

Frank Stiegler

Generative KI und Recht: Herausforderungen von LLMs für Versicherer

Künstliche Intelligenz kann Arbeitsprozesse erheblich effizienter gestalten, stellt Versicherungsunternehmen jedoch auch vor rechtliche Hürden, wenn Daten in Large Language Models (LLMs) eingebracht, daraus generiert und im Geschäftsalltag verwendet werden. Der folgende Beitrag beleuchtet verschiedene Herausforderungen für Versicherungsunternehmen mit Blick auf den EU AI Act, Datenschutz sowie Digital Operational Resilience Act (DORA).

LLMs nutzen fortschrittliche Deep-Learning-Techniken und gewaltige Datensätze, um menschliche Sprache und komplexe Daten zu interpretieren und zu generieren. Durch ihre besondere Fähigkeit zur Mustererkennung können diese generativen KI-Systeme bereits heute in der Versicherungsbranche Trends in Daten analysieren sowie einen Vergleich von Dokumenten auf inhaltlicher Ebene durchführen, um Änderungen allgemeiner Versicherungsbedingungen zu erkennen oder Leistungen mit dem aktuellen Marktangebot zu vergleichen.

Jedoch sind LLMs nicht fähig, verlässlich faktische Wahrheiten zu beurteilen, sondern können nur im Rahmen ihres Trainingsmaterials sowie ihrer Funktionalität und Vektorentiefe Aufgaben ausführen und Inhalte wiedergeben. Zudem sind KI-Systeme aufgrund der Intransparenz ihrer Funktionsweise im Vergleich zu traditionellen Algorithmen schwer kontrollierbar, da sie keine Software sind, in deren Code kleinteilig eingegriffen werden kann, falls unerwünschte Ergebnisse auftreten.

Da einmal eingespeiste Trainingsdaten in der Regel nicht mehr einfach aus dem System gelöscht werden können, sind LLM auch eine Herausforderung beim Umgang mit personenbezogenen Daten. Ebenso kann die Intransparenz des Ergebnisproduktionswegs in Bezug auf Auskunftspflichten im Bereich Datenschutz und Informationssicherheit Probleme aufwerfen. Im Folgenden wird beleuchtet, worauf Versicherungsunternehmen achten sollten, um mögliche Rechtshürden hinsichtlich der Nutzung von LLMs aufzudecken.

EU AI Act: Hohe Bußgelder bei Verstoß

Das europäische KI-Gesetz (EU AI Act), das im August 2024 in Kraft getreten ist und die verantwortungsvolle Entwicklung und Verwendung von KI in der EU fördern

soll, ordnet KI-Anwendungen in drei Risikokategorien ein. Verboten werden Systeme mit inakzeptablem Risiko wie staatlich betriebenes Social Scoring. Zu Anwendungen mit hohem Risiko gehören zum Beispiel KI-Tools im Bereich kritische Infrastruktur, Polizei sowie auch Anwendungen, die den Zugang zu Dienstleistungen im Banken- und Versicherungsbereich beeinflussen. Weitgehend unreguliert bleiben KI-Anwendungen, die nicht explizit verboten oder als risikoreich eingestuft sind.

Bei Verstoß in Bezug auf verbotene Anwendungen sind Strafzahlungen von bis zu 7% oder 35 Mio. Euro des Jahresumsatzes möglich. Bei allen anderen Verstößen werden in der Regel bis zu 3% oder 15 Mio. Euro fällig. Bei inkorrekten Auskunftsn gegenüber Behörden drohen zudem Bußgelder von bis zu 1,5% des Jahresumsatzes oder 7 Mio. Euro.

KI-Folgeabschätzung (KIFA) nach EU AI Act

Eine KI-Folgeabschätzung nach Art. 27 AI Act ist für Versicherer dann zu erstellen, wenn KI-Systeme eingesetzt werden sollen, die für die Risikobewertung und Preisgestaltung in Bezug auf natürliche Personen im Falle von Lebens- und Krankenversicherungen verwendet werden sollen (Anhang III Nr. 5). Werden Vertragskonditionen also nicht nur in der Entwurfsphase mithilfe von KI vorbereitet, sondern KI-gestützt festgelegt, ist eine KIFA notwendig.

In der KIFA müssen Unternehmen die Verfahren sowie Zeitraum und Häufigkeit beschreiben, in denen das KI-System entsprechend seinem vorgesehenen Zweck eingesetzt wird, sowie die von der Verwendung betroffenen Personen(gruppen). Darüber hinaus muss die KIFA die spezifischen Schadensrisiken, Überwachungsmaßnahmen sowie Maßnahmen im Fall des Eintritts von Risiken, gepaart mit Re-

gelungen für die interne Steuerung und Beschwerdemechanismen umfassen.

KI und Datenschutz: Rechtsgrundlagen und Hürden

Sollen personenbezogene Daten in einem KI-System verarbeitet werden, gibt es für Versicherungsunternehmen basierend auf Art. 6 DSGVO regelmäßig nur drei Möglichkeiten als *Rechtsgrundlage*. Entweder man holt die *Einwilligung der betroffenen Personen* ein; die Datenverarbeitung ist Teil der Leistungsbeschreibung von Verträgen mit den Betroffenen, wodurch sie *erforderlich zur Vertragserfüllung* ist; oder man begründet die Datenverarbeitung mit einer *Interessenabwägung*, was umso schwerer umzusetzen ist, je stärker die Auswirkungen für die Betroffenen sind.

Eine Zusatzhürde stellt das *Profiling* nach Art. 22 DSGVO dar, wonach eine betroffene Person das Recht hat, „nicht einer ausschließlich auf einer automatisierten Verarbeitung – einschließlich Profiling – beruhenden Entscheidung unterworfen zu werden, die ihr gegenüber rechtliche Wirkung entfaltet oder sie in ähnlicher Weise erheblich beeinträchtigt.“

In diesem Fall müssen Unternehmen Betroffene darüber aufklären, wie diese automatisierte Verarbeitung inklusive der „involvierten Logik“ (Art. 13 DSGVO) eingesetzt wird. Aufgrund der Intransparenz von LLMs ist es jedoch schwierig, ausreichend zu erklären, wie das System zu seinen Entscheidungen kommt. Rechtlich gibt es zwei Stellschrauben, durch die diese Auskunftspflicht nicht greift:

– Dient die KI nur zur Vorbereitung von Entscheidungen und werden diese

Frank Stiegler

Senior Legal Counsel und Datenschutzbeauftragter der Convista-Gruppe

Abbildung: Essentials bei der Verwendung von LLM

		
LLM sind wegen Intransparenz schwer zu kontrollieren	Wenn Vertragskonditionen KI-gestützt festgelegt, KIFA notwendig	Bei LLM meist DSFA nötig (flankierend zu KIFA)
		
Auskunft und Löschung tendenziell schwierig	DORA: Zusatzhürden nur, wenn KI die Risiken beeinflusst	LLM-Lokalbetrieb bringt Kontrollierbarkeit, Kosten (und Staubanfälligkeit)

Bildquelle: Convista

durch einen Menschen verifiziert, trifft der Umstand einer „ausschließlich automatisierten Verarbeitung“ nicht zu.

- Ist die KI nicht in Vorgänge wie etwa Vertragsbedingungen involviert, die deutliche Konsequenzen für Kunden haben, werden diese auch nicht „erheblich beeinträchtigt“.

Auch die Einhaltung von Betroffenenrechten (Art. 12 bis 23 DSGVO) von *Auskunft und Löschung* rund um personenbezogene Daten, die in einem LLM verarbeitet oder in diesem Zusammenhang produziert wurden, kann aufgrund der Intransparenz dieser KI-Systeme tendenziell schwierig sein, jedenfalls dann, wenn man davon ausgeht, dass personenbezogene Daten in einem LLM gespeichert sind. Denn: Verlangt eine Person die Löschung, kann es sein, dass das ganze Modell eliminiert und neu trainiert werden muss, wenn die Daten so tief in das Modell verwoben sind, dass eine Löschung nicht mehr möglich ist.

Datenschutz-Folgenabschätzung (DSFA) nach DSGVO

Eine Datenschutz-Folgenabschätzung ist laut Art. 35 DSGVO zu erstellen, wenn „[...] eine Form der Verarbeitung, insbesondere bei Verwendung neuer Technologien, aufgrund der Art, des Umfangs, der Umstände und der Zwecke der Verarbeitung voraussichtlich ein hohes Risiko für die Rechte und Freiheiten natürlicher Personen zur Folge [hat]“.

Da eine DSFA wie die KIFA Risiken für die Betroffenen betrachtet, können beide Abschätzungen „parallel“ erstellt werden. So kann z. B. die DSFA die Basis der KIFA bilden und umgekehrt. Wie die KIFA, nur aus Datenschutzsicht, beschreibt die DSFA nämlich das eingesetzte KI-System und die damit verbundene Datenverarbeitung inklusive der Verantwortlichkeiten, der mit den Maßnahmen auftretenden Risiken und den risikomitigierenden Maßnahmen.

Digital Operational Resilience Act (DORA) bei Verwendung von KI-Systemen

DORA soll die digitale operationale Resilienz des europäischen Finanzsektors stärken. Hierbei stehen vor allem Risikomanagement, das Überwachen und Testen der IKT-Infrastruktur sowie die Behandlung, Klassifizierung und Berichterstattung IKT-bezogener Vorfälle im Vordergrund. Beim Einsatz von LLMs bestehen in Bezug auf DORA nur dann Zusatzhürden, wenn die KI Risiken, in Form einer geschäftskritische Entscheidung, beeinflusst.

Beispiel: Ein Versicherer setzt ein LLM-basiertes System ein, das Schadenmeldungen aus eingegangenen Textdokumenten, aufgezeichneten Kundenanrufen und Schadenfotos auswertet. Das Modell bewertet, ob der gemeldete Schaden versicherungsrelevant ist und damit eine Schadenzahlung fällig wird. Es muss nach den

DORA-Richtlinien sichergestellt werden, dass das verwendete LLM-basierte System nachvollzieh- und prüfbar ist, insbesondere da es in diesem Fall an finanziell kritischen Entscheidungen beteiligt wird. Zudem fordert DORA robuste Sicherheitsmaßnahmen, um Betrug und Datenmanipulation zu verhindern. Es muss eine Implementierung von zusätzlichen Prüfschritten zur Erkennung verdächtiger Schadenmeldungen und zur Absicherung gegen betrügerische Eingaben erfolgen.

Eine rechtliche Herausforderung ist die Risikobewertung hinsichtlich Bias, Fehlinformationen und Fehlern in der Entscheidungsfindung aufgrund der Intransparenz von LLMs. Dies kann ebenso ein Problem bei Meldungen an Aufsichtsbehörden bei Informationssicherheitsvorfällen darstellen, da es schwierig sein kann zu klären, wie der Vorfall überhaupt zustande kam und wie oder ob das Modell beteiligt war.

Wird die KI durch einen IKT-Dienstleister betrieben, ist es wichtig, diesen miteinzubeziehen, um Risiken – etwa durch Penetrationstests – zu entdecken und so kontrollierbar zu machen sowie ein Meldewesen zwischen Anbieter und Unternehmen zu etablieren. Um die Haftungslücke zu schließen, müssen auch Risiken vertraglich mit dem Dienstleister adressiert werden. Eine vertragliche Weitergabe von DORA-Pflichten kann sich allerdings insbesondere bei großen KI-Anbietern wie etwa aus den USA schwierig gestalten.

Ein LLM-Lokalbetrieb bringt hier deutlich mehr Kontrollierbarkeit in puncto Sicherheit, aber auch höhere Kosten. Darüber hinaus entwickelt sich ein lokales Modell im Gegensatz zu einem gemieteten nicht weiter, wodurch es weniger zukunftsgerichtet ist.

Hilfreiche Fragen zur Bewertung von Rechtshürden

Die rechtlichen Herausforderungen beim Einsatz von LLMs sind vielfältig. Daher sollten sich Anwenderunternehmen stets folgende grundlegenden Fragen stellen:

- Welche Daten wurden/werden zum Training des LLMs genutzt?
- Welche Daten wandern vom Anwenderunternehmen in das LLM?
- Welche Daten liefert das LLM dem Anwenderunternehmen und wofür sollen

- sie genutzt werden?
- Welche Daten nutzen der Betreiber des Modells und gegebenenfalls sonstige Personen?
- Welche Konsequenzen kann der LLM-Output für Betroffene haben?
- Wie sind die Systeme, auf denen das LLM läuft, gesichert und wo, wie und von wem werden sie betrieben?
- Hierdurch können Unternehmen bereits in der Planungsphase für den Einsatz eines LLMs mögliche Rechtshürden ermitteln und bewerten sowie einhergehende Risiken adressieren.