

企業が直面する生成AIリスク： 保険の役割に関する考察

調査研究概要 | 2025年10月

Ruo (Alex) Jia, Director Digital Technologies, Geneva Association
Associate Professor of Insurance, Peking University

寄稿者：

Martin Eling, Director of the Institute of Insurance Economics and Professor of Insurance Management, University of St. Gallen

Tianyang Wang, Professor of Finance, Colorado State University

ジェネレーティブAI（以下「生成AI」）とは、ユーザープロンプトに回答してAI独自のコンテンツ（テキスト、画像、コード、音声など）を生成する高度なAIシステムを指します。2022年後半以降、様々な業界で生成AIの導入が急速に進み、OpenAIのChatGPTをはじめとするツールが前例のない普及を見せています。企業はイノベーションと業務効率の向上を目的に、自社製品や社内プロセスに生成AIを組み込んでいます。しかし、この導入は新たなリスクを生むだけでなく、従来型AIの利用に伴う既存のリスクを増幅します。生成AIモデルは予測不能な回答を返す場合があります。つまり、時として（誤った情報や誤解を招く情報を自信満々に出力する）「ハルシネーション」を生み出すほか、著作権で保護されたコンテンツを意図せずに複製してしまうこともあります。生成AIのこうした挙動は、歴史的にもほとんど前例のない問題を引き起こす一方で、偏見の拡大、エラーや事実誤

認の発生、セキュリティ上の脆弱性の増大といった従来からある懸念を増幅します。端的に言えば、生成AIは計り知れないメリットをもたらすと同時に、適切な管理を必要とする新たなリスクも生み出しているのです。

生成AIに関する事業リスクは、7つのカテゴリーに分類できます（表1）。製品面では、テクノロジー企業が開発した生成AIツールを使用する事業者が財務的損失を被る可能性があり、その結果としてツール提供者であるテクノロジー企業が潜在的な賠償責任を負うリスクが生じます。事業運営の面では、生成AIを活用する事業者は、誤った意思決定に加え、業務の非効率化や財務的損失を被るリスクにも直面しています。さらに、生成AIシステムは、サイバー攻撃に対してより脆弱である可能性があります。つまり、サイバーセキュリティリスクが増大するということです。

表 1：生成 AI が誘発する各種リスク

	リスクカテゴリー	具体例
自社起因のオペレーショナルリスク	オペレーショナルリスク	アルゴリズムエラー、安定性の喪失、低信頼性
		ブラックボックス問題、悪意のある攻撃
	サイバーセキュリティとプライバシー	AI 駆動型サイバー攻撃、データ・プライバシー侵害
		顧客の信頼の喪失、ブランドイメージの低下
	風評リスクとマーケットリスク	依存リスク、競争リスク
自社起因のオペレーショナルリスクと第三者製品リスク	人材面の課題	配置転換
		AI スキル要件
	規制／コンプライアンス	進化する AI 規制
		説明責任と賠償責任の増大
	バイアスおよび倫理的懸念	差別とバイアス
		倫理的意思決定
	ESG	環境およびエネルギー関連リスク

出典:Geneva Association

生成 AI 関連保険に対する需要：企業顧客調査

ジュネーブ協会は、世界 6 大保険市場（中国、フランス、ドイツ、日本、英国、米国）で、保険の意思決定者 600 名を対象にアンケート調査を実施し、生成 AI を利用する企業におけるリスク認識度と、そのリスクを補償する保険に対する需要動向を評価しました。調査の結果、生成 AI は広く導入されているものの、その有用性に対する認識は国や地域によって異なることが分かりました。具体的には、米国と中国では有用性が高く評価されている一方で、日本、フランス、ドイツではより慎重な評価にとどまっています。これは、デジタル分野の成熟度と組織文化の違いを反映しているものと考えられます。

生成 AI 導入に際して主要な課題として挙げられたのは、人材不足に加え、データ品質の低さやデータ統合の不十分さです。対象企業の約 3 分の 1 が、これらを大きな障害と回答しています。ドイツとフランスでは、生成 AI に対する社内の抵抗（従業員や顧客の懐疑的な姿勢）が、もう一つの主要な障壁となっています。一方、米国とアジアの企業は生成 AI 導入に前向きですが、導入や管理を担える十分な専門人材の確保に苦労しています。

多くの企業はすでに生成 AI に関する問題や失敗を経験しており、リスクの高まりを強く意識しています。例えば、回答者の多くが、生成 AI が不正確で誤解を招くアウトプットを出すことや、既存システムへの統合が困難であることを報告しています。このことは、アウトプットを厳格に検証する必要性や、業務に導入する際の慎重なチェンジマネジメントの重要性を強調しています。

生成 AI に関して主な懸念事項について尋ねたところ、多くの企業が一番に挙げたのはサイバーセキュリティでした。半数を超える回答者が、生成 AI の導入によりハッキング、データ侵害、悪意ある AI 駆動型サイバー攻撃に対する脆弱性が高まることを懸念しています。次いで第三者賠償責任、即ち、生成 AI システムの誤りや不具合が顧客やパートナー企業に損害を与え、訴訟につながるリスクが 2 番目にランクインしています。第 3 位は業務の混乱で、生成 AI の停止やエラーによって事業継続が妨げられることへの不安を反映しています。注目すべき点として、風評被害は比較的低い順位にとどまっており、企業は現状、ブランドへの無形の悪影響よりも、目に見える具体的な財務リスクや法的リスクを重視していることがうかがえます。

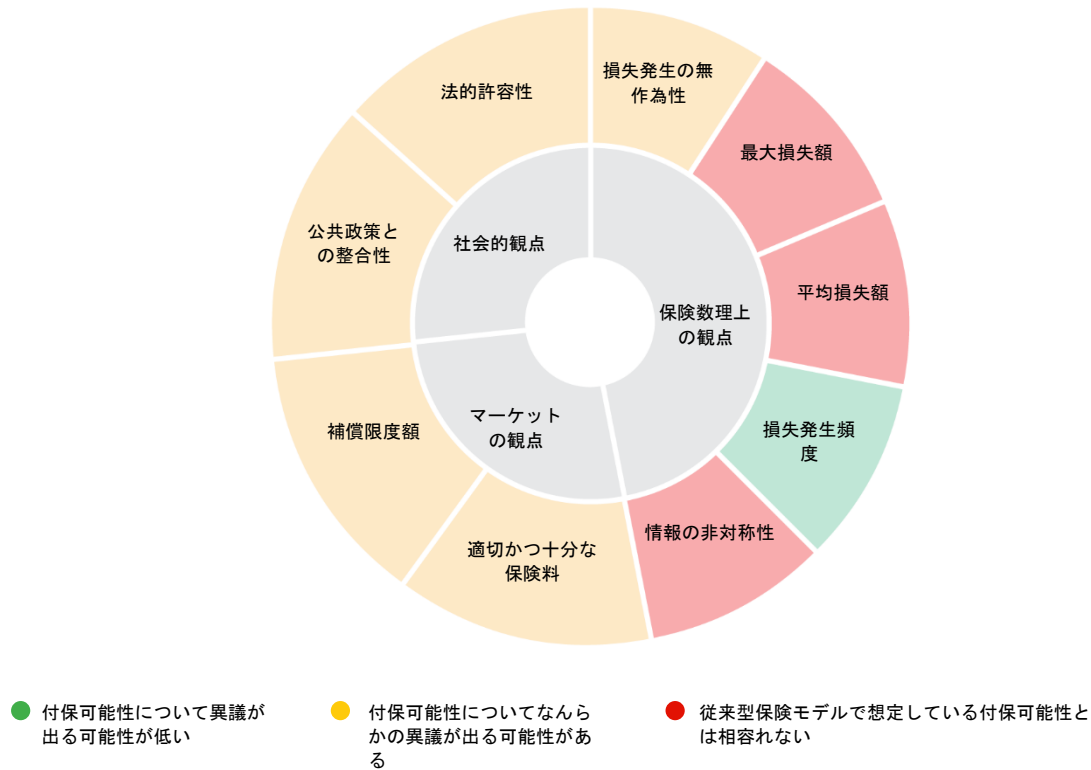
こうした懸念に呼応するように、関連リスクを補償する保険に対する明確な需要が確認できます。調査対象企業の 90% 以上が、AI 関連リスクに備えるための保険が必要だと考えており、3 分の 2 以上が、現在の保険料に 10% 上乗せされても加入したいと回答しています。需要が最も高いのは大企業や中堅企業、そしてテクノロジーや金融といった業種です。地域別に見ると、関心の高さは導入状況と一致しています。米国と中国（生成 AI 導入における二大先進国）の回答者は、AI 保険の需要が最も高く、日本、ドイツ、フランスの回答者はより慎重な姿勢を示しています。英国はその中間に位置しています。

本調査は、逆選択の力学についても示唆しています。生成 AI を広範に活用している企業や、すでに深刻な AI 関連の事故を経験した企業は、その他の企業と比べて保険への関心が高い傾向があります。

付保可能性に関する課題とマーケットの反応

生成 AI の登場は、Berliner の古典的な保険適格性基準に対して複数の点で課題を突きつけています。中でも、最大損失額が過大、平均損失額が大きい、情報の非対称性、の三つの項目が際立っています（図 1）。¹

図 1：生成 AI 関連リスクの付保可能性



出典: Geneva Association

こうした逆風にもかかわらず、保険会社は生成 AI リスクに対応する新たな保険商品の導入に向けた取り組みを始めています。市場の初期的な反応は次のとおりです。

- **既存保険契約の補償拡張。**多くの保険会社は、サイバー保険や専門職賠償責任保険（E&O 保険）などの従来の補償範囲を拡張し、生成 AI に関連するリスクを明示的に含めることで対応しています。例えば、サイバー保険では、AI 駆動型サイバー攻撃やデータ漏洩を補償範囲に含めるケースがあり、E&O 保険では、AI が生成したコンテンツに起因する誤りを補償する場合もあります。こうした補償拡張は多くの場合、特約として提供され、補償範囲を制限するためのサブリミットや条件が付されるのが一般的です。
- **引受方針の調整。**保険会社は新たな引受戦略を試行しています。一部の保険会社では、（特定の AI 障害イベントが発生した際に事前に定められた金額を支払う）パラメトリック型トリガーを導入し、不確実な環境下

での請求手続きを簡素化しようとしています。また、他の保険会社では、補償を提供する前に、被保険者の AI システムやガバナンス体制を精査する（技術監査に類似）ことで、引受基準を厳格化しています。こうした対応は、情報の非対称性の軽減に役立っています。

- **スタンドアロン型 AI 保険。**一部の保険会社は、複数の AI 関連リスクをまとめて補償する AI 専用保険を試験的に導入しています。例えば、ある保険会社は、AI 開発企業のアルゴリズムの誤りや、知的財産権侵害に対する賠償責任を一つのパッケージで補償する保険を提供することになるかもしれません。これらの商品はまだ初期段階にあり、多くの保険会社は慎重に対応を進めていますが、オーダーメイド型の AI 専用リスクソリューションを提供しようとする動きがみられます。スタンドアロン型の AI 保険が主流となるのか、あるいは既存保険の補償範囲を拡張する形で生成 AI リスクに対応していくのかは、今後の市場の動向次第です。

¹ Berliner 1982.

保険会社は可能な範囲で補償を拡張していますが、多くの場合、一定の制限を設けたり、保険料を引き上げたり、データ収集に重点を置いたりしています。これは、数十年前のサイバーリスクなど、他の新興リスクへの取り組みと類似しています。

結論と提言

生成 AI リスクに効果的に対応するために、保険会社は以下の点を考慮して、取り組む必要があります。

- **積極的に行動し、実践を通して学ぶ。** 保険会社は、完璧なデータが揃うのを待つのではなく、生成 AI リスクの境界を定義し、今すぐ試験的な補償提供を始めるべきでしょう。これは、AI リスクに対する限定的な補償拡張や試験的な保険商品の導入を通じて、経験を蓄積していくことを意味します。サイバー保険と同様に小さく始め、試行錯誤を繰り返すことで、アンダーライターは損失発生パターンと顧客ニーズをリアルタイムで把握することができます。保険会社が早期に取り組むことで、生成 AI リスクの状況が成熟するにつれて、補償をより賢く拡大していくことが可能になります。
- **リスク評価とガバナンスにおける協働。** 保険会社は、生成 AI リスクに単独では対処できません。AI 開発企業、顧客、規制当局と協力し、バイアスの検証、出力の妥当性確認、データ保護、説明責任を含むガバナンス基準を確立する必要があります。業界共通の基準や事故データを共有することで、不確実性を軽減し、賠償責任の明確化と付保可能性の向上につながります。
- **リスク軽減と備えの促進。** 保険は、強固な AI リスク管理と組み合わせることで初めて効果を発揮します。保険会社は、人による監視、バイアスのチェック、サイバーセキュリティ対策、コンティンジェンシープランなどの安全対策を顧客に求めていく必要があります。また、AI リスク監査のような付加価値サービスを提供することもできます。予防策と個別に設計された補償を組み合わせることで、企業が生成 AI の恩恵を受けつつ、リスクを適切に管理し、レジリエンス（回復力）を高めることができます。

生成 AI リスクの保険は、今後数年で進化していくことでしょう。今この分野に慎重かつ果敢に取り組む保険会社は、需要の高まりを取り込むだけでなく、この変革的なテクノロジーが生み出すリスクを、社会全体としていかに適切に管理していくかという枠組みづくりにも寄与することになります。臨機応変に対応できる態勢を維持し、ナレッジへの投資を行い、協力して取り組むことで、保険業界は生成 AI の恩恵が安全に実現されるよう支援できるのです。そのためにも、生成 AI の発展を支える強固なリスク移転メカニズムの整備が不可欠です。

参考文献

Berliner, B.1992. *Limits of Insurability of Risks*. Prentice Hall.