

保険事業における AI 利用に関する 海外主要規制・監督機関等の主な取組み

主席研究員 佐藤 智行

目 次

1. はじめに

2. AIに関する海外主要機関等の主な直近の公表文書

3. 国際機関

- (1) 国際連合「AIに関する画期的な決議」
- (2) 欧州評議会「AI国際条約」
- (3) G7「広島 AI プロセス」
- (4) 経済協力開発機構「AI原則」改訂
- (5) 保険監督者国際機構「保険における AI と機械学習の規制と監督」
- (6) ジュネーブ協会「保険における AI の規制」

4. イギリス

- (1) 政府「AI規制に対するイノベーション推進のアプローチ」
- (2) 健全性監督機構と金融行為規制機構の文書
- (3) イギリス保険協会「AIガイド」
- (4) AI Code of Conduct「保険金請求対応の AI 使用等に関する行動規範」

5. 米国

- (1) ホワイトハウス「AIの安全・安心・信頼できる開発と利用に関する大統領令」
- (2) 全米保険監督官協会「保険会社による AI システムの利用」
- (3) NY 金融サービス局「AI システムと外部消費者データと情報源の利用」
- (4) 米国損害保険協会の声明

6. EU

- (1) EU「AI法」

- (2) 欧州保険・企業年金監督機構の報告書等
- (3) 保険ヨーロッパ「キーマッセージ」

7. ドイツ

- (1) ドイツ連邦金融監督庁の見解
- (2) ドイツ保険協会の AI 法への見解

8. わが国における AI 利用に関する行政動向

9. おわりに

要旨

生成 AI の技術を用いた ChatGPT が 2022 年 11 月に登場して以降、世界で広く AI 関連のルール作りや AI 取扱いの検討が行われている。海外主要国の保険事業の規制・監督当局も AI 利用上の規制のあり方の検討、規制ガイダンスや考え方の公表を進めており、またこれらに呼応する形で保険協会等が当局による規制ガイダンスへの意見表明や AI 関連のガイド作成を行う構図となっている。

わが国においては、2024 年 4 月に総務省と経済産業省から「AI 事業者ガイドライン」が公表されたが、2024 年 5 月末現在、保険事業を直接に対象にした AI 関連の規制や指針は示されていない。

保険事業における AI 利用について、海外主要国の保険事業の規制・監督当局が現在検討している、または策定した規制ガイダンスや考え方に概ね共通して見出すことができるのは、リスクガバナンスの確立が求められている、ということである。そのリスクガバナンスの確立を担保する手段や方法について、既存の法規制で対応するのか、または新たな規制ガイダンスや考え方の公表により対応するのかについては、各国政府の方針もあって一様ではない。それらが、今後のわが国の保険事業における AI 利用に関する規制・監督動向に影響してくる可能性を見据え、それら主要国における規制ガイダンスや考え方の検討、策定の取組みなどについて、2024 年 5 月末現在の最新の状況を報告する。

1. はじめに

2022年11月の米国 OpenAI 社の ChatGPT の登場以降、人工知能（以下「AI」）、特に「生成 AI」という言葉に触れる機会が多くなってきた。従来の AI が機械学習に基づく認識や分析による特定の機能を果たす自動化であったのに対して、生成 AI は任意の命令により学習データから新たなコンテンツを生成することを特長とする。このような生成 AI の登場は、ビジネスの効率化、生産性の向上、教育や研究の進歩などのメリットをもたらすため、多くの期待が寄せられている。他方で、権利・利益の侵害、間違った情報の利用による偏見の生成とその結果もたらされる不公正な取扱い、ディープフェイクやフェイクニュースの生成などのデメリットも指摘されているところである。

このようなメリット・デメリットを併せ持つ生成 AI 技術の急速な進展に対し、特にデメリットが及ぼす影響を考慮して、国際機関や各国政府レベルでのルール作りが急務として行われてきた。そして、これらに連なる形で、海外の保険事業の規制・監督当局も生成 AI を含めた AI 利用上の規制ガイダンスや考え方を公表し、さらにこれらに呼応して保険業界等がそれらに対する意見表明や会員会社向けのガイドを作成、公表するなどしている。

わが国の保険事業でも既に AI の利用が進展している中、今のところ、規制・監督当局から保険事業を直接に対象にした AI 関連の規制や指針は示されておらず、また保険協会からも保険事業における AI 利用上の指針や刊行物は公表されていない。

海外の保険・規制監督当局による AI 利用上の規制ガイダンスや考え方は、国際保険監督者機構が AI・機械学習に特化したテーマ別レビューで参照している。同機構がそれらをもとにして保険・規制監督当局宛に統一的に推奨するアプリケーション・ペーパーの作成を表明していることから、わが国当局の今後の AI 利用に関する監督動向に影響する可能性も考えられる。また、海外の保険協会の AI 利用に関する先行した意見表明やガイドは、今後のわが国の保険業界における AI 利用にあたっての指針作りなどで参考になると考えられる。

そこで、本稿では、保険事業と AI に関する文書、報告書、ガイドなどを中心に、主として国際機関、イギリス、米国、EU、ドイツにおける政官民の主要機関による最新の取組みを報告のうえ、最後にわが国保険事業が AI を利用していくにあたって把握しておくべき行政動向にも触れる。

なお、本稿における意見・考察は筆者の個人的見解であり、所属する組織を代表するものではないことをお断りしておく。

2. AI に関する海外主要機関等の主な直近の公表文書

図表 1 と 2 は、主な国際機関、政府、保険規制・監督当局、保険協会等が 2019 年以降に公表してきた AI に関連した主要な文書、報告書である。

図表 1 は国際レベルの機関による公表文書等の一覧である。これらの中には、「AI に関する OECD 原則」のように、文書中の「原則」や考え方が各国政府や保険規制・監督当局の公表物に引用、準拠されているものもあるため、また、AI に関する国際的な取組みの趨勢を一定程度は把握しておく必要もあることから、取り上げている。図表 2 はイギリス、米国、EU、ドイツで中央政府、保険規制・監督当局、保険協会等による公表物であり、本稿ではこれらを中心として説明する。

なお、本稿で取り上げるのは、ChatGPT という生成 AI が 2022 年 11 月に出現した時期を起点として、それ以降に公表された文書や報告書などとする。生成 AI を含めて急速な AI 技術の進展がもたらす影響が認識され、それが一部に反映されているからである。

図表 1 AI 関連の国際レベルにおける主要な公表文書・報告書

公表年月	機関	文書等の名称	説明箇所
2019.5	経済協力開発機構	AI に関する OECD 原則	—
2019.6		G20 AI 原則	—
2020.1		保険分野におけるビッグデータと AI のインパクト	—
2020.1	ジュネーブ協会	保険における責任ある AI の推進	—
2023.9		保険における AI の規制	3.(6)
2023.12	G7	広島 AI プロセス	3.(3)
	保険監督者国際機構	保険における AI と機械学習の規制と監督	3.(5)
2024.3	国際連合	AI に関する画期的な決議	3.(1)
	経済協力開発機構	AI に関する OECD 原則改訂	3.(4)
2024.5	欧州評議会	AI 国際条約	3.(2)

(出典：各種資料をもとに作成)

図表 2 イギリス・米国・EU・ドイツにおける AI 関連の主な公表文書・報告書

国など	機関	公表年月	文書等の名称	説明箇所
イギリス				
政府	科学・イノベーション・技術省	2023.3	AI 規制に対するイノベーション推進のアプローチ	4.(1)
当局	健全性監督機構・金融行為規制機構	2022.10	AI と機械学習（討議文書）	—
		2023.3	AI と機械学習（回答声明）	4.(2) a
	健全性監督機構	2024.4	AI アップデート	4.(2) b
	金融行為規制機構	2024.4	書簡（(DSIT・HMT letter)	4.(2) b
協会等	イギリス保険協会	2024.2	AI ガイド	4.(3)
	AI Code of Conduct	2024.1	保険金請求対応における AI の開発、実装、使用に関する行動規範	4.(4)
米国				
政府	ホワイトハウス	2019.2	AI における米国のリーダーシップ維持のための大統領令	—
	国立標準技術研究所	2019.9	説明可能な AI の 4 原則	—
	科学技術政策局	2021.1	2020 年国家 AI イニシアチブ法	—
		2022.10	AI 権利章典のための青写真	—
	国立標準技術研究所	2023.1	AI リスク管理フレームワーク	—
	ホワイトハウス	2023.10	AI の安全・安心・信頼できる開発と利用に関する大統領令	5.(1)
当局	全米保険監督官協会	2020.8	AI 原則	—
		2023.12	保険会社による AI システムの利用	5.(2)
	ニューヨーク州金融サービス局	2024.1	保険引受と保険料設定における AI システムと外部消費者データと情報源の利用	5.(3)
協会等	米国損害保険協会	2020.10	説明可能な AI の 4 原則への意見	—
		2021.6	AI システム分類の枠組みへの経済協力開発機構宛コメント	—
		2024.2	AI に関する法律案への声明	5.(4)
EU				
政府	EC ^(注)	2019.4	信頼できる AI のための倫理指針	—
	EU ^(注)	2024.5	AI 法	6.(1)
当局	欧州保険・企業年金監督機構	2021.6	AI ガバナンス原則：欧州保険分野における倫理的で信頼できる AI に向けて	—
		2023.4	AI の機会と課題に対処するための分野別アプローチ	6.(2) b
		2023.9	EIOPA のデジタル戦略	6.(2) a
		2023.9	保険業界における AI の導入と規制の動向	6.(2) b
		2024.2	AI 法と欧州金融セクターへの影響	6.(2) b
		2024.4	欧州保険分野のデジタル化に関する報告書	6.(2) c
協会等	保険ヨーロッパ	2023.9	AI 法に関する現在進行中の三者協議へのキーメッセージ	6.(3)
ドイツ				
当局	ドイツ連邦金融監督庁	2021.6	ビッグデータと AI：意思決定プロセスにおけるアルゴリズムの使用に関する原則	—
		2023.9	イノベーションへの取組み	7.(1)
協会等	ドイツ保険協会	2023.6	AI 法案に関する三者協議の開始について	7.(2)
		2024.3	AI 法案承認を受けたコメント	

（注）EU（欧州連合）は、欧州委員会（EC）、欧州理事会、欧州議会から構成される。EC は、行政役・執行役を担う機関である。

（出典：各種資料をもとに作成）

3. 国際機関

本項目では、2023 年 9 月以降に主な国際機関から公表された 6 つの文書を取り上げる。保険事業の観点では、保険監督者国際機構の動向が最も重要である。

(1) 国際連合「AI に関する画期的な決議」

国際連合は 2024 年 3 月に開催した総会で、「AI に関する画期的な決議」¹を採択した。この決議は、AI 利用を巡る安全性の確保や国家間のデジタル格差の是正などを求めているほか、加盟国に対して AI を貧困や気候変動など世界規模の課題解決に役立てていくことや、国際的な人権の法規を無視した AI 技術の運用停止などを求めている。

なお、今回の決議案は米国が主導し、日本を含む 120 以上の国・地域が共同提案した。国際連合が AI の規制に関する決議案を採択したのは今回が初めてである。

(2) 欧州評議会「AI 国際条約」

人権、民主主義、法の支配などに主軸を置く欧州評議会²は 2024 年 5 月、AI の使用に関し人権保護や法の支配、民主主義の尊重を目的とした初の国際条約「欧州評議会の AI と人権、民主主義、法の支配に関する枠組み条約」³を採択した。2024 年 9 月に署名が予定されており、その後各国の批准を経て発効する。

条約は、AI を巡る人権保護など締約国の義務を明記し、リスクの特定・軽減を定めている。民主主義や法の支配と相いれないと判断された AI については、使用の一時停止や禁止を検討しているほか、人権侵害に対する救済措置も確保することとしている。

(3) G7「広島 AI プロセス」

広島 AI プロセスは、2023 年 5 月に開催された G7 広島サミットの結果を踏まえ、急速な発展と普及が国際社会全体の重要な課題となっている生成 AI について議論するために立ち上げられた。同年 12 月、安全、安心で信頼できる高度な AI システムの普及を目的とした指針と行動規範からなる初の国際的政策枠組みとして「広島 AI プロセス包括的政策枠組み」⁴がまとまり、G7 首脳に承認された。

¹ 正式名称は、「持続可能な開発のための安全、安心で信頼できる AI システムに係る機会活用」(Seizing the opportunities of safe, secure and trustworthy artificial intelligence systems for sustainable development)である。

² 欧州評議会は EU (欧州連合) とは別組織である。フランスやドイツといった EU 各国やイギリスなど 46 か国が加盟しており、日本や米国などがオブザーバーとして参加している。批准は欧州以外の非加盟国にも門戸が開かれているため、締約国は世界各地に広がる可能性があると考えられている。

³ 原題は、“Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy and the Rule of Law”である。

⁴ 次の 4 点を内容とする。①生成 AI に関する G7 の共通理解に向けた OECD レポート、②すべての AI 関係者向けおよび AI 開発者向け広島プロセス国際指針、③高度な AI システムを開発する組織向けの広島プロセス国際行動規範、および④偽情報対策に資する研究の促進等のプロジェクトベースの協力

(4) 経済協力開発機構「AI 原則」改訂

経済協力開発機構（Organisation for Economic Co-operation and Development：以下「OECD」）は 2019 年 5 月に AI 原則⁵を採択していたが、2024 年 5 月にその改訂版を採択した^{6,7}。生成 AI の出現など、近年の AI 技術の進歩を踏まえて改訂されたこの原則は、プライバシー、知的財産権、安全性、情報の完全性など、AI 関連の課題に今まで以上に踏み込んで対応している。

より具体的には、生成 AI の利用の急拡大に対応して、AI 開発者らに対し、偽・誤情報への対処を求める項目が新たに盛り込まれたほか、AI の透明性や説明責任に関する項目が見直され、AI 開発者らに対し AI の能力や限界に関する情報や AI の学習データや生成過程に関する情報の開示が求められるようになっている。

(5) 保険監督者国際機構「保険における AI と機械学習の規制と監督」

保険監督者国際機構（International Association of Insurance Supervisors：以下「IAIS」）は、2023 年 12 月に公表した文書「保険における AI と機械学習の規制と監督：テーマ別レビュー」⁸で、初めて「生成 AI」という言葉を使用した。

この文書によると、加盟各国の監督当局は、AI と機械学習がもたらす大きな機会（メリット）と新たなリスク（デメリット）を認識し、後者への懸念についてはハイレベルの原則、より詳細な基準やガイダンスを公表することにより対応してきたことが説明されている。

IAIS は、さらに、効果的かつ国際的に整合的な保険監督を促進する立場から、AI・機械学習の分野に関する IAIS ガイダンスの必要性があるかどうかについてテーマ別レビューにより調査した。その結果、既存の保険基本原則（Insurance Core Principle：以下「ICP」）⁹がリスク管理を目的とした設計思想とされていることから、各国監督当局に対して引き続きこの ICP が AI・機械学習モデルの使用によるリスク管理の監督についても適用可能であり適切であると考えたとしている。しかし一方で、AI・機械学習に特化した政策対応は、拘束力のある法律や基準、ハイレベルの AI 原則、拘束力のない監督指針、討議文書の発行など、世界各国で様々な監督手法を組合わせて行われていることが浮き彫りになったとしている。そして、明らかになった最も重要なことは、調

⁵ 2019 年 5 月、42 カ国の間で、AI に関する初の国際的な政策ガイドラインとして正式に採択され、AI システムが健全、安全、公正かつ信頼に足るように構築されることを目指す国際標準を支持することで合意した。

⁶ 2019 年の OECD 閣僚理事会による採択時当初から、5 年経過後の改訂が求められていた経緯にある。

⁷ 今回の AI 原則の改訂に先立つ 2023 年 12 月、OECD が使用する「AI システム」という言葉の定義については、先行して改訂が行われていた。

⁸ 原題は、“Regulation and supervision of artificial intelligence and machine learning (AI/ML) in insurance: a thematic review”である。

⁹ ICP は、健全な保険セクターの促進および保険契約者の適切な保護のために必要な保険監督にあたっての基本原則を定めた監督文書で、IAIS のメンバーである国は ICP に則った監督制度を実施することが推奨されている。また、ICP は原則としてすべての保険会社、グループに適用される。

査に応じたすべての国・地域が、生成 AI や大規模言語モデル¹⁰から生じるリスクを含め、急速に進化するリスクに対応するために、AI・機械学習モデルの分野で継続的な監督強化の必要性を期待している、ということであった。

このような背景から、IAIS は 2024 年中に AI・機械学習モデルの分野に関するアプリケーション・ペーパー¹¹を作成する方針を表明している。

(6) ジュネーブ協会「保険における AI の規制」

ジュネーブ協会¹²が 2023 年 9 月に公表した報告書「保険における AI の規制：消費者保護とイノベーションのバランス」^{13,14}は、偏見、差別、引受拒否など保険と AI との関係で問題となりやすい事項を取り上げ、保険実務の視点で分かりやすく説明されている¹⁵。保険実務家が今後の保険業務への AI の導入や利用にあたって深く考えるうえで参考になるものと思われる¹⁶。

この報告書は、生成 AI を含めたイノベーションが急速に進展する中、消費者保護とイノベーションのバランスを探究する目的で、保険事業における AI の規制に関する調査結果をまとめたものであり、規制・監督当局や保険業界などの利害関係者に貴重な洞察と推奨事項を提供している。報告書冒頭の重要な論点において、「生成 AI の台頭は、政策立案者や規制・監督当局に対応を迫り、その結果、(原則策定など) 分野横断的な取り組みや保険に特化した取り組みが行われるようになった。」と述べられている。

報告書の内容は図表 3 のとおりであり、特に第 6 章「AI の規制のあり方」では各国の保険当局において AI 関連リスクが既存の法規制でどのように関連付けされているのかを横断的に集約して比較し (図表 4 参照)、その分析のもとに第 7 章の端的な「結論

¹⁰ 「大規模言語モデル」とは、「計算量」「データ量」「モデルパラメータ数」の 3 要素を大量に使ってトレーニングされた自然言語処理のモデルのことを指す。

¹¹ IAIS が作成する文書の一類型である。IAIS の監督文書の特定のテーマにおいて、原則や基準の一律な解釈や適用が難しい場合に、事例の提供により助言や推奨の考え方を示している。

¹² ジュネーブ協会 (The Geneva Association) は、世界の保険会社約 70 社の CEO で構成される保険業界のシンクタンクである。会員企業、学術機関、多国間組織との協力による厳密な調査に基づいて、保険業界に影響を与える可能性のある主要なリスク分野を調査し、提言を作成し、関係者が議論するための場を提供している。

¹³ 原題は、“Regulation of Artificial Intelligence in Insurance: Balancing consumer protection and innovation”である。

¹⁴ 報告書の概要版に日本語版があり、次の URL から入手可能である。

https://www.genevaassociation.org/sites/default/files/2023-10/Regulation%20of%20AI%20report_summary_JP.pdf

¹⁵ ジュネーブ協会は、AI 関連について、2023 年 9 月に公表した報告書のほかに、2020 年 1 月にも「保険における責任ある AI の推進 (Promoting Responsible Artificial Intelligence in Insurance)」を公表している。この 2020 年 1 月の報告書では、政府機関、非政府組織、民間企業が発行した多数の倫理ガイドラインの一部を分析し、特定された AI の責任ある使用に関する 5 つの中核原則のうち、保険会社にとって解釈と実装が特に複雑な 2 つ、①透明性と説明可能性、②公平性について詳しく取り上げられている。また、これらの原則を適用する際に生じる主要な相反事項への対処方法についても検討されている。

¹⁶ 報告書の作成にあたっては、次の機関、団体、保険会社などが助言をしたり、意見を述べたりしている。OECD、規制・監督当局のイギリス金融行為規制機構、欧州保険・企業年金監督機構、シンガポール金融監督庁、民間保険事業者団体のイギリス保険協会、米国損害保険協会、および保険会社のゼネラル・AIG、スイス再保険、平安保険など。

と提言」（図表 5 参照）に結び付けている。

図表 3 「保険における AI の規制：消費者保護とイノベーションのバランス」の内容

1. はじめに
2. AI の定義
 - 2.1 定義
 - 2.2 高度分析 vs AI
 - 2.3 機械学習
 - 2.4 既存の定義の発展
3. 保険における AI のメリット
 - 3.1 AI が保険を変える様相
 - 3.2 AI のメリット
4. AI とアンダーライティングにズームイン
 - 4.1 アンダーライティング・プロセスと AI：何が新しいのか？
 - 4.2 意思決定の変化：AI 対人間
5. AI に関するリスクと懸念
 - 5.1 透明性と説明可能性
 - 5.2 偏見、間接差別、公正さへのリスク
 - 5.3 データ入力
 - 5.4 過度の保険料差別化と排除
 - 5.5 保険会社が AI 関連のリスクと懸念に対処し、軽減させる方策
 - 5.6 AI リスクは保険にとって新しいリスクか？
 - 5.7 AI リスクと保険のビジネスモデル
6. AI の規制のあり方
 - 6.1 管轄地域特有の動き
 - 6.2 既存の規制が AI リスクにどう対処するかの方策
 - 6.3 イノベーションと顧客保護のバランス
7. 結論と提言

（出典：The Geneva Association, “Regulation of Artificial Intelligence in Insurance: Balancing consumer protection and innovation”（2023.9））

図表 4 AI 関連リスクの EU、中国、イギリスおよび米国における既存の法規制への取込み

状況	法規制	説明
偏見、差別、および公平性	○EU：人種平等法 ○イギリス：平等法 ○米国：各州の法規制	民族などによる差別の禁止
	○EU：ジェンダー指令 ○イギリス：平等法 ○米国：各州の法規制	性差などによる差別の禁止
	○EU：GDPR 5 条 ○イギリス：UK GDPR 13-21 条 ○米国：カリフォルニア州消費者プライバシー法	個人データの合法的、公正、かつ透明性のある利用および処理を保証
	○EU：保険販売業務指令（IDD）20 条 ○米国：NAIC 不公正取引慣行モデル法 ○中国：インターネット規制に関する措置 17 条	保険商品に消費者の需要とニーズを満たすことを要求
透明性およびデータガバナンス	○EU：保険販売業務指令（IDD）20 条	保険会社が顧客に客観的な商品情報を提供するよう義務付け
	○EU：GDPR 5 条、13 条、14 条 ○イギリス：UK GDPR ○中国：個人情報保護法 5 条	データの使用と処理におけるオープン性と透明性の義務付け
	○EU：GDPR 5 条	データの妥当性、関連性、および

状況	法規制	説明
	○イギリス：UK GDPR ○米国：グラム・リーチ・ブライリー法 ^(注) 、公正信用報告法、カリフォルニア州消費者プライバシー法	正確性など、データ処理に関連する原則を概説
	○EU：GDPR 30 条 ○イギリス：UK GDPR 35 条	処理活動の記録を保持することを要求
人間による監視	○EU：GDPR 22 条 ○イギリス：UK GDPR ○中国：個人情報保護法 24 条	自動化された意思決定に異議を唱える権利を提供
	○EU：ソルベンシー II 枠組指令 41 条 ○中国：保険法 5 条	健全かつ慎重な事業運営のための効果的なガバナンスシステムの要求

(注) 1999 年 11 月に成立した金融制度改革法を指し、金融持株会社が子会社方式により銀行業務、証券業務、保険引受業務、および金融関連業務を営むことが可能になった。

(出典：The Geneva Association, “Regulation of Artificial Intelligence in Insurance: Balancing consumer protection and innovation” (2023.9))

図表 5 ジュネーブ協会の「保険における AI の規制」の提言

1. AI の慎重な定義

規制目的での AI の定義をめぐる議論が続いている。実践可能な定義として、AI を自己学習アプリケーションに限定し、保険分野で確立された規制が過剰になることを避けるため、機械学習に焦点を当てるべきである。

2. 既存の規制の適用

AI 関連のリスクについては、規制当局が既存の技術に偏りし過ぎない中立的な枠組みを活用し、AI の文脈でこれらの規制を適用するためのガイダンスを更新することが極めて重要である。

3. 原則に基づく規制の策定

急速に進化する AI の性質上、その規制は複雑で移り変わりの激しいものとなる。現行の規制を基礎とした原則に基づく規制アプローチは、イノベーションと競争を阻害することなく AI のリスクを管理する最も有望なアプローチとなる。

4. 保険における AI の特徴への配慮

保険分野においては、既存の規制の枠組みが有効であることが証明されている。分野横断的な規制は、テクノロジーなどの規制が緩い分野や、AI による意思決定が不可逆的で深刻な結果をもたらす可能性がある分野に比べて、はるかに効果が低い。

5. 顧客の成果への焦点

データガバナンス^(注)の枠組みは、保険数理上の公平性を確保し、差別を防止する上で重要な役割を果たすが、リスクを評価し、保険料を決定するために使用される個々の料率決定要素の規制を過度に重視しないことが重要である。顧客の成果に焦点を当てたバランスの取れたデータガバナンスのアプローチは、公正かつ非差別的な方法でイノベーションを促進するのに役立つ。

6. 国際的な協力

各法域は、保険における AI の活用事例に特化したガイダンスを策定するために協力すべきである。規制やガイダンスが各法域で統一されることで、保険会社は AI がもたらす課題や機会をより効果的に乗り越えることができるようになる。

(注) データガバナンスとは、企業において、データの取扱いを全社横断で監督し、効果的かつ安全に活用できる状態に治めることをいう。

(出典：The Geneva Association, “Regulation of Artificial Intelligence in Insurance: Balancing consumer protection and innovation” (2023.9))

4. イギリス

イギリスでは、政府が関係者の意見を十分に聞き取りながら、AIに関する規制のあり方の検討を行っている。金融サービス事業者の規制・監督当局である健全性監督機構と金融行為規制機構は、政府の方針を尊重し、連携した対応を行っている。保険業界も政府の定めた原則に立脚したガイドや行動規範を策定している。

(1) 政府「AI 規制に対するイノベーション推進のアプローチ」

政府は 2023 年 3 月、科学・イノベーション・技術省が中心となってまとめた「AI 規制に対するイノベーション推進のアプローチ」と題する白書¹⁷（以下「AI 白書」）を公表した。この白書は、AI 技術の急速な進化に伴って一定の規制の必要性が認識されるようになったが、過剰な規制によって技術進歩を阻害するのではなく、技術を発展させつつ、同時に倫理的な懸念やリスクを適切に取り扱うための枠組みを政府として提供することを主眼として作成されたものである。AI 技術の懸念やリスクに対して立法措置を講じると、企業に対して過度な負担を強いる危険性があることに鑑み、政府として大きな枠組みを示したうえで、事業分野別の規制当局に主導権を与えて規制する構図となっている¹⁸。

a. 5 原則

AI 白書では、大きな規制枠組みを構成する元になる 5 つの原則が図表 6 のとおり示されている。これらは、OECD の AI 原則を基礎としており、5 原則中の 4 原則が OECD の AI 原則とほぼ同じ内容となっている。この点について、AI 白書に意見を述べた者からは、多国間で合意された OECD の AI 原則を準用したことは国際的な連携と相互運用性を促進するものであるとして、支持されている。AI 白書では、この 5 原則について、責任ある AI の設計、開発、利用の重要な要素を定めるものとして、すべての事業分野に横断して解釈、適用することが提案されている。

図表 6 AI 規制に対するイノベーション推進のアプローチの 5 原則（「AI 規制の 5 原則」）

原則	原則の説明
安全性・セキュリティ・堅牢性	○AI システムは堅牢で安全に機能し、リスクが継続的に特定、評価、管理されるべきである。 ○規制当局は、規制対象に対して、AI システムが設計どおりに信頼性を保持し、意図したとおりに機能するための措置を導入する必要がある。
適切な透明性と説明可能性	○AI システムは適切に透明であり、説明可能でなければならない。 ○透明性とは、AI システムに関する適切な情報（例えば、AI システムがいつ、どのように、どんな目的で使用されるかに関する情報）が関係者に伝達されることを指す。

¹⁷ 原題は、“A pro-innovation approach to AI regulation”である。

¹⁸ なお、イギリスでは科学・イノベーション・技術省を中心とする政府による対応のほかに、AI 規制に関する機能を専門に担当する AI 機関の設立などを内容とする議員立法が提出されており、2024 年 5 月現在、貴族院で審議が行われている。

原則	原則の説明
	○説明可能性は、関係者が AI システムの意思決定プロセスにアクセスし、解釈し、理解することが可能である程度を指す。 ○適切な透明性と説明可能性のレベルは、公衆の信頼を高め、AI の採用を促進する重要な要因である。 ○AI システムを十分に説明できない場合、AI の提供者や利用者は法律を無意識に破る、権利を侵害する、害を引き起こす、または AI システムのセキュリティを危険にさらす可能性がある。
公平性	○AI システムが個人や組織の法的権利を損なわず、個人を不当に差別せず、不公正な市場結果を生み出さないようにすべきである。 ○公平性とは、平等や人権、データ保護、消費者や競争法、公法や一般法、および弱者を保護する規則など、法規制の多くの分野に根ざしている概念である。 ○規制当局は、担当する規制分野について、AI システムに適用される公平性に関する説明を考案し、関連法規制、技術基準などを考慮に入れたガイダンスを作成する必要がある。 ○規制当局は、担当する規制分野で AI システムがこれら公平性の説明に基づいて設計、展開、使用されていることを確認する必要がある。 ○公平性の概念が交錯する規制分野でそれが重要である場合、規制当局同士が共同ガイダンスを開発することが優先事項となる場合もある。
説明責任とガバナンス	○AI システムの提供と使用の効果的な監督を確保するために、AI ライフサイクル^(注)全体にわたって明確な責任のラインが確立されるべきである。 ○AI ライフサイクルの当事者は、原則を考慮し、AI ライフサイクルのすべての段階で原則の実施に必要な措置を導入すべきである。 ○規制当局は、AI サプライチェーン内の適切な当事者に規制遵守と適切な慣行に関する明確な期待を示す方法を模索する必要がある。
異議申立と救済	○AI の意思決定や結果に異議を唱えることができるべきであり、その結果が有害であるか、有害のリスクを生じる可能性がある場合には、異議を唱えることができる仕組みを確立すべきである。

(注) AI ライフサイクルとは、データ収集、モデル作成、モデルトレーニング、展開、評価など AI プロジェクトにおいて必要なデータを提供するためのデータ管理のプロセスのことをいう。

(出典：GOV.UK, “A pro-innovation approach to AI regulation” (2023.3) ほかをもとに作成)

b. AI を利用する保険会社への政府の規制アプローチ

AI 白書では、AI を利用する事業者が直面する可能性のある法的リスクに効果的に対処しつつ、自信をもって AI イノベーションに投資し発展させることができるようにするため、政府が担当する法規制分野で積極的な介入、支援を行っていくことが述べられている。

その例として説明に用いられているのが図表 7 に示す保険事業の事例である。AI を利用した保険料設定について、AI 白書の掲げる 5 原則のうちの 1 つである「公平性」を担保しようとするだけでも、既存の法規制でどれが適用となり、どれが適用とならないか、錯綜する傾向にあることが政府による検討で明らかになった¹⁹。こうした問題点を改善し、法規制の不確実性と差分に対処するために、規制当局が適切な方法で介入する一方、既存の規制を活用することにより、イノベーションを推進する規制の枠組みを

¹⁹ 例えば、AI システムによる保険料設計において、平等法が規定する保護されるべき特性である人種や宗教を根拠にした保険料設計は平等法違反となるが、そのほかにもこのような AI による保険料設計が消費者保護法やデータ保護法といった法的枠組みの対象にもなりうるような状態をいう。

実現し、規制当局、業界との協力を促進することを表明している。

AI 白書では、保険事業のもう一例が図表 8 のとおり取り上げられている。こちらは規制当局の具体的な名称を挙げて、それら規制当局が保険会社による適切なコンプライアンス行動が可能となる共同規制ガイダンスを作成、提示のうえ、AI を利用した保険商品の導入を後押ししようとするものである。

図表 7 保険事業における公平な AI 利用と規制上の問題点

- | |
|---|
| <p>○X 保険会社</p> <ul style="list-style-type: none">・顧客のリスクを正確に反映した保険料を設定するために、AI を活用した新しいアルゴリズムを設計している。・公平な保険料を設定し、消費者の信頼を構築することが、同社のブランドの重要な要素になっている。・事業運営にあたり、関連する法律とガイダンスに確実に準拠することを優先事項としている。 <p>○保険料設定に AI を使用するときの考慮事項</p> <ul style="list-style-type: none">・データ保護法^(注1)、平等法^(注2)、一般消費者保護法^(注3)を含む様々な法的枠組みの対象となる可能性がある。・金融サービス・市場法のような分野別規制の対象となる可能性もある。 <p>○問題点</p> <ul style="list-style-type: none">・AI を利用する X 保険会社にとって、どのようなルールが関連するのかを特定したうえで、個々の利用事例に自信をもって適用することは難しい。・X 保険会社のような企業に対して規制状況を案内するための支援が不足している。・分野横断的な原則がなかった。また、規制システム全体の調整も限られてきた。 |
|---|

(注1) Data Protection Act 2018 を指す。

(注2) Equality Act 2010 を指す。

(注3) Consumer Protection from Unfair Trading Regulations 2008 を指す。

(出典：GOV.UK, “A pro-innovation approach to AI regulation” (2023.3) ほかをもとに作成)

図表 8 保険事業における規制当局との連携による AI 利用推進のイメージ

- | |
|---|
| <p>○Y 保険会社</p> <ul style="list-style-type: none">・新開発していた AI 利用のアプリの導入を延期した。・延期の理由は、AI 主導の保険商品について、現行の関連規制要件がつかはざ状態のため、適切なコンプライアンス行動を特定することが困難だったからである。 <p>○今後の予想</p> <ul style="list-style-type: none">・AI 規制に対するイノベーション推進の政府アプローチのもと、情報コミッショナーオフィス^(注1)、平等人権委員会^(注2)、金融行為規制機構、その他の関連規制当局が共同でガイダンスを作成することが予想される。・AI に関連する規制要件がより明確になるとともに、保険や消費者向け金融サービスとの関連でこれらの要件を満たす方法のガイダンスが提供されるようになる。 <p>○AI 白書が提案する規制枠組み環境</p> <ul style="list-style-type: none">・AI 白書の 5 原則に対応するために規制当局が発行する新規または更新されたガイダンスにより、事業継続の道筋を立てることができる。・規制当局間の協力の結果として発行される共同規制ガイダンスに従うことも可能であり、リスク評価や透明性措置のテンプレート、関連する技術基準（透明性や偏見緩和のための国際基準など）など、規制当局が提供する一連のツールを利用することもできる。・規制当局間の協力とガイダンスに基づく実践により、Y 保険会社は規制の状況を把握しながら業務運営を行い、差別などの特定のリスクに対処することが容易になる。・Y 保険会社は、AI を利用した保険商品について責任をもって展開できるようになる。 |
|---|

(注1) Information Commissioner's Office を略して、ICO と通称されている。日本の個人情報保護委員会に相当する機関で、イギリスにおけるデータ保護機関である。

(注2) 平等人権委員会 (Equality and Human Rights Commission) は、2006 年平等法に根拠を有する独立した国家機関として、2007 年 10 月に設立された。人権と平等について、広範な問題を取り扱う組織である。

(出典 : GOV.UK, “A pro-innovation approach to AI regulation” (2023.3) ほかをもとに作成)

(2) 健全性監督機構と金融行為規制機構の文書

a. PRA&FCA「AI と機械学習」

健全性監督機構 (Prudential Regulation Authority : 以下「PRA」)²⁰の「ルールブック」や、金融行為規制機構 (Financial Conduct Authority : 以下「FCA」)²¹の「ハンドブック (Handbook)」には、現在のところ、AI に直接的に触れた規制事項はない。

保険会社を含む金融サービス事業者が AI を利用することが多くなってきた背景に鑑み、PRA と FCA の両金融監督当局は共同して 2022 年 10 月、監督当局が金融サービス事業者による AI の安全かつ責任ある導入を支援するうえで果たすべき役割に関する意見を求めるため、今後の AI へのアプローチのあり方の概要を説明した討議文書「AI と機械学習」²²を発出して、利害関係者からの意見を募った。

そして、1 年後の 2023 年 10 月、寄せられた意見を集約した回答声明「AI と機械学習」²³を公表した。上記討議文書発出時点では、生成 AI の ChatGPT が世間に登場していなかったため、生成 AI 関連の記載はなかったが、その後の ChatGPT の登場を受け、生成 AI の使用に伴う新たなリスクを課題認識する必要がある旨の意見があったことが紹介されている。

この回答声明の位置付けとして、政策提案は含まれておらず、また AI と機械学習に関する現在または将来の規制提案を監督当局がどのように明確化、設計、実施することを検討しているのかを示唆するものではないことが明記されている。しかし、回答声明で提起された図表 9 の内容から、消費者保護策の必要性や委託先等第三者との関係明確化等の点で、監督当局としての AI に対する将来的なアプローチの方向性を推察することができる。

²⁰ PRA は、金融サービス市場法を根拠法として設置された規制・監督当局であり、金融サービス事業者の健全性の確保を担う。

²¹ FCA は、金融サービス市場法を根拠法として設置された規制・監督当局であり、金融サービス事業者の業務上の行為の監督を担う。

²² PRA&FCA, “DP5/22 - Artificial Intelligence and Machine Learning” (2022.10)

²³ PRA&FCA, “FS2/23 - Artificial Intelligence and Machine Learning” (2023.10)

図表 9 フィードバック声明「AI と機械学習」で提起された主な事項

項目	内容
AI 規制の原則ベースアプローチ	○規制により AI を定義することは有用ではない。 ○多くの回答者は、AI の特定の特性や AI によってもたらされる、あるいは増幅されるリスクに焦点を当てている。また、AI の定義について代替的な原則ベースまたはリスクベースのアプローチを使用している。
進化する AI に対する「最新の」規制ガイダンス	○AI の急速な進化を考えると、規制当局が「最新の」規制ガイダンス、すなわち定期的に更新されるガイダンスやベストプラクティスの例を提供することが有益である。
AI 官民フォーラム ^(注) による継続的な業界関与	○業界の継続的な関与が重要である。 ○AI 官民フォーラムなどの取組みは有益であり、継続的な官民連携のテンプレートとして機能する可能性がある。
AI 規制における国内外の調整と連携	○AI に関する規制の状況は複雑かつ細かい。 ○国内外を問わず、規制当局間の調整と協調を強化することが有用である。
データリスク管理における規制の連携	○特にデータ規制が細分化されており、データリスク、中でも公平性、偏見、保護された特性の管理に関連するデータリスクに対処するために、規制をより連携させることが有益である。
消費者保護を中心とした規制の焦点	○偏見、透明性、弱い立場の消費者からの搾取といったリスクを考えれば、規制の重要な焦点は消費者の成果と保護であるべき。
第三者に関する AI ソリューションの規制ガイダンス	○第三者によるモデルとデータの使用の増加が懸念事項であり、より多くの規制上のガイダンスが役立つ分野である。 ○AI ソリューションを提供する第三者は、自社の AI 製品の責任ある開発、独立した検証、継続的なガバナンスを裏付ける証拠を提供する必要がある。
事業部門間の連携による AI リスク軽減	○AI システムは複雑で、企業全体の多くの分野に関与する可能性がある。 ○AI リスクを軽減するためには、事業部門や機能間の連携したアプローチが有効であろう。特に、データ管理チームとモデル・リスク管理チームとの緊密な連携は有益であろう。
企業ガバナンス構造による AI リスクの管理	○AI リスクへの対処について、既存の会社のガバナンス構造（および上級管理職・認証制度（SM&CR）などの規制的枠組み）で十分である。

(注) 金融分野における AI の利用と影響について理解を深めるために、FCA とイングランド銀行が 2020 年 10 月に設立した官民両部門からなるフォーラムのことをいう。

(出典：PRA&FCA, “FS2/23 - Artificial Intelligence and Machine Learning” (2023.10) ほかをもとに作成)

b. FCA「AI アップデート」と PRA「書簡」

FCA と PRA は 2024 年 4 月、科学・イノベーション・技術省を中心とする政府が方針提示している事業分野別の AI 規制原則の導入²⁴について、意見²⁵を提出した。もともと、金融分野の規制・監督当局として、どのような規制アプローチを講じるのかについて、計画案を公表するように要請されていたものへの対応である。

²⁴ イギリスにおける AI 戦略の司令塔の役割を果たしている科学・イノベーション・技術省は 2024 年 2 月、保健医療、通信、金融などの規制当局に対して、「イギリスの AI 規制原則の導入：規制当局向けの初期ガイダンス」(Implementing the UK's AI regulatory principles: initial guidance for regulators)を発出した。同省が今後、上記規制当局向けの AI 規制に関するガイダンスを作成するにあたり、2023 年 3 月の AI 白書で述べた AI 5 原則を基礎として、各分野の最終的な規制・監督の実施は各規制当局の裁量に委ねる方向のガイダンス内容とすることを示していた。

²⁵ FCA の意見は「AI アップデート (AI Update)」に、PRA の意見は「科学・イノベーション・技術省、財務省宛書簡 (DSIT-HMT letter)」にそれぞれまとめられている。

回答文書である FCA「AI アップデート」と PRA「書簡」によると、いずれも政府の AI 原則に基づいた事業分野主導のアプローチを歓迎していることが示されている。特に、AI 白書で示された事業分野横断的な AI の 5 原則が両規制当局の規制アプローチと一致しており、それぞれの規制の枠組みが AI 規制の 5 原則を満たす、またはそれに対処することができるとして、図表 10 のような関連付けを行っている。そして、既存の FCA と PRA の規制枠組みが、AI 規制の 5 原則に沿って、リスクに対処しながら金融事業と経済全体に利益をもたらす方法で AI イノベーションを支援するのに適切であるという見解が示されている。

しかしながら、FCA は、監督対象の金融サービス事業者に対する今後の規制アプローチについて、AI の進化スピード、規模、複雑性に適応する必要があり、AI モデルのテスト、検証、理解により重点を置くとともに、イギリス政府の掲げる AI 原則の 1 つである「説明責任」原則を適用させる必要があると指摘している。

他方、PRA は金融サービスにおける AI の継続的な導入が金融安定性に影響を及ぼす可能性があるとしており、2024 年中にこれに関する分析を深め、その結果の取扱いを検討するとしている。

現地の法律事務所、Hogan Lovells²⁶によると、金融サービス事業者の AI 利用に関する両規制当局の方針が必ずしも明確ではなく、また、金融サービス事業者が AI システムの導入でベンダーと交渉する場合、ベンダーが特定の規制の適用関係を慎重に見定めることは少ないため、金融サービス事業者にとってはむしろ、ベンダーに提示可能なガイダンスの存在を歓迎し、明確な実践例の提示が役に立つだろうと指摘している。

図表 10 AI 5 原則と FCA・PRA による既存規制枠組みとの主な関連付け

原則	FCA が所管する既存の規制枠組み	PRA が所管する既存の規制枠組み
安全性・セキュリティ・堅牢性	○業務行為プリンシプル ・十分な技能、注意力および勤勉さをもって業務を遂行する。 ・責任を持って、効果的かつ適切なリスク管理体制で業務を組織し、管理するために合理的な注意を払う。 ○基準条件 ・特定の規制対象業務を行う会社は、ビジネスモデルが適切であることを保証する必要がある。 ○上級管理職における計画、システム、および統制 ・事業継続に関する要求 ・組織上の一般的な要求 ・リスクコントロール ・重要な機能のアウトソーシング ・オペレーショナル・リスク ・オペレーショナル・レジリエンス	○SS2/21^(注1)（アウトソーシングと第三者リスク管理） ・PRA の方針は、サード・パーティ・ビジネス・サービスが重要なビジネス・サービスをサポートしている場合、規制対象会社はサプライヤーからのリスクを管理する義務を負う。

²⁶ Hogan Lovells, “UK: The FCA, Bank of England and PRA issue their strategic approach to regulating AI in response to the government’s AI White Paper” (2024.4)

原則	FCA が所管する既存の規制枠組み	PRA が所管する既存の規制枠組み
適切な透明性と説明可能性	○消費者への義務 ・誠実に行動する必要がある、これは顧客に対する誠実さ、公正かつオープンな取引を特徴とする。 ・顧客の情報ニーズを満たし、顧客が効果的かつタイムリーかつ適切な情報に基づいた意思決定を行えるようにする必要がある。 ○業務行為プリンシプル ・消費者への義務が適用されない場合、顧客の情報ニーズに十分な配慮を払い、明確かつ公正で誤解を招かない方法で顧客とコミュニケーションをとることを求める。	○DP5/22^(注1) (AI と機械学習) ・イングランド銀行と PRA は、FCA とともに、AI と機械学習のモデルの透明性と説明可能性の欠如が金融システムにもたらす潜在的リスクを指摘している。
公平性	○消費者への義務 ・会社は、顧客により結果をもたらすために、より積極的な役割を果たすことが求められる。 ・会社は誠実に行動し、予見可能な損害を引き起こすことを回避し、顧客が金融上の目標を追求できるように支援することが求められる。 ○業務行為プリンシプル ・会社は、利益相反の管理に関するプリンシプル 8 と、助言と裁量的決定の適合性に関するプリンシプル 9 を考慮する必要がある。 ・会社が小売市場での事業を行っておらず、消費者への義務が適用されない場合、会社は顧客の利益を十分に配慮し、公正に扱い（プリンシプル 6）、明確かつ公正で誤解を招くことのない方法で情報を伝達する必要がある（プリンシプル 7）。 ○平等法^(注2) ・この分野横断的な法律（金融部門と非金融部門の区別なく適用される）は、保護されるべき特性に対する差別を禁止する。	○PRA は、公平性が PRA の権限に関連するとみなされる場合、会社が正当な理由とともに公正性を自ら定義することを期待する。
説明責任とガバナンス	○行為プリンシプル ・プリンシプル 3 は、会社が適切なリスク管理システムを使用して、責任を持って効果的に業務を組織し、管理するために合理的な注意を払うことを求めている。 ○上級経営陣における計画、システム、および統制 ・会社は、明確で透明性のある一貫した責任系統を持つ組織構造、会社がさらされている、またはさらされている可能性のあるリスクを特定、管理、監視、報告するための効果的なプロセス、健全な管理・会計手続きや情報処理システムの効果的な管理・保護措置を含む内部統制の仕組みを含む、強固なガバナンス体制を持たなければならない。 ○上級管理職と認定制度	○上級管理職と認定制度 ・規制対象会社は、1 名以上の上級管理職（会社内の主要な意思決定者）が、会社の主要な活動、事業分野、および経営機能に対する全体的な責任を有することを保証する必要がある。 ・会社の活動、事業分野、経営機能に関連する重要業務に AI を利用することは、上級管理職の責任範囲に含まれるものとして設定されなければならない。 ・上級管理職個人は、その責任範囲内で規制違反があった場合、また、それを防止するための合理的な措置を講じなかった場合、責任を問われる可能性がある。 ○PRA ルールブック

原則	FCA が所管する既存の規制枠組み	PRA が所管する既存の規制枠組み
	・この規制は上級管理職の説明責任を強調しており、AI の安全かつ責任ある使用に関連する。 ○消費者への義務 ・会社は、顧客により結果をもたらすために行動する義務が戦略、ガバナンス、リーダーシップに反映されていることを確認する必要がある。	・「一般的な組織要件」と「業務管理態勢」は、銀行と保険会社それぞれのガバナンスルールの概要を説明している。 ・リスク管理に関する PRA ルールブックのセクション（2.1A）は、リスク管理の効果的な手順に関する高レベルのガバナンス要件を焦点としている。
異議申立と救済	○紛争処理 ・会社が苦情にどのように対処すべきかを詳述する規則とガイダンスが含まれている。	○この原則は、消費者を対象とする規制当局の領域であるため、イングランド銀行と PRA は、現時点ではこの原則を実施するための措置を講じていない。

（注１）PRA が発出する、目的別の文書の記号と番号を表している。SS は Supervisory statement（監督声明）の略語で、新旧の期待を織り込んだ柔軟な枠組みを企業に設定する。DP は Discussion paper（討議文書）の略語で、ルールの策定や期待事項の設定を検討している問題について、議論を喚起するために使用される。

（注２）Equality Act 2010 を指す。

（出典：FCA, “AI Update”（2024.4）および BOE&PRA, “DSIT-HMT letter”（2024.4）ほかをもとに作成）

（３）イギリス保険協会「AI ガイド」

イギリス保険協会（Association of British Insurers：以下「ABI」）が 2024 年 2 月に公表した「AI ガイド」²⁷は、保険会社に対して、イギリス政府の AI 白書に示された AI 規制の 5 原則²⁸を適用するための実践的なアプローチを提供し、AI の責任ある利用を支援する内容となっている。

項目として、「責任ある使用を支援するための質問」、「想定されるリスクと軽減措置」、「AI ソリューションに関する好事例集」、および「適用される法規制の概要」を取り上げ、それぞれについて AI 白書で挙げられた 5 原則を説明する体系的構成となっている。タイトルの副題に「責任ある AI を始めるための実践的アイデア」と付けられており、AI を保険事業に利用し始めるにあたっての実務的なヒントが多数盛り込まれている（全体の仮訳について、後掲参考図表 1-1-1 ～ 1-4-5 を参照）。

（４）AI Code of Conduct「保険金請求対応の AI 使用等に関する行動規範」

JEL Consulting²⁹の創設者兼取締役である Eddie Longworth 氏が中心となっている AI Code of Conduct は 2024 年 1 月、「保険金請求対応における AI の開発、実装、使

²⁷ ABI, “AI Guide Practical ideas for getting started with responsible AI”（2024.2）

²⁸ 前記 4.(1) a を参照願う。

²⁹ JEL Consulting は、保険金請求対応関連のコンサルタント業務を行っているイギリスの民間会社である。

用に関する行動規範」³⁰を公表した。これは、主に保険会社やロスアジャスターなどを対象として、保険金請求対応場面における AI の使用等について策定した自主行動規範である（図表 11 参照）。

この行動規範の策定の背景には、2023 年に米国の医療保険会社 Cigna が被保険者らから提訴された事例があった。具体的には、Cigna が被保険者からの保険金請求の対応にあたり、効率化のために導入した保険金請求対応のシステムがわずか 2 カ月間に 30 万件的の保険金請求をアルゴリズムで瞬時に却下した（平均して 1 件あたり 1.2 秒の審査時間）ため、保険金請求への対応にあたって州規則で規定されている「徹底的に、公正に、客観的に」審査すること³¹に違反したとして、集団訴訟に発展したものである³²。

このような裁判例に鑑み、イギリスで保険金請求対応を行う事業者が AI を使用するにあたり、保険金請求者から不審な目を向けられることなく、AI を適切に使用し、十分な配慮を行っていることを自主的な行動規範として定めていることを提示している。

技術革新のペースは速いため、規制・監督当局からのガイダンス等を待つことなく、政府の AI 規制の 5 原則に準拠する形により、保険金請求対応業務における AI 利用に当てはめて Eddie Longworth 氏が自主的に策定したという。

この行動規範を掲載しているウェブサイト³³では、行動規範への署名者を募るとともに、保険会社や大手のロスアジャスターなどが既に賛同を表明したことを紹介している³⁴。イギリスで保険に関する資格・教育・研修を行っている勅許保険協会（CII）のウェブサイトにおいても、「すべての事業者が署名することを推奨する。この規範は、透明性と保険金請求対応業務における AI の適切な使用に対するシンプルな自主的な取り組みであり、業界にとっては顧客の懸念を理解し、確実に対処するために率先して取り組んでいることを示す絶好の機会である。」として紹介されている³⁵。

図表 11 保険金請求対応における AI の開発、実装、使用に関する行動規範

項目	行動規範の内容
安全性・セキュリティ・堅牢性	○保険金請求対応の AI システムは、安全性、セキュリティ、堅牢性に揺るぎない重点を置いて開発、適用されなければならない。 ○保険金請求対応の管理と支払可否の決定に使用されるすべての情報は、検証可能で、監査可能で、サイバーセキュリティの厳格な基準に適合していなければならない。

³⁰ 原題は、“Code of Conduct for the Development, Implementation, and Use of Artificial Intelligence in Claims”である。

³¹ カリフォルニア州規則（California Code Of Regulations）Title. 10 § 2695.7 は、保険会社が保険金請求対応において徹底的で公正かつ客観的な調査を行うことを求めている。

³² Richard Nieva, “Cigna Sued Over Algorithm Allegedly Used To Deny Coverage To Hundreds Of Thousands Of Patients” (2023.7)

³³ <https://www.aicodeofconduct.co.uk/>

³⁴ 保険会社として Allianz、ロスアジャスター関係として Verisk、勅許ロス・アジャスター協会（Chartered Institute of Loss Adjusters）、Crawford、McLarens、carpenters group などの名称のロゴがある。

³⁵ <https://www.cii.co.uk/learning/learning-content-hub/articles/voluntary-code-of-conduct-for-the-use-of-ai-in-claims/108945>

項目	行動規範の内容
適切な透明性と説明可能性	○保険金請求対応の管理と支払可否の決定における AI の適用には、透明性があり、説明可能で、精査可能な意思決定ロジックがなければならない。 ○支払可否決定の結果は、すべての利害関係者、特に請求者本人から信頼できると認識されなければならない。
公平性	○AI の影響を受ける保険金請求対応システムは、回避可能な偏見を排除して公平に運用されなければならない。 ○AI の意思決定パターンと保険金請求対応に使用されるトレーニングデータを定期的に監査し、すべての人口統計と地域にわたって公平性を維持すべきである。 ○AI の公平性に対処するための技術基準が不可欠である。
説明責任とガバナンス	○コンプライアンス、説明責任、AI を活用した保険金請求対応の管理システムおよびプロセスの設計と実装を定義し、管理するためのガバナンス・メカニズムを確立すべきである。 ○保険金請求対応における AI の使用に関するすべての決定とその後の行動が記録され、監査可能で、倫理的に正当化すべきである。
異議申立と救済	○異議申立人は、AI による決定に異議を唱え、救済を求める手段を持たなければならない。 ○受け取ったすべての苦情について、意思決定プロセスや結果に熟練した人間の意見を取り入れながら、AI が下した結果に異議を唱えるための明確な手続きを設け、周知すべきである。

(出典：AI Code of Conduct ウェブサイトをもとに作成)

5. 米国

米国においては、連邦政府がこれまでに AI について大統領令の発布や種々の文書化を行ってきた。また、米国における保険監督者である全米保険監督官協会は、保険会社の AI 利用に関するモデル告示を独自に策定している。

(1) ホワイトハウス「AI の安全・安心・信頼できる開発と利用に関する大統領令」

米国のバイデン大統領は 2023 年 10 月、「AI の安全・安心・信頼できる開発と利用に関する大統領令」³⁶に署名して発効させた³⁷。これは、AI の能力向上が国民の安全や安心に及ぼす影響から国民を守るために、「安全性とセキュリティの新基準」、「米国民のプライバシーの保護」、「公平性と市民権の推進」など全部で 8 つの項目について、法的拘束力を含む対応を連邦政府機関や民間企業などに対して求める内容となっている。

(2) 全米保険監督官協会「保険会社による AI システムの利用」

全米保険監督官協会（National Association of Insurance Commissioners：以下

³⁶ 原題は、“Executive Order on the Safe, Secure, and Trustworthy Development and Use of Artificial Intelligence”である。

³⁷ AI 関連の他の大統領令として、2019 年当時のトランプ大統領が署名した「AI における米国のリーダーシップ維持のための大統領令」(the Executive Order on Maintaining American Leadership in Artificial Intelligence)がある。これは、AI の研究開発を促進させ、経済と安全保障の面で米国の優位性を保持させようとする目的のもとに策定された。

「NAIC」) は 2023 年 12 月、「保険会社による AI システムの利用」³⁸というモデル告示³⁹を採択して公表した⁴⁰。内容は、大きく分けると以下の 4 つのパートから構成される。

a. 第 1 パート：序論、背景、告示策定

この告示が採択された背景には、保険会社の AI 利用事例が増加し、それに伴って生じ得る重要な課題やリスクから、消費者を保護したり適切な規律を確保したりするための枠組みの必要性が認識されてきたことがあったとされている。AI システムの使用による消費者への潜在的な影響を最小限に抑えるためのガイダンスとして、保険会社に期待する行動を規制・監督当局として示すことで、適切な管理と透明性の確保を保険会社に求める内容となっている。

告示策定については、この告示中の第 3 パートと第 4 パートを規定するにあたり、NAIC がこれまでに策定してきた各種モデル法のうち、AI 利用に関して直接に参照すべきモデル法として、不公正取引慣行モデル法など⁴¹に立脚して今回の告示を策定していることを示している。

b. 第 2 パート：定義

第 2 パートは、「消費者に不当な結果」、「アルゴリズム」、「AI システム」、「AI」、「消費者に対する潜在的な損害の程度」など、告示で用いる用語を定義している。

ChatGPT 登場後に策定された告示のため、「生成 AI」という用語も取り上げており、「データ、テキスト、画像、音声、映像の形でコンテンツを生成する AI システムの一種で、既存のデータやコンテンツの直接的なコピーではないが、類似したものを生成するものを指す。」と定義している。

³⁸ 原題は、“Model Bulletin on the Use of Artificial Intelligence Systems by Insurers”である。

³⁹ NAIC におけるモデル法とモデル告示の相違について、モデル法には全米各州が州別規制を敷いている中、保険規制の法的枠組みの最低基準をすべての州で調和させる効果がある。これに対して、モデル告示は、法的拘束力を持たせておらず、各州の保険規制・監督当局が保険会社に対して提示して期待を示すガイダンス的位置付けとして、各州の保険規制当局間で統一を図ったものである。

⁴⁰ NAIC は、2020 年 8 月に「AI 原則 (Principles on Artificial Intelligence)」を公表している。この AI 原則は、保険会社による AI 使用の一般的なガイダンスを提供したものである。これに対して、今回のモデル告示は、保険会社が AI システムを実装する際の具体的なガイダンスとして、関係するこれまでの NAIC の各種モデル法に準拠させたものであり、保険規制当局による監督要素の強い内容となっている。また、NAIC の 2020 年 8 月の AI 原則にしても、今回のモデル告示にしても、米国連邦レベルで制定された「2020 年米国 AI イニシアチブ法」(National AI Initiative Act of 2020)を直接に参照している箇所はないため、紐づくものではない。なお、NAIC による 2020 年 8 月の「AI 原則」については、渡部美奈子「欧州・米国の保険業界における AI の活用事例と AI 原則等の動向」損保総研レポート第 140 号(損害保険事業総合研究所、2022.8)に詳しい。

⁴¹ 他に、不正請求支払慣行モデル法、コーポレート・ガバナンス年次開示モデル法、損害保険モデル料率法、および市場行為監視モデル法が挙げられている。

c. 第3パート：規制のガイダンスと期待

AIシステムを利用する保険会社が重視すべきと考えられるのがこの第3パートである（第3パート全体の仮訳について、後掲参考図表2-1を参照）。AIシステムの利用により行われる業務上の決定は、不公正取引慣行モデル法の対象となるため、不正確でないこと、恣意的でないこと、ムラのないこと、不当に差別的でないことを求めていることを説明している。そして、こうしたモデル法の規定は、保険会社がどのような手段や方法を用いて業務上の決定を下そうとも遵守されなければならない、適切なコントロールがない場合、AIは消費者にとって不正確、恣意的、ムラが生じやすく、不当な差別的結果をもたらすリスクを高める可能性があるとしている。

リスク対策として適切なコントロールの役目を果たすために保険会社に対して期待する事項がこの第3パートで述べられているAIシステム・プログラムの策定である。保険事業の認可を受けた保険会社が保険業務に関する意思決定や支援を行うAIシステムを利用する場合、このAIシステム・プログラムで述べられていることを策定、実施、維持することが期待されている。

d. 第4パート：規制監督と審査における考慮事項

第4パートは、各州の保険規制当局がAIシステムを利用する保険会社に対して監督上の調査を行う際に要求する可能性のある情報や文書についてのガイダンスを示しており、いわば規制・監督当局における監督上の着眼点のようなものである（第4パート全体の仮訳について、後掲参考図表2-2を参照）。

保険会社は、一般的に、ガバナンスの枠組み、リスク管理、内部統制のような事項について規制・監督当局から質問されることが想定される。また、保険会社によるAIシステムの開発、展開、使用、あるいは特定の予測モデル、AIシステム、アプリケーション、そしてそれらのAIシステムの使用による結果（消費者に悪影響を及ぼす結果を含む）に関して、保険規制当局が関連性のあると判断したその他の情報や文書についても、当局から質問されることが想定されるとしている。

(3) NY 金融サービス局「AIシステムと外部消費者データと情報源の利用」

ニューヨーク州金融サービス局（New York Department of Financial Services：以下「NYDFS」）は2024年1月、「保険引受と保険料設定におけるAIシステムと外部消費者データと情報源の利用」⁴²と題するサーキュラー（通達案）を公表して、市中協議に付した（全体の仮訳について、後掲参考図表3を参照）。

NYDFSがこのサーキュラーで述べていることは、保険会社がAIシステム・外部消費者データ・情報源を利用して保険引受と保険料設定をする場合、顧客となる消費者に

⁴² 原題は、“Use of Artificial Intelligence Systems and External Consumer Data and Information Sources in Insurance Underwriting and Pricing”である。

対して差別的な取扱いをしないことを保険会社に求めている。

サーキュラーの中で、イノベーションをもたらす新しいテクノロジーの用語として、次の 2 つを挙げて定義するとともに、その利用に伴う懸念を述べている。

○ AI システム

- ・定義：推論、学習、自己改善により、通常の人間の知能に関連する機能を実行するように設計された機械ベースのシステムであって、従来の医療保険、損害保険、傷害保険の引受や保険料設定を補完するために、または保険申込者の引受もしくは保険料設定の評価に寄与する可能性のある指標を設定するために、全体的もしくは部分的に使用されるものをいう。
- ・懸念：AI システムの自己学習行動は、不正確な、恣意的な、ムラのある、あるいは不当に差別的な結果をもたらす危険性を増大させ、不釣り合いな影響を与える可能性がある。

○ 外部消費者データ・情報源（External Consumer Data and Information Sources：以下「ECDIS」）

- ・定義：従来の医療保険、損害保険、傷害保険の引受または保険料設定を補完するために、または保険加入申込者の引受もしくは保険料設定の評価に資する「ライフスタイル指標」を設定するために、全体的もしくは部分的に使用されるデータもしくは情報をいう。
- ・懸念その 1：システムの偏見を反映する可能性があり、その使用は不公平等を助長し、悪化させる可能性がある。このことは、不当な悪影響や差別的な意思決定の可能性について重大な懸念を生じさせる。
- ・懸念その 2：精度や信頼性にばらつきがあり、規制監督や消費者保護の対象となっていない事業体から提供されている可能性がある。

保険会社がこれら AI システムや ECDIS のような新しい技術の利用を発展させて管理するにあたっては、利用に内在するリスクの全体的な複雑さと重要性を考慮して、消費者に及ぼす可能性のある損害を軽減し、関連するすべての法的義務を遵守するために、サーキュラーで説明している「公平性の原則」「ガバナンスとリスク管理」「透明性」に立脚して、合理的で適切なアプローチを取るべきであるとしている。

なお、NAIC が 2023 年 12 月に採択したモデル告示をそのまま、または微修正して自らの州のガイダンスとしている各州保険規制当局が多い中、NYDFS のサーキュラーの構成や内容は、NAIC のモデル告示と異なっている点で独自性が認められるが、AI を利用する保険会社に向けて期待を表明するガイダンス文書という点で、思想的には NAIC 告示と整合的なものと言える。

(4) 米国損害保険協会の声明

米国損害保険協会（American Property Casualty Insurance Association：以下「APCIA」）による AI 規制対応の状況については開示資料が少ない。コネチカット州が州法として策定しようとしている「AI に関する法律案」⁴³に対して、次のような声明を発表している資料⁴⁴があり、これらは APCIA における AI 法規制への考えを代表しているものと考えられることから、引いて紹介する。

- 保険会社と消費者の双方に有益な効用をもたらしている AI の利用について、これを妨げないことが重要である。
- コネチカット州当局による保険事業に対する法規制は、消費者への公正な取扱いを含めて、包括的な保険規制システムを提供しており、保険会社による AI の利用を含めて審査・執行の手段を完備している状態にある。
- 2023 年 12 月、NAIC は AI がもたらす特有の問題に対処するために、保険会社による AI の利用を監督するための包括的な枠組みとしてモデル告示を策定した。この中には、適用される法的基準の再表明、厳格なリスク管理とコーポレート・ガバナンス体制への期待などが含まれている。
- NAIC のモデル告示は、既に多くの州で、そのまま、または一部修正して採用されている。コネチカット州もそのモデル告示を採用し、州内保険会社が AI システムを使用する際の重要な規制の枠組みを整えた。
- AI を開発、導入する業種や企業を対象とした新しい AI 規制が必要で優先事項であることは十分に認識している。しかし、保険会社は AI の利用を含め、既に完全に州保険規制当局によって規制されている。
- 「AI に関する法律案」が成立して施行されれば、州保険規制当局による規制と重複し、相反する可能性もあることに懸念を抱いている。この事態は、保険会社のコンプライアンスを混乱、複雑化させ、保険会社や消費者にとって非生産的なコストを増やす可能性がある。

APCIA は結論として、コネチカット州保険規制・監督当局が既に NAIC モデル告示を採用しているため、「AI に関する法律案」が成立して施行されると、保険会社にとっては屋上屋を架された規制状況となることに鑑み、こうした状況への配慮を州議会の審議で強く求める意見表明としている。

⁴³ 原題は、“An Act Concerning Artificial Intelligence”である。包括的なリスクベースアプローチを敷いており、米国における州法レベルでの主要な法的枠組みになると見られている。2024 年 4 月現在、コネチカット州の議会で審議されているが、成立すれば EU の AI 法と同等の規模で民間部門における AI の開発と導入を管理する米国で初めての法律となると見込まれている。

⁴⁴ Statement American Property Casualty Insurance Association (APCIA) S.B. No. 2 – An Act Concerning Artificial Intelligence General Law Committee February 29, 2024

6. EU

EU 域内各国で重要となるのは、EU が定めた AI 法であり、保険事業に関連する部分を中心に説明する。また、AI 法の成立を受け、EU の保険規制・監督当局である欧州保険・企業年金監督機構の動向などについても説明する。

(1) EU「AI 法」

a. 概要

AI 法⁴⁵は 2024 年 5 月に EU 議会で承認され、成立した法規制であり、2026 年中に全面施行の予定である。EU 加盟国に直接適用となり、AI の開発や運用を包括的に規制する⁴⁶。AI が人々の健康、安全、基本的権利にもたらすリスクに対処するための安全性を確保する一方、AI への投資とイノベーションを促進するための法的確実性を確保させることも目的としている。そして、これらの目的を実施するため、次のような特徴を有している。

- リスクの程度に応じて規制の強度を 4 つに区分したリスクベースアプローチ
- 「指令」ではなく「規則」により加盟国に直接適用させる統一ルール⁴⁷
- 規制違反の場合の巨額制裁の賦課

上記のうち、最も特徴的な 4 区分のリスクについては、図表 12 に掲げる内容となっている。そもそも AI の利用を禁止するリスクを挙げるほか、高リスクと限定的リスクについては、一定の規制のもとに利用を認めるという建付けとしている。

なお、当初の草案は、2021 年 4 月に公表されたが、2022 年 11 月の ChatGPT のような生成 AI の登場を機に途中で見直しが行われ、生成 AI に対する規定の追加などが行われた。

図表 12 AI 法におけるリスク区分、対応・条件、対象事例

区分	対応・条件	対象となる AI の事例
許容できない リスク (禁止)	安全、セキュリティ、 基本的権利の観点で 容認できないリスク をもたらすため、使 用禁止	○潜在意識に働き掛ける操作的または欺瞞的な技術(サブリミナル効果狙い) ○ソーシャルスコアリング(個人の信用スコア) ○個人が罪を犯すリスクの評価 ○顔認証データベースの生成 ○職場や教育現場での感情の推測 など

⁴⁵ 正式名称は、「AI に関する調和の取れたルールを定める規則」(Regulation of the European Parliament and of the council laying down harmonised rules on artificial intelligence (artificial intelligence act) and amending certain Union legislative acts) である。

⁴⁶ AI 法は、EU の欧州委員会 (EC) が 2019 年 4 月に策定した「信頼できる AI のための倫理指針」を発展させた、より具体的な法的拘束力を有する法規制と言える。

⁴⁷ 「指令 (Directive)」は達成するための手段と方法を加盟国に委ねるのに対して、「規則 (regulation)」はすべての加盟国を拘束して一律に適用する。

区分	対応・条件	対象となる AI の事例
高リスク	市場投入前の適合性評価など、厳格な法の要件の遵守を条件に使用可	○自然人の生体識別・分類 ○教育・職業訓練 ○雇用、労働者管理 ○必要不可欠な公共および民間サービスへのアクセスと享受 ○法執行 など
限定的リスク	透明性確保など限定的な条件のもとに可	○人と直接対話する AI システム（チャットボット） ○ディープフェイクに関するシステム など
最小リスク	一定の自主的な行動規範の作成を推奨、強制される措置はない	○上記以外のすべての AI システム ・写真編集 ・商品の推薦 ・スパムフィルタリング ・スケジュール管理のソフト など

（出典：各種資料をもとに作成）

b. 基本的権利影響評価

保険事業の視点で留意しておくべき点は、AI 法における 4 区分のリスクのうち、高リスクに関する区分である。生命保険や医療保険のような保険商品に関するリスク評価や保険料設定で AI システム⁴⁸を導入して利用する場合、前掲図表 12 にある高リスクの事例として掲げた「必要不可欠な公共および民間サービスへのアクセスと享受」に該当し、市場投入前の適合性評価として基本的権利影響評価（Fundamental Rights Impact Assessment：以下「FRIA」）の実施が求められる⁴⁹ことになるからである。

FRIA の目的は、影響を受ける人々の基本的権利に対する具体的なリスクと、そのリスクが顕在化した場合に取りべき措置を特定するためにあり、次の事項が含まなければならないとされている。

- 高リスクの AI システムがその意図された目的に沿って使用される導入者のプロセスの説明、および各高リスクの AI システムが使用される予定の期間の説明
- 特定の文脈で、AI システムの使用によって影響を受ける可能性のある自然人および集団のカテゴリー
- 前項に従って特定された自然人または集団のカテゴリーに影響を及ぼす可能性のある具体的な危害のリスク
- 内部ガバナンスと苦情処理体系の取り決めを含む、リスクが顕在化した場合に講じられる措置

EU 所在の法律事務所である Hogan Lovells⁵⁰によると、FRIA への対応にあたって

⁴⁸ 付属書 III の 5 号 (c)に「生命保険や医療保険の場合における自然人に関するリスク評価や価格設定に使用することを目的とした AI システム」とある。

⁴⁹ AI 法第 27 条（高リスク AI システムの基本的権利への影響評価）

⁵⁰ Hogan Lovells, “Artificial intelligence in the insurance sector: fundamental right impact assessments”（2024.4）

は図表 13 に掲げるような事項が対応要領となる。

図表 13 保険会社における FRIA への対応要領

項目	内容説明
対象となる AI システム	○自然人の信用度を評価する、または信用スコアを確立する AI システム（金融詐欺の検出に使用されるシステムは除かれる） ○生命保険と医療保険で、自然人に関するリスク評価と保険料設定に使用される AI システム
実施時期	○初めて導入する前に FRIA を実施する。 ○AI システムの構成要素に変更や古さがあるときは、FRIA を更新する必要がある。 ○以前に実施された同様の FRIA や、システムのプロバイダーが実施した既存の影響評価に依存することもできる。 ○FRIA は、EU 一般データ保護規則（GDPR）に基づくデータ保護影響評価（DPIA） ^{（注1）} の一部となり、これを補完する可能性がある
当局通知	○導入者は、FRIA の結果を市場監視当局 ^{（注2）} に通知する。 ○加盟国 AI 当局が開発する自動ツールを通じてアンケートに回答する。
実際の対応上のヒント	○FRIA を DPIA と同時に実施することが合理的である。データ保護担当者やプライバシーチームも関与することにより、多くの相乗効果を活用できる可能性がある。 ○保険会社は既に DPIA の実施体制を整備しているため、FRIA を同じプロセスの一部に組み込むと、問題が少なくなり、必要な資源を少なくできる可能性がある。 ○FRIA は保険会社の AI ガバナンスプログラムと連携させる。個人に対する偏見や差別などのリスクへの対応は、AI ガバナンスプログラムに含まれていることが多いためである。

（注1）DPIA（Data Protection Impact Assessment）とは、データ取扱いに用いるコンピューター・システムで新たな技術や新たな取扱手法を導入する場合等において、そのデータ取扱いが自然人の権利と自由に対して高いリスクを生じさせる可能性がある場合、データ取扱いを行う前に、データ保護影響評価を実施しなければならないことをいう（GDPR 第 35 条 1 項）。

（注2）市場監視当局とは、市場で広く利用されている AI システムが AI 法に準拠していることを確認する責任を負う、加盟各国に設置される規制当局をいう。

（出典：Hogan Lovells, “Artificial intelligence in the insurance sector: fundamental right impact assessments”（2024.4）をもとに作成）

（2）欧州保険・企業年金監督機構の報告書等

EU の保険規制・監督当局である欧州保険・企業年金監督機構（European Insurance and Occupational Pensions Authority：以下「EIOPA」）は 2021 年 6 月に「AI ガバナンス原則: 欧州保険分野における倫理的で信頼できる AI に向けて」⁵¹を公表したが、その後の EU における AI 法案の審議や生成 AI の登場に対応して、以下のような対応を行っている。

⁵¹ 原題は、“Artificial intelligence governance principles: towards ethical and trustworthy artificial intelligence in the European insurance sector”である。詳しくは、渡部美奈子「欧州・米国の保険業界における AI の活用事例と AI 原則等の動向」損保総研レポート第 140 号（損害保険事業総合研究所、2022.8）を参照願う。

https://www.sonposoken.or.jp/reports/wp-content/uploads/2022/08/sonposokenreport140_1.pdf

a. 「EIOPA のデジタル戦略」

EIOPA が 2023 年 9 月に公表した「EIOPA のデジタル戦略」⁵²は、保険・年金セクターにおけるデジタル変革への対応にあたり、消費者、市場、EU 加盟各国監督当局への支援として、今後 3 年間の重点取組みを概説したものである。

この文書が作成された主な目的には、「デジタル化による機会と影響の認識」、「戦略的方針の明確化」、「関係者への情報提供」、および「消費者保護と金融安定性の促進」が含まれており、EU の保険・年金セクターにおけるデジタル化に関する EIOPA の将来的な戦略的アプローチ、方向性を提示している。

「デジタル化による機会と影響の認識」の中では、AI と保険との関連を取り上げ、データ分析の改善により、リスク管理の改善、保険料設定の精度向上、効率性の向上など様々な面で機会（メリット）の享受があることを挙げる一方、次のような影響（デメリット）に対する懸念を示している。

- 生成 AI により、ガバナンスと説明可能性に関する新たな問題が生じている。
- より細分化されたリスクは、保険の公平性の原則を損ない、一部の消費者やリスクを排除する可能性がある。
- データ革命は、データ保護の必要性を一段と喚起させるものであり、意図しない偏見や差別から生じる新たな倫理的問題や危険を引き起こしている。

そして、デジタル化による将来の変革の方向性やそれがもたらすメリットとデメリットを正しく理解して対応するため、EIOPA 自身と加盟各国の保険規制・監督当局自身が評価、管理できるだけの能力を強化する必要があると述べられている。

b. AI に関する考え

EIOPA は、急速に進展する AI 技術への対応に関する考え方をウェブサイトに公表している。ここでは、その掲載内容を図表 14 のとおり、「AI への現状認識」、「保険会社の AI への対応」、「AI 法の影響と対応」、および「政策対応・監督方針」の 4 つに区分して、その主なものを紹介する。

図表 14 AI に関する考え（抜粋）

項目	考え方
AI への現状認識	○（ChatGPT のような）大規模な言語モデルがもたらした一歩は、AI がこの分野や社会全体に対してできることの、間違いなくほんの始まりしか見えていないことを示している（注¹）。 ○AI は多くのメリットをもたらす可能性を秘めているが、AI 固有の特性によりリスクを悪化させる可能性がある。主なリスクの 1 つは偏見と差別に

⁵² 原題は、“EIOPA’s Digital Strategy – Support consumers, markets and the supervisory community through digital transformation.”である。

項目	考え方
	関連するものである。AI システムは、学習訓練の対象となるデータから偏見を引き継ぐ可能性があり、それが特定のグループに対して不利な、または差別的な結果をもたらす可能性がある ^(注1) 。
保険会社の AI への対応	○AI の導入による保険業務の自動化と高度化は、保険会社の従業員にとって、業務支援や課題軽減につながるが、利用のためには適切な訓練が必要である^(注1)。 ○保険会社での生成 AI の普及は、まだ初期段階にあるが、保険会社は消費者へのアドバイスの提供、保険契約者への保険金請求手続きの案内、保険料設定や引受プロセスの強化など、その潜在的な用途を積極的に模索している^(注2)。
AI 法の影響と対応	○AI 法のような新法は、既存の分野別の規制枠組みに新法の規定を統合させるという複雑なタスクを伴う。保険分野は既に高度に規制されており、AI 法の適用により多少の摩擦が生じる可能性がある^(注1)。 ○生命保険や医療保険における保険料設定とリスク評価は、高リスクの AI 利用事例とみなされるため、(AI 法で規定されている) 高度な要件に準拠する必要がある^(注3)。
政策対応・監督方針	○AI の普及は、規制当局や監督者に多くの疑問を生じさせる。規制の枠組みは、AI による技術の進歩に適応させる必要があるか、あるとして、どのように行われるべきか^(注1)。 ○複雑な AI システムによってもたらされる課題と、AI の責任ある使用を促進する必要性を認識しており、さらなるガイダンスを作成する準備をしている。ただし、重複することが多い新しいルールを作成するのではなく、ガバナンス、リスク管理、事業運営、保険商品の監督と管理に関する既存の要件に基づき、ガイダンスを作成すべきであると強く信じている^(注1)。 ○保険金請求管理、マネーロンダリング対策、不正行為検出などの事例で AI の使用が既にかなり広範に行われていることを踏まえ、監督者は特定の利用事例で既存のルールがどの程度十分であるか、また追加のガイダンスが必要な場合があるか、比例性、公平性、説明責任などの観点から評価する必要がある^(注3)。

(注1) 「AI の機会と課題に対処するための分野別アプローチ (A sectorial approach to address the opportunities and challenges of AI)」(2023.4)

(注2) 「保険業界における AI の導入と規制の動向 (AI in the insurance sector industry adoption and regulatory developments)」(2023.9)

(注3) 「AI 法と欧州金融セクターへの影響 (AI Act and its impacts on the European financial sector)」(2024.2)

(出典：EIOPA ウェブサイトをもとに作成)

c. 「欧州保険分野のデジタル化に関する報告書」

この報告書⁵³は、EIOPA 加盟各国の規制・監督当局を通じて保険会社を実施した、デジタル化をテーマとするモニタリング調査の結果⁵⁴であり、2024 年 4 月に公表された。調査目的は、AI、ブロックチェーン・暗号資産、IoT (モノのインターネット)、サイバー保険、レグテック⁵⁵などのデジタルイノベーションが EU 保険分野にもたらす変

⁵³ 原題は、“Report on the digitalization of the European insurance sector”である。

⁵⁴ 2023 年の 4 月から 6 月にかけて調査実施され、EU 加盟 22 カ国から 209 の保険会社が回答した。

⁵⁵ レグテック (RegTech) とは、規制 (Regulation) と技術 (Technology) を合わせた造語であり、2015 年頃からイギリスと米国を中心に使われ始めたと言われている。

化について、保険会社から実証的証拠を収集するために行われた。実施結果は、EIOPAの今後のデジタル戦略の指針に反映するとしている。

調査の結果は、図表 15 のとおりである。これによると、AI と保険の市場動向では、損害保険業界の方が生命保険業界よりも AI 利用が進んでおり、今後は生成 AI の自動車保険での利用増加が予想されている。AI ガバナンスの観点では、AI 導入や業務自動化に伴うオペレーショナル・リスクの増加が認識されているものの、ORSA 報告書⁵⁶への記載やガバナンス原則・リスク管理原則の採用までには至っていないとの回答が多い。また、AI を利用している保険会社に共通した回答として、特に影響の大きい AI を利用する場合に、その最終承認に経営陣と取締役会が深く関与していることを強調していることが挙げられる。これは、AI アルゴリズムが保険事業における意思決定に使用される場合に不正確な予測を生成したときでも、その説明責任は経営陣と取締役会にもあることを述べており、ガバナンスを担保する仕組みが機能しているものと言える。

報告書の結論では、規制当局として AI 関連に今後取り組むステップとして、EU の AI 法が保険分野にどのように影響するかを評価し、また保険会社や保険仲介業者による AI システムの利用をどのように監督すべきかについて、加盟各国規制・監督当局と議論を継続してベストプラクティスを目指すとしている。さらに、AI 法のもとで高リスクに分類されていない保険種目や AI の利用について、監督上の期待を明確に表明していく方針を掲げている⁵⁷。

図表 15 「欧州保険分野のデジタル化に関する報告書」の AI 関連の調査結果（抜粋）

OAI と保険の市場動向

- ・(EU 加盟 22 カ国の 209 の保険会社のうち) AI を既に利用している比率は、損害保険会社が 50%、生命保険会社が 24%である。
- ・生成 AI のような開発を考慮すると、AI の利用は今後数年間で増加し、特に損害保険、中でも自動車保険で顕著と予想される。
- ・テキストからデータを抽出して要約するための自然言語処理 (NLP) ^(注) の使用実績が認められている。
- ・保険バリューチェーンの中で、保険会社による AI 利用の事例が多いのは、保険料設定と保険金請求対応管理の分野であり、これに僅差で募集・販売、不正検知の分野が続いている。
- ・一部の保険会社がリスク評価や保険引受リスクの管理に AI を活用しているほか、損害引当金や保険金支払いの (半) 自動化など、よりニッチな業務に AI を活用していることも注目に値する。

OAI ガバナンス

- ・ある保険会社は、実施義務のあるデータ保護影響評価をする際に、AI システムの使用についても説明し、文書化している。
- ・保険会社は、AI 導入や業務自動化に伴うオペレーショナル・リスクの増加を認識している。一部の先進的な保険会社は、AI 関連リスクを ORSA 報告書で評価しているが、多くの保険会社はまだ具体的な対応を取っていない。
- ・国際および欧州レベルでのいくつかの AI 方針の取組み (文書化されたもの) は、公正性、説明可

⁵⁶ 保険会社がリスクとソルベンシーの自己評価を行って規制・監督当局に提出する報告書のことをいう。ORSA は、Own Risk and Solvency Assessment の略である。

⁵⁷ EIOPA は、2024 年 9 月 30 日、「保険と年金における AI の規制と監督」をテーマとして、公開討論を開催予定である。

能性、ヒトの監督、責任、記録保持、堅牢性、およびパフォーマンスなどの高レベルのガバナンスおよびリスク管理原則を共通して述べている。いくつかの保険会社は既に同様のガバナンスおよびリスク管理原則を実施していると回答したが、大多数の保険会社はまだこれらの原則を採用していないと回答した。

- ・リスク評価アンケートやチェックリストを用いて AI システムのリスクを評価し、それに見合ったガバナンスとリスク管理措置を実施し、包括的な文書化と記録管理を実施している保険会社もある。
- ・ほとんどの保険会社は、インパクトの大きい AI 利用の最終承認に経営陣と取締役会が深く関与していることを強調した。これは、例えば、AI アルゴリズムが意思決定に使用される不正確な予測を生成した場合の説明責任という点で、ガバナンスと関連していると思われる。

(注) 自然言語処理 (Natural Language Processing : NPL) とは、人間が日常的に使っている自然言語をコンピューターに処理させる一連の技術をいう。

(出典 : EIOPA ウェブサイトをもとに作成)

(3) 保険ヨーロッパ「キーマッセージ」

欧州の保険会社・再保険会社の連合体である保険ヨーロッパ (Insurance Europe) ⁵⁸ は 2023 年 9 月、EU 議会で審議されていた AI 法案の中に、保険事業に直接に関連する事項が含まれていたことから、「AI 法に関する現在進行中の三者協議へのキーマッセージ」⁵⁹を公表した。このメッセージの中で取り上げられていたのは、次の 3 項目である。

- ① AI システムに関する定義事項の修正提案
- ② 高リスク AI システムに生命保険と医療保険が含まれていることに関する意見
- ③ 高リスク AI システムがもたらすリスクレベルの評価で免除可能性のある規定が盛り込まれたことの歓迎表明

①の提案内容は、AI システムに関する定義について、多国間で採択された OECD の AI 原則による定義と整合的にしてグローバル対応できるようにしてはどうかというものであったが、2024 年 3 月に承認された AI 法の定義では OECD による包括的な定義内容よりもやや具体的になっている。

また、②の意見は、高リスク AI システムから生命保険と医療保険を除外することを要望した内容であったが、承認された AI 法では含めることが明記されている。

今後の焦点は、同じ②に関連して意見表明していた保険ベースの投資商品 (Insurance-Based Investment Products : 以下「IBIPs」) の取扱いである。IBIPs は、投資の要素と、死亡や病気などのリスクに対する保障の要素を組み合わせものである

⁵⁸ 保険ヨーロッパは、欧州にある 37 の加盟団体 (各国の保険協会) を通じて、あらゆる種類と規模の保険会社・再保険会社を代表している。ブリュッセルに本部を置き、欧州の保険料収入の約 95%を占める保険会社を代表している。

⁵⁹ 原題は、“Key messages in view of the ongoing trilogue discussions on the Artificial Intelligence (AI) Act”である。

が、保険ヨーロッパの要望はこの IBIPs が AI 法に定められる高リスク AI システムに分類されないことを明記すべきということだった。結局のところ、2024 年 3 月に承認された法条文では明記はされておらず、今後の協議に委ねられる見込みである。

7. ドイツ

ドイツは EU の加盟国であり、保険事業における AI 利用に関する規制・監督は AI 法の直接の適用となるため、ドイツの規制・監督当局も今後国内対応を図ることになる。本項目では、ドイツにおける規制・監督当局および保険協会の現状を報告する。

(1) ドイツ連邦金融監督庁の見解

ドイツ連邦金融監督庁（以下「BaFin」）は 2021 年 6 月にビッグデータと AI に関する監督原則⁶⁰を発出しているが、生成 AI まで含んだ公文書は 2024 年 5 月末時点で発出されていない。これは、前記の EIOPA が今後、生成 AI を含む AI に関するガイダンスを公表する方針を示していることから、そのガイダンス公表を待ってドイツ国内の保険会社への監督対応を図るものと思われる。

AI に関する現在の BaFin の見解は図表 16 のとおりである。

図表 16 BaFin の AI に関する見解（抜粋）

○AI 使用上の問題

保険会社にとって、AI の使用によりこれまでよりも正確な個人のリスクプロファイルを作成することができるになれば、個人情報保護に要する費用が大幅に増大し、または保険の母集団の構成対象から除外（引受拒否）したりする可能性が発生する。特に後者は大数の法則により運営される保険事業の精神に反し、極めて政治的、倫理的問題である。

○AI 使用者に求める対応

- ・ AI によるすべての決定は、透明性があり、理解しやすく、説明可能でなければならない。
- ・ 顧客への不当な差別があってはならない。データは適切に保護される必要がある。
- ・ 責任を持って新しいテクノロジーを使用し、慎重なガバナンスを確保することを期待する。
- ・ AI ツールがどれほど強力であるかを定期的に確認する必要がある。意思決定プロセスを管理し、介入し続ける必要がある。

○生成 AI 使用上の留意事項

- ・ 生成 AI の言語モデルは、ハルシネーション^(注)を起こすことが知られており、誤ったことを伝えている可能性がある。また、誤った情報を AI に提供することにより、その誤った情報を AI が学習して生成することもある。
- ・ 法的なリスクとして、生成 AI 型の言語モデルによって生成されたものが知的財産権を侵害する可能性や、生成されたデータを信頼し過ぎることによって誤った決定を下すことによるレピュテーション・リスクもある。

(注) 例えば、生成 AI が出力した文章が正確性や正当性を欠いているにもかかわらず、あたかももっともらしく出力することを指し、日本語では一般的に「幻覚」と訳されている。

(出典：BaFin ウェブサイトをもとに作成)

⁶⁰ 原題は、“Big Data und künstliche Intelligenz: Prinzipien für den Einsatz von Algorithmen in Entscheidungsprozessen”（ビッグデータと AI：意思決定プロセスにおけるアルゴリズムの使用に関する原則）である。

(2) ドイツ保険協会の AI 法への見解

ドイツ保険協会（以下「GDV」）は、EU による AI 法案の審議が行われていた 2023 年 6 月にポジションペーパー⁶¹を公表して、AI 法案に関する保険業界の立場からの見解表明を行った。

その内容は、前記の保険ヨーロッパが行った「AI システムに関する定義事項の修正提案」と「高リスク AI システムに生命保険と医療保険が含まれていることに関する意見」に加えて、AI 法案で項目として挙げられているソーシャルスコアリング、リスク管理システム、個人情報の使用、ガバナンス規制など全部で 13 の項目⁶²にわたった。

その後、2024 年 3 月に AI 法案が承認されると、GDV はコメントをウェブサイトに掲載した⁶³。「すべての人にとって拘束力のある AI 使用ルールの合意はよい。AI は、保険業界にとって次の大きな発展の推進力である。」として歓迎したうえで、「既存の規制によって保険業界における顧客保護のレベルは既に非常に高いにもかかわらず、基本的権利影響評価（FRIA）のような新たな義務が課される。EU にはこの点でもっと先見の明を持ってほしかった。」として批判を述べている。

GDV は今後、AI 法の完全実施が見込まれる 2026 年に向けて、AI 法に規定された権限の設計、規制要件の仕様、現行規制との関係付けなど、まだまだ明確にする必要があるとして、議論の場⁶⁴を設けて対応していく方針を示している。

8. わが国における AI 利用に関する行政動向

わが国において、金融庁が AI 関連で公表している文書は 2 つある。1 つが「AI モデルに関する利用規約」⁶⁵であり、もう 1 つが「モデル・リスク管理に関する原則」⁶⁶であ

⁶¹ 「AI 法案に関する三者協議の開始について（zum Beginn der Trilogverhandlungen zur KI-VO）」

⁶² 例えば、ガバナンス規制に関する意見表明の内容は、AI 法案第 17 条が高リスク AI システムの提供者に求める品質管理システムに関する規定が、ソルベンシー II 指令第 41 条が保険会社に求める効果的なガバナンス体制の構築と重複するため、このような不必要な負担や潜在的な矛盾は回避すべきである、というものであった。

⁶³ 「保険会社は EU による AI 法の承認を歓迎するが、高リスク分類を批判（Versicherer begrüßen Zustimmung des EP zum AI Act – kritisieren jedoch Hochrisikoeinstufung）」

⁶⁴ 2024 年 6 月 11 日、ベルリンで「欧州 AI 規制の国内導入に関するディスカッションラウンド」が開催される予定である。

⁶⁵ 金融機関による効果的・効率的な事業者支援の取組みを後押しするため、「AI や ICT 技術を活用した経営改善支援の効率化に向けた調査・研究」の実施により、AI 技術を活用して各種データを分析することを目的に作成された利用規約である。

⁶⁶ 金融機関がモデルを使用する際に生じるリスクを包括的に管理するためのアプローチを期待目線で示しており、モデル・リスク管理態勢の構築と強化を求めている。想定されるモデル・リスク管理は、「特定のモデルのカテゴリーに限定せず、広範なモデルを対象としている。例を挙げると、プライシング・モデル、市場リスク・信用リスク等のリスク計測モデルのほか、AML で使われるモデルや市場監視モデル等も含まれ得るが、これらに限定されるものではない。本文書は、モデルがリスクをもたらし得る限り、そのリスクを管理すべきという考え方に基づいている。」（II.適用＞（2）「モデル」の範囲）。生成 AI によるコンテンツ生成も、広義には生成 AI モデルが使用されているため、このモデル・リスク管理に関する

る。いずれも、保険会社が事業・業務に AI を利用する際の直接的なガイダンスや指針を示しているものではない。

このような状況の中、現在、わが国保険会社が AI 使用に関して依拠することができる公的文書は、2024 年 4 月に総務省と経済産業省から広く一般的な事業者向けに公表された「AI 事業者ガイドライン」⁶⁷である。このガイドラインは、AI の開発・提供・利用にあたって必要な取組みについての基本的な考え方を示している。実際の AI 開発・提供・利用でこのガイドラインを参考の 1 つ⁶⁸としながら、AI 活用に取り組むすべての事業者が自主的に具体的な取組を推進することが重要としている。

その後、わが国政府の AI 戦略会議⁶⁹は、「AI 事業者ガイドライン」という事業者の自主的な取組みに委ねるという従来の基本方針を軌道修正し、AI、特に生成 AI について国際的に規制強化の方向で議論が進むのに合わせた対応をしようとしている⁷⁰。このため、わが国でも EU の AI 法に準じるような立法措置や法規制化が行われる可能性はある。

保険事業に関しては、IAIS が 2024 年中に AI・機械学習モデルの分野に関するアプリケーション・ペーパーの作成を表明していることから、わが国の保険行政においても今後、AI に関する何らかの法令や指針が示される可能性がある。このことは、前記 3. (6) のジュネーブ協会の報告書が「保険という特殊性と、その確立された規制の枠組みは、分野横断的な規制よりもむしろ、分野固有の規制を必要としている。」と述べている⁷¹こととも符合する。

9. おわりに

保険事業においても、今後ますます AI 技術の利用は増加していくことが予想される。総務省と経済産業省から公表された AI 事業者ガイドラインは、AI 活用に取り組むすべ

原則の対象になり得るが、本原則は銀行と証券会社等への適用を念頭に置いて策定されている（「コメントの概要及びコメントに対する金融庁の考え方」）。

⁶⁷ AI を活用するすべての人が守るべき原則として「人間中心」、「安全性」、「公平性」、「プライバシー保護」、「セキュリティ確保」、「透明性」、「アカウントビリティ（説明責任）」、「教育・リテラシー」、「公正競争確保」、「イノベーション」の 10 原則を掲げている。

⁶⁸ AI 事業者ガイドライン案の策定中から AI、特に生成 AI がもたらす負の要素への対策として、拘束力を有する法規制の立法化が一部で要望されていた。自由民主党デジタル社会推進本部「AI の進化と実装に関するプロジェクトチーム」のワーキンググループ有志は、「政府は、極めて大きなリスクがある AI モデルに対し、必要最小限の法的枠組みを整備すること」を提言し、「責任ある AI 推進基本法（仮称）」の制定を提案している。

⁶⁹ 2023 年 5 月、広島で開催された G7 サミットで「広島 AI プロセス」が設置されたことを受け、わが国における AI 政策の司令塔として立ち上げられた。

⁷⁰ 読売新聞オンライン「生成 AI 法規制議論ようやく 海外動向受け 政府、なお慎重論も」（2024.5.23）

⁷¹ このほか、イギリスの保険業界紙にも、「保険会社が変化するビジネス環境に対応し、自らの立場を確立するには、AI テクノロジーに関する強固な規制が唯一の方法である。」という論調（Frances Stebbing, “UK develops secure AI guidelines”（Insurance POST, 2023.11））はある。

ての事業者を対象としている点で極めて汎用的なものであり、保険事業に AI を利用する場合の分野固有の考え方は示されていない。保険事業における AI 利用の当面の外部に向けた対応としては、例えば個人情報保護法制におけるプライバシーポリシーの公表と同じように、AI 事業者ガイドラインに準拠した AI 倫理ポリシーのようなものを作成して公表することなどが考えられる。

内部における対応としては、本稿が取り上げたジュネーブ協会の報告書、イギリス ABI の AI ガイド、米国の NAIC 告示モデル、および NYDFS のサーキュラーなどの内容も参考にしつつ⁷²、AI 利用に関するグループ内・社内向け文書を作成し、AI 利用に伴うリスクへの体制整備を図ることが AI ガバナンスの確立につながるものと考えられる。

⁷² または、金融データ活用推進協会（Financial Data Utilizing Association : FDUA）が策定している「生成 AI ガイドライン」（2024 年夏に公表予定）を参照することも選択肢である。なお、同協会には、国内の損害保険会社のほかに、生命保険会社、大手銀行、通信会社などが参画している。

参考図表 1-1-1 ABI「AI ガイド」安全性・セキュリティ・堅牢性を担保する質問

質問の観点	責任ある使用を支援するための質問例
信頼性	○システムのパフォーマンスが落ちるのはどんなときか？ ○AI システムが期待どおりに機能しなかった例を教えてください。 ○AI システムで特定された問題に対処するプロセスはどうなっているか？ ○パフォーマンスや精度の低下をどのように測定し、特定しているか？
障害からの復旧	○障害からの復旧と事業継続計画（BCP）について教えてください。 ○AI システムが安全で、サイバー脅威から保護されていることをどのように保証するか？ ○災害からの復旧計画や事業継続計画をどのくらいの頻度でテストしているか？ ○災害復旧および事業継続計画の有効性を測定するために、どのような重要業績評価指標（KPI）を使用しているか？ ○AI システムを提供（または一部提供）する可能性のある委託先等第三者との契約に、不測事態への対応方針について、どのように組み込まれているか？ ○障害を特定し是正するために、AI に関する専門知識と理解が組織全体にあるか？
モニタリング	○AI システムのパフォーマンスをどのように維持するか？ ○AI システムのパフォーマンスの測定に、どのような指標を使っているか？ ○AI システムが期待どおりに機能していることをどのように確認するか？ ○AI システムのパフォーマンスを継続的に改善するために、どのようなプロセスを導入しているか？

（出典：ABI, “AI Guide”（2024.2）をもとに作成、以下の参考図表 1-4-5 まで同じ）

参考図表 1-1-2 ABI「AI ガイド」適切な透明性と説明可能性を担保する質問

質問の観点	責任ある使用を支援するための質問例
説明可能性	（AI システムが意思決定を支援する場合、分かりやすい結果を出力するように設計されていることを前提として） ○どのように決定されるか？ ○出力結果の通じやすさ、理解しやすさを判断する基準は何か？ ○特に業務委託先等第三者の AI システムを使用している場合、誰がその説明を行うのか？ ○説明のしやすさとは何か？ ○説明が明確で誤解を招かず、最終利用者に理解されているか？
利用者への開示	（データ、機械学習、AI の活用に関する関連情報の提供） ○利用者はどのように考えているか？ ○AI システムの提供する説明が、正確で、適切で、意図する利用者にとって理解しやすいものであることをどのように保証するか？ ○利用者が交流するかどうかを決めるために、どのようなタイミングで透明性を示すべきか？
目的	（事業目的達成のための AI 利用とその必要性に関する明確な理解のあることを前提として） ○この AI アプリケーションで使用するデータは、この目的のために収集されたものかどうか？ ○目的を適切に保つために、どのように利用を監視（制限）すればよいか？

参考図表 1-1-3 ABI「AI ガイド」公平性を担保するための質問

質問の観点	責任ある使用を支援するための質問例
サービスの公平性	(AI システムが異なる集団に比較可能なサービス品質を提供することを前提として) ○サービス品質は集団によってどのように異なるのか？ ○異なる集団間で比較可能なサービス品質を提供していることを確認するため、AI システムのパフォーマンスをどのように監視・評価しているか？ ○AI システムが、様々な層に対して比較可能なサービス品質を提供している例を教えてほしい。 ○消費者は、自動的な決定に異議を申し立てる権利について、いつ、どの時点で、どのように通知されるか？ ○同意がない場合、AI による自動的な決定に関して公正かつ平等な結果を提供する上でどのような影響があるか？例えば、利用者が同意しないことを選択した場合、価格決定から除外されるのか？
偏見の最小化	(不当な偏見を特定し、理解し、最小限に抑えるために最善を尽くすことを前提として) ○不公平な偏見を最小限にするには？ ○AI システムに不当な偏見や差別がないことを保証するための措置は？ ○AI システムにおける不公平な偏見をどのように特定しているのか、例を挙げてほしい。
一貫性	(可能な限り、利用しているモデルが一貫した再現性のある結果を出すようにしていることを前提として) ○結果や決断はどのように行われ、そしてなぜ異なるのか？ ○AI モデルの下した判断が、期待された結果と矛盾していた例を挙げてほしい。 ○AI モデルのトレーニングに使用されるデータが、その AI モデルが意図する母集団を代表するものであることを、どのようにして確認しているか？ ○AI システムの結果の矛盾を検出し、緩和するために、どのような対策を講じているか？

参考図表 1-1-4 ABI「AI ガイド」説明責任とガバナンスを担保するための質問

質問の観点	責任ある使用を支援するための質問例
影響	(AI システムが人々や組織、環境にどのような影響を与えるかを留意することを前提として) ○この AI システムはどのような影響をもたらすか？ ○AI システムの効果を確実にするために、具体的にどのような対策を講じているか？ ○AI システムが人、組織、環境に与える影響をどのように測定しているか？ ○AI システムが人、組織、環境に及ぼす悪影響に対処するために、どのような手段を講じているか？ ○規制体制が異なる場合、国境を越えた利用はどのような影響を及ぼすか？
副作用	(AI システムが人、組織、環境に悪影響を及ぼす場所の特定) ○何が問題となっているか？ ○あなたの組織が過去に確認した悪影響の例を挙げてほしい。 ○AI システムの悪影響を確実に軽減、あるいは排除するにはどうすべきか？ ○原則が適切かつ効果的であり続けるために、原則を見直し、更新するプロセスはどのようなものか？ ○AI を導入する企業か、またはソリューションを提供する第三者の AI プロバイダーであるか、どちらか？
監督レベル	(可能な限り、適切なレベルの人的監視を可能にすることを前提として) ○どうすれば決定に異議を申し立てることができるのか？ ○AI システムの限界を理解しているか？ ○人的監督のレベルが適切であることを保証するために、具体的にどのような対策を講じているか？

質問の観点	責任ある使用を支援するための質問例
	○人間の監督レベルが適切であることをどのように保証しているか？ ○AI システムの悪影響を特定し、緩和するために、人間の監視が効果的であることをどのように保証するか？ ○AI システムの決定に異議を唱えるためのプロセスはどのようなもので、どのようにして透明性と公平性を確保するか？ ○モデルの最終的な責任は誰が負うのか？
法的根拠	(常に、関連する法律および規制の枠組みのガイドラインと制限の範囲内で行動すること、また、自動的な決定に関して、明確で文書化され、証拠となる同意の立場を取っていることを前提として) ○法的根拠をどのように明確に示したか？ ○AI の利用が合法的であることを保証するために、具体的にどのような法律や規制の枠組みに従っているか？ ○AI の利用が関連する法律や規制の枠組みに準拠していることをどのように確認しているか？ ○AI 利用の透明性と説明責任を確保するために、どのような対策を講じているか？
リスク対策の実行主体（当事者）	(AI システムを利用するときに責任を負うのは誰であることを知っていることを前提として) ○このシステムの RACI（レイシー）^(注) は、どのようにになっているか？ ○リスク対策の実行主体（当事者）はどのように組織全体に分散されているか？ ○AI の利用を導く原則に従わなかった場合、どのような結果になるのか？

(注) RACI（レイシー）とは、プロジェクトにおいて、役割や責任を明確化するために用いられる概念である。タスクの実行に関連する役割である「Responsible（実行責任者）」「Accountable（説明責任者）」「Consulted（相談先）」「Informed（報告先）」の頭文字を並べたものである。

参考図表 1-1-5 ABI「AI ガイド」異議申立と救済を担保するための質問

質問の観点	責任ある使用を支援するための質問例
明確でアクセスしやすい窓口	(利用者や影響を受ける当事者が、不利な結果に対する救済を求めるための、単一の明確な窓口や仕組みがあることを前提として) ○不利な結果や影響に対して、あるいは不正確なデータや有害なデータを修正・削除するために、不服申し立てや苦情、救済を求めるにはどうすればよいか。 ○有害な AI の決定や結果を迅速かつ効果的に解決するためのプロセスは何か？ ○潜在的に有害な AI の結果や決定によって影響を受ける可能性のある人々が、不利な結果に対する認識を高め、異議を唱え、救済を求めるにはどうすればよいか？ ○正しい結果や決定が、不公平または不正確であると利用者が認識した場合、その不一致にどのように折り合いを付けるか？（コンピューターが正しく「ノー」と判定した場合など） ○誤用、誤作動、欠陥のあるモデル、機械学習による偏見、または機械学習による偏見に起因する有害なアウトプットの責任は誰が負うのか？不正確なデータ、特にデータが他から入手されたものである場合はどうなるのか？

参考図表 1-2-1 ABI「AI ガイド」リスクと軽減措置：安全性・セキュリティ・堅牢性

想定リスク	軽減措置
○不十分なデータ品質 AI システムが堅牢でなかったり、不正確な出力を生成したりするリスク	○人間による介入 ○データガバナンスの方針と手順、ライフサイクルを通じたデータ品質の維持 ○アルゴリズムの選択と検証 ○異常検知 ○モニターとフィードバック

想定リスク	軽減措置
	○データ品質指標
○内部と外部の不正 AI システムの脆弱性が悪意を持って悪用されるリスク	○倫理的な AI 開発 ○規制と政策に合わせる ○AI の開発者と利用者の教育と意識向上
○モデルドリフト ^(注) ・時間が経つにつれて、モデルの性能が低下するリスク ・運用データの分布や環境要因の変化による正確性の低下	○連続モニタリング ○データ品質保証 ○更新されたデータセットによるモデルの定期的再トレーニング ○モデルの性能を評価するための自動テスト ○フィードバック・ループ：末端利用者からの反応の収集
○データセキュリティ 不正アクセス、データ漏洩、データ操作により AI システムの完全性が損なわれるリスク	○データ・セキュリティ・ポリシー ○訓練データの整合性チェック ○データアクセス監査証跡 ○安全なコーディングの実践 ○安全な通信チャネル ○定期的なセキュリティ・アップデート ○訓練と意識向上 ○インシデント対応計画
○サイバーセキュリティの脅威	○AI サイバー攻撃のリスクと脅威について、利用者へのセキュリティ意識向上訓練 ○堅牢な認証・認可メカニズム ○定期的なセキュリティ監査と侵入テスト ○侵入検知・防止システムの導入 ○脅威インテリジェンスに関する協力

(注) モデルドリフトとは、モデルの訓練に使用されたデータと、モデルが使用されるデータとの間の定義、分布、統計的特性などの根本的な変化から生じる、時間の経過に伴うモデルの性能の減衰を指す。

参考図表 1-2-2 ABI「AI ガイド」リスクと軽減措置：適切な透明性と説明可能性

想定リスク	軽減措置
○解釈可能性と説明可能性の欠如 AI のモデルやシステムは複雑で、決定がどのように行われるかを示すには理解が難しいというリスク	○モデルを単純化するテクニックを使う。 ○解釈可能なモデル、決定木分析、線形回帰を代理として使用する。 ○モデルの検証とテストを行う。 ○モデル・アーキテクチャと決定プロセスに関する文書化を維持する。 ○人間による監督と介入を行う。 ○連続モニタリングを行う。 ○モデル使用に関する利用者訓練を行う。

参考図表 1-2-3 ABI「AI ガイド」リスクと軽減措置：公平性

想定リスク	軽減措置
○AI 偏見 AI システムが学習データの偏見を増幅し、特定のコミュニティやグループに対して不当な差別を行うような形で適用されるリスク	○多様な代表データ ○偏見の検出と評価 ○データ倫理フレームワーク ○多様な開発チーム ○監査とモニタリング ○明確な説明責任
○プライバシー侵害 AI システムによる個人データの収集、悪用する可能性	○データの最小化 ○匿名化と仮名化 ○データガバナンス方針 ○説明責任と監督

想定リスク	軽減措置
	○データ倫理フレームワーク ○定期監査 ○継続的なトレーニング

参考図表 1-2-4 ABI「AI ガイド」リスクと軽減措置：説明責任とガバナンス

想定リスク	軽減措置
○第三者およびサプライヤーの管理 第三者が AI を開発し、それが当社の認識なしに当社の代理として行動し、規制当局に暴露されるリスク	○信頼できるベンダー選定 ○契約書 ○リスクと管理監督 ○明確なガイドライン要件 ○第三者に対する監督と監査 ○法務およびコンプライアンスの監督
○第三者プロバイダーへの依存 第三者が AI 技術やサービスに依存することで、サービスの品質、信頼性、継続性に影響を及ぼすリスク	○信頼できるベンダーの選定 ○契約書 ○リスクと管理監督 ○第三者に対する監督と監査 ○サード・パーティ・サプライヤーの枠組み ○顧客サービス基準
○スキル・ギャップ 従業員のトレーニングや専門知識が不十分なため、AI のアウトプットの誤用や誤った解釈につながるリスク	○継続的な学習と開発 ○多様な学際的チーム ○監査とレビュー ○責任ある AI に関するガイドラインと文書

参考図表 1-3-1 ABI「AI ガイド」好事例集：安全性・セキュリティ・堅牢性

好事例	説明
1.1 AI がセキュリティに与える影響の考慮	○AI ライフサイクルの様々な段階で適用される可能性のあるセキュリティ上の脅威を特定し、緩和するために、データシステムは厳格なサイバーセキュリティプロセスの対象とすべきである。 ○機械学習や AI の各使用に関する RACI には、サイバーセキュリティ管理および関係者の詳細を含めるべきである。 ○機械学習モデルを保護するための NCSC (National Cyber Security Centre) の原則や、安全な AI システム開発のためのガイドラインなど、検討可能な情報は既にある。
1.2 データ保護とプライバシー	○AI システムの設計、訓練、使用に関わる個人データの処理は、イギリス一般データ保護規則 (GDPR) およびデータ保護法 2018 に基づく要件、特にもっぱら自動化された意思決定に関する要件に準拠すべきである。 ○リスクの高い処理には、AI の使用がもたらすデータ保護法非準拠のリスクを最小化するためのデータ保護影響評価が必要となる。
1.3 プライバシーに重点を置くべきデータ環境	○データ環境は分析者に最小限のデータを公開すべきである。 ○ID、ハッシュ、その他の技術を駆使して、データへのアクセスを制限し、適切なアクセス制御を行う。
1.4 プライバシーを強化する技術の検討	○データサイエンス環境に含まれるデータの粒度を検討する。 ○例えば、郵便番号の中には世帯が 1 つしかないものがあり、この郵便番号で個人を特定することができる。 ○また、より狭い地域であれば、郵便番号 1 つと年齢で個人を特定することができる。 ○このような場合は、直接個人を特定できるようなデータの使用を避けるために、集計技法を検討する。
1.5 スタッフ教育	○AI は私たちの生活を向上させる驚くべき力を持っているように見える。 ○しかし、その便利さゆえに、人は AI に頼りすぎたり、その結果を受け入れ

好事例	説明
	過ぎたり、物事がうまくいかない状況を見逃してしまうこともある。 ○そのため、スタッフは AI の限界と人間の判断を維持する必要性を認識するための定期的な訓練が必要である。 ○企業は、物事が基準外であった場合にそれを検知するプロセスを構築する必要がある。AI を使用するスタッフは、その結果について責任を持つべきである。

参考図表 1-3-2 ABI「AI ガイド」好事例集：適切な透明性と説明可能性

好事例	説明
2.1 十分な技術の使用の推奨	○社内モデルを構築する際には、手法や変数の面で最も単純なモデルの使用が推奨される。 ○精度と説明可能性は、逆相関の関係にあることが多いため、この 2 つの考慮事項のバランスが最適になるように、様々な手法やモデルを試す必要がある。 ○適切なバランスの保持を確認するために、独立した技術的検証者（モデルを構築した個人やチームではない）を活用すべきである。 ○第三者サプライヤーがそのモデルの全詳細を開示する可能性は低い、そのアプローチに関する一般的な情報を要求することが推奨される。
2.2 ブラックボックス的な機械学習の可視化	○視覚化は、予測モデルの結果を完全に理解するのに役立ち、公正な決定が行われることを保証する。 ○ブラックボックス技術（例：ニューラルネットワーク、ツリー・アンサンブル・モデルなど）を使用した機械学習モデルは、解釈手法によって完全に視覚化されるべきである。 ○検証者やガバナンス部門は、ブラックボックス技術を使用した機械学習モデルについて、一般的な文言での説明をできるようにすべきである。 ○ブラックボックス技術を使用した機械学習モデルの内部理解を深めることで、ガバナンス機能が適切なレベルの課題を提供できるようになる。 ○外部委託先等第三者のモデルは、適切な理解レベルを確立できるようテストされるべきである。
2.3 顧客の現実的な期待の範囲内での取り組み	○データは、顧客の信頼を損なうような方法で使用されるべきではない。 ○そのため、顧客の期待や、顧客が予測・予想しないような影響を受けないかどうかを考慮する必要がある。 ○データの利用がどのように顧客の期待に応えるかを、モデルの開発や AI の利用事例に盛り込むことを推奨する。
2.4 機微情報の慎重かつ控えめな使用	○データを使用する正当な理由を文書化し、それが合法的かつ比例的であることを確認するための評価を行うべきである。
2.5 適切な説明	○顧客データを取得し使用する際には、顧客に対して適切なレベルの透明性を確保すべきである。 ○個人情報担当部門は、各利用事例に必要な透明性のレベルについて助言すべきである。

参考図表 1-3-3 ABI「AI ガイド」好事例集：公平性

好事例	説明
3.1 顧客利益の確保	○顧客は、そのデータが使用される際に、目に見える見返りを受け取るべきである。 ○AI の利点は文書化されるべきであり、リスク管理も必要である。 ○AI が生み出しうる潜在的なリスクと機会のバランスを取るために、評価が実施される。
3.2 学習訓練用データの多様性の確保	○データの偏りを考慮する。データにおける偏見は、代表性の欠如（サービスへのアクセスや利用の欠如など）、設計の不備、人的偏見、歴史的または社会的偏見が繰り返され、強化されるなど、様々な理由によって引き起こ

好事例	説明
	される可能性がある。 ○その結果、有害な結果を招きかねない。偏見を検出するプロセスは容易ではないため、モデル構築に直接関与していない専門家が、どのような偏見が存在する可能性があるかの検討・検証を導入することを推奨する。 ○偏見が検出された場合は、それを軽減するために何ができるかを検討するための手段を講じる。 ○利害関係者は、潜在的な偏見と、偏見を制限するために講じられたあらゆる措置について認識されるべきである。
3.3 高水準のデータ品質の確保	○モデルは、正確で公正な結果を保証するために、信頼できるクリーンなデータに基づいて構築される必要がある。 ○モデルで使用されるデータは、正確性、適時性、完全性、一貫性、妥当性を厳密にテストされるべきである。 ○データの質に関する問題はすべて報告され、文書化され、管理されるべきである。 ○質の低いデータが取得され、採点に使用されることを防止するためのプロセスを導入すべきである。 ○第三者のモデルを使用する場合は、そのモデルがクリーンで信頼できるデータセットの上に構築されていることを保証するために、適切な管理が行われているかどうかを検討する。
3.4 保護されるべき特性（LGBTQ やマイノリティなど）の慎重な取り扱い	○限られた状況において、モデルは保護された特性（例えば、医療リスクモデルにおける年齢）を使用することができるが、保護された特性およびそれに近い代理の使用が各文脈において適切かどうかを検討する。 ○保護特性またはその代理の使用の正当化については、法的根拠および比例性の検討を含め、文書化する。 ○第三者モデルを使用する場合、保護特性またはその代理の使用を特定し、正当化するために適切な管理が行われているかどうかを検討する。
3.5 誰が影響を受けるかの理解	○意思決定に情報を提供するため、また誰が影響を受けるかを理解するために、機械学習モデルを視覚化する必要がある。 ○可視化することで、誰が最も高いスコアを持っているか、誰が最も多く支払わなければならないか、どの顧客グループが不正確なスコアを持っているか、などの情報を提供する必要がある。 ○異なる顧客グループのスコアとエラー率は、どのグループも不利益を被らないように検証されるべきである。
3.6 テスト計画の策定	○AI や機械学習モデルを構築または導入する前に、テストプロセスを検討し、長所と短所を文書化する必要がある。 ○テストの結果は、AI の限界とどこでミスをおかすかに焦点を当てるべきであり、テストは社内と社外の両方のモデルやツールに対して行うべきである。 ○生成 AI には、偏見（性別、人種など）や幻覚（不正確さ）など、いくつかの既知の限界がある。このため、生成 AI を使用する場合は、人間の監視と定期的なテストの両方を行うことが重要である。
3.7 単なる相関関係ではなく、因果関係の重視	○モデル変数とモデルの結果の間に合理的な因果関係があるかどうかを検討する。 ○使用するモデルの関係が論理的で因果関係があることを確認すべきである。 ○機械学習モデルの各変数の方向性と因果関係を記述した書面を作成し、モデルを展開する前に独立したレビューを行うべきである。
3.8 適切な精度レベルの検討	○要求される精度のレベルは、AI システムの構築前または購買段階で確立されるべきであり、この精度レベルを満たすモデルのみが使用されるべきである。 ○そうすることで、不公平な結果を招きかねない、性能の低いモデル（オーバーフィッティングやアンダーフィッティング）を使用するリスクを軽減することができる。 ○したがって、機械学習モデルは、ライフサイクルを通じて、高いパフォーマンスを確保し、ドリフトを検出するためにモニタリングされるべきである。

好事例	説明
	る。 ○第三者のモデルや生成 AI については、ランダム化テストを行い、高いパフォーマンスを維持できるようにすることを推奨する。
3.9 人間による検証と独立したテスト	○公平性は人間的なものであり、純粋に数学的な問題ではない。 ○人間による適切な監視とモニタリングの仕組みが導入されるべきであり、すべてのモデルが技術専門家による独立した挑戦の対象となることを推奨する。 ○技術専門家は、顧客の独立した声として機能すべきである。

参考図表 1-3-4 ABI「AI ガイド」好事例集：説明責任とガバナンス

好事例	説明
4.1 AI の使用が必要かどうかの検討	○すべての問題がモデルやアルゴリズムで解決できるわけではないため、最も公正なアプローチは予測分析や AI の使用を完全に回避することである。 ○プロジェクトを始める前に、その文脈で AI を使用する利点と限界を検討する。 ○AI の使用が特定の目的を達成するための最も適切かつ公平な方法であることを示すために、その正当性を文書化する。
4.2 AI カタログ／インベントリ	○機械学習の使用はすべて、インベントリ／商品管理システムの中で完全に文書化されるべきである。 ○文書化には、別のデータサイエンティストや開発者がゼロからモデルを再現できるような十分な詳細が必要である。 ○個別の場所を持つことで、組織はこの文書に簡単かつ効率的にアクセスでき、より効果的な監査とガバナンスが容易になる。
4.3 データ・カタログ	○機械学習モデルや生成 AI で使用されるすべてのデータは、データ・カタログに文書化されるべきである。
4.4 文書化された役割と責任	○AI を構築または使用する場合、AI のライフサイクル全体を通じて AI リスクの管理に対する説明責任と効果的な監督を確保するために、文書化された役割と責任を持つことが重要である。 ○生成 AI と第三者による AI については、調達、内部実装、テスト、活用に関わるすべての個人（データサイエンティスト、データエンジニア、データオーナー、マネジメント、テスター、ステークホルダー、ガバナンスなど）をリストに含める必要がある。
4.5 トレーニングと意識向上	○スタッフと利害関係者は、意図しない結果を招く可能性を減らし、責任ある AI の開発とガバナンスを確保するために、AI に関連する倫理的配慮とリスクについて教育を受けるべきである。 ○モデルを開発するスタッフと、モデルの検証とその承認を担当するスタッフには、専用のトレーニングを実施すべきである。
4.6 意思決定者に対する技術的支援	○独立した技術専門家は、要請に応じて、関連する作業部会や委員会に対して意見具申や勧告を行うことができるようにすべきである。
4.7 エスカレーションとコミュニケーション・チャンネル	○機械学習と AI のプロジェクトでは、認識、説明責任、管理を確実にするために、上級管理職への明確なエスカレーションとコミュニケーション・チャンネルが必要である。
4.8 信頼できる第三者サプライヤーとのみ取引を行う	○第三者からの顧客情報を使用する場合、その情報がどのようなプロセスを経ているのかについて、不完全な知識しか持たない、またはデータが自らに課している倫理基準で収集されていない可能性がある。 ○したがって、その慣行について透明性があり、公正な顧客成果を確保するための積極的な措置を講じているサプライヤーと協力することが重要である。

参考図表 1-3-5 ABI「AI ガイド」好事例集：異議申立と救済

好事例	説明
5.1 問合せ対応チームの強化	○AI のアウトプットを顧客や規制当局に説明するために必要なツールや訓練を提供する。AI ソリューションを説明し視覚化する技術（SHAP、LIME など）が開発されるべきである。 ○問合せ対応チームには、予測モデルの操作に関する完全な訓練が行われるべきである。
5.2 異議申立の受付手続き	○顧客に悪影響を及ぼす AI の決定や結果に異議を唱えるためのプロセスを、顧客が利用できるようにすべきである。
5.3 データ修正プロセス	○AI による意思決定で使用された不正確な情報にフラグを立てることができるプロセスを顧客が利用できるようにすべきである。 ○不正確な情報は、正確な結果を保証するために修正されるべきである。

参考図表 1-4-1 ABI「AI ガイド」適用される法規制：安全性・セキュリティ・堅牢性

法規制	説明
一般データ保護規則（UK GDPR）の関連規定	○適切なセキュリティレベルを保証する適切な技術的・組織的な対策を実施する義務（第 32 条） ○データ侵害通知義務がある場合、当局への通知義務（第 33 条） ○特に新技術を使用した処理が自然人の権利と自由に対して高いリスクをもたらす可能性がある場合、データ保護影響評価を実施する義務（第 35 条）
情報コミッショナーオフィス（ICO）の対応	○データ保護責任者向けに、AI のセキュリティとデータ最小化に関するガイダンス ^{（注1）} を公表している。これには、AI モデルについて、個人情報の侵害というリスクを管理するためにどのような措置を講じるべきか、AI システムのデータ最小化と個人情報保護技術に関するガイダンスが含まれている。 ○組織が AI ソリューションを開発する際に完全なリスクを検討し、リスクを軽減するためにどのような実用的な手段を講じるかを支援するために、有用な AI とデータ保護のリスクツールキット ^{（注2）} を提供している。
FCA ハンドブックの「消費者への義務」	○組織は強固なガバナンス・モデルを導入することが求められる。「消費者への義務」において、AI の安全性、セキュリティ、堅牢性については言及されていないが、このガバナンス・モデルの性質上、安全で堅牢な AI 性能を維持する必要がある。
その他	○政府は、AI を活用したシステムの設計、開発、導入、調達に携わるすべての人のために、AI 保証技術一式 ^{（注3）} を公表した。

（注 1）このガイダンスの英文正式名称は、“Guidance on AI and data protection”である。

（注 2）このガイダンスの英文正式名称は、“AI and data protection risk toolkit”である。

（注 3）公表された AI 保証技術一式の中で、保険事業に関連するものとして、“Credo AI Governance Platform”がある。再保険会社がアルゴリズム偏見の評価と報告を行うことを主目的として作成されている。

参考図表 1-4-2 ABI「AI ガイド」適用される法規制：適切な透明性と説明可能性

法規制	説明
FCA の「消費者への義務」に関するガイダンス ^{（注1）}	○次の事例は、誠実に行動すべき「消費者への義務」に反する。 ・機械学習や AI を含むアルゴリズムを、消費者被害につながる可能性のある方法で製品やサービス内で使用すること。 ・これは、アルゴリズムが偏見を埋め込んだり増幅したりし、一部の顧客グループにとって組織的に悪い結果をもたらす。
情報コミッショナーオフィス（ICO）のガイ	○ICO は、AI による意思決定および AI とデータ保護に関する手引きについて説明するための詳細なガイダンスを作成した。

法規制	説明
ダンス ^(注2)	○このガイダンスの説明は、「プロセスベースと成果ベース」、「合理性」、「責任」、「データ」、「公平性」、「安全性とパフォーマンス」、「影響」に分類されている。 ○イギリスの GDPR とデータ保護法に言及し、AI モデルが個人データを使用して決定する場合には、そのことを説明することの重要性を指摘している。
平等人権委員会のガイダンス ^(注3)	○平等人権委員会は、金融サービス事業者向けに、違法な差別に関するガイダンスとして「平等法—銀行およびその他の金融サービス事業者」を公表した。 ○意思決定プロセスに AI システムを使用する場合、それが違法な差別につながることを確認し、それを示すことができるようにする必要がある。

(注1) FCA のハンドブックでは規定しなかったが、ガイダンスレベルで知らせることを述べている。

ガイダンスの英文正式名称は、“FG22/5 Final non-Handbook Guidance for firms on the Consumer Duty”である。

(注2) ガイダンスの英文正式名称は、“Explaining decisions made with AI”である。

(注3) ガイダンスの英文正式名称は、“Equality law - Banks and other financial services providers”である。

参考図表 1-4-3 ABI「AI ガイド」適用される法規制：公平性

法規制	説明
FCA の消費者保護	○ハンドブックの「消費者への義務」は、金融サービス事業者が顧客をどのように扱うべきかについて、高い基準を設けている。 ○社会的弱者や保護されるべき特性を持つ顧客を含め、顧客の多様なニーズを考慮するよう求めることで、差別的被害に対処しようとしている。 ○異なる顧客グループを区別する AI 由来の価格戦略は、特定の顧客グループにとって悪い結果をもたらす場合、要件に違反する可能性がある。 ○自社の AI モデルが、顧客層によって価格や価値に差異をもたらすかどうかを監視し、説明し、正当化する必要がある。
2010 年平等法	○企業が意思決定プロセスに AI を活用する場合、9 つの保護されるべき特性^(注1) に対して違法な差別につながらないようにしなければならない。 ○企業は平等法に配慮しなければならず、脆弱性^(注2) の特徴の多くは保護されるべき特徴と重なり合っていると、FCA の「脆弱な顧客の公正な扱いに関する企業向けガイダンス」^(注3) は指摘している。 ○(AI システムによる差別的判断などの) 平等法違反は、FCA 規則に違反する可能性があり、規制当局による措置の対象となる。
情報コミッショナーオフィス (ICO) — データ保護	○企業が個人データの処理に AI を使用する場合、UK GDPR とデータ保護法 2018 のもとで定められた規制義務を遵守しなければならない。 ○ICO は、データ保護措置の遵守を執行する責任を負う。 ○ICO は、AI の文脈で適用される UK GDPR の公正原則の解釈方法に関するガイダンス^(注3) を公表した。このガイダンスは、自動化された意思決定とプロファイリングのみに関する UK GDPR 第 22 条のセーフガードの公正さに対する重要性に言及している。また、データ保護の要件は、公平性や差別が概念として存在する他の法律や規制と合わせて読まれるべきであると指摘している。

(注1) 「9 つの保護されるべき特性」とは、年齢、障害、性自認（性別再指定）、結婚している異性カップルと結婚していない同性カップル、妊娠・出産、人種、宗教・信条、性別、および性的志向をいう。これら特性は、差別から保護されるべきことが平等法で規定されている。

(注2) 認知判断能力の低下した顧客、識字や計算に関する能力の低下した顧客などが想定されてい

る。

(注3) ガイダンスの英文正式名称は、“How do we ensure fairness in AI?”である。

参考図表 1-4-4 ABI「AI ガイド」適用される法規制：説明責任とガバナンス

法規制	説明
FCA の「消費者への義務」に則した考え方	○会社は、消費者の利益を守るために、予見可能な損害を特定し、技術が善意で使用されることを確認する必要がある。 ○開発および展開の各段階で、企業は AI の使用によって顧客が経験する不良結果を検出、特定、修正するために、管理情報を収集し分析する必要がある。 ○会社は、消費者に関連する 4 つの指標（製品とサービス、価格と価値、消費者理解、消費者サポート）のもとで潜在的な損害を考慮する必要がある。 ○例えば、AI を使用して作成されたコンテンツが顧客に理解されることを確認するために、どのようなガバナンスとテストが行われているか。 ○チャットボットの展開により、特に弱者とされる顧客のサポートニーズを満たすために、どのような管理情報が収集されているか。 ○会社は、年次報告書と証明書のプロセスにおいて、これらのリスクを効果的に管理し、顧客の不良結果の予防と是正を提示できるようにしておくべきである。
FCA ハンドブック ＞基準条件	○企業は、認可を受けて規制対象の業務を行うためには、FCA の基準条件を満たす必要がある。 ○その中で、十分な非金融資源を保有することが求められている。 ○具体的には、AI の使用に対するガバナンスと監督を行うために、必要な人員数や適切なスキルを持っているかどうかを検討するべきである。
FCA ハンドブック ＞システムと統制	○「SYSC」は、企業の事業運営対して、適切なシステムと統制を設立し、維持するために合理的な注意を払うことを求めている。 ○実装する必要がある既存のシステムと統制は、AI の使用についても拡張して適用される。 ○企業は、技術のリスクを特定し軽減するために適切な方針、手続き、およびコントロールを実施する必要がある。 ○企業は、技術に対する十分な理解とアクセスを持ち、効果的な監督を提供する必要がある。 ○AI ライフサイクルのすべての段階で、SYSC が要求するリスク管理と監督を活用することで、効果的なモデル・リスク管理を実施することができる。
上級管理職と認証制度	○「SMCR」は、規制当局が企業内で特定の領域や機能に誰が責任を持つかを確立するための主要なツールである。 ○規制当局は、金融サービス事業者による AI の使用について、指定された責任と説明責任を割り当てるために、SMCR の規定を活用することを確認している。 ○企業は、AI 使用の監督の説明責任を誰に持たせるかを検討する必要がある。その個人が役割を遂行するために十分なトレーニングと能力を持っていることを確認し、その内容を企業の「責任マップ」に文書化するべきである。 ○上級管理職機能を有する個々の責任者以外の責任者も AI の展開と使用に関与する可能性がある。このような場合、これらのシステムの使用にも責任を持つ人々のために責任声明書を修正することを検討するべきである。

参考図表 1-4-5 ABI「AI ガイド」適用される法規制：異議申立と救済

法規制	説明
FCA の 「消費者への義務」	○「消費者への義務」は、企業が消費者により結果をもたらすような商品やサービスを設計しなければならないと定めている。 ○カスタマージャーニーを通じて、サプライチェーンのすべての部分がどの

法規制	説明
	ようにこれらを提供するかを示さなければならない。
情報コミッショナーオフィス (ICO) –UK GDPR–AI 向けガイダンス ^(注1)	○UK GDPR は、組織が次を行う必要があると規定している。 ・個人に対し、関係するロジックやその意義、想定される結果について、有意義な情報を積極的に提供すること。 ・少なくとも個人には、管理者側からの人的介入を得る権利、自分の意見を表明する権利、および決定に関連するコンテキストを得る権利が与えられる。
情報コミッショナーオフィス (ICO) –AI システムにおける個人の権利	○ICO は、AI システムにおける個人の権利を定め、是正の権利に関するガイダンスを提供している。 ○保険会社は、AI の成果に関連する訓練データに使用された個人データの正確性を確認し、修正する可能性がある。
情報コミッショナーオフィス (ICO) の開発ツール	○ICO は、組織が効果的に対応できるような方法で個人からの要請を支援するため、個人情報の開示請求用のツールを導入している。このツールは、人々が次を特定するのに役立つ。 ・どこに要請を送ればよいか、何を期待すべきかを説明する。 ・受け取った組織は ICO から情報を受け取り、迅速かつ簡単に対応できるようにする。
2018 年データ保護法 第 3 部・第 4 部	○データ主体に不利な法的影響または重大な影響を及ぼす、および法執行目的で実施される、もっぱら自動化された決定に対する保護。 ○個人は人間の介入を受け、自分の意見を表明し、決定についての説明を受け、それに異議を唱えることができる。
FCA ハンドブック –DISP ^(注2) 苦情処理規定	○DISP では、会社は苦情を効果的に特定し、調査し、管理し、適時に解決するための処理手続きと統制を整備することが義務付けられている。 ○苦情を申し立てる個人は、金融オンブズマン・サービス (FOS) に訴えることもできる。 ○苦情を処理するための資格およびその他の要件は、DISP にさらに概説されている。
2010 年平等法	○意思決定プロセスに AI システムを使用する場合、それが差別につながらないことを保証し、また、それを示すことができるようにする必要がある。 ・保護されるべき特性の一つを理由として、決定された意思の受領者が他の誰かよりも不利な扱いを受ける。 ・その結果、保護されるべき特性を持つ人が、そうでない人よりも悪い影響を受ける。

(注1) このガイダンスの英文正式名称は、“Guidance on AI and data protection”である。AI システムでの情報の使用に UK GDPR の原則を適用する方法について、「AI における説明された意思決定」「AI における透明性の確保」「AI における合法性の確保」「正確性と統計的正確性について知るべきこと」などに分けて説明している。

(注2) DISP (Dispute Resolution: Complaints (紛争解決：苦情)) は、FCA ハンドブックのうち、「救済 (Redress)」の配下にある項目である。

参考図表 2-1 NAIC モデル告示 第 3 パート：規制のガイダンスと期待

項目	内容
1.0 一般ガイドライン	
1.1	AI システム・プログラムは、保険会社による AI システムの利用が消費者に不利益をもたらすリスクを軽減するように設計されるべきである。
1.2	AI システム・プログラムは、ガバナンス、リスク管理統制、内部監査機能に取り組むべきである。
1.3	AI システム・プログラムは、AI システム・プログラムの開発、実施、モニタリング、監督、および AI システムに関する保険会社の戦略策定の責任を、取締役会または取締役

項目	内容
	会の適切な委員会に責任を負う経営幹部に負わせるべきである。 1.4 AI システム・プログラムは、保険会社による AI や AI システムの利用、依存に適応、見合ったものでなければならない。管理と手続きは、消費者に悪影響を及ぼす結果を軽減することに重点を置くべきであり、ある AI システムの利用事例に適用される管理と手続きの範囲は、その利用事例に関する消費者への潜在的な危害の程度を反映し、それに沿ったものでなければならない。 1.5 AI システム・プログラムは、保険会社の既存のエンタープライズ・リスク・マネジメント（ERM）プログラムから独立したものであっても、その一部であってもよい。AI システム・プログラムは、米国標準技術局（NIST）の AI リスク・マネジメント・フレームワーク第 1.0 版のような、公的な第三者標準化機関が開発した枠組みや標準の全部または一部を採用したり、取り入れたり、依拠したりすることができる。 1.6 AI システム・プログラムは、商品開発・設計、マーケティング、利用、引受、料率・保険料設定、訴訟管理、保険金管理・支払い、不正検知などの分野を含む、保険のライフサイクル全体にわたる AI システムの利用に取り組むべきである。 1.7 AI システム・プログラムは、設計、開発、検証、導入（システムとビジネスの両方）、使用、継続的な監視、更新、廃棄を含む、AI システムのライフサイクルの全段階に対応すべきである。 1.8 AI システム・プログラムは、保険会社が開発したものであれ、第三者ベンダーが開発したものであれ、規制された保険業務に関して使用される AI システムに対応すべきである。 1.9 AI システム・プログラムは、影響を受ける消費者に AI システムが使用されていることを通知し、AI システムが使用されている保険ライフサイクルの段階に応じた適切なレベルの情報へのアクセスを提供するプロセスと手順を含むべきである。
2.0 ガバナンス	
	2.1 AI システムのライフサイクルの各段階（開発提案から廃止まで）において従うべき、リスク管理および内部統制を含む方針、プロセス、手順。 2.2 AI システム・プログラム方針、プロセス、手順、基準への準拠を文書化するために保険会社が採用する要件。文書化の要件は、セクション 4 を念頭に置いて策定されるべきである。 2.3 保険会社内部の AI システム・ガバナンスの説明責任構造： a) 事業部門、商品スペシャリスト、保険数理、データサイエンスおよびアナリティクス、引受、クレーム、コンプライアンス、法務など、保険会社内の適切な分野や部門の代表者で構成される集中型、連合型、またはその他の方法で構成される委員会の設置。 b) 責任と権限の範囲、指揮系統、意思決定階層。 c) AI システムのライフサイクルの各段階における意思決定者と防衛ラインの独立性。 d) 監視、監査、エスカレーション、報告のプロトコルと要件。 e) 継続的なトレーニングの開発と実施、および人員の監督。 2.4 特に予測モデル：予測モデルの設計、開発、検証、配備、使用、更新、監視のための保険会社のプロセスと手順（予測モデルの使用に起因するエラー、パフォーマンス上の問題、異常値、保険業務における不当な差別を検出し、対処するために使用した方法の説明を含む）。
3.0 リスク管理と内部統制	
	3.1 AI システムの開発、採用、買収の監督と承認プロセス、および自動化と設計に関する制約と管理の特定を行い、機能とリスクのバランスを取る。 3.2 データの最新性、系統性、品質、完全性、偏りの分析と最小化、適切性を含む、データの実践と説明責任手順。 3.3 予測モデル（そこで使用されるアルゴリズムを含む）の管理と監督： a) 予測モデルの目録と説明。 b) 予測モデルの開発と使用に関する詳細な文書。 c) 解釈可能性、再現性、頑健性、定期的な調整、再現性、トレーサビリティ、モデルドリフト、および適切な場合にはこれらの測定の見直し可能性などの評価。 3.4 モデルの開発、訓練、検証、監査に使用したデータの適合性を含め、AI システムの実装時の出力の一般化を評価するために、必要に応じて検証、テスト、および再テストを行う。検証は、モデル開発時に利用可能であった未知のデータにおけるモデルの性能と、

項目	内容
	実装後のデータで観察された性能を比較する、専門家のレビューに対する性能を測定する、またはその他の方法を探ることができる。
	3.5 予測モデル自体への不正アクセスを含む、非公開情報、特に消費者情報の保護。
	3.6 データと記録の保持。
	3.7 予測モデル：モデルの意図された目標と目的、および当該モデルに依存する AI システムが当該目標と目的を正確かつ効率的に予測または実施することを保証するために、モデルがどのように開発され、検証されたかを説明すること。
4.0 第三者の AI システムとデータ	
	4.1 消費者に不利益な結果をもたらす可能性のある、当該 AI システムによる意思決定やそのサポートが、保険会社自身に課された法的基準を満たすことを保証するために、保険会社が第三者やそのデータ、第三者から取得した AI システムを評価するために採用したデューデリジェンスやその方法。
	4.2 適切かつ利用可能な場合には、第三者との契約に次の条項を盛り込む： a) その第三者に対して、監査権を提供する、および／または、保険会社が適格監査機関による監査報告書を受け取る権利を与える。 b) 保険会社による第三者の製品またはサービスの使用に関する規制当局の照会や調査に関して、保険会社に協力するよう第三者に求める。
	4.3 第三者が契約上および該当する場合は規制上の要件を遵守していることを確認するための監査および／またはその他の活動に関する契約上の権利の履行。

(出典：NAIC, “Use of Artificial Intelligence Systems by Insurers” (2023.12) ほかをもとに作成)

参考図表 2-2 NAIC モデル告示 第 4 パート：規制監督と審査の考慮事項

項目	内容
1. AI システムのガバナンス、リスク管理、使用プロトコルに関する情報と文書	
	1.1. 保険会社の AI システム・プログラムに関連する、あるいはそれを証明する情報と書類： a) 書面化された AI システム・プログラム b) AI システム・プログラムの採択に関する情報および証拠書類 c) 保険会社の AI システム・プログラムの範囲（AI システム・プログラムに含まれない、あるいは AI システム・プログラムで扱われない AI システムや技術を含む） d) AI システム・プログラムが、保険会社による AI システムの使用と依存、消費者に不利益な結果をもたらすリスク、消費者への潜在的な損害の程度にどのように適合し、比例しているか。 e) 保険会社の AI システム・プログラムの採用、実施、維持、監視、監督に関する方針、手順、ガイダンス、トレーニング資料、その他の情報： i. 次のような AI システムの開発、採用、買収のプロセスや手順： (1)自動化と設計に関する制約と制御の特定。 (2)データガバナンスと管理、データの系統性、品質、完全性、偏りの分析と最小化、適合性、およびデータの通貨価値性に関する慣行。 ii. モデルや AI システムの開発、検証、監督において保険者が採用または使用する測定値、基準値を含む、予測モデルの管理・監督に関するプロセスと手順。 iii. 予測モデル自体への不正アクセスを含む、非公開情報、特に消費者情報の保護。
	1.2. 第三者が開発したデータまたは AI システムの、保険会社による取得前／使用前デリジェンス、監視、監督、監査に関する情報および文書。
	1.3. 保険会社が AI システム・プログラムを実施し、遵守していることを証明する情報や文書。これには、遵守に関する保険会社のモニタリングや監査活動に関する文書などが含まれる： a) AI システムの開発、使用、監督に関する保険者の調整機関の設立と継続的な運営に関する文書、またはそれを証明する文書。 b) データの系統、品質、完全性、偏りの分析と最小化、適合性、データの通貨性など、データの実務と説明責任手続きに関連する文書。 c) 予測モデルと AI システムの管理・監督：

	i. 消費者に不利益な結果をもたらす可能性のある意思決定を行うため、あるいはそれをサポートするために保険会社が使用する予測モデル、AI システムの一覧とそれらの説明。 ii. 調査または検討の対象となる特定の予測モデルまたは AI システム： (1)保険会社が導入した予測モデルと AI システムの開発、使用、監督において、適用されるすべての AI プログラムの方針、プロトコル、手順を遵守していることを文書化すること。 (2)データソース、出所、データ系統、品質、完全性、偏見分析と最小化、適合性、データの最新正確性を含む、特定のモデルまたは AI システムの開発および監督に使用されたデータに関する情報。 (3)保険者が使用する技術、測定、基準値、同様の管理に関する情報。 d) 予測モデルに基づく意思決定に影響を与える出力の信頼性を評価するためのモデルドリフトの評価を含む、検証、テスト、監査に関する文書。検証、テスト、および監査の性質は、予測モデルに基づくか生成 AI に基づくかにかかわらず、AI システムの基礎となる構成要素を反映したものでなければならないことに留意すること。
2. 第三者の AI システムとデータ	
	2.1 第三者およびそのデータ、モデル、AI システムに対して実施されるデューデリジェンス。
	2.2 第三者の AI システム、モデル、またはデータベンダーとの契約（表明、保証、データセキュリティとプライバシー、データソーシング、知的財産権、秘密保持と開示、および／または規制当局との協力に関する条項を含む）。
	2.3 第三者が契約上および該当する場合は規制上の義務を遵守していることに関する監査および／または確認プロセス。
	2.4 モデルドリフトの評価を含む、検証、試験、監査に関する文書。

（出典：NAIC, “Model Bulletin on the Use of Artificial Intelligence Systems by Insurers”（2023.12）ほかをもとに作成）

参考図表 3 NYDFS「AI システムと外部消費者データと情報源の利用」

項目	内容
Ⅱ. 公平性の原則	9. 保険会社は、適用法に従って、そのデータソースやモデルが、保険法第 26 条に従って保護されているいかなる階級も使用しておらず、いかなる形でもそれに基づいていないことを証明できない限り、引受や価格設定の目的で ECDIS や AI システムを使用すべきではない。さらに、保険会社は、ECDIS や AI システムを引受や価格設定の目的で使用することが、不公正な差別をもたらすか、許容するか、あるいは保険法やその下で公布された規制に違反する場合は、使用すべきではない。
A. データの数理的妥当性	10. 保険会社は、保険引受や料率設定に用いられる他の変数と同様に、ECDIS が一般に認められた保険数理実務の標準によって裏付けられ、統計的研究、予測モデリング、リスクアセスメントを含めて、これらに限定されない、実際または合理的に予測される経験に基づくものであることを証明できない限り、基礎となる分析は、使用される変数と被保険者の関連リスクとの間に、明確で、経験的、統計的に有意で、合理的で、不当に差別的でない関係があることを実証するものでなければならない。 11. 保険会社は、保険引受と保険料設定のために採用された ECDIS が、保険法またはそれに基づき公布された規制によって禁止されていないことを証明できなければならない。また、不当または違法な差別を与える可能性のある保護されるべき人々の代理として機能しないことを証明できなければならない。
B. 不公正で違法な差別	12. 州法および連邦法は、保険会社が特定の保護対象者を違法に差別すること、また、保険会社が特定の基準に基づいて引受を行うことを含め、不公正な差別を行うことを禁じている。保険会社は、ECDIS や AI システムが不当・不法な差別や不公正な取引慣行となるような基準を収集・

	使用していないと判断しない限り、ECDIS や AI システムを引受や保険料設定に使用してはならない。 13. 保険業務の一環として ECDIS や AI システムを使用する場合、保険会社自身がデータを収集し、直接消費者の引受を行うか、直接引受や価格設定の一部または全部の代替となることを意図した外部業者の ECDIS や AI システムに依存するかにかかわらず、保険会社はこれらの識別的差別禁止法を遵守する責任がある。保険会社は、ECDIS や AI システムを使って、保険会社が直接収集・利用することを禁じられている情報を収集・利用してはならない。保険会社は、差別禁止法の遵守を判断するために、識別的でないというベンダーの主張や、独自のサードパーティ・プロセスだけに頼ってはならない。識別的差別禁止法を遵守する責任は、常に保険会社にある。 14. 保険会社は、包括的な評価を通じて、引受または保険料設定ガイドラインが保険法に違反する不当または違法な差別的行為でないことを立証できない限り、引受または保険料設定において ECDIS または AI システムを使用すべきではない。ECDIS または AI システムに由来する引受または保険料設定ガイドラインが、同様の立場にある個人間で不当な差別を行うか、または保護されるべき階級に対して違法な差別を行うかどうかの包括的な評価は、少なくとも以下のステップを含むべきである：(略)
C. 不公正または違法な差別の分析	15. 文書化 保険会社は、ECDIS や AI システムの使用状況、ECDIS や AI システムの複雑さ、重要性に見合った、不当または違法な差別のテスト方法と分析のプロセスと理由を適切に文書化すべきである。保険会社は、要請があれば、そのような文書を保険局に提供する用意がなければならない。 16. テストの頻度 不当または違法な差別的テストと分析は、ECDIS または AI システムのいずれかに重要な更新または変更が加えられるたびに、AI システムを本番稼働させる前、およびその後も定期的の実施されなければならない。 17. 定量的評価 保険会社は、包括的な理解と評価を確実にするため、データおよびモデルの出力を評価する際に、複数の統計的指標を使用することが推奨される。このような指標には、特に以下のものが含まれる：(略) 18. 定性的評価 定量的分析に加え、保険会社の包括的な評価には、不当または違法な差別の質的評価を含めるべきである。これには、保険会社の AI システムがどのように機能するかを常に説明できること、ECDIS や他のモデル変数と被保険者または潜在的被保険者のリスクとの間の直感的な論理的関係を明確に説明できることが含まれる。
Ⅲ. ガバナンスとリスク管理	19. 11NYCRR § 90.2 は、保険会社に、保険会社の性質、規模、複雑性に見合ったコーポレート・ガバナンスの枠組みを持つことを求めている。11NYCRR § 90.1(c)は、「コーポレート・ガバナンスの枠組み」を「保険会社またはシステムの監督、指示、管理、マネジメントのために使われ、法律上、規制上の要件を確実に遵守するための構造、プロセス、情報、関係」と定義している。保険会社は、保険法およびその下で公布された規制の遵守を確保するため、保険会社による ECDIS および AI システムの利用を適切に監督するコーポレート・ガバナンスの枠組みを持つべきである。
取締役会および 上級管理職 の監督	20. 保険会社の取締役会またはその他のガバナンス団体の役割は、取締役会またはその他のガバナンス団体の戦略的ビジョンを実行し、事業体のリスク選好度を監視するための効果的なガバナンスの枠組みを提供することを含め、保険会社の活動を監督することである。 21. 取締役会または他の統治機関は、ECDIS と AI システムの開発・管理を含め、保険会社の活動を監督するための特定の任務と権限を、取締役会または他の統治機関の委員会および上級管理職に委譲すること

	ができる。特定の任務と権限を委譲する場合、認可団体は、理事会または他の統治機関の情報ニーズを満たすため、定期的で質の高い報告とともに、適切な報告システムを確保すべきである。これには、理事会または他の統治機関が、保険会社の ECDIS や AI システムの使用に関連する重要な活動やリスクを理解するための、すべての適時かつ関連する事実を含めるべきである。 22. 上級管理職は、取締役会または他の統治機関の戦略的ビジョンとリスク選好に合致した、保険会社による ECDIS と AI システムの開発とマネジメントの日々の実施に責任を負う。これには、適切な方針と手続の確立、有能なスタッフの配置、モデル・リスク・マネジメントの監督、効果的なチャレンジと独立したリスクアセスメントの確保、内部監査の結果のレビュー、必要な場合の迅速な改善措置などが含まれる。 23. 保険会社が ECDIS と AI システムの利用を効果的に実施するための機能を果たすにあたり、上級管理職は、法務、コンプライアンス、リスクマネジメント、商品開発、保険引受、保険数理、データサイエンスなど、主要な機能分野の代表者を適宜集めた横断的なマネジメント委員会などを通じて、すべての関連業務分野が適切に関与していることを確認すべきである。
方針、手順、文書化	24. ECDIS や AI システム を使用する保険会社は、ECDIS や AI システムの開発および管理を、本サーキュラーレターと整合性のある、文書化された方針および手順で正式に行うべきである。 25. 保険会社の取締役会、もしくはその他のガバナンス団体、または委任された権限を持つ上級管理職は、保険会社の ECDIS や AI システム関連の方針と手続を少なくとも年 1 回見直し、承認し、保険会社の ECDIS や AI システムの使用における変化や、業界におけるベストプラクティスを常に最新のものにしておくべきである。 26. 方針と手続には、明確に定義された役割と責任、およびモニタリングと上級管理職への報告要件が含まれていなければならない。 27. 方針および手続は、ECDIS および AI システムの責任ある合法的な使用に関する、社員の責任に適切に合わせた関連職員のための訓練を含むべきである。さらに、訓練プログラムには、新入社員に対する迅速な訓練と、その後の定期的な訓練、および適時に訓練を完了するための説明責任を含むべきである。 28. 保険会社は、社内開発されたものであれ、第三者から提供されたものであれ、11 NYCRR 243 に準拠した ECDIS を含む、すべての AI システムの使用に関する包括的な文書を保持し、要請があった場合には、当該文書を同局に提供できるように準備しなければならない。当該文書には以下が含まれる：(略) 29. 保険会社は、AI システムや ECDIS の使用に関する消費者からの苦情や問合せに対応するため、苦情を受け取り、対応する手順を実施しなければならない。保険会社は、AI システムや ECDIS に関する苦情の記録を 11 NYCRR 243 に従って保持し、要請があればそのような記録を保険局に提供する準備をしておかなければならない。
リスク管理と内部統制	30. 保険会社は、AI システムのライフサイクルの各段階において、関連するリスクを管理すべきであり、個々の AI システムモデルと全体からリスクを検討すべきである。保険会社は、保険法で義務付けられているように、既存の ERM (エンタープライズ・リスク・マネジメント) 機能の中で AI システムのリスクを管理することも、独立したプログラムの一部として個別に管理することもできる。 31. 保険会社はモデルの開発、導入、使用、検証に関する標準を含むべきであり、保険会社の ECDIS と AI システムの開発とリスクマネジメントに関するリスク分析、検証、テスト、開発、その他のプロセスに対する独立した評価と効果的な取組みを促進すべきである。 32. 保険会社は、役割と責任を明確に定義し、適切なアカウンタビリティの手段をもって、AI システム・リスクマネジメントを実行し、監督す

	る有能で資格のある要員を持つべきである。 33.11 NYCRR § 89.16 は、資産を保護し、コントロールの有効性と効率性を評価し、方針と規制の遵守を評価するために必要な全般的および特定の監査、レビュー、テストを行う内部監査機能を持つことを保険会社に求めている。保険会社は、内部監査機能が、財務上、業務上、コンプライアンス上のリスクと整合性を保ちながら、保険会社の ECDIS と AI システムの利用に適切に関与していることを確認すべきである。このような監査は、AI システムと ECDIS のリスクマネジメントの枠組みの全体的な有効性を評価すべきであり、これには以下が含まれる：(略)
第三者ベンダー	34. 保険会社は、保険引受と保険料設定に使用されるツール、EDCIS、AI システムが第三者によって開発され、または導入されたものであることを理解し、そのようなツール、EDCIS、AI システム が適用されるすべての法、規則、規制に準拠していることを確認する責任を負う。 35. 第三者の適切な監督を確保するため、保険会社は、第三者が開発・配備した ECDIS や AI システムの取得、使用、依存について、文書化された標準、方針、手順、プロトコルを策定すべきである。さらに、保険会社は、誤った情報を第三者に報告し、必要に応じ、さらなる調査や更新を求める手続きを設けるべきである。さらに、保険会社は、保険会社が識別した、あるいは第三者に報告された AI システムの不正確な情報を修正し、排除するための手順を策定すべきである。
IV. 透明性	
情報開示と通知	36. 通達書簡第 1 号（2019 年）で述べたように、透明性は、保険引受と保険料決定における ECDIS の利用において重要な考慮事項である。保険法第 3425 条および第 3426 条は、具体的な根拠や理由が被保険者に書面で提供されない限り、非商業用および特定の商業用損害保険契約を解約、不更新、条件付き更新してはならないと規定している。さらに、保険法第 4224 条(a)(2)および(b)(2)は、ニューヨーク州で事業を行う生命保険会社または傷害保険会社および医療保険会社は、被保険者または被保険者となる可能性のある者の身体的または精神的な障害、傷害、疾病、またはその既往歴のみを理由として、保険を拒否したり、保険の継続を拒否したり、個人が利用できる保険の金額、範囲、種類を制限したり、同じ保険に対して異なる料金を請求したりしてはならないと定めている。ただし、その拒否、制限、料率差が法律や規則で認められており、健全な保険数理原則に基づいている場合、または実際の経験や合理的に予想される経験に関連している場合はこの限りではない。この場合、保険会社は被保険者または被保険者となる可能性のある者に、そのような拒否、制限、または料率差の具体的な理由を受け取る権利、またはその説明をする医療専門家を指定する権利を通知しなければならない。さらに、被保険者または被保険者となる可能性のある者に対し、拒否、制限、または料率差のその他の具体的な理由または理由を適切に開示しなかった場合は、保険業務の遂行における不公正または欺瞞的な行為および慣行とみなされる可能性があり、保険法第 2402 条(c)に定義される、決定された違反を構成する取引慣行とみなされる可能性があり、その場合は保険法第 2403 条の違反となる可能性がある。 37. 保険会社が ECDIS または AI システムを使用している場合、被保険者もしくは被保険者となりうる者、または医療専門家が指名した者に提供される理由には、保険会社が引受拒否、制限、料率差、またはその他の不利な引受判断の根拠としたすべての情報についての詳細が含まれていなければならない（保険会社が不利な引受判断または価格決定の根拠とした具体的な情報源を含む）。 38. (i)保険会社が保険引受または保険料設定プロセスにおいて AI システムを使用しているかどうか、(ii)保険会社が外部業者から入手した個人に関するデータを使用しているかどうか、(iii)当該個人は、保険引受または保険料設定の根拠となった具体的なデータに関する情報を請求する権利を有していること（その請求のための連絡先を含む）。

		39. 保険会社は、不利な保険引受または保険料設定に関する具体性の欠如を正当化するために、第三者のアルゴリズム・プロセスの独自性に依拠してはならない。 40. AI システムの重要な要素、および AI システムが依拠する外部データソースを消費者に適切に開示しないことは、保険法第 24 条に基づき、不公正な取引方法を構成する可能性がある。
--	--	---

(出典 : NYDFS, “Use of Artificial Intelligence Systems and External Consumer Data and Information Sources in Insurance Underwriting and Pricing” (2024.1) ほかをもとに作成)

＜参考資料＞

- ・ 欧州連合日本政府代表部「EU AI 規則案の最新動向」(2023.9)
- ・ 三部裕幸「EU の AI 規制法案の概要」(2022)
- ・ 自民党 AI の進化と実装に関する PT WG 有志「責任ある AI の推進のための法的ガバナンスに関する素案」(2024.2)
- ・ 隅山正敏「AI (人工知能) のガバナンス～何を制御すべきか～」(SOMPO インスティテュート・プラス、2024.3)
- ・ 総務省「広島 AI プロセスについて」(2023.12)
- ・ 損害保険事業総合研究所『主要国における個人情報保護規制の動向と保険業界の対応』(2017.9)
- ・ 野村哲郎、山本英生「金融業界における生成 AI 活用の可能性とリスク」損害保険事業総合研究所 2024 年度特別講座 (2024.5)
- ・ 殿村桂司、丸田颯人「日本の AI ガバナンスの基本となる「AI 事業者ガイドライン (第 1.0 版)」の概要」テクノロジー法ニュースレター (長島・大野・常松法律事務所、2024.5)
- ・ 福岡真之介、松下外『生成 AI の法的リスクと対策』(日経 BP、2023.10)
- ・ 渡部美奈子「欧州・米国の保険業界における AI の活用事例と AI 原則等の動向」損保総研レポート第 140 号 (損害保険事業総合研究所、2022.8)
- ・ ABI, “AI Guide Practical ideas for getting started with responsible AI” (2024.2)
- ・ BOE&PRA, “DSIT-HMT letter” (2024.4)
- ・ Department for Science, Innovation and Technology, “Implementing the UK’s AI regulatory principles: initial guidance for regulators” (2024.2)
- ・ EIOPA, “EIOPA’S DIGITAL STRATEGY – Support consumers, markets and the supervisory community through digital transformation.” (2023.9)
- ・ EIOPA, “Report on the digitalisation of the European insurance sector” (2024.4)
- ・ European Parliament, “Artificial Intelligence Act” (2024)
- ・ FCA, “AI Update” (2024.4)
- ・ Frances Stebbing, “UK develops secure AI guidelines” (Insurance POST, 2023.11)
- ・ GDV, “zum Beginn der Trilogverhandlungen zur KI-VO” (2023.6)
- ・ GOV.UK, “A pro-innovation approach to AI regulation” (2023.3)
- ・ Geneva Association, “Promoting Responsible Artificial Intelligence in Insurance” (2020.1)
- ・ Geneva Association, “Regulation of Artificial Intelligence in Insurance: Balancing consumer protection and innovation” (2023.9)
- ・ Hogan Lovells, “Artificial intelligence in the insurance sector: fundamental right impact assessments” (2024.4)
- ・ Hogan Lovells, “UK: The FCA, Bank of England and PRA issue their strategic approach to regulating AI in response to the government’s AI White Paper” (2024.4)
- ・ IAIS, “Regulation and supervision of artificial intelligence and machine learning (AI/ML) in

- insurance: a thematic review” (2023.12)
- ・ insurance europe, “Key messages in view of the ongoing trilogue discussions on the Artificial Intelligence (AI) Act” (2023.9)
- ・ NAIC, “Model Bulletin on the Use of Artificial Intelligence Systems by Insurers” (2023.12)
- ・ NYDFS, “Use of Artificial Intelligence Systems and External Consumer Data and Information Sources in Insurance Underwriting and Pricing” (2024.1)
- ・ OECD, “The Impact of Big Data and Artificial Intelligence (AI) in the Insurance Sector” (2020.1)
- ・ Richard Nieva, “Cigna Sued Over Algorithm Allegedly Used To Deny Coverage To Hundreds Of Thousands Of Patients” (2023.7)
- ・ PRA&FCA, “DP5/22 - Artificial Intelligence and Machine Learning” (2022.10)
- ・ PRA&FCA, “FS2/23 - Artificial Intelligence and Machine Learning” (2023.10)

<参考ウェブサイト>

- ・ 金融庁 <https://www.fsa.go.jp/>
- ・ 金融データ活用推進協会 <https://www.fdua.org/>
- ・ 国際連合広報センター <https://www.unic.or.jp/>
- ・ 国際連合日本政府代表部 <https://www.un.emb-japan.go.jp/>
- ・ 総務省 <https://www.soumu.go.jp/>
- ・ 損害保険事業総合研究所 <https://www.sonposoken.or.jp/>
- ・ 平将明衆院議員公式サイト <https://www.taira-m.jp/>
- ・ デロイト トーマツ <https://www2.deloitte.com/jp/>
- ・ 内閣府 <https://www.cao.go.jp/>
- ・ 日本経済新聞 <https://www.nikkei.com/>
- ・ 日本損害保険協会 <https://www.sonpo.or.jp/>
- ・ 日本貿易振興機構（ジェトロ） <https://www.jetro.go.jp/>
- ・ まるちゃんの情報セキュリティ気まぐれ日記 <http://maruyama-mitsuhiko.cocolog-nifty.com/security/>
- ・ 読売新聞オンライン <https://www.yomiuri.co.jp/>
- ・ SOMPO インスティテュート・プラス <https://www.sompo-ri.co.jp/>
- ・ イギリス保険協会（ABI） <https://www.abi.org.uk/>
- ・ イングランド銀行 <https://www.bankofengland.co.uk/>
- ・ 欧州議会 <https://www.europarl.europa.eu/>
- ・ 欧州評議会 <https://www.coe.int/>
- ・ 欧州保険・企業年金監督機構（EIOPA） <https://www.eiopa.europa.eu/>
- ・ 金融行為規制機構（FCA） <https://www.fca.org.uk/>
- ・ 経済協力開発機構（OECD） <https://www.oecd.org/>
- ・ 健全性監督機構（PRA） <https://www.bankofengland.co.uk/prudential-regulation/>

- ・全米保険監督官協会（NAIC） <https://content.naic.org/>
- ・勅許保険協会（CII） <https://www.cii.co.uk/>
- ・ドイツ保険協会（GDV） <https://www.gdv.de/>
- ・ドイツ連邦金融監督庁（BaFin） <https://www.bafin.de/>
- ・ニューヨーク州金融サービス局（NYDFS） <https://www.dfs.ny.gov/>
- ・米国コネチカット州議会 <https://www.cga.ct.gov/>
- ・米国損害保険協会（APCIA） <https://www.apci.org/>
- ・保険ヨーロッパ <https://www.insuranceeurope.eu/>
- ・保険監督者国際機構（IAIS） <https://www.iaisweb.org/>
- ・AI Code of Conduct <https://www.aicodeofconduct.co.uk/>
- ・EU Artificial Intelligence Act <https://artificialintelligenceact.eu/>
- ・Forbes <https://www.forbes.com/>
- ・Hogan Lovells <https://www.engage.hoganlovells.com/>
- ・Insurance POST <https://www.postonline.co.uk/>
- ・PwC <https://www.pwc.com/>
- ・Shaping Europe's digital future <https://digital-strategy.ec.europa.eu/en>
- ・The Geneva Association <https://www.genevaassociation.org/>
- ・The White House <https://www.whitehouse.gov/>
- ・UK Parliament <https://www.parliament.uk/>
- ・UN News <https://news.un.org/>