. Der zweiten Unterabschnitt ist den größten Herausforderungen für einen weitergehenden Einsatz von KI gewidmet. Das letzte Unterkapitel schließt das Kapitel mit Einschätzungen ab, wie der KI-Einsatz weiter verbessert und erleichtert werden kann.

4.4.1 Zukünftiges Potenzial

Auf die Frage »Wie schätzen Sie das Potenzial durch KI in der Betrugserkennung für die Zukunft ein?« äußerten die Befragten übereinstimmend, dass das Potenzial für KI sehr hoch ist. Insbesondere im Zusammenhang mit der stetig fortschreitenden Digitalisierung und der

zunehmenden Bedeutung des Onlinehandels müsse auch die Effizienz in der Betrugserkennung skalieren, und daher sei eine KI-basierte weitergehende Automatisierung und Digitalisierung der Betrugserkennung zwingend erforderlich.

Ein Teilnehmer sagte, dass der KI-Einsatz mittlerweile »mehr oder weniger ein Standard« sei, wenn die benötigten Daten und Kontextdaten verfügbar seien. Einige der identifizierten Vorteile einer KI-Integration waren eine weitergehende Prozessautomatisierung, höhere Trennungsgenauigkeit und die Integration weiterer Daten. Zusätzliches Potenzial liege in der Weiterentwicklung der eingesetzten Algorithmen, z. B. für die Zeitreihenanalyse und das Verstehen komplexerer Texte mit neuen Sprachmodellen. Auch wurde das Potenzial angesprochen, mithilfe von KI eine Vielzahl von Merkmalen aus verschiedenen Bereichen wie den eingesetzten Geräten, Verhaltensweisen und Personendaten in Beziehung setzen zu können, und diese dann mit weiteren Dimensionen wie Identitäts- oder Bonitätsprüfungen zu verbinden. Dies könne dabei helfen, sich dem kreativen und sich schnell ändernden Verhalten der Betrüger*innen besser stellen zu können. Allerdings betonte ein Befragter auch, dass KI – zumindest gegenwärtig – nicht für alle Angriffs- und Einsatzszenarien dasselbe hohe Potenzial habe.

Insgesamt ergibt sich, dass KI nicht nur schon heute für die Identitätsbetrugsprävention wichtig ist, sondern in Zukunft eine noch größere Rolle spielen wird.

»KI in der Betrugserkennung wird deutlich zunehmen. Aufgrund der immer stärkeren Digitalisierung nimmt der persönliche Kontakt zum Verbraucher ab. Gleichzeitig muss die Effizienz erhöht werden, um zu skalieren und damit werden automatisierte KI-Verfahren immer wichtiger.«

Prof. Dr. Gjergji Kasneci, CTO, SCHUFA AG

4.4.2 Die größten Herausforderungen

Die Expert*innen wurden gebeten, ihre Einschätzung zu der Frage »Was sind die größten Herausforderungen, mit denen Sie derzeit konfrontiert sind, wenn es um die Integration von KI-basierten Ansätzen in Ihren Produkten geht?« zu geben. Die Studie ergab, dass es die folgenden zentralen Herausforderungen gibt, mit denen der Einsatz von KI in der Praxis konfrontiert ist:

»Betrugsmuster sind sehr vielfältig. Diese können gut klassifiziert werden, dass wird jedoch selten gemacht. Damit werden verschiedene Zielstellungen vermisch. Selbstlernende Verfahren leiden unter dieser schlechten Datenqualität und verlieren an Trennschärfe.«

Dirk Mayer, Fraud Expert & Catalyst, RISK IDENT GmbH

Datenqualität für Supervised Learning

Die Datenqualität ist oft problematisch insbesondere in Bezug auf die Qualität der Label. Erstens ist das Dunkelfeld relativ hoch und es bleibt für Kunden oft unklar, ob eine Anomalie wirklich ein Betrugsfall ist. Zweitens ist nicht immer klar, welche Betrugsarten bei einem Vorfall dahinterstehen, eine genaue Differenzierung ist schwierig, auch weil es an guten und anerkannten Definitionen mangelt.

Datenqualität für Unsupervised Learning

Für die Erkennung neuer Betrugsmuster muss man oft Unsupervised-Modelle einsetzen, die ohne vorgegebene Label arbeiten. Dafür braucht man sehr viele Trainingsdaten, die oft fehlen bzw. inkonsistent sind. Der Datenzugang und die -zusammenführung ist ebenfalls oft schwierig, da verschiedene Datensilos und veraltete IT-Systeme angebunden werden müssen.

Akzeptanz der Verfahren durch Regulierungsbehörden

Man muss die Regulatoren wie z. B. BAFIN, BSI, Bundesnetzagentur davon überzeugen, dass die KI keine Blackbox, sondern vertrauenswürdig und nachvollziehbar bzw. überprüfbar ist und so die regulatorischen Anforderungen erfüllt. Kompetente Ansprechpartner*innen seien nicht immer in ausreichender Zahl verfügbar und es mangle noch an klaren regulatorischen Anforderungen für den KI-Einsatz in diesem Bereich. So sei es manchmal schwierig, Sachverhalte zügig klären zu können.

Im Bereich Datenschutz sei die Zusammenarbeit einfacher, weil sich relativ gut erklären lässt, dass der ethische, nicht-missbräuchliche Einsatz von KI im Vergleich zu klassischen Datenbanken bzw. manuellen Kontrollen zu einer Verbesserung führt.

Um Datenminimalität zu erreichen, sollten nur die wirklich statistisch relevanten Informationen verwendet werden, d. h. man möchte im Prinzip einen kausalen Zusammenhang mit der Vorhersage anstreben.

Insgesamt wird die Wichtigkeit der Zusammenarbeit mit den Regulierungsbehörden als sehr hoch eingeschätzt, und eine weitere regulatorische Harmonisierung sei – auch über die EU hinaus – sehr wünschenswert.

Kontinuierliche Qualitätssicherung

Die kontinuierliche Qualitätssicherung ist unentbehrlich, um die Fundiertheit und Robustheit zu garantieren. Allerdings wären ressourcenschonendere Lösungsansätze wünschenswert.

Kosten

Um KI-Modelle zu trainieren und auszuführen, braucht man GPU-basierte Serversysteme, deren Anschaffung und Betrieb (hohe Energieaufnahme) immer noch sehr kostspielig sind. In Produkten, in denen die Betrugserkennung eine Funktionalität unter anderen ist, stehen die hohen Kosten einem weitergehenden Einsatz von KI teilweise im Wege.

Darüber hinaus wurden den Befragten eine Reihe von Aussagen zum Potenzial und den verschiedenen Herausforderungen im Zusammenhang mit der Anwendung von KI in der Betrugserkennung vorgelegt. Die Befragten sollten ihre Ablehnung bzw. Zustimmung auf einer Skala von 1 (starke Ablehnung) bis 5 (starke Zustimmung) angeben. Die Durchschnittswerte sind in Abbildung 5 zu sehen.

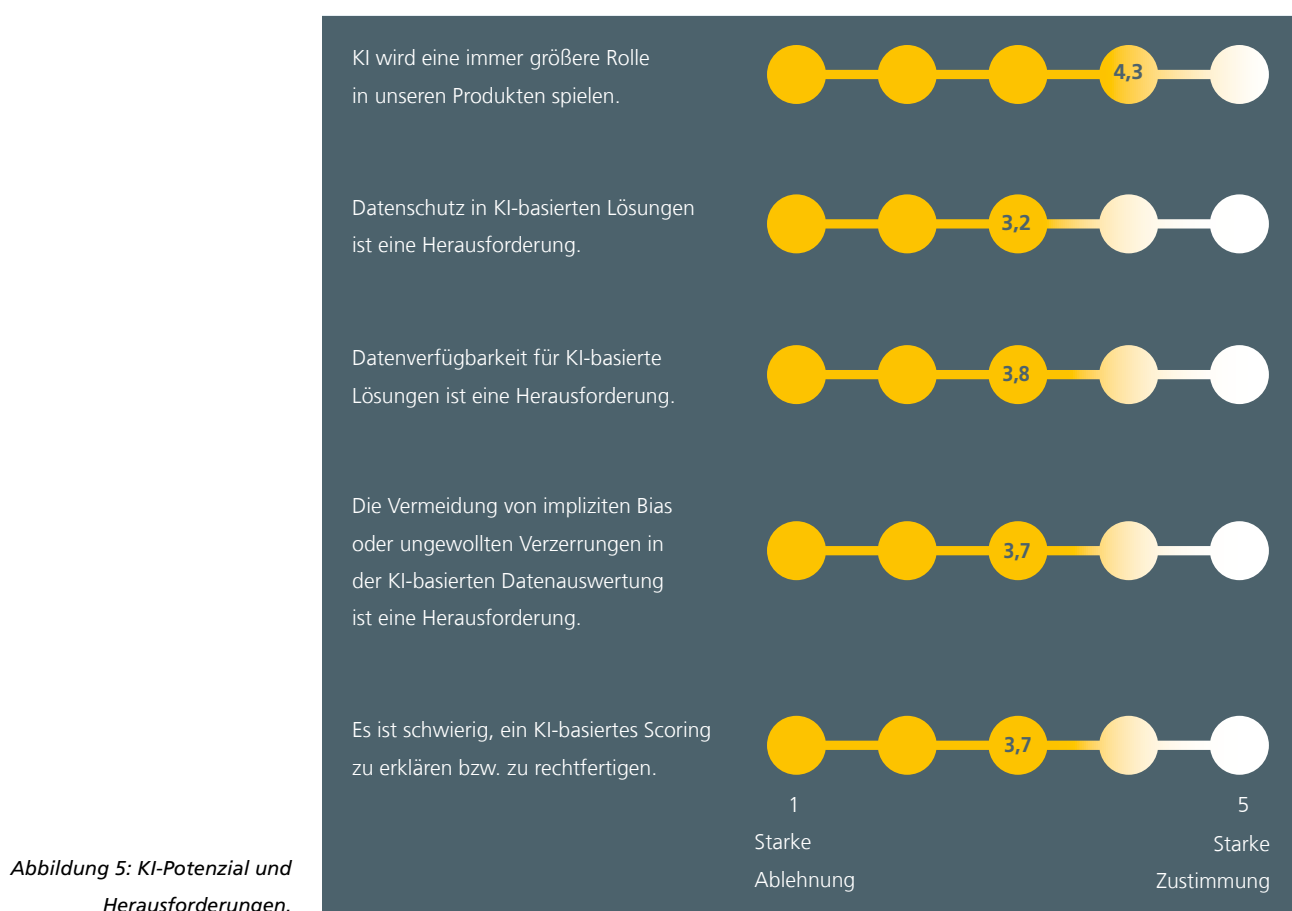
Die Herausforderungen, nach denen gefragt wurde, waren: Datenschutz, Datenverfügbarkeit, implizites Bias und die Erklärbarkeit einer KI-basierten Bewertung. Die Befragten stimmten darin überein, dass die Rolle von KI in ihren Produkten größer werden wird (durchschnittliche Bewertung: 4,3). In Bezug auf den Datenschutz äußerten die Befragten, dass es zwar Herausforderungen gebe, diese aber relativ gut beherrschbar seien. Ein Teilnehmer kommentierte, dass es eigentlich keine Problem gebe, solange die Daten ausschließlich zur Betrugserkennung verwendet werden (durchschnittliche Einschätzung: 3,2).

Die Datenverfügbarkeit wurde als ein Problem für KI-Lösungen eingestuft, wie an einer durchschnittlichen Bewertung von 3,8 zu sehen ist. Einige Befragte nannten als bestehende Herausforderungen bei der Datenverfügbarkeit eine schlechte Datenqualität, die Datenstrukturierung und die »labels«.

Für Herausforderungen im Zusammenhang mit implizitem Bias oder unbeabsichtigten Verzerrungen bei der KI-Datenanalyse gaben die Befragten einen Wert von 3,7 an. Dies unterstreicht die Bedeutung einer guten Qualitätssicherung.

Auch die Herausforderung, KI-basierte Bewertungen erklären und rechtfertigen zu können, wurde von den Befragten im Durchschnitt mit 3,7 von 5 Punkten bewertet. In Bezug auf erklär-bare KI drückten einige Befragte aus, dass dies im Moment noch ein sehr wichtiges und komplexes Thema sei. Es gebe aber aus der Forschung immer bessere Lösungsansätze, und tendenziell werde die Bedeutung dieser Herausforderung abnehmen. Im Wesentlichen bestehe die Herausforderung auch nur im Hinblick auf Regulierungsbehörden. Forscher*innen würden die Modelle im Detail verstehen und könnten den entsprechenden Argumenten gut folgen. Für das Marketing und die Kunden gehe es im Kern darum, dass die Modelle funktionieren und nicht

wirklich, wie sie funktionierten, und warum. Im Umgang mit den Regulierungsbehörden bestehe die Herausforderung, die verwendeten Modelle und Ansätze im Detail erklären zu müssen.



4.4.3 Verbesserungspotenzial

Zum Abschluss fragten wir die Expert*innen, wo sie zukünftiges Verbesserungspotenzial für den Einsatz von KI in der Betrugserkennung sehen. Dies steht im engen Bezug zu den identifizierten Herausforderungen:

Allgemeingültige, trennscharfe Definitionen für verschiedene Betrugsarten

Eine bessere Klassifizierung von Betrugsarten kann dazu beitragen, bessere und eindeutige »label« zur Verfügung zu haben und so die Erkennung von Verdachtsfällen zu verbessern. Zudem kann dies helfen, geeignetere und klarer abgegrenzte Reaktionen auf Verdachts- und Betrugsfälle zu triggern.

Datenqualität

Hier gibt es nach Ansicht einiger der Befragten ein klares Verbesserungspotenzial hinsichtlich Datensammlung, Datenbeschaffung, Datenaufbereitung und -berichtigung. E-Commerce-Player in der Wirtschaft und die Finanzwirtschaft sollten ihre tatsächlichen Betrugsfälle zusammenbringen und der Gemeinschaft zur Verfügung stellen, damit sich auf Basis dieser Fälle immer bessere Verfahren antrainieren lassen.

Einfachere Regulatorik und weitergehende regulatorische Harmonisierung

Die DSGVO habe einen wichtigen Beitrag zur Harmonisierung von Datenschutzerfordernungen geleistet. Eine weitergehende Harmonisierung sei auch im Hinblick auf den Einsatz von KI wünschenswert, auch wenn dies eine Herausforderung für alle Beteiligten sei. Reallabore und agile Regulatorik könnten hier einen sinnvollen Beitrag leisten.

Kostenreduzierung und bessere Hardware

Es wäre wünschenswert, wenn KI-Modelle kosteneffizienter erstellt werden könnten, einerseits in Bezug auf Infrastruktur und Hardware. Andererseits wäre eine bessere Verknüpfung zwischen Softwarearchitektur und Data Science wünschenswert, um Skalierbarkeit sicherstellen zu können, z. B. durch eine größere Verfügbarkeit von entsprechend geschultem Personal.

- ▶ **Allgemeingültige, trennscharfe Definition für verschiedene Betrugsarten**
- ▶ **Datenqualität**
- ▶ **Einfachere Regulatorik und weitergehende regulatorische Harmonisierung**
- ▶ **Kostenreduzierung und bessere Hardware**

Abbildung 6: Verbesserungspotenzial.

5 ZUSAMMENFASSUNG

Die vorliegende Studie zeigt, dass schon heute die effektive Prävention von Identitäts- und Verbraucherbetrug ohne den unterstützenden Einsatz von KI nicht mehr denkbar ist und KI eine immer größere Rolle in den auf dem Markt verfügbaren Dienstleistungen und Produkten spielen wird.

Die auf dem Markt verfügbaren Lösungen spiegeln die Vielfalt der Betrugsmuster wieder und werden oft in Kombination eingesetzt. Dementsprechend ist auch der Einsatz von KI sehr vielfältig und reicht von Text- über Bild- bis zur Anomalieerkennung und der Optimierung von Regelsystemen.

Für die Lösungsanbieter sind mittlerweile die meisten Herausforderungen durch den KI-Einsatz in den Bereichen Datenschutz, Biasreduktion und Erklärbarkeit durch intensive Qualitätssicherung gut beherrschbar. Vor allem ein systematischerer und wissenschaftlicherer Ansatz zur Klassifikation von Betrugsmustern sowie eine weitergehende regulatorische Harmonisierung können den Einsatz von KI in der Zukunft immer mehr vereinfachen. Um das Potenzial von KI noch umfassender ausschöpfen und um neue Betrugsmuster leichter erkennen zu können, werden unternehmensübergreifende branchenspezifische Datenpools an Bedeutung gewinnen.

Unternehmen, die entsprechende Produkte einsetzen möchten, sollten darauf vorbereitet sein, dass neben einmaligen Aufwänden zur Erschließung der notwendigen Daten und der Prozessintegration kontinuierlich an der Qualitätssicherung gearbeitet werden muss. Nur so kann das Verbesserungspotenzial durch den KI-Einsatz voll ausgeschöpft werden. Dafür lässt sich dann ein deutlicher Anstieg der Trennschärfe in der Beurteilung von Verdachtsfällen und damit ein signifikanter Anstieg an neuem, sicherem Geschäft erwarten.

6 STECKBRIEFE DER UNTERNEHMEN

Für den Umgang mit Identitätsbetrug gibt es eine Vielzahl von Lösungen auf dem Markt von standardisierten Produkten bis zu hoch kundenindividuellen Dienstleistungen. Viele Interessenten entscheiden sich für eine Kombination von verschiedenen Angeboten, um den eigenen Anforderungen gerecht werden zu können.

Im Folgenden stellen wir mit Steckbriefen die Unternehmen aus verschiedenen Marktsegmenten vor, die sich an dieser Studie beteiligt haben. Gezielt haben wir dazu Anbieter angesprochen, die mit verschiedenen Lösungsansätzen arbeiten. Wir betonen, dass diese Liste keinen Anspruch auf Vollständigkeit erhebt und es selbstverständlich weitere hoch qualifizierte Anbieter auf dem Markt gibt, die man bei der Auswahl einer geeigneten Lösung berücksichtigen sollte.

Die folgenden Fragen sollten sich Unternehmen stellen, um für sie geeignete Anbieter zu identifizieren:

- Wird eine Identitäts- und Accessmanagementlösung gesucht oder ist dafür schon eine Lösung vorhanden?
- Muss die Zuverlässigkeit der Kundenidentifikation regulatorischen Anforderungen genügen?
- Müssen die Identitäten von natürlichen und/oder juristischen Personen überprüft werden?
- Wie lassen sich die verschiedenen Lösungen in die bestehende IT-Landschaft und -Prozesse integrieren?
- Kann man eher von standardisierten Lösungen und dem Austausch bzw. dem Poolen von Daten über Unternehmensgrenzen hinweg profitieren oder gibt es besondere Anforderungen, Betrugsmuster bzw. Datengrundlagen, für die eine individuelle Lösung vorteilhafter ist?
- Ist eventuell ein hybrider Ansatz, also eine Kombination von standardisierten und individuellen Lösungen am besten geeignet?

6.1 CRIF Bürgel GmbH



CRIF Bürgel unterstützt Finanzinstitute und Unternehmen als Partner für Informationsmanagement mit integrierten und internationalen B2B2C-Lösungen aus einer Hand – für gesteigertes Unternehmenswachstum mittels datengestützter Entscheidungen. CRIF Bürgel gehört zur weltweit tätigen Gruppe CRIF mit Hauptsitz in Bologna. 35+ Länder, 70+ Unternehmen, 5500+ Expert*innen.

Produkte und Dienstleistungen

Produkte	Lösungen für integriertes Identitäts-, Kreditrisiko- und Betrugsmanagement über eine Schnittstelle sowie innovative Ansätze in den Bereichen Digitalisierung, Compliance, Adressermittlung, Marketing Services, Analytics und Consulting. DS-Portal: Eine Gemeinschaft von Unternehmen der Kreditwirtschaft tauscht über CRIF Bürgel Informationen aus, um sich gegenseitig gegen Betrug zu schützen.
Verfügbarkeit	■ Klassische Schnittstellenkonzepte mit einer technischen Integration ■ On premise ■ Apification (Benutzung von Microservices über APIs)
Abrechnung	Kein allgemeingültiges Prinzip. Die Gestaltung hängt von den Bedürfnissen und den Anwendungsbereichen des Kunden ab.
Reichweite des Unternehmens	Weltweit, unter anderem mit 40 Ländergesellschaften und darüber hinaus über Partnernetzwerke
Zielkunden	Alle Unternehmen mit einer Vielzahl von Beziehungen zu Konsument*innen und/oder weiteren Unternehmen, insbesondere: ■ E-Commerce ■ Banken ■ Telekommunikation ■ Energie ■ Versicherungen ■ Automobilindustrie

Besonderheiten	■ HYBRIGHT: Integrierte Identitäts- und Bonitätsprüfung und Betrugsabwehr über eine einzige Schnittstelle. ■ B2B2C-Konzept (d. h. sowohl für Unternehmen, deren Kunden Konsument*innen und/oder weitere Unternehmen sind).
Einsatz von KI	In allen Produkten, insbesondere dort, wo eine Vielzahl von Daten verarbeitet werden müssen.
Vorteile des KI-Einsatzes	■ Trennschärfezugewinn und damit die verbesserte Prognosefähigkeit ■ Fähigkeit, auf sich schnell verändernde Phänomene reagieren zu können ■ Unterschiedliche Anwendungsfelder in gemeinsamen Bezug zu bringen und so eine insgesamt bessere Entscheidung treffen zu können.
Kundeninteraktion mit dem Produkt	Technisch: Der Kunde ist über die Schnittstelle integriert und bekommt in Echtzeit Entscheidungsempfehlungen geliefert. Fachlich: Enger und regelmäßiger, teilweise wöchentlicher Austausch zur weiteren Verbesserung und Optimierung der Prognosen mit dem Kunden. Der Kunde entscheidet unabhängig, auf Basis der fachlichen Beratung und der Scores und in Abwägung seiner Policies und seiner individuellen Situation.
Grundlegende Informationen	
Standort	Leopoldstr. 244, 80807 München, Deutschland
Gründung	1885
Mitarbeiter	300 in Deutschland
Webseite	www.crifbuergel.de
Telefon	+49 4089803-0

6.2 d-fine GmbH

d-fine

d-fine ist ein europäisches Beratungsunternehmen mit Fokus auf analytisch anspruchsvollen Themen, die von naturwissenschaftlich geprägten Mitarbeiter*innen mit einem hohen Maß an Verantwortung für zukunftsfähige Lösungen und ihrer nachhaltigen technologischen Umsetzung bearbeitet werden.

Produkte und Dienste	
Dienstleistungen	Entwicklung von hoch personalisierten Lösungen in den Bereichen: ■ Know Your Trade (KYT) ■ Know Your Customer (KYC) ■ Anti-Geldwäsche ■ Betrugserkennung
Verfügbarkeit	In der Regel werden die Lösungen direkt in den Kundensystemen implementiert.
Abrechnung	Individuell vereinbart
Reichweite des Unternehmens	Weltweit
Zielkunden	■ Bankwesen & Kapitalmärkte ■ Versicherung & Vermögensverwaltung ■ Chemie, Energie & Fertigung ■ Pharma- & Gesundheitswesen ■ Mobilität & Transport ■ Öffentlicher Sektor ■ E-Commerce ■ Telekommunikation
Besonderheiten	■ Kundenindividuelle Softwareentwicklung ■ Erschließung des Datenuniversums des Unternehmens
Einsatz von KI	In allen Bereichen

Vorteile des KI-Einsatzes	■ Automatisierung, Prozessautomatisierung, d. h. Arbeits- erleichterung ■ Verbesserte Erkennungsgenauigkeit als mit nur regelbasierten Systemen
Kundeninteraktion mit dem Produkt	Funktionalitäten werden über dedizierte Dashboard-Anwendungen zugänglich gemacht oder direkt in bestehende Monitoring-Lösungen des Kunden integriert. In der Regel werden Verdachtsfälle mit den zugehörigen Hintergrundinformationen dem Kunden angezeigt, der diese dann im Detail analysieren und bewerten kann.
Grundlegende Informationen	
Standort	An der Hauptwache 7, 60313 Frankfurt, Deutschland
Gründung	2002
Mitarbeiter	1000
Webseite	d-fine.com
Telefon	+49 69 90737-0

6.3 NECT GmbH



Die NECT GmbH veröffentlichte in 2018 die NECT-App im Apple® und Android App Store. Die NECT-App ermöglicht die schnelle und sichere Identitätsprüfung von Personen mittels der Robo-Ident® Technologie. So können sich Kunden dank modernster Technologie auf Basis Künstlicher Intelligenz völlig autark identifizieren – egal wo, egal wann.

Produkte und Dienste	
Produkte	NECT-App zur Fernidentifizierung einer natürlichen Person
Verfügbarkeit	■ Für Kunden: Einbindung ähnlich PayPal®. Vorgangsanlage via Back-end API – Weiterleitung Kunden zur NECT Lösung via »NECT Button« ■ Für Konsument*innen: ■ Apple® App Store ■ Android App Store
Abrechnung	Transaktionsbasiert
Reichweite des Unternehmens	Deutschland
Zielkunden	■ Alle regulierten Unternehmen ■ Versicherungen ■ Banken ■ Telekommunikationsunternehmen
Besonderheiten	Pricing und Nutzerfreundlichkeit
Einsatz von KI	Zur Erkennung der Echtheit von Hologrammen und zur Liveness Detection.
Vorteile des KI-Einsatzes	Gleichbleibende oder bessere Qualität im Vergleich zur menschlichen Kontrolle und das 24/7.

Kundeninteraktion mit dem Produkt	Der Kunde wird durch »flags gewarnt. Im Verdachtsfall wird ein Bericht mit detaillierten Informationen als PDF an den Kunden übermittelt.
Grundlegende Informationen	
Standort	Großer Burstah 21, 20457 Hamburg, Deutschland
Gründung	2017
Mitarbeiter	66
Webseite	nect.com
Kontakt	nect.com/de/contact

6.4 RISK.IDENT GmbH



RISK IDENT ist ein schnell wachsendes Softwareunternehmen mit Sitz im Herzen Hamburgs. Im Jahr 2012 als Tochterunternehmen der Otto Group gegründet, hat sich das Unternehmen innerhalb kürzester Zeit zu einem der Marktführer im Bereich der Betrugsprävention entwickelt.

Produkte und Dienste	
Produkte	DEVICE IDENT Identifizierung und Risikobewertung von Geräten ■ Realtime-Analyse und Scoring der Geräte von Endkund*innen ■ Erstellung eines individuellen Risikoprofils der Geräte ■ Abgleich mit globalem Gerätepool FRIDA Intelligente Betrugspräventionssoftware ■ Automatische Antrags- und Transaktionsbewertung ■ Betrugsmustererkennung auf unterschiedlichen Ebenen: Antrag, Kund*innen, Vermittler, Third Party
Verfügbarkeit	Von RISK IDENT betriebene Cloudlösung als Backend
Abrechnung	DEVICE IDENT: Transaktionsbasiert, FRIDA: Lizenzmodell
Reichweite des Unternehmens	DEVICE IDENT: Weltweit, FRIDA: EU; Hauptmärkte: DACH
Zielkunden	■ E-Commerce ■ Finanzdienstleister ■ Telekommunikation ■ Mobility ■ Ticketing

Besonderheiten	DEVICE IDENT: Größter Pool in D-A-CH. FRIDA: Vollständige Fallbearbeitung, prozessoptimierte Oberflächen, Visualisierung von Betrugsnetzwerken, Erkennung neuer Muster, die außerhalb menschlicher Wahrnehmung liegen.
Einsatz von KI	In allen Produkten
Vorteile des KI-Einsatzes	Identifizierung neuer Betrugsmuster, Optimierung false/positives: Regeloptimierung und Algorithmen, Identifizierung trennscharfer Daten zur Betrugsprävention.
Kundeninteraktion mit dem Produkt	DEVICE IDENT: Nutzung von Hinweisen direkt zur manuellen Prüfung oder in Entscheidungssystemen, wie z. B. FRIDA. FRIDA: automatisiertes Decisioning und manuelle Prüfung von Verdachtsfällen. Technisch: API; Organisatorisch: Integration in Prozesse des Kunden.
Grundlegende Informationen	
Standort	Am Sandtorkai 50, 20457 Hamburg, Deutschland
Gründung	2012
Mitarbeiter	75
Webseite	riskident.com
Telefon	+49 4060945259-0

6.5 SCHUFA AG



Die SCHUFA ist ein führender Lösungsanbieter von Auskunft- und Informationsdienstleistungen für Unternehmen und Verbraucher in Deutschland. Seit 1927 werden die kreditrelevanten Informationen der SCHUFA rund um Bonität, Identität und Betrugsprävention genutzt, um Vertrauen zwischen Geschäftspartnern zu bilden.

	Produkte und Dienste
Produkte	Produktbereiche: ■ Bonität ■ Identifikation ■ Betrugsprävention ■ Compliance Identifikation: ■ IdentitätsCheck: Abgleich mit SCHUFA-Datenbestand ■ GiroIdent: PSD2-basierte Lösung Betrugsprävention: ■ SCHUFA-FraudPreCheck: Prädiktive Lösung für den E-Commerce, die Abhängigkeiten zwischen Transaktionen auswertet und in Echtzeit entscheiden kann, ob eine Transaktion auffällig ist. ■ SCHUFA-FraudPool: Finanzdienstleister informieren sich über diese Lösung gegenseitig über Vorfälle.
Verfügbarkeit	In der Regel »on premise«-Lösungen
Abrechnung	Hauptsächlich transaktionsbasiert
Reichweite des Unternehmens	Deutschland

Zielkunden	■ E-Commerce ■ Finanzdienstleistungen ■ Telekommunikation ■ Versicherungen ■ Öffentlicher Sektor
Besonderheiten	Hohe Datenqualität auf Grund detaillierter historischer Informationen; stark fundierte und abgesicherte KI-Verfahren
Einsatz von KI	In allen Bereichen
Vorteile des KI-Einsatzes	Die KI-Integration führt zu einem signifikanten Zugewinn an neuem, sicherem Geschäft für den Kunden
Kundeninteraktion mit dem Produkt	Integration der von der SCHUFA bereitgestellten Informationen in die eigenen analytischen Verfahren; Intensive Qualitätssicherung wird gemeinsam mit dem Kunden betrieben
Grundlegende Informationen	
Standort	Kormoranweg 5, 65201 Wiesbaden, Deutschland
Gründung	1927
Mitarbeiter	900
Webseite	www.schufa.de/en/
Telefon	+49 6119278-0

6.6 WidasConcepts GmbH



WidasConcepts ist ein globales professionelles Dienstleistungsunternehmen, das eine breite Palette von Dienstleistungen und Lösungen in den Bereichen Big Data, Internet of Things, Mobile, IT-Strategie, Managed Services sowie IAM/CIAM Beratung anbietet.

Produkte und Dienste	
Dienstleistungen und Produkte	Cidaas: Ein Cloud Identity und Access Management, die eine Identität pro Nutzer über all Ihre Kanäle schafft. cidaas übernimmt die Kundenidentifikation, Authentifizierung und Autorisierung. Es basiert auf den Standards OAuth2.0 und OpenID Connect, sorgt für eine starke API-Security und mit der integrierten Bot Net Detection und Betrugserkennung für eine sichere Reise über alle Geschäftskanäle. ■ Features ■ Single Sign On ■ Multi-Faktor-Authentifizierung ■ Einwilligungsmanagement ■ Social Login, eIDAS konforme Identifizierung & Real World Identification ■ Integrierte Betrugserkennung ■ API-Security
Verfügbarkeit	Cloud
Abrechnung	Cidaas: Subscription, Andere Lösungen: Individuell vereinbart
Reichweite des Unternehmens	Weltweit
Zielkunden	Alle Branchen, die digitale Identitäten benötigen.
Besonderheiten	Mehr als 14 verschiedene Multifaktor-Authentifizierungsverfahren
Einsatz von KI	Zur starken Nutzeridentifikation, zur Bot Net Detection und verhaltensbasierten Betrugserkennung

Vorteile des KI-Einsatzes	Zuverlässige, sichere Nutzeridentifikation, insbesondere bei sensiblen, kritischen Anwendungen und in regulierten Branchen + integrierte Betrugserkennung als willkommenes Extra
Kundeninteraktion mit dem Produkt	Das gesamte Identitäts- und Zugangsmanagement inkl. Betrugs-erkennung wird as a Service dem Kunden bereitgestellt für eine standardisierte Authentifizierung und Autorisierung über alle digitalen Kanäle hinweg.
Grundlegende Informationen	
Standort	Maybachstraße 2, 71299 Wimsheim, Deutschland
Gründung	1997
Mitarbeiter	120
Webseite	www.widas.de
Telefon	+49 704495103-0

LITERATUR

- [1] LexisNexis® Risk Solutions, "LexisNexis® Risk Solutions 2018 True Cost of FraudSM Study," Aug. 2018.
- [2] Bundeskriminalamt, "Cybercrime - Bundeslagebild Cybercrime 2019," 62, Sep. 2020. Accessed: Feb. 22, 2021. [Online]. Available: <https://www.bka.de/SharedDocs/Downloads/DE/Publikationen/JahresberichteUndLagebilder/Cybercrime/cybercrimeBundeslagebild2019.html;jsessionid=7B9CBC303554D45E30E04AA52E8A6CB6.live0602?nn=28110>.
- [3] P. H. Pisani and A. C. Lorena, "A systematic review on keystroke dynamics," J Braz Comput Soc, vol. 19, no. 4, pp. 573–587, Nov. 2013, doi: 10.1007/s13173-013-0117-7.
- [4] P. S. Teh, A. B. J. Teoh, and S. Yue, "A Survey of Keystroke Dynamics Biometrics," The Scientific World Journal, Nov. 03, 2013. <https://www.hindawi.com/journals/tswj/2013/408280/> (accessed Feb. 24, 2021).
- [5] C. Phua, V. Lee, K. Smith-Miles, and R. Gayler, "A Comprehensive Survey of Data Mining-based Fraud Detection Research." May 18, 2013.
- [6] A. Dey, "Machine Learning Algorithms: A Review," International Journal of Computer Science and Information Technologies., no. 7, pp. 1174–1179, 2016.
- [7] D. Choi and K. Lee, "An Artificial Intelligence Approach to Financial Fraud Detection under IoT Environment: A Survey and Implementation," Security and Communication Networks, vol. 2018, p. 15, 2018.
- [8] J. West and M. Bhattacharya, "Intelligent financial fraud detection: A comprehensive review," Computers & Security, vol. 57, pp. 47–66, 2016.
- [9] S. Bhattacharyya, S. Jha, K. Tharakunnel, and J. C. Westland, "Data mining for credit card fraud: A comparative study," Decision Support Systems, vol. 50, no. 3, pp. 602–613, Feb. 2011, doi: 10.1016/j.dss.2010.08.008.

- [10] E. W. T. Ngai, Y. Hu, Y. H. Wong, Y. Chen, and X. Sun, "The application of data mining techniques in financial fraud detection: A classification framework and an academic review of literature," *Decision Support Systems*, vol. 50, no. 3, pp. 559–569, Feb. 2011, doi: 10.1016/j.dss.2010.08.006.
- [11] M. Cihan Sorkun and Toraman, "Fraud Detection on Financial Statements Using Data Mining Techniques," Sep. 30, 2017. https://www.researchgate.net/publication/320132308_Fraud_Detection_on_Financial_Statements_Using_Data_Mining_Techniques (accessed Feb. 24, 2021).



KI-FORTSCHRITTSZENTRUM

Das KI-Fortschrittszentrum »Lernende Systeme« unterstützt Firmen dabei, die wirtschaftlichen Chancen der Künstlichen Intelligenz und insbesondere des Maschinellen Lernens für sich zu nutzen. In anwendungsnahen Forschungsprojekten und in direkter Kooperation mit Industrieunternehmen arbeiten die Stuttgarter Fraunhofer-Institute für Arbeitswirtschaft und Organisation IAO sowie für Produktionstechnik und Automatisierung IPA daran, Technologien aus der KI-Spitzenforschung in die breite Anwendung der produzierenden Industrie und der Dienstleistungswirtschaft zu bringen. Finanzielle Förderung erhält das Zentrum vom Ministerium für Wirtschaft, Arbeit und Wohnungsbau Baden-Württemberg.

Europas größte Forschungsk Kooperation auf dem Gebiet der KI

Das KI-Forschungszentrum ist Forschungspartner des Cyber Valley, einem Konsortium aus den renommierten Universitäten Tübingen und Stuttgart, dem Max-Planck-Institut für intelligente Systeme und einigen führenden Industrieunternehmen. In gemeinsamen Forschungslabors werden Grundlagenforschung und anwendungsorientierte Entwicklung zu aktuellen wie auch zukünftigen Bedarfen behandelt und vorangetrieben.

Menschzentrierte KI

Alle Aktivitäten des Zentrums verfolgen das Ziel, eine menschzentrierte KI zu entwickeln, der die Menschen vertrauen und die sie akzeptieren. Nur wenn Menschen mit neuen Technologien intuitiv interagieren und vertrauensvoll zusammenarbeiten, kann deren Potenzial optimal ausgeschöpft werden. Daher konzentrieren sich die Forschungsaktivitäten unter anderem auf die Themen Erklärbarkeit, Datenschutz, Sicherheit und Robustheit von KI-Technologien.

Studienreihe »Lernende Systeme«

Die Studienreihe »Lernende Systeme« gibt Einblick in die Potenziale und die praktischen Einsatzmöglichkeiten von KI. Nähere Informationen und die aktuellen Versionen der Studien finden Sie unter: www.ki-fortschrittszentrum.de/studien

FRAUNHOFER-GESELLSCHAFT

Forschen für die Praxis ist die zentrale Aufgabe der Fraunhofer-Gesellschaft. Die 1949 gegründete Forschungsorganisation betreibt anwendungsorientierte Forschung zum Nutzen der Wirtschaft und zum Vorteil der Gesellschaft. Vertragspartner und Auftraggeber sind Industrie- und Dienstleistungsunternehmen sowie die öffentliche Hand.

Die Fraunhofer-Gesellschaft betreibt in Deutschland derzeit 74 Institute und Forschungseinrichtungen. Mehr als 28 000 Mitarbeiterinnen und Mitarbeiter, überwiegend mit natur- oder ingenieurwissenschaftlicher Ausbildung, erarbeiten das jährliche Forschungsvolumen von mehr als 2,8 Milliarden Euro. Davon entfallen mehr als 2,3 Milliarden Euro auf den Leistungsbereich Vertragsforschung. Rund 70 Prozent dieses Leistungsbereichs erwirtschaftet die Fraunhofer-Gesellschaft mit Aufträgen aus der Industrie und mit öffentlich finanzierten Forschungsprojekten. Rund 30 Prozent werden von Bund und Ländern als Grundfinanzierung beigesteuert, damit die Institute Problemlösungen entwickeln können, die erst in fünf oder zehn Jahren für Wirtschaft und Gesellschaft aktuell werden.

Internationale Kooperationen mit exzellenten Forschungspartnern und innovativen Unternehmen weltweit sorgen für einen direkten Zugang zu den wichtigsten gegenwärtigen und zukünftigen Wissenschafts- und Wirtschaftsräumen.

Mit ihrer klaren Ausrichtung auf die angewandte Forschung und ihrer Fokussierung auf zukunftsrelevante Schlüsseltechnologien spielt die Fraunhofer-Gesellschaft eine zentrale Rolle im Innovationsprozess Deutschlands und Europas. Die Wirkung der angewandten Forschung geht über den direkten Nutzen für die Kund*innen hinaus: Mit ihrer Forschungs- und Entwicklungsarbeit tragen die Fraunhofer-Institute zur Wettbewerbsfähigkeit der Region, Deutschlands und Europas bei. Sie fördern Innovationen, stärken die technologische Leistungsfähigkeit, verbessern die Akzeptanz moderner Technik und sorgen für die Aus- und Weiterbildung des dringend benötigten wissenschaftlich-technischen Nachwuchses.

Ihren Mitarbeiterinnen und Mitarbeitern bietet die Fraunhofer-Gesellschaft die Möglichkeit zur fachlichen und persönlichen Entwicklung für anspruchsvolle Positionen in ihren Instituten, an Hochschulen, in Wirtschaft und Gesellschaft. Studierenden eröffnen sich aufgrund der praxisnahen Ausbildung und Erfahrung an Fraunhofer-Instituten hervorragende Einstiegs- und Entwicklungschancen in Unternehmen.

Namensgeber der als gemeinnützig anerkannten Fraunhofer-Gesellschaft ist der Münchner Gelehrte Joseph von Fraunhofer (1787–1826). Er war als Forscher, Erfinder und Unternehmer gleichermaßen erfolgreich.

Fraunhofer IAO

Mensch und Technik in der digitalen Arbeitswelt, Wirtschaft und Gesellschaft

Digitale Technologien verändern unsere Arbeitswelt und haben tiefgreifende Auswirkungen auf Wirtschaft und Gesellschaft. Lang etablierte Methoden und Prozesse werden in kurzer Zeit modernisiert und revolutioniert. Das Fraunhofer IAO kooperiert eng mit dem Partnerinstitut IAT der Universität Stuttgart und entwickelt gemeinsam mit Unternehmen, Institutionen und Einrichtungen der öffentlichen Hand wirksame Strategien, Geschäftsmodelle und Lösungen für die digitale Transformation.

Die digitale Transformation und neue IT-Technologien eröffnen für Unternehmen viele Chancen: innovative Produktangebote für neue Zielgruppen, bessere und kostengünstigere Prozesse, eine »intelligentere« Kundenkommunikation und höhere Automatisierung. Dafür kommen innovative, vernetzte IT-Lösungen auf Basis von Big Data, Künstlicher Intelligenz, Cloud und Internetplattformen zum Einsatz.

Die richtige Strategie und IT sind eine wesentliche Grundlage für den Erfolg und die Wettbewerbsfähigkeit von Unternehmen. Voraussetzung für erfolgreiche Anwendungen ist ein klarer Nutzen für das Unternehmen, seine Kund*innen und seine Partner.

Unsere Leistungen basieren auf fundierter Technologie- und Marktkenntnis sowie branchenübergreifenden Erfahrungen. Durch den Einsatz unserer praxiserprobten Methoden und erfahrenen Mitarbeitenden sichern wir den Projekterfolg. Unser Fraunhofer-Netzwerk ermöglicht uns den Zugriff auf ein umfassendes Kompetenzspektrum.

Das Fraunhofer IAO und das IAT der Universität Stuttgart beschäftigen gemeinsam mehr als 650 Mitarbeitende und verfügen über rund 15 000 Quadratmeter Büroflächen, Demonstrationen sowie Entwicklungs- und Testlabors.

Forschungsbereich »Mensch-Technik-Interaktion«

Wir gestalten Technik für die Menschen, damit sie sich intuitiv und komfortabel bedienen lässt. Neue Technologien können Prozesse digitalisieren, die Produktivität steigern und unser Leben bereichern. Sie bieten viele Chancen für neue Produkte und attraktive Services. Doch letztlich hängt ihr Erfolg davon ab, wie sie von den Menschen wahrgenommen und angenommen werden.

In unserer Forschung arbeiten wir an Lösungen, die ein effizientes Zusammenspiel von Menschen und intelligenter Technik ermöglichen: bei der Arbeit, im Fahrzeug, zu Hause oder unterwegs. Wir gestalten menschenzentrierte Innovationsprozesse und entwickeln bedarfsgerechte Konzepte für die IT-Sicherheit.

Die Fokusthemen des Forschungsbereichs »Mensch-Technik-Interaktion« sind:

- Technologien und Gestaltungsansätze für User Interfaces und ergonomische Produkte
- Mensch-Fahrzeug-Interaktion, Fahrassistenzsysteme und Automatisierung des Fahrens
- Identitätsmanagement und Personalisierung
- User Experience und Usability Engineering
- Strategien für die menschengerechte Digitalisierung
- Neuroarbeitswissenschaft
- Intuitive und nutzerorientierte IT-Sicherheitslösungen

Team »Identitätsmanagement«

Da IT-Systeme inzwischen das Rückgrat praktisch aller wichtigen Geschäftsprozesse darstellen, ist klar, wie wichtig deren sichere und privatsphärenfreundliche Gestaltung ist, um den Abfluss von Wissen, direkte Angriffe oder Datenskandale zu vermeiden. Dass dies häufig scheitert, liegt unter anderem daran, dass IT-Sicherheit häufig zu kurz gedacht und eine rein technische Herangehensweise bevorzugt wird, statt reale Marktbedürfnisse, die Wünsche von Nutzer*innen sowie den wirtschaftlichen Kontext zu beachten.

Identitätsmanagement ermöglicht die Kontrolle, Verwaltung und Nachvollziehbarkeit von Zugriffen auf Informationen sowie Ressourcen und leistet somit einen zentralen Beitrag zur Informationssicherheit.

Wir unterstützen Unternehmen darum mit unserem Ansatz der tragfähigen Sicherheit. Hierunter verstehen wir Lösungen, die am Markt erfolgreich sind und im Praxiseinsatz tatsächlich die Sicherheit verbessern. Bearbeitet werden wirtschaftliche, technische und organisatorische Fra-

gestellungen. Dabei stehen insbesondere die ökonomische Tragfähigkeit, die organisatorische Einbindung der Systeme sowie deren Zusammenspiel mit den Ökosystemen der Unternehmen im Fokus.

Das Team betreibt mehrere Labore und Demonstrationszentren, insbesondere das Labor für tragfähige Sicherheit (viable security – VS-Lab). Im VS-Lab unterstützen wir Unternehmen bei der Entwicklung von bedarfsgerechten, an den Bedürfnissen von Nutzer*innen orientierten und sicheren IT-Lösungen.

Die folgenden Themenfelder beschreiben die aktuellen Forschungsschwerpunkte des Teams:

Identity- und Accessmanagement

Der zuverlässige Schutz vor unberechtigt Zugriff auf Firmendaten ist Voraussetzung für organisationsübergreifende Kooperationen. Wir unterstützen Ihre Organisation bei der Auswahl und Einführung sicherer und zugleich nutzerfreundlicher Systeme.

Integratives Sicherheitsmanagement

Integratives Sicherheitsmanagement umfasst nicht nur Auswahl, Einbindung und Implementierungsbegleitung, sondern auch die betriebswirtschaftliche Gestaltung (Budgetierung, Prozessgestaltung) sowie die Kommunikation der Entscheidungen. So finden sich hier unsere Leistungen zur Awarenesssteigerung, Policy-Compliance-Sicherung, Risiko-Governance und Infrastrukturberatung wieder.

User-oriented Security-Technologien

Lösungen zur IT-Security werden häufig noch rein aus technischer Perspektive betrachtet. Dabei gelten Nutzer*innen immer noch als der »weakest link«. Social Engineering stellt dabei nur ein Problem dar. IT-Security-Lösungen, die ohne die Einbindung von Nutzer*innen entwickelt werden, weisen häufig eine schlechte Usability auf, was wiederum zu kritischen Sicherheitslücken führen kann. Wir beschäftigen uns damit, wie IT-Security-Lösungen nutzerfreundlich gestaltet werden können, um die Sicherheit und Nutzerakzeptanz zu steigern.

Blockchain und Distributed-Ledger-Technologien

Auch nach dem Abflauen des Hypes um Bitcoin und andere Kryptowährungen bestehen die Potenziale der Blockchain weiterhin fort. Wir untersuchen, wie sich die Technologie in der praktischen Anwendung schlägt, wo ihre Grenzen liegen und an welchen Stellen sich perspektivisch Entwicklungspotenziale ergeben.

Sozioökonomische Aspekte von IT-Sicherheit

Datenschutzfreundliche und sichere IT-Systeme müssen sich am Markt bewähren. Darum unterstützen wir Sie bei der Analyse von Use Cases, entwickeln abgestimmte Geschäftsmodelle, untersuchen das Marktgeschehen und für Ihre Produkte relevante gesellschaftliche sowie politische Entwicklungen.

Data Protection and Compliance

Rechtliche Vorgaben zum Datenschutz wie die der EU-Datenschutz-Grundverordnung sind häufig komplex und stellen Unternehmen vor Herausforderungen. Hierzu bieten wir praxisgerechte Lösungen und unterstützen Unternehmen bei einer effizienten Umsetzung.

AUTORINNEN UND AUTOREN



Dr. Christian Schunck

Fraunhofer-Institut für Arbeitswirtschaft und Organisation IAO, Stuttgart

Dr. Christian Schunck promovierte im Jahr 2008 in Physik am Massachusetts Institute of Technology, USA. Anschließend arbeitete Dr. Schunck für die Boston Consulting GmbH in München, Deutschland, als Berater mit dem Schwerpunkt Risikomanagement für Finanzinstitute. Im Jahr 2010 zog er nach Rom, Italien, und arbeitete bei Nestor Scarl und Fondazione INUIT an der Forschung und Entwicklung in den Bereichen IT-Risiko und -Sicherheit, digitales Identitätsmanagement und der Validierung von verteilten elektronischen Transaktionen. Danach wechselte er als Assistenzprofessor an die Abteilung für Wirtschaftsingenieurwesen »Mario Lucertini« an der Universität Rom »Tor Vergata«. Derzeit ist Christian Schunck Senior Researcher in der Gruppe für digitale Identität am Fraunhofer IAO. Christian Schunck war wissenschaftlicher Koordinator des thematischen Netzwerks SSEDIC (»Scoping the Single European Digital Identity Community«) und Innovationsmanager im H2020-Projekt PICASO, wo er auch das Arbeitspaket zu »Integrated Care Orchestration with Risk Assessment and Decision Support« leitete. Als Forscher ist er Co-Autor von mehr als 40 Artikeln, darunter sieben in Nature und Science. Seine Publikationen haben mehr als 5000 Zitate erhalten.



Rachelle Sellung, M. Sc.

Fraunhofer-Institut für Arbeitswirtschaft und Organisation IAO, Stuttgart

Rachelle Sellung ist Senior Scientist im Kompetenzteam Identity Management am Fraunhofer IAO in Stuttgart. Rachelle Sellung hat einen Master of Science in Wirtschaftswissenschaften von der Universität Hohenheim in Stuttgart, Deutschland sowie einen Bachelor of Business Administration in Marketing von der University of Mississippi, USA. Sie forscht an einer Vielzahl von Technologien im Bereich der IT-Sicherheit und des Identitätsmanagements aus einer wirtschaftlichen Perspektive. In dem groß angelegten EU-FP7-Projekt FutureID, das eine Identitätsmanagement-Infrastruktur für Europa entwickelte, analysierte sie die sozioökonomischen Auswirkungen. Insbesondere betrachtete sie die Integration von bestehenden eID-Technologien, Vertrauensinfrastrukturen und föderierten Identitätsmanagement-Diensten. Auch im EU-Horizon-2020-Projekt LIGHTTest, das eine globale, transparente, und domänenübergreifende Vertrauensinfrastruktur zur Bewertung elektronischer Transaktionen aufbaute, leitete sie die beglei-

tende sozioökonomische Analyse. Ihre Forschungsinteressen umfassen die Bereiche Sicherheit, Identitätsmanagement, Vertrauensmanagement und die Einflüsse von neuen Technologien wie Artificial Intelligence, Machine Learning und Blockchain.

Dr. Heiko Roßnagel

Fraunhofer-Institut für Arbeitswirtschaft und Organisation IAO, Stuttgart

Dr. Heiko Roßnagel ist Leiter des Competence Teams Identitätsmanagement am Fraunhofer IAO. Er verfügt über umfangreiche Erfahrung aus zahlreichen nationalen und internationalen Forschungsprojekten wie den EU-geförderten Projekten SECUR-ED, SSEDIC, WiTness und FIDIS oder den national geförderten Verbundprojekten VerSiert, VERTRAG, SkIDentity, SANDRA, CUES, Industrial Data Space und ENTOURAGE. Er koordinierte die EU-geförderten Projekte von Horizon2020 LIGHTest (2016–2019) und FP7 FutureID (2012–2015). Heiko Roßnagel hat an der Technischen Universität Darmstadt Informatik studiert und an der Goethe-Universität im Fachbereich Wirtschaftswissenschaften promoviert.

Seine aktuellen Forschungsinteressen liegen im Bereich Datenschutz, IT-Sicherheit und Identitätsmanagement mit einem Fokus auf menschliche Faktoren, Wirtschaftlichkeit und Marktakzeptanz.



Kontaktadresse

Fraunhofer-Institut für Arbeitswirtschaft und Organisation IAO
Nobelstraße 12, 70569 Stuttgart

Autorinnen und Autoren

Christian H. Schunck
Telefon +49 711 970-2430
christian.schunck@iao.fraunhofer.de

Rachelle Sellung

Telefon +49 711 970-2403
rachelle.sellung@iao.fraunhofer.de

Heiko Roßnagel

Telefon +49 711 970-2145
heiko.rossnagel@iao.fraunhofer.de

Herausgeber

Wilhelm Bauer, Oliver Riedel, Thomas Renner, Matthias Peissner

Satz und Gestaltung

Franz Schneider, Fraunhofer IAO

URN-Nummer

[urn:nbn:de:0011-n-6306863](https://nbn-resolving.org/urn:nbn:de:0011-n-6306863)

Online verfügbar als Fraunhofer-ePrint

<http://publica.fraunhofer.de/dokumente/N-630686.html>

Gefördert durch das Ministerium für Wirtschaft, Arbeit und
Wohnungsbau Baden-Württemberg

Alle Rechte vorbehalten

© Fraunhofer IAO, 03/2021

Gefördert durch



Baden-Württemberg

MINISTERIUM FÜR WIRTSCHAFT, ARBEIT UND WOHNUNGSBAU

CyberValley

