

AI 利活用における法的リスク管理とガバナンス態勢の構築 －金融庁 AI ディスカッションペーパーを踏まえて－

金融 & ロボット/AI ニュースレター

2025 年 8 月 8 日号

執筆者:

[福岡 真之介](#)

s.fukuoka@nishimura.com

[山本 俊之](#)

to.yamamoto@nishimura.com

AI 技術の急速な進展に伴い、金融機関においても AI の導入が進みつつある中で、AI 利活用におけるリスク管理・ガバナンス態勢の構築が重要な課題となっている。金融庁は「AI ディスカッションペーパー（第 1.0 版）」（以下「AI ディスカッションペーパー」）を 2025 年 3 月に公表し、金融分野における AI の健全な利活用を促進するための初期的な論点整理を示すに至った。その中で、金融庁は、AI のもたらし得る負の影響だけではなく、金融機関が技術革新に取り残されて中長期的に良質な金融サービスの提供が困難になる「チャレンジしないリスク」を指摘し、金融機関による AI 活用の取り組みを後押しするとしている。金融機関には AI 導入への積極的な取り組みとともに、それに伴った AI のガバナンス態勢・リスク管理態勢が求められるよう。

I AI 導入の阻害要因と対策

AI ディスカッションペーパーは、金融機関における従来型 AI と生成 AI 共通の課題として、①データ整備・品質管理、②サードパーティーリスク管理（外部ベンダー管理）、③投資対効果（ROI）の不確実性が挙げられている。これらの対策としては、①についてはデータ収集・管理態勢の整備、②については契約上の手当てや管理フレームワークの採用、③についてはデジタル部門での予算化や KPI の設定とモニタリングなどが考えられる。

II 生成 AI により特に難化した課題（1）

以上に加えて、AI ディスカッションペーパーは、生成 AI により特に難化した課題として、以下の①から⑥の項目を挙げている。

①説明可能性の確保：生成 AI による生成は膨大なパラメータによるプロセスを経ており、その特性上、人間が理解・説明することが困難である。他方で、金融機関には、顧客に影響を及ぼす意思決定などについて一定の説明責任を果たすことが求められるため、説明責任を果たすことができないリスクがある。

②公平性・バイアスのリスク：AI の学習データやアルゴリズムに偏りが含まれると、特定の属性を持つ顧客層に不当な不利益を与えるリスクがある。例えば、融資判断において、他の条件は同一であるにもかかわらず、性別や人種などによって、融資の可否が異なるといったリスクが生じる可能性がある。

③AI システムの開発・運用及びモデル・リスク管理：AI モデルについても、他の金融モデルと同様にモデルの誤り又は不適切な使用に基づく意思決定によって悪影響が生じるリスクがある。このリスクに対しては、金融庁「モデル・リスク管理に関する原則」（2021 年 11 月公表）などに基づいたリスクベース・アプローチがある。対策としては、社内で使用している主要な AI モデルを一覧で把握するモデル・インベントリーの

整備が挙げられる。

④個人情報保護と情報セキュリティ・サイバーセキュリティ：生成 AI では、IT 企業（ベンダー）が提供するクラウド上の大規模モデルを利用するケースも多く、モデルにデータを入力する際の個人情報漏えいや、出力から個人情報が推測されるリスクがある。また、第三者が提供する AI サービスを使う際に、そのサービス自体がサイバー攻撃を受けて改ざん・悪用されるリスクや、自社システムとの連携部分にセキュリティホールが生まれるリスクがある。もっとも、個人情報保護法違反リスクについては、個人情報保護法に従えば、生成 AI に入力することは可能であり、問題はどのようにして個人情報保護法を遵守するかであると言える。

⑤ハルシネーション：生成 AI はもっともらしいが誤った情報を作り出す「ハルシネーション」と呼ばれる現象がある。そのため、例えば、AI が作成したレポートや回答に誤情報が含まれ顧客の意思決定を誤らせるリスクがある。対策として、人間が途中で関与し、内容をチェックすることが挙げられる。しかし、すべてについて人間が関与するのであれば AI を導入した意味がなくなってしまう。生成 AI の利用方法としては、（i）社内利用、（ii）対顧客サービスへの間接的な利活用、（iii）対顧客サービスへの直接的な利活用があり、社内利用と対顧客への間接的な利活用については従業員教育によりかなりの部分が対応可能である。対顧客への直接的な利活用については、人間が関与するか、生成 AI であることを明示してハルシネーションの危険性を警告する対策が考えられる。

⑥金融犯罪への悪用：生成 AI は、巧妙なフィッシングメールやディープフェイクと呼ばれる本人なりすまし音声や画像を簡単に生み出せるため、詐欺やサイバー犯罪に悪用されるリスクがある。金融機関としては、システム的な防御に加えて、自らが騙されないように態勢整備や従業員教育をするとともに、騙された顧客による送金等を検知する態勢を整備し、顧客の被害や不安を抑えることが求められる。また、生成 AI が生成する情報が、SNS 投稿などを通じて、誤情報に基づく風説の流布や、AI 主導のアルゴリズム取引が予期せぬ連鎖反応を引き起こすことで金融市場のボラティリティが高まるリスクがある。リスク管理部門はマーケットリスクやレピュテーションリスクの観点からもこの問題に対応する必要がある。

Ⅲ 生成 AI により特に難化した課題（2）

生成 AI の課題としては、上記の AI ディスカッションペーパーが取り上げた項目に加えて以下のものが挙げられる。

⑦秘密情報の漏えい・守秘義務違反：個人情報と同様に、秘密情報を生成 AI に入力した場合における情報漏えい、金融機関の守秘義務違反、秘密保持契約違反のリスクがある。営業秘密を漏えいした場合には不正競争防止法違反を問われるリスクもある。秘密情報については、生成 AI への入力が必要しも禁止されるものではなく、どのような場合に入力を認めるかのルールの整備や外部に情報が開示されないシステムの構築によって対応が可能な部分もある。

⑧著作権侵害：AI を開発する段階では AI 学習のためのデータの複製が著作権侵害となる可能性がある。また、AI を使った生成・利用段階ではアウトプットの生成・利用が著作権侵害になる可能性がある。前者については著作権法 30 条の 4 により、原則として著作権者に無許諾で著作物を利用することが可能であるが、学習元著作物の全部及び一部の創作的表現の出力を目的とした学習をする場合には、学説上争いのあるものの、同法 30 条の 4 の適用がない可能性が高い。特に RAG については、仕様によってはウェブサイトなどで検索した第三者の著作物の全部又は一部を出力するものもあり注意が必要である。生成・利用段階においては、ライセンスを受けていない限り、著作権侵害の現実的可能性が高まることから、他人の著作物のプロンプト入力を避けることも対策の 1 つとして挙げられる。また、他人の著作物と類似したアウトプットの利用

は避けるべきである。

IV 金融機関の AI ガバナンス態勢構築のポイント

1. AI ガバナンスとは

新たな技術の利活用には新たなリスクを伴うし、リスクをゼロとすることもできない。他方で、それに伴うリターンや便益もある。そこで AI については、ガバナンス態勢を構築して、リスクをコントロールしながら利活用を進めるべきという考え方が一般的である。

その際、鍵となる概念の1つとして、総務省＝経済産業省「AI 事業者ガイドライン（第 1.1 版）」（2025 年 3 月）にもあるように、AI の利活用にあたっては、予め事前に当該利用分野における利用形態に伴って生じ得るリスクの大きさ（危害の大きさ及びその蓋然性）を把握したうえで、その対策の程度をリスクの大きさに対応させる「リスクベース・アプローチ」が重要となる。このリスクベース・アプローチは、金融機関に適用されるガイドライン、例えば、前記Ⅱ記載の金融庁「モデル・リスク管理に関する原則」における指摘とも親和的である。

したがって、法令遵守を前提としながら、リスクベース・アプローチの採用により、リスク低減策を講ずるためのガバナンス態勢構築を検討するのがよいと考えられる。また、AI ガバナンスは、明示的に法令にその内容が定められているものではないため、各金融機関が各社の実情にあわせて構築していかざるを得ない。なお、AI のように進歩が早い分野においてはアジャイル・ガバナンスの考え方も重要であろう¹。

2. 金融機関特有の論点

AI ガバナンス自体は、業種を問わず、AI を利活用する企業であればその導入を考えるだろうが、金融機関特有の論点として、例えば以下の事項が考えられる。

- | |
|---|
| (i) 規制業種であること、公共性、社会的信用の高さ |
| (ii) 金融機関における 3 つの防衛線・3 線モデルの適用 |
| (iii) 関連部署が多い中で、誰・どの部署がリーダーシップを取るか |
| (iv) 営業・ビジネス部門、IT・開発部門の「やる気」を損なわない仕組み作り |

(i) については、AI ディスカッションペーパーが「技術革新に取り残されて中長期的に良質な金融サービスの提供が困難になる『チャレンジしないリスク』も十分に認識されるべき」とするように、AI 利活用にあたっては「攻め」の姿勢を取ることが必要となってくる反面、金融機関は規制業種であり、公共性、社会的信用の高さがそのビジネスの根幹・特徴となっていることから、「守り」の姿勢や、「攻め」と「守り」のバランスが重要な考慮要素となってくるであろう。

(ii) については、金融機関では、第 1 線の営業・ビジネス部門、第 2 線のリスク管理部門やコンプライ

¹ PDCA（Plan-Do-Check-Act）を内包しつつも、その前提となる環境分析やゴール設定を常に見直し続けるとともに、外部に対する透明性やアカウンタビリティを確保するモデル。詳細は、Society5.0 における新たなガバナンスモデル検討会『アジャイル・ガバナンスの概要と現状 GOVERNANCE INNOVATION Vol.3』（2022 年 8 月）16～18 頁参照。

アンス部門、第 3 線の内部監査部門を通じて、リスク管理態勢を構築し、全社的なリスク管理を行っているのが通例と思われるところ、この仕組み自体が AI ガバナンスにおいても役立つものであり、むしろ、既存の 3 つの防衛線・3 線モデルに、AI に関するガバナンス態勢をどう組み込むか、という点が検討事項となるべきである。

(iii) については、AI の導入プロジェクトでは、上記 (ii) 記載の部門に加えて、IT・開発部門や企画部門といった金融機関における数多くの部門が関与するのが通常である。そのため、各部門がバラバラに導入プロジェクトを進めてもうまいかないことから、関係部門を主導・統括する部門が AI ガバナンス導入にあたって必要であり、重要になると考えられる。最近では、金融機関でも DX に関する企画立案を行う部門が設けられ、当該 DX 部門が主導するケースも多いと思われるが、これも各金融機関の組織や考え方による。さらに、AI に詳しい知見を有している者がまだまだ少ない分野では、関係部門だけでなく、特定の個人によるプロジェクトのリードも重要になるのが実情である。

(iv) については、(i) の「攻め」の姿勢とも共通するが、慎重になるあまり、例えば、やたらと分厚いマニュアルや緻密なチェックリストが実務的に機能するのかという問題である。チェックリストを埋めること自体が目的となつては本末転倒であり、形式的なチェックはかえってリスクに対する認知を埋没させることになりかねない。リスクベース・アプローチの下で、メリハリをつけ、「攻め」「守り」のバランスを考慮した仕組み作りが求められる。

上記に加えて、AI ガバナンスの導入自体が法令遵守を必ずしも保証するわけではないので、法務部門による確認プロセスも重要となろう。

3. AI ガバナンス態勢構築において考慮すべきポイント

下記に限定されるわけではないが、以下の表では、AI ガバナンス態勢構築において考慮すべきポイントをまとめている。AI ガバナンス態勢構築においては、FDUA（一般社団法人金融データ活用推進協会）の「金融機関における生成 AI の開発・利用に関するガイドライン（第 1.1 版）」第 4 章も参考となる。

【AI ガバナンス態勢構築において考慮すべきポイント】

項目	考慮すべきポイント
1. 経営陣の関与	① 経営陣が AI 活用の目的と戦略を明確に示し、全社方針としてガバナンス態勢の構築を主導 ② リスクアペタイト（許容できるリスクの範囲）を明文化し、AI に対する取り組みの方向性を統一 ③ 取締役会・経営会議に対する定期的な報告、モニタリングやレビューの実施
2. AI ポリシー・社内規則の制定	① AI 活用に関する基本方針（AI ポリシー）を策定・公表 ➢ 基本原則、利用範囲、説明責任、公平性、倫理など ② より詳細な社内規則、実務指針やマニュアル整備 ➢ 所管部署 ➢ AI に関する部署間の役割 ➢ 取締役会・経営会議への報告ルール ➢ リスクの特定・評価手法

項目	考慮すべきポイント
	➤ AI の導入や利用に関する承認プロセス ➤ 法務面では、例えば、個人情報・個人データ、顧客情報、機密情報、秘密情報の取扱い・ルール、著作権等の知的財産権の取扱い・ルールなど
3. 組織の整備	① 主導・統括部門、責任者の取り決め ② 営業・ビジネス部門、企画・DX 部門、リスク管理部門、法務・コンプライアンス部門、IT・開発部門などが連携する横断的な委員会・ワーキンググループの設置 ➤ 3つの防衛線・3線モデルの意識 ③ 必要に応じて外部アドバイザーを招聘し、第三者目線で、最新技術やガバナンス手法を取り入れる（アドバイザリーボード・外部委員会の設置）
4. リスクの特定・評価、承認プロセスの整備	① 「リスク」の定義 ② リスクの特定や評価の目線 ➤ 「内側」（社内利用・業務効率化等）か、「外側」（対顧客サービスへの利用）か ➤ 個人情報、顧客情報、機密情報、秘密情報の取扱い状況 ➤ 人間の関与度合い ➤ モデル自体の複雑さ、説明可能性 ➤ 公平性 ➤ モデルの内製・外製の別（サードパーティーリスク） ③ ハイリスクな AI の取扱い ④ 外部ベンダー利用時のサードパーティーリスク ⑤ AI の開発・利用前におけるリスク評価と承認プロセスの明確化 ⑥ 性能・精度・バイアス・データ品質等に関するチェックリストの運用、承認
5. モニタリング	① 稼働中の AI モデルについて定期的なモニタリングを実施（予測精度、異常検知、エラー率等） ② リスクが高いユースケースについてはリアルタイム監視や人間の判断（Human in the Loop）を組み合わせたことも検討 ③ モニタリング結果を経営陣・関連部門と共有し、改善アクションを迅速に実施 ④ AI ポリシーや社内規則の見直し ⑤ AI 導入後も継続的なレビューと性能確認（チェンジ管理、バージョン管理など）を実施
6. 従業員教育、カルチャーの醸成	① 全社員に対する AI リテラシー教育と、役割に応じたリスク対応研修の実施 ② ケーススタディ、具体的事案による教育プログラム提供 ③ AI 活用における留意事項や禁止事項の明示 ④ 説明可能性や透明性を重視し、AI 判断の過信を避ける組織文化

項目	考慮すべきポイント
	の醸成 ⑤ インシデントやヒヤリハットを共有し、組織全体での学習と再発防止を促進 ⑥ 経営陣から現場まで一体となって、リスクと責任を自覚する風土を築く

「1. 経営陣の関与」の下、どれも重要な項目とはなるが、とりわけ「4. リスクの特定・評価、承認プロセスの整備」については、評価基準の策定から実際の運用まで、態勢構築に非常に時間を要するのではないと思われる。

例えば、「リスク」の定義 1 つ取ってみても難しい点がある。「AI 事業者ガイドライン」のように、リスクを危害の大きさとその蓋然性とした場合であっても、金融機関で一般に用いられている、一定の確率分布を想定して過去データを用いながら数値結果を算出する計量モデルとは異なり、AI のリスクについては計量化が必ずしも容易とは言えず、定性的に評価せざるを得ないのだと思われる。したがって、リスクを特定できた後であっても（前記Ⅱ及びⅢ参照）、「危害の大きさ」を実際に計量化することには困難を伴い、「その蓋然性」についても AI のような新規分野では過去データから推定することは難しいので、定性評価も加味しながら、一定の類型化を行うことで対応していかざるを得ないのではなかろうか。その意味で表中の「②リスクの特定や評価の目線」に記載の通り、考えられる評価の目線を数多く持ちながら検討していくことが重要になると考えられる。さらにその際、異なる営業・ビジネス部門で用いられる各種の AI を、一定の基準をもっていわば「横串」を刺す形でリスク評価することもやはり困難であり、「走りながら考える」アジャイル的な発想も重要となろう。

ハイリスク AI については、デジタル庁「行政の進化と革新のための生成 AI の調達・利活用に係るガイドライン」（2025 年 5 月）において記載されている高リスク判定シートや、当該判定シートにおけるリスク判定ロジックのような考え方も参考になろう。リスク軸として、利用者の範囲・種別、生成 AI 利活用業務の性格、要機密情報や個人情報等の学習等の有無、出力結果の政府職員による判断を経た利活用、という視点が示されており、民間分野でもアレンジすることで利用可能となるかもしれない。

外部ベンダー利用時のサードパーティーリスクについては、外部委託に関する既存の法令に加えて²、FISC（公益財団法人金融情報システムセンター）「金融機関等コンピュータシステムの安全対策基準・解説書（第 13 版）」（2025 年 3 月）における AI の安全対策基準も参考となろう³。また、サードパーティーリスクについては、管理・監督、モニタリングに加えて、サービス提供停止・終了・契約解除時の対応策についても予め検討が必要になるとと思われる。

AI ガバナンスについては、いわゆるハードローではないため、各金融機関自らが、利活用する AI についてリスク評価等を行いながら、自身で態勢構築を行う必要があるところに難しさがあると言えるし、逆に創意工夫の余地もあると言える。

² 例えば、銀行法 12 条の 2 第 2 項、同施行規則 13 条の 6 の 8。金融庁「主要行等向けの総合的な監督指針」III-3-3-4、同「中小・地域金融機関向けの総合的な監督指針」II-3-2-4。

³ 同安全対策基準実 150、統 20～23。

V おわりに

AIの利活用には「チャレンジしないリスク」もあることから、AIリスクに対応しながらも、技術革新を前向きに取り入れる姿勢が重要であると考えられる。金融庁は、「金融庁 AI 官民フォーラム」などを通じて、業界と対話しつつ、柔軟で実効的なガバナンス態勢の整備を後押しするとしている。そのため、金融機関としては、自社の状況に即した態勢やルールの構築とともに、今後に変化する技術環境を的確に捉え、AI を安心して活用できる枠組みを構築していくことが求められる。

当事務所では、クライアントの皆様のビジネスニーズに即応すべく、弁護士等が各分野で時宜に合ったトピックを解説したニュースレターを執筆し、随時発行しております。N&A ニュースレター購読をご希望の方は [N&A ニュースレター 配信申込・変更フォーム](#) よりお手続きをお願いいたします。

また、バックナンバーは [こちら](#) に掲載しておりますので、あわせてご覧ください。

本ニュースレターはリーガルアドバイスを目的とするものではなく、個別の案件については当該案件の個別の状況に応じ、日本法または現地法弁護士の適切なアドバイスを求めているいただく必要があります。また、本稿に記載の見解は執筆担当者の個人的見解であり、当事務所または当事務所のクライアントの見解ではありません。

西村あさひ 広報課 newsletter@nishimura.com