

生成AIの台頭で大きく注目されるAIガバナンス ～欧米のAI・デジタル動向の考察と示唆～

【後編】

2025年4月10日

株式会社日本総合研究所 先端技術ラボ

株式会社三井住友フィナンシャルグループ

シリコンバレー・デジタルイノベーションラボ

本資料は、作成日時点で弊社が一般に信頼出来ると思われる資料に基づいて作成されたものですが、情報の正確性・完全性を保証するものではありません。また、情報の内容は、経済情勢等の変化により変更されることがあります。本資料の情報に基づき起因してご閲覧者様及び第三者に損害が発生したとしても執筆者、執筆にあたっての取材先及び弊社は一切責任を負わないものとします。尚、本資料の著作権は株式会社日本総合研究所に帰属します。

<本件に関するご照会先>

田谷洋一 株式会社日本総合研究所 シニアエキスパート 兼 JRI America, Inc. Executive Director (yoichi.taya@jri-america.com)

緒方雄二 株式会社三井住友フィナンシャルグループ 部長代理 (yuji.ogata@smbcgroup.com)



- 近年、グローバルにAIの社会実装が進み、さまざまな業界におけるAI活用が進んでいる。日本国内でも生成AIの利用に向けた議論がなされているが、実際の利用率や検討状況などを見ると、わが国でAI活用が大きく進展しているとは言い難い。
- 日本企業がAI導入を進める際の障壁として挙げられるのは、AIのリスク管理やプライバシー保護等のガバナンス施策の推進である。AIは多分野にわたる活用が期待される反面、安全性や信頼性など、AIが社会に与える影響を多角的に検証することも併せて求められており、組織が部門横断的に施策や戦略を立案する重要性が指摘されている。
- AI実装で先行する欧米ではAI法規制の整備が進みつつある。欧州では産業横断的に厳格なAIリスク管理を行う法規制が採択された一方、米国では産業の特性に応じて、より効果的なガバナンスの整備とイノベーションの促進を実現するための政策がとられている。
- 欧米においても法規制の整備は発展途上の段階であるため、現時点でAI施策を評価するのは時期尚早である。しかし、過去のデジタル進展を振り返ると、米国ではイノベーションとともに大きな経済成長がもたらされた事例も多く、AIのリスクを適切にコントロールしながら、イノベーションを強く促す方法を模索することはAIの発展においても有効だと考えられる。
- 前編で整理した内容を基に後編では、主に米国の事例や政策などの紹介を中心に、企業に求められるAIガバナンスの要素と日本のAI活用の進展に向けた考察をする。



第3章

AI実装が進む米国の動向と AIガバナンス施策に取り組む企業の先進事例



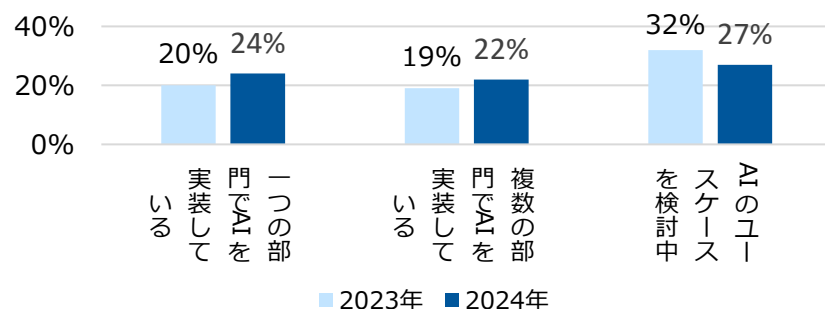
1. グローバルに活用が進むAIと求められるガバナンス整備

- AI導入を推進する企業やAI施策の専用予算を確保する企業の割合がグローバルに増加している
- 一方、責任あるAIの実装がAI導入における課題となっており、AIガバナンス態勢の整備が企業に求められている

AIの導入はグローバルに増加

- 調査会社Omdiaの「2024 AI Market Maturity Survey Analysis」によると、**AIを導入している企業の割合は2023年の20%から2024年の24%に増加しており、グローバルにAIの導入が拡大している**（図表1）。一方で、AI導入を進める組織の50%以上はまだ導入プロセスの初期段階にある。
- AI施策の予算を確保する企業の割合も2023年の50%から2024年の60%に増加しており、多くの組織がAIへ強いコミットメントを示している**。特に、AIに100万ドル以上の投資をする企業の割合が大きく増加しており、テクノロジー予算の5%以上をAIに投資する企業はROIを得ていると回答している。

図表1：AI導入の状況（グローバル）

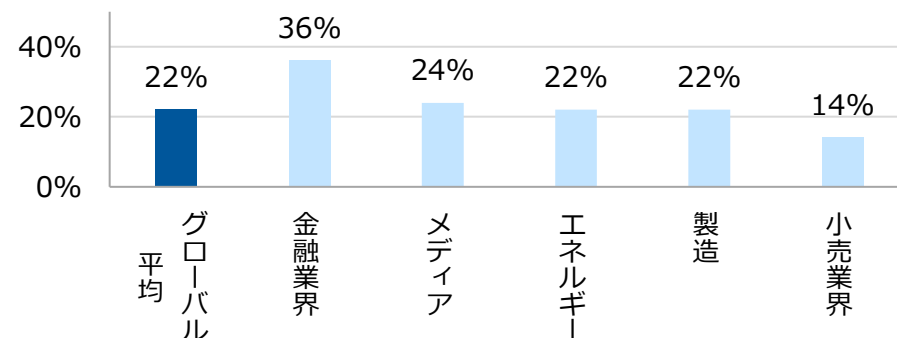


（参考）The AI Summit New York 2024 “Informa TechTarget Market Trends Report (2024/12/11)”を基に作成

【参考】AI導入状況に関するデータ

- AIの導入では金融業界が他業界をリードしており、AIを導入している企業の世界平均22%に対して36%と高い割合**である（図表2）。金融業界は膨大な数値データ（顧客、取引、市場、等）を扱うため、データの傾向分析や異常検知など、広範な領域でAIを活用する大きなメリットがある。
- 企業のAI導入がグローバルに進むなか、**責任あるAIの実装が、企業のAI推進における主な障壁**となっている。特にEUでは、5月に採択されたAI規制法によって企業のコンプライアンスや規制への順守が強く求められており、同法が企業のAI導入に大きな影響を及ぼしている。

図表2：AIを導入している企業の割合（グローバル）



（参考）The AI Summit New York 2024 “Informa TechTarget Market Trends Report (2024/12/11)”を基に作成

2. AIの社会実装とともに重要性を増す法規制・ガバナンスの整備

- AIの実用化が進む米国では、AIサービスがユーザーや消費者の生活や思想に与える影響も大きくなりつつある
- 企業がAIを適切に扱うためには、法規制でのコントロールに加え、組織が安全性や信頼性などの多角的な観点でガバナンスやリスク管理を強化するための施策を推進することが肝要となる

AI実装が進む中で発生した米国内での事件

- ・ 近年、米国ではAIの実用化が進み、**AIが消費者の生活や思想に与える影響も大きくなりつつある**。米国では、すでにAIの利用に関する訴訟や事件が多数発生しており（下表）、**AIが一層浸透する社会に向けて、AIモデルを適切に管理し、透明性を証明する必要性が指摘されている**。

事件	内容
AIチャットボットとの会話にのめり込んだ少年が自殺をした事件（フロリダ州）	フロリダ州の14歳の少年はAIが生成したキャラクターとコミュニケーションができるチャットボットアプリ（Character.ai）を利用していた。少年はチャットボットと会話をするうちに依存症となり、恋に落ちたとされ、最終的には現実世界で生きる気力を失い、2024年2月に自殺に至ったと報道されている。
UnitedHealthcareのCEOが銃殺された事件（NY州）	2024年12月UnitedHealthcareのCEOがNY市内で射殺された。同社は他社の保険対比、保険請求への却下率が高く、保険サービスへの不満が本事件の一因になったと指摘される。同社の保険審査にはAIが利用されており、患者のリハビリ期間や介護施設の利用期間を予測し、保険適用を判断していたと報道されている。

（参考）<https://www.nbcnews.com/tech/characterai-lawsuit-florida-teen-death-rcna176791>,
The AI Summit New York 2024 “The Enterprise AI Legal Guide: Insights on Navigating Compliance and Regulatory Challenges (2024/12/11)”などを基に作成

AIの社会実装とともに重要性が高まるガバナンス整備

- ・ AIの社会実装は一層拡大することが予想されており、AIモデルを適切に管理していくことが組織や企業にとっての大きな責任に
- ・ AIの実装で先行する米国では、AIに関連する法規制の整備や変更が進み、企業のAI管理態勢の強化も拡大しつつある。

1. 法規制の観点

- ・ 米国では、**トランプ大統領がバイデン前政権のAIの安全性に関する大統領令を撤回し、規制緩和を指示する大統領令を発表**した。一方、**州法レベルではAIの規制強化が進んでいる（P6, 8参照）**。また、規制の厳しい金融業界においては、AIに対する米国当局の動きが活発化しており、財務長官がAIを金融システムにおける最重要課題だと強調した（P9参照）。

2. 組織の対策の観点

- ・ AIを活用する大手企業は、**説明可能性や透明性、法令順守、AI人材の強化、データ整備など、さまざまな角度からAIの安全性や信頼性を担保するガバナンス強化関連施策を推進している（P10～P15にて欧米企業の事例を紹介）**。



3. トランプがAIの安全利用に関する大統領令を撤回

- 2025年1月20日、トランプ大統領がバイデン前政権のAIの安全性に関する新基準などの大統領令を撤回した
- 米国連邦政府レベルでは規制緩和が進むものの、EUのAI規制法や州法レベルで規制強化が図られる状況を鑑みると、グローバル企業には引き続きAIガバナンスへの対応が強く求められると予想される

バイデン政権のAIの安全利用に関する大統領令を撤回

- 2025年1月20日 **トランプ大統領がバイデン前政権のAIの安全性に関する新基準などの大統領令（Executive Order 14110）*を撤回**した
- また、**2025年1月23日には、AIに対する規制緩和を指示する大統領令を発表**した。同令では、経済競争力や国家安全保障等を促進するために、**米国のAI競争優位性を確立することが重要であり、イノベーションの障壁となる既存のAI政策を無効化し、米国がAIのグローバルリーダーシップを維持するための行動を取る**ことなどが示された。
- トランプ大統領は、大統領補佐官らに対して、バイデン元大統領が発表した**AIの安全性に関する新基準などの大統領令に基づいて取られた政策や指令、規制などを見直す**とともに、米国のAI優位性を確立する上で障壁となっている場合は、**それらを修正、撤回などの提案をするように指示**した。

（参考）<https://www.jetro.go.jp/biznews/2025/01/5e89056b1f856822.html>
を基に作成

* <https://www.jri.co.jp/page.jsp?id=108778>の前編P20～P23参照

企業のAIガバナンスに関する今後の影響は

- **バイデン前政権のAI規制に関する大統領令の撤回が、AI業界に与える影響が懸念**されている。同令については、マイクロソフトが「AIガバナンスにおける重要な前進」、Googleが「我々のAIの責任慣行が原則に合致する」などの評価をしていた。
- しかし、同令の撤回によってAIガバナンスの規制の枠組みが失われることで、**AI開発における安全性の担保が企業の自主性に一層委ねられることになる**と予想される。
- 連邦レベルで規制緩和が進む一方、**州レベルでは新たな規制強化が進む可能性**も指摘されている。カリフォルニア州やコロラド州などが独自のAI規制の整備を進めており（P8参照）、**企業には各州の規制に対応する必要が生じる**。
- また、**国際的な観点では、EUのAI規制法との整合性も十分に考慮する必要がある**。米国で規制緩和が進んだとしても、グローバルに事業展開する企業は、より厳格なEUの基準に合わせざるを得ない状況が予想される。**結果として、企業はAIガバナンスへの対応をとっていくことが不可欠**となろう。

（参考）<https://innovatopia.jp/ai/ai-news/46920/>,
<https://tech.yahoo.com/ai/articles/trump-dumps-biden-ai-safety-093617185.html>
などを基に作成

【参考】DeepSeekの登場と米国政策への影響（米国ベンチャーキャピタルの見解等）

- DeepSeekのようなオープンソースLLMが台頭すると、AI開発のスタイルが大きく変化する可能性がある。
- LLMのコモディティ化により、米中のAI開発競争がさらに激化するとともに、トランプ政権はAIの開発やイノベーション促進を一層強化するための政策を取るだろう

DeepSeekの登場により米中AI開発競争が激化

- シリコンバレーではOpenAIやAnthropicのようにクローズドなAI開発が主流*であるため、**DeepSeekのようなオープンソースのLLMが台頭するようになると、AI開発のスタイルが大きく変化する可能性**がある。
- 従来は、高性能なLLMとその技術を独占することに企業の競争優位性があったが、オープンソースによって**LLM開発がコモディティ化すると、LLMのカスタマイズやアプリケーションの競争力が重要になる**。実際に投資家の注目は、**LLMの性能優位性からLLMを活用したアプリケーションへと移りつつある**。
- LLMのコモディティ化により、米中のAI開発競争がさらに激化することも予想され、トランプ政権はAIの開発やイノベーションを一層強化するための政策を取ると考えられる。
- とりわけAIは軍事的にも広く利用される技術であるため、**連邦政府の政策がイノベーションの促進に偏ると、AIのガードレールとなるガバナンス施策が十分に議論されない可能性もある**。今後の米国のAI政策には注視が必要な状況である。

（参考）Geodesic Capital投資家Jon Rezneck, Brent Wu, Arvind Ayyala, Will Horyn, Justin Yueからヒアリングした内容（2025年2月4日）を基に作成

*Metaが提供しているLlamaはオープンソースであるが、学習データは公開していない。一方、DeepSeekは学習データも公開しており、学習プロセスも透明性が高いとされる。

DeepSeekとは

- 中国のLLM開発企業およびAIモデルの名称、2023年5月設立。
- GPTシリーズやClaude、Llamaなどと同じくテキスト生成や自然言語処理を行うAIモデルを提供。
- GPTなどの超大規模GPUクラスタを必要とするLLMと比べ、**コスト効率良く強力なAIをトレーニングできるように最適化されたアルゴリズムを採用**。これにより、**比較的低コストなGPUを利用し、高性能なAIモデルを開発することが可能**に。
- 2024年12月にリリースされたDeepSeek-V3は**OpenAIのGPT-4oやAnthropicのClaude3.5 Sonnetと同等の性能**を持つとされる。
- DeepSeekは**学習データも含めたオープンソースのモデルを公開しており、誰でも利用・カスタマイズが可能**。
- 従来は、GPT-4クラスなどの高性能なLLMを保有すること（技術の独占）に競争優位性であったが、**DeepSeekの登場により、多くの開発者が高性能なLLMを開発できるようになり、LLM開発のコモディティ化やLLM市場の激化が進む**と予想されている。

（参考）<https://venturebeat.com/ai/deepseek-v3-ultra-large-open-source-ai-outperforms-llama-and-qwen-on-launch/>などを基に作成

【参考】米国のAIに関する主な州法

- カリフォルニア州やコロラド州では、AIの管理やAIサービスの評価等に関する先進的な法案が制定されている
- 最新の州法では、人間の生活や社会に与える影響が大きいAIを介在させた判断について、企業や組織にAIモデルの透明性や説明責任を果たすことが求められている

カリフォルニア州 (Assembly Bill 2013 (AB 2013))

- カリフォルニア州では、Assembly Bill 2013 (以下、AB 2013) が成立し、2026年に施行される予定である。本法案の目的は、**AIの最大の課題の1つであるブラックボックス問題を解決すること**である。
- AB 2013は、生成AIモデルの開発者に対して、**モデルのトレーニングに使用するデータに関する詳細な情報を開示し、学習の仕様やアルゴリズムを明らかにすることを求めている**。
- 同法は米国の州法の中でも特に先進的な事例であり、カリフォルニア州がAIの説明責任を強く働きかけ、**AI企業に対して、AIシステムの透明性を示すように要求している**。
- AB 2013は、主にAIをトレーニングする際のデータに焦点を当てており、トレーニングデータデータ内の隠れたバイアスや問題がモデルの出力結果に大きく影響する点に着目している。
- AB 2013に則り、企業はデータのソースや種類、著作権で保護された情報や機密情報の有無など、トレーニングデータに関する広範な情報を開示する必要がある。

(参考) <https://www.bakerdonelson.com/californias-new-generative-ai-law-what-your-organization-needs-to-know>を基に作成

コロラド州(Colorado Artificial Intelligence Act)

- 2024年5月、コロラド州は米国初の包括的なAI法である Colorado AI Actを制定し、2026年2月に施行する予定。
- 同法では、**アルゴリズムによる差別が発生しうる「リスクの高いAIシステム」の検知や、予測可能なリスクから消費者を保護するために、AI開発者が十分な注意を払う義務を課している**。
- 「リスクの高いAIシステム」とは、利用者に対するサービス提供や拒否、またはサービスの費用や利用条件等に対して、重大な影響を及ぼすシステムを指す（以下）。
 - ✓ 教育や雇用に関するシステム
 - ✓ 金融または融資サービス
 - ✓ 行政サービス
 - ✓ ヘルスケア、ハウジングサービス、等
- **リスクの高いAIシステムの導入者は、同法に基づくシステムの影響評価をすることが義務付けられる**。影響評価には、差別が発生しうるリスクやAIの入出力データ、AIシステムに導入した透明性の措置の内容などを含める必要がある。

(参考) <https://www.whitecase.com/insight-alert/newly-passed-colorado-ai-act-will-impose-obligations-developers-and-deployers-high>を基に作成

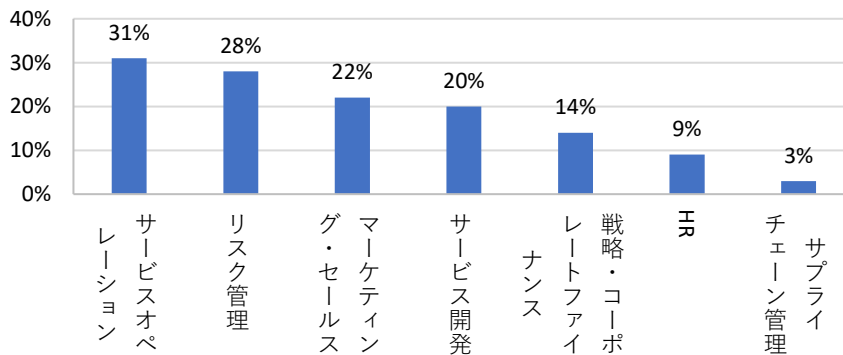


- 生成AIのユースケースは金融業界で急増しており、金融機関は広範な業務領域での生成AI活用を検討している
- 金融機関におけるAI導入の拡大と共にAIに対する米国金融当局の動きも活発化しており、イエレン元米財務長官はAIを金融システムにおける最重要課題と強調した

金融機関で活発化するAI活用

- **生成AIのユースケースは金融業界で急増している**
- Goldman Sachsは社内業務だけでなく、**投資や融資領域にも生成AIを活用する可能性を示唆**している。特に開発者の生産性を高め、品質、セキュリティ、管理の水準を高く維持しながら業務効率を高めることに注力する姿勢を示している。
- JPMorganも**400以上の生成AIのユースケースを実運用**しており、ソフトウェアエンジニアリングやカスタマーサービス、生産性向上など幅広い分野での生成AIの活用を検討している。

図表3：金融業界におけるAI活用領域（2023）



（参考）<https://www.constellationnr.com/blog-news/insights/financial-services-firms-see-genai-use-cases-leading-efficiency-boom#:~:text=Generative%20AI%20use%20cases%20are,data%20and%20controls%20in%20order.>を基に作成

【参考】AIに対する米国当局の動き

- 2024年6月、米金融安定監視評議会（以下、FSOC）にて、イエレン元米財務長官は**AIが金融システムの重大なリスクになり得る**と指摘した
- FSOCは2023年の年次報告書で、AIが資産運用や不正検知など金融領域での利用が拡大している点に触れ、**AIの判断を説明できない点などにリスクが含まれている**とコメント。
- イエレン氏はAIの複雑さや不透明性に加え、多くの市場参加者が同じデータやモデルに依存するリスクなどに言及した。また、犯罪が疑われる取引の発見などにAIが貢献している点にも触れ「**AI活用に伴う多大な機会と重大なリスクは米財務省とFSOCの最重要課題になっている**」と強調した。



（参考）<https://www.nikkei.com/article/DGXZQOGN0704U0X00C24A6000000/>を基に作成

- Wells Fargoは責任あるAIの実装のため、従業員のAI倫理教育やAIの透明性をチェックする体制等を整備
- Goldman Sachsは、AIの安全管理が施された環境で、アプリケーション開発ができる全社プラットフォームを構築

Wells Fargoの取組事例

- Wells Fargoは、スタンフォード大学とともにAI教育プログラムを創設し、**4,000人以上の従業員を対象にAI倫理の教育プログラムを実施している。**
- 同社は、AIの透明性や安全性を高めるため、AIが利用するデータやアルゴリズムを検証するための運用も実施している。
- 例えば、潜在的なバイアスを取り除くために、複数のAIモデルを比較して、AIのアルゴリズムのチェックと修正を行っている。また、**独立した開発チームがAIを使った製品や体験を複数の開発ポイントでレビューし、顧客に提供できる準備が整っているかどうかを確認している。**

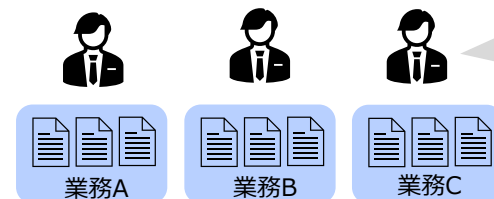


（参考）<https://stories.wf.com/how-wells-fargo-builds-responsible-artificial-intelligence/>を基に作成

Goldman Sachsの取組事例

- Goldman Sachsは、OpenAIのGPT-4やGoogleのGemini、MetaのLlamaなど**複数のLLMをユースケースに応じて安全に利用できるAIプラットフォームを活用中（図表4）。**
- 本プラットフォームでは、**安全性の管理を中央集権化（共通化）することで、規制に準拠したデータを使って開発者がモデルを調整することができる。**また、アクセス権限のない開発者にデータを提供しないような制御も組み込まれている。
- プラットフォーム上に制御機能を設けることで、**開発者が個別に安全対策を取る手間が削減され、アプリケーション開発に専念することができる。**結果として、開発期間が大幅に削減された。

図表4：Goldman Sachsの生成AIプラットフォームのイメージ図



開発者は、社内の管理が施された安全な環境で、用意されたLLMを基にアプリケーションを開発することができます。

生成AIプラットフォーム
(OpenAI、Google、MetaなどのLLMを安全に利用可能)

（参考）<https://www.wsj.com/articles/goldman-sachs-deploys-its-first-generative-ai-tool-across-the-firm-cd94369b>を基に作成

- ドイツ銀行は、生成AIを活用して規制文書を分析するアプリケーションをスタートアップと共同で開発
- 生成AIは金融業界の規制対応業務において、規制文書の要約やコンプライアンス状況の分析、リスク評価の自動化など、さまざまな場面での活用が期待されている

ドイツ銀行の取組事例

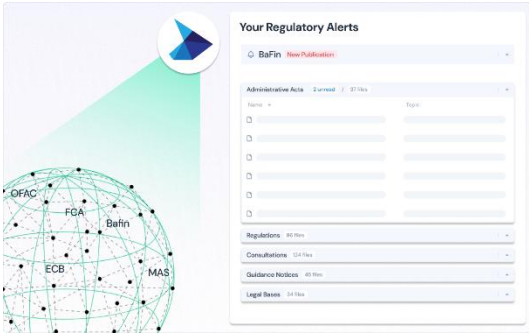
- ドイツ銀行は、**生成AIを活用して規制文書を分析するアプリケーションをKodex AI（右記）と共同で開発**した。
- ドイツ銀行の証券サービスチームは、同アプリケーションを使用して、**顧客の業界に関する規制文書や法的テキストを迅速に解析し、要約や影響分析等の情報を提供することで、クライアントの規制対応をサポート**している。
- また、ドイツ銀行とKodex AIは、共同でホワイトペーパーを作成し、生成AIが金融サービス業界においてどのように規制分析を強化できるかを検討している（以下、活用事例）。

生成AIの活用事例	内容
規制文書の自動要約	新たな規制が発行された際、主要な変更点や影響を特定する
コンプライアンスチェックリストの生成	規制要件に基づき、必要な手続きをリスト化する
リスク評価の自動化	規制違反のリスクがある領域を特定、評価する

（参考）<https://corporates.db.com/publications/White-papers-guides/adopting-generative-ai-in-banking>,
https://www.db.com/news/detail/20240531-best-custody-services-provider-in-emerging-markets?language_id=1を基に作成

Kodex AI（規制対応支援ソリューション）

- ドイツ（ベルリン）本社、2022年設立
- 金融機関向けに生成AIを活用したコンプライアンス管理ソリューションを提供
- 世界中の規制を継続的に監視するとともに、最新の規制と社内ポリシーを相互に参照して、コンプライアンスのギャップを特定し、迅速なAI規制対応を企業に促すプラットフォーム**
- 主なサービスには、自動規制スキャン（新しい規制の検出）、規制ギャップ分析（規制とコンプライアンス体制との比較）、リアルタイムレポート生成（コンプライアンス文書の作成）、AIチャットアシスタント（規制関連の問い合わせ対応）などがある。



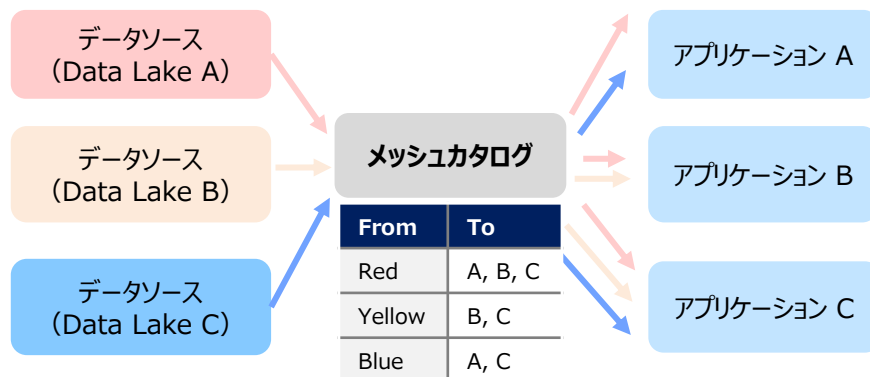
（参考）<https://www.kodex-ai.com/>を基に作成

- JPMorganは、AIを広範に活用するために、データを安全かつ全社的に利用できる基盤の整備を実施している
- データの利用にはデータ保護も併せて対策することが望ましく、データマスキングソリューションなどが注目されている

JPMorgan Chaseの取組事例

- AIを広く組織内で活用するため、**データを安全かつコンプライアンスに準拠した形で企業全体で共有できるように設計したデータメッシュアーキテクチャを構築した。**
- 本アーキテクチャでは、アプリケーションがデータを利用する場合、メッシュカタログを基にデータソース（Data Lake）から必要なデータを探し出して、アプリケーションにデータが連携される。
- これにより、**データをコピーすることなく、さまざまなアプリケーションでデータの利活用ができるようになり、迅速かつ安全なデータ管理を実現することができた（図表5）。**

（図表5）データメッシュアーキテクチャ



（参考）<https://www.jpmorgan.com/technology/technology-blog/evolution-of-data-mesh-architecture>を基に作成

【参考】データ保護対策、マスキング

- データ整備においては、**データの保護を併せて対策することも望ましいとされる。**近年、LLMに取り込むデータを自動でマスキングするデジタルソリューションも登場している。

データ保護ソリューションPrivateAI

- 従業員がChatGPTなどにインプットする個人情報を特定し、自動的にマスキングすることが可能。
- PDFや画像、音声に対しても適用可能。また、画像に含まれる個人識別情報（顔の輪郭など）を特定し、該当箇所だけを切り取るとなどのプライバシー保護対策も可能。
- GDPRやCPRAなどのプライバシー規制に準拠しており、大手金融機関でも導入が進んでいる。



（参考）<https://private-ai.com/>などを基に作成



7. 米国金融機関のAI施策（AI人材の強化）

- JPMorganは、AI専門人材の獲得と育成が銀行のAI競争力を高めるうえで決定的に重要と強調している
- BNY Mellonは、AI人材育成のプログラムに積極的に投資しており、AIの倫理的利用への対応力を強化している

JPMorgan Chaseの取組事例

- JPMorganは、2023年にテクノロジー投資に153億ドルを費やし、900名以上のデータサイエンティスト、600名以上の機械学習エンジニア、200名以上のAI研究者を抱え、AI開発に注力している。
- AI人材の増員は、顧客のデジタル体験を向上させるだけでなく、AIガバナンスの強化にも大きく寄与すると考えられる。

（図表6）AI人材の強化で期待されるガバナンス効果

項目	内容
専門知識の強化	AIリスクや倫理的課題を理解する専門家が増えることで、適切なガイドラインやポリシーの策定が可能になる
透明性と説明責任の向上	AI開発・運用プロセスを適切に監視・管理できる人材が増えることでAIの透明性や説明可能性を高めることができる
リスク管理の強化	AIモデルのバイアスやセキュリティリスクを早期発見し、適切な対策を講じる体制が構築できる
規制対応の迅速化	AI関連の法規制が進化するなか、専門人材がいることで企業が適切かつ迅速なコンプライアンス対応を実施できる

（参考）<https://www.constellationnr.com/blog-news/insights/jpmorgan-chase-digital-transformation-ai-and-data-strategy-sets-generative-ai>などを基に作成

Bank of New York Mellonの取組事例

- BNY Mellonは、金融サービスの提供において責任ある倫理的なAIの推進にコミットする点を強調している。
- 従業員に対してAIラーニングコンテンツ（約130時間）を提供しており、2,700名以上がAIトレーニングを受講している
- 全従業員がアクセス可能なAIプラットフォームを構築し、従業員のAIスキルの習得を促進しているほか、1,000名以上がAIハッカソンに取り組むなど、全社的なAI人材育成が進められている。
- また、MITやカーネギーメロン大学とのAI研究に関する情報なども発信しており、将来のAI人材へのアピールも行っている。
- ルールやシステム化によるガバナンス整備と併せて、個々の従業員がAIリテラシーを高めていくことが、AIの倫理的利用やAIリスクへの対応力の向上に大きく寄与すると考えられる。

（図表7）BNY MellonのAI人材育成への投資

育成コンテンツ	規模（時間、人数）
AIラーニングコンテンツ	130時間
AIトレーニング	2,700人以上
AIプラットフォーム利用	12,500人以上
AIハッカソン	1,100人

（参考）Evident社のレポートを基に作成

- 本ページでは、金融業界以外でAIガバナンス強化施策を積極的に推進する注目すべき2社の事例を紹介する
- 両社の事例に共通するのは、AIガバナンスを強化する上で社内外を含めた多角的な視点を持つ管理体制を構築しているという点である

マクドナルドの取組事例

- マクドナルドはアクセンチュアと戦略的パートナーシップを締結し、**世界中の店舗に生成AIソリューションを導入し、顧客体験の向上や従業員のデジタルスキルを強化**すると発表した。
- 本パートナーシップでは、マクドナルドがデジタルテクノロジーチームを組成し、アクセンチュアがマクドナルド従業員のAIやデータ管理等に関するスキルの育成をサポートしている。
- また、マクドナルドはAIを活用する上で、**経営やビジネスリーダー、人事、テクノロジー、法務など、全社的なクロスファンクショナルなチームを組成することが重要**だと強調している。
- 同社は、先進のAI技術の導入と管理体制の強化を推進することで、**技術革新とAI倫理的な運用を両立させるビジネスモデルの構築を目指している**。

クロスファンクショナルチームの組成による技術革新とAIの倫理的な運用の両立

経営

ビジネス
リーダー

人事

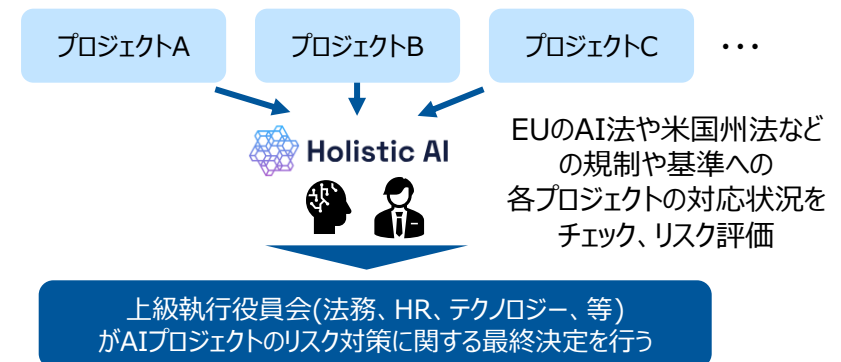
テクノロジー

法改正

広報

【参考】Unileverの取組事例

- UnileverはAIガバナンスの強化と規制遵守を目的に、スタートアップHolistic AI（P23詳述）との協業を進めている。
- 同社はEUのAI Actに対応するため、**Holistic AIのソリューションを活用してAIプロジェクトのリスクを評価し、リスクの軽減策を講じている**。また、**法務やHR、テクノロジー部門の代表者からなる上級執行役員会が同社の法規制対応について最終決定を行う体制を整備**している。
- 実際に、Unileverは300以上のAIプロジェクトにおけるバイアスや説明可能性、風評などの問題について、リスクを50%軽減することに成功している。



（参考）<https://newsroom.accenture.jp/jp/news/2024/release-20240116>, GTC 2024（2024/3/18-21）を基に作成

（参考）<https://www.unilever.com/news/news-search/2024/the-eu-ai-act-has-arrived-how-unilever-is-preparing/>を基に作成



9. グローバル企業のAI施策（AIセキュリティの強化）

- AIガバナンスの一環として、AIモデルへのサイバー攻撃や脆弱性を悪用した攻撃への対策も重要な要素となる
- AIやLLM固有の新たな脅威を認識するとともに、米国などで利用が拡大する先進のAIセキュリティソリューションの活用事例の調査やPoCを進めることもAIガバナンスの有効な施策の一つとなるだろう

AIに関する主な脅威

攻撃名	内容
敵対的攻撃	AIシステムを操作して不正な出力を生成したり、機密情報を抽出したりする攻撃。 【データポイズニング】 モデルを侵害するためにトレーニングデータを破損または変更する 【プロンプトインジェクション】 悪意のあるプロンプトをAIにインプットすることで、意図しない、または潜在的に有害な出力を生成するようにモデルをだます
AIを対象にしたサイバー攻撃	基盤モデルの開発やAIのアクセス制御など、AIを標的とする攻撃。 機械学習を侵害するように設計されたマルウェアやソフトウェアサプライチェーン攻撃、APIの脆弱性の悪用など。
AIに内在するリスク	機密データの漏洩や誤った情報の生成（ハルシネーション）、不適切または偏ったコンテンツの生成 など。

【参考】AIセキュリティソリューションの例

HiddenLayer

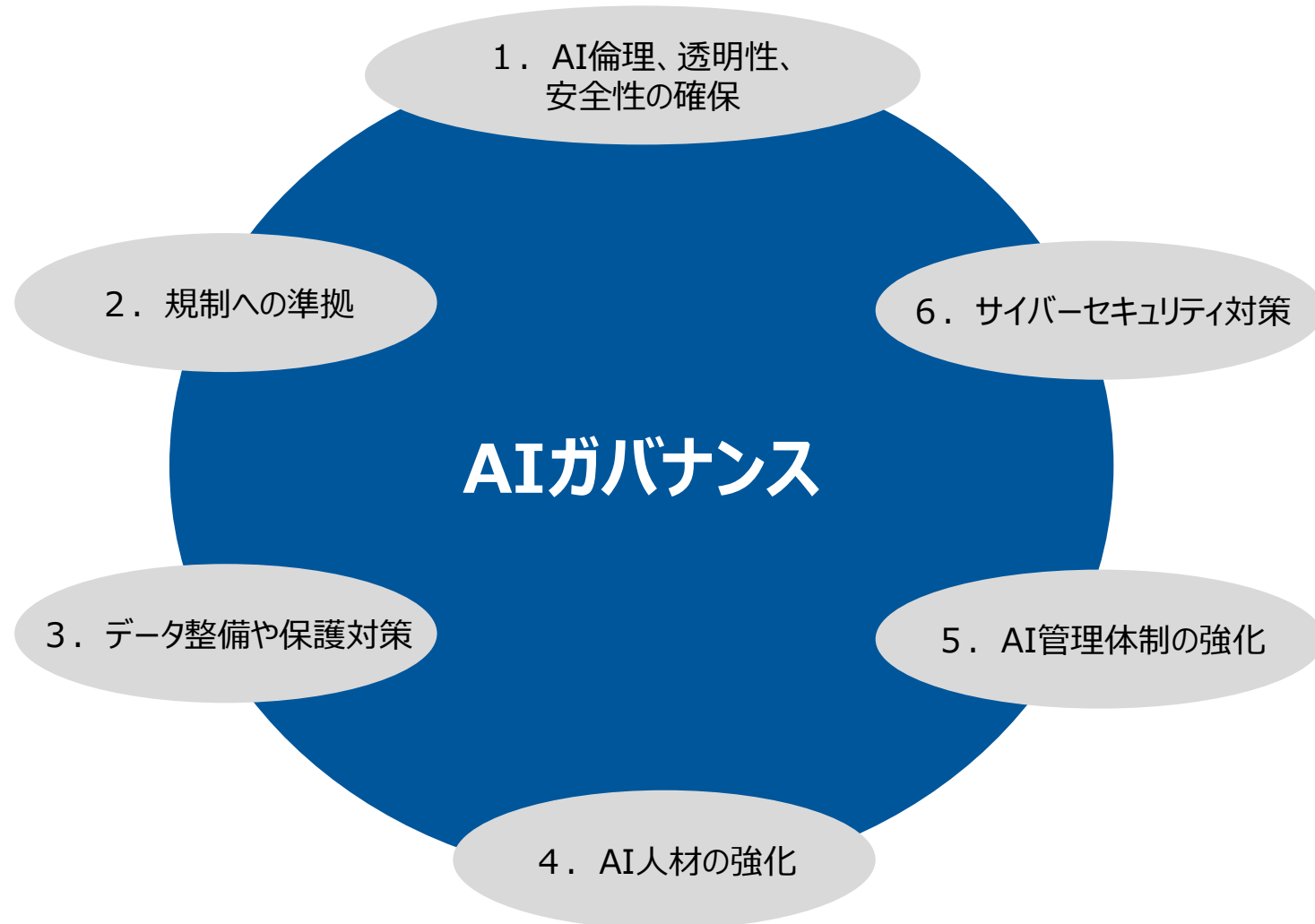
- <https://hiddenlayer.com/solutions/finance/>
- AIのアルゴリズムや入力・出力データを常時監視し、**AIモデルへの悪意のある攻撃を検出して防御**する。モデルの盗難や改ざん、データポイズニング、プロンプトインジェクションなど、敵対的攻撃への防御ソリューションを提供。
- グローバル金融企業などで導入実績あり

Lakera

- <https://www.lakera.ai/>
- **AIに対するプロンプトや応答を全て監視し、AIへの脅威が判明した際に対策を実行**する。個人情報の漏洩や偏見、セキュリティ脅威などの基準に基づいてAIの応答のスコアを算出し、しきい値を超えた場合に応答をブロックする。
- 危険なプロンプトを検知し、即時に対応するとともに、AIの応答に関するスコアの傾向を確認することが可能。
- Fortune 500企業などで導入実績あり。

（参考）<https://app.cbinsights.com/research/ai-security-trends-2024/>を基に作成

- 先進企業の事例（P10～15）を基に組織が取り組むべきAIガバナンス施策を整理すると6つの施策に大別できる
- AIの管理態勢を整備するためには、組織がこれらの施策を横断的に実施するための体制の構築や施策の立案が必要となる



- 近年グローバルにAIの導入が進み、AIがユーザーの生活や思想に与える影響も大きくなりつつある。AIを適切に扱うためには、法規制でのコントロールに加え、組織がAIの安全性や信頼性を確保するための多角的なガバナンス施策を進めることが肝要になるだろう。
- AIを積極的に活用する企業の先進事例から、AIガバナンスの主要施策を整理すると、6つ（1. AI倫理、説明可能性、安全性の確保、2. 規制への準拠、3. AI人材の強化、4. AI管理体制の強化、5. データ整備、6. サイバーセキュリティ対策。）に大別できる。AIガバナンスを整備するためには、特定の部署が単一の施策を行うのではなく、さまざまな施策を有機的に繋げて、全社的な活動として発展させていくことが求められる。
- AIの法規制に関するグローバルの動向にも継続的な注視が必要である。米国では、トランプ大統領がバイデン前政権のAIの安全性に関する新基準などの大統領令を撤回し、イノベーションを促進するための大統領令を新たに発表した。これによって、米国では、AI開発における安全性の担保が企業の自主性に一層委ねられることになった。
- 米国の連邦政府が規制緩和を進める一方、州法レベルでは新たな規制強化が進んでいる。また、欧州ではEUのAI規制法によって、グローバルに事業を展開する企業に厳格なリスク管理基準への適合が求められている。欧米のAI政策の方向性は大きく異なるものの、多くの企業にとっては、引き続きAIガバナンスの強化が最優先課題となることが予想される。



第4章

日本企業がAI活用を進めるために
一企業に求められるAIガバナンス態勢の整備と
企業をサポートする政策の立案を一

1. AIガバナンスの強化はグローバル企業にとっての優先事項

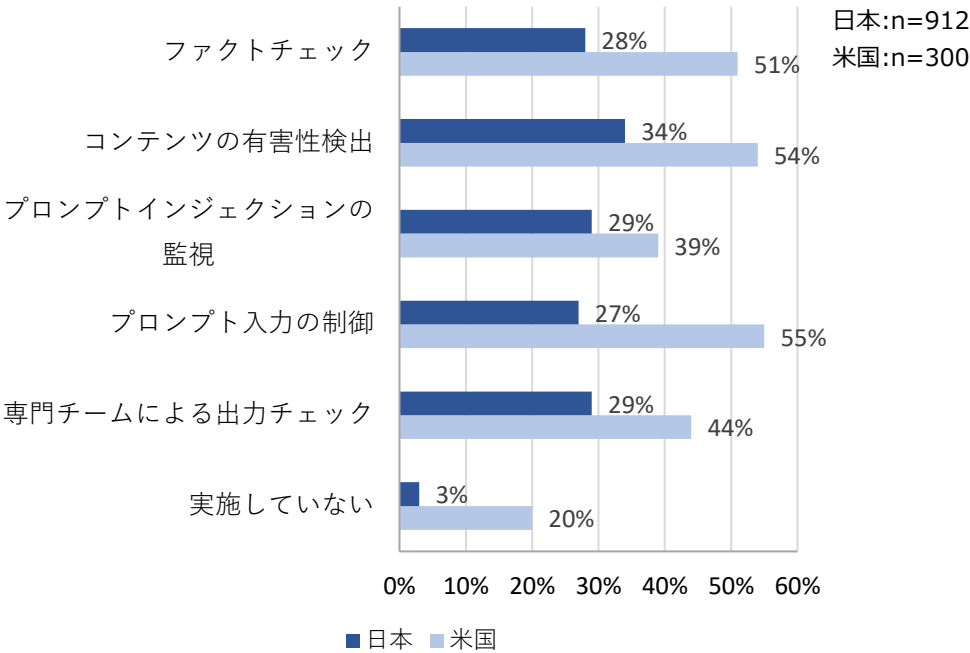
- 米国では連邦政府の規制の枠組みがなくなり、AI開発の安全性確保は企業の自主性に委ねられることに
- ただし、AIが社会に与える影響の大きさやEUのAI規制法への対応などを考慮すると、グローバル企業のAIガバナンス施策は緩められることはなく、むしろ強化される方向に進むことも予想される
- AI規制を取り巻く動向に注視するとともに、日本企業においてもAIリスク対策を整理していくことが求められるだろう

不安定な情勢でも求められるAIリスク管理

- ・前章で触れたとおり、米国ではバイデン前政権のAI規制に関する大統領令が撤回されて、**連邦政府としてのAIガバナンス規制の枠組みがなくなった**。これにより、**AI開発における安全性の担保が企業の自主性に一層委ねられることに**。
- ・一方、米国の州法レベルでは規制強化が進む兆しも見られる。また、グローバルの観点では、企業にはEUのAI規制法への対応も求められる。米連邦政府が規制緩和を進めても、**グローバルに事業を展開する企業は、より厳格な基準に合わせざるを得ない**。
- ・実際に米国企業は生成AIリスクへの対応策を既に進めており（図表8）、AIが社会に与える影響やEUのAI規制法への対応などを考慮すると、**グローバル企業のAIガバナンス施策はむしろ強化されることも予想される**。
- ・欧米のAI政策アプローチが大きく異なる状況のなか、AI実装を進める日本企業には、**AIを取り巻くグローバルの動向を継続的に注視するとともに、先進企業の事例を参考にAIリスクへの方針と対策を十分に整理していくことが求められる**だろう。

（図表8）日米における生成AIリスクへの対応策（2024年）

- ・米国企業ではガバナンス関連施策が優先事項として捉えられ、日本と比較すると生成AIへのリスク対策も進む傾向がある



（参考）<https://www.pwc.com/jp/ja/knowledge/thoughtleadership/generative-ai-survey2024-us-comparison.html>を基に作成



2. 先進事例を参考にAIガバナンス施策への着手・推進を

- 先行する欧米企業等の事例から、AIガバナンスを整備・強化する企業の主要施策として以下の6つがあげられる
- AI活用を進める企業には、AI管理責任や法規制への順守、AI管理組織の構築やサイバーセキュリティ対策など、部門横断的にAIを継続的かつ的確にコントロールしていく態勢の整備が求められる

(図表9) AIガバナンスの主要要素

#	項目	内容、取り組み例	参考ページ
1	AI倫理、透明性、安全性の確保	➤ 従業員のAI倫理教育やAIの説明可能性、透明性確保などのツール活用。 ➤ 安全性の管理を中央集権化（共通化）したAI開発環境の構築。	P10
2	規制への準拠	➤ デジタルツールなどの活用により、世界の法規制のアップデートをタイムリーに把握し、規制準拠対応漏れなどのリスクを最小限に抑える。	P11
3	データ整備やデータ保護対策	➤ 安全かつコンプライアンスに準拠した形で企業全体でデータを共有する設計。 ➤ プライバシー法規制等に準拠した形でデータを利用する対策。	P12
4	AI人材の強化	➤ AIに精通した人材を集め、組織的にAI管理を整備する。	P13
5	AI管理体制の強化	➤ 経営、ビジネス企画、人事、テクノロジー、法務等で構成されるクロスファンクショナルチームを組成。さまざまな観点でAIの出力や動作を監視・管理。	P14
6	サイバーセキュリティ対策	➤ AI固有の新たな脅威を認識するとともに、米国などで利用が拡大する先進のAIセキュリティソリューションの調査や、ソリューションのPoC・導入などを行う。	P15

* 本ページで挙げている項目はAIガバナンスの主要要素であるため、実際に施策を推進するためには、経営層・AIシステムの開発者・AIシステムを利用するユーザーなど、それぞれの観点で具体的な施策をブレイクダウンして策定するとともに、責任者を明確にすることが必要になる。また、全てのAIシステムを一元的に管理できることが理想ではあるが、システムのリスク（安全性や権利・公平性、司法、政治・選挙等に関する高リスクシステムか否か）や重要度、影響度に応じて優先順位をつけて管理・対策を整備していくことも必要になる。また、サードパーティや外部ベンダーによって、AIガバナンスに関する取り組みを第三者の目線で公平に評価してもらうことも有効な手段となる。



- 日本企業がAIガバナンス施策を強く推進するためには、政策によるサポートやインセンティブも欠かせない
- しかし、欧米と比較すると、日本には拘束力のある法律はなく、AI政策については慎重な検討が行われている
- また、日本ではリスク管理等に関する評価基準も不足しており、企業には自主的な取組みが求められる状況である

（図表 1 0）AIガバナンスに関する政策動向（欧米と日本の比較）

項目	EU	米国	日本
法規制の整備	2024年に「AI規制法（AI Act）」を可決し、リスクベースのアプローチを導入。高リスクAIシステムには厳格な規制を設ける。	バイデン前政権のAIの安全性に関する新基準などの大統領令が撤回される。トランプ政権がAIに対する規制緩和を指示する大統領令を発表。	AI事業者ガイドラインが策定され、AIシステムの開発におけるリスク管理やガバナンス体制の構築に関する指針が示されるものの、拘束力のある法律はまだない。
リスク管理やフレームワーク	AI Actでは、開発・運用プロセスにおけるリスク評価を義務付け	NIST（米国国立標準技術研究所）が「AIリスクマネジメントフレームワーク（AI RMF）」を発表し、リスク評価を推奨	AIのリスク評価や監査に関する明確な基準が不足しており、企業の自主的な取組みに依存
AI倫理や透明性の強制力	AI開発企業に対し、透明性確保や倫理的ガイドラインの遵守を義務付ける	連邦政府が規制緩和を示すものの、州法レベルではAIの透明性と公平性に関する規制強化が進行中	AI倫理に関する指針はあるものの、実際の監査や検証プロセスは具体化されていない
国際競争力と政策対応のスピード	生成AIの台頭に合わせて迅速に政策対応を実施	生成AIの台頭に合わせて迅速に政策対応を実施	AI政策は慎重な検討が続く

4. イノベーションとガバナンスの両立を実現する効果的なAI政策を

- AIを取り巻く国際情勢は不安定であり、日本企業がAIガバナンス施策を自主的に進めるのは困難な状況になりつつある。**日本のAI発展には、企業のAI施策を後押しする有効な政策の立案が併せて必要となろう。**
- AI政策の立案においては、AIが発展途上の技術である点を鑑み、初めから完璧な政策を作り上げるのではなく、技術の発展や企業の状況なども十分に検証しながら、法規制を柔軟に変更していくことも検討すべきであろう。

AI開発とガバナンス強化を促進するAI政策の立案を

- 欧米のAI政策アプローチは大きく異なり、AIを取り巻く国際情勢が刻々と変化する状況のなか、**日本企業にとって、適切なAIガバナンス施策を立案することが困難な状況になりつつある。**
- AI競争力を維持しながら、グローバルスタンダードにも対応できる柔軟な開発・運用体制の構築をするためには、**企業の判断に一任をするだけでなく、インセンティブが働く法規制やAIリスクを評価するための明確な基準などを示す必要がある**だろう。
- また、欧米の状況や先行事例を十分に検証し、企業がAI開発とガバナンス強化を両立して進めていくための有効な政策を立案することも併せて必要になるだろう。
- また、**政策の立案をゴールにするのではなく、AIを活用する企業の課題や状況を継続的に確認するとともに、政策が柔軟に変更できるような余地を考慮しておくことも、日本のAI産業の発展にとって重要な要素**となるだろう。

【参考】米国カリフォルニア州AI安全法案の拒否

- カリフォルニア州は米国のAI産業をリードしており、**同州の動向がAI業界へ与える影響が大きい**ため、**業界の成長を鑑み、モデル開発者が不利益を被る法案を拒否**した（以下）。
- **AI政策の変更に対する柔軟な姿勢***は、**日本の政策立案にとっても参考になる**だろう。
 - ✓ 米カリフォルニア州のギャビン・ニューサム知事が2024年9月「最先端AIシステムのための安全で安心な技術革新法（SB1047）」の法案を拒否した。
 - ✓ 本法案はモデルの開発者に安全性テストを義務付けるとともに、**AIによる壊滅的な事象が発生した場合、第三者が当該モデルをベースに作り上げた別のAI製品であっても、当該モデルの開発者が全責任を負う必要があった。**
 - ✓ 本法案は発展途上の技術に対して過度な負担を開発者に課すとして、AI企業だけでなく、オープンモデルを利用するスタートアップやVC、学術関係者も多くの反対を表明していた。

（参考）<https://www.jetro.go.jp/biznews/2024/10/73761c6c009dcecf.html>を基に作成
* SB 1047が否決された一方、同時期に提案されていたAIシステムの透明性を企業に求めるAB 2013は成立している（P8参照）。同州がイノベーションとガバナンスを両立させるための政策を模索する様子が見られる。

- 近年は、規制への準拠やAIの安全性の確認、品質向上などさまざまなAI関連ソリューションが登場している
- もっとも、P16やP20で挙げたようなAIガバナンスの主要施策を網羅的にカバーするデジタルソリューションはないため、企業には戦略やポリシー、優先度などに応じて適切にソリューションを選択していくことが求められる
- 本編で紹介したように、欧米では大手企業がスタートアップと協業してAIガバナンスを強化するケースも多く見られる

Holistic AI



- [Holistic AI - AI Governance Platform](#)
- 本社ロンドン、2020年設立
- 企業が利用している**AIのガバナンス管理やリスク管理、バイアス評価、法令順守等をサポートするプラットフォーム**を提供。
- **EUのAI法やニューヨーク市のバイアス監査、米国州法、NISTのAIリスク管理フレームワーク、ISO／IEC 42001などグローバルの主要な規制や基準への対応をサポート。**
- 全世界のAI規制やインシデントをまとめたデータベースを活用し、AIリスクに精通したスペシャリストが、AIシステムのリスク評価やリスク低減策の提案、監視等を行う。人材採用や銀行業務、消費者アプリなど幅広い分野で導入実績あり。
- UnileverはAI監査業務で同ソリューションを活用。100件以上のAIプロジェクトのリスクマネジメントを実施し、3割のプロジェクトにおいて、リスクが高レベルから低レベルまで低減された。

Acuvity



- <https://acuvity.ai/>
- 本社カリフォルニア州、2023年設立
- ユーザーがAIのアプリケーションやチャットボット、サービスなどを安全に利用できるソリューションを提供
- **機能は主に、①組織内で生成AIを利用するアプリケーションの検知 ②AI利用状況の可視化 ③AI活用のポリシー設定 ④防御 の4つ**
- ②AI利用状況の可視化では、どのような種類のデータ (e.g. 個人情報, クレデンシャル情報) が誰によって、どの外部サービス (e.g. Chatgpt, GitHub) に送られているのかを把握することなどが可能。
- ③AI活用のポリシー設定では、「個人情報が含まれている画像やテキストを生成AIサービスはクラウドにアップロードしない」などの制御を行うことが可能。

- 本編で触れたように、企業がAIガバナンスの整備を進めるためにはさまざまな施策を多角的に行うことが必要になる
- 技術の進展が著しいAIを的確かつ継続的にコントロールしていくためには、マニュアルを前提にした（人手・担当者の判断が多い）運用を立案するのではなく、先進のデジタルソリューションを効果的に活用していくことも肝要となろう

Shelf



- <https://shelf.io/>
- 本社コネティカット州、2015年設立
- 生成AIを活用する際、**信頼性の高い回答を得るために、企業内のさまざまな情報源やデータベースを統合的に管理して高品質なデータ管理をサポートする**ソリューション。
- データ品質の向上を向上させることにより、企業が効果的に生成AIを活用し、AIシステムのパフォーマンスを最適化させることが目的。
- SharePointやSalesforce、boxなど企業内のさまざまなデータソースと統合することが可能。組織内の情報のサイロ化を防ぎ、AIが必要なデータへ的確かつ容易にアクセスすることを可能にする。
- また、**AIが利用したファイルを特定したり、ファイルに不正確な情報や矛盾がないかをチェックし、生成AIが誤った回答を生成するリスクを低減する。**

Reality Defender



- <https://www.realitydefender.com/>
- 本社ニューヨーク州、2018年設立
- **AIが生成した偽のコンテンツ、ディープフェイクから企業や政府機関を守るためのソリューションを提供。**
- **音声、動画、画像、テキストなど多様なメディア形式のAI生成コンテンツをリアルタイムで検出することが可能。**コールセンターの音声詐欺やメディアコンテンツの信頼性確認など、さまざまな分野でのリスク軽減に寄与。
- 2024年5月にRSAカンファレンスのInnovation Sandboxコンテストで最も革新的なスタートアップとして表彰された。
- 2024年10月にはアクセンチュアが同社への戦略的投資を発表した。この提携により、金融、メディア、ハイテク分野の顧客に対するディープフェイク詐欺への防御が強化されると期待されている。

- AIが社会に与える影響の大きさやEUのAI規制法への対応などを考慮すると、グローバル企業のAIガバナンス施策は今後強化される方向に進むと予想される。AI規制を取り巻く動向に注視するとともに、日本企業においてもAIリスク対策を整理していくことが求められるだろう。
- AI活用を積極的に進める企業の先進事例からは、企業のAIガバナンスにとって、透明性の確保や法規制への順守、AI管理組織の構築など、部門横断的にさまざまな観点でAIを継続的かつ的確にコントロールする態勢整備の重要性が示唆される。
- 日本企業がAI競争力を強化しつつ、グローバルスタンダードにも対応できるAIガバナンス体制を構築するためには、企業の自主的な判断に一任するだけでなく、企業のインセンティブが働く法規制や、AIリスクを評価するための明確な基準を示すことも併せて必要となる。
- 欧米の法規制や先進企業の取り組みなどを参考に、日本国内で応用が見込める事例や課題を多角的に検証し、企業がAI開発とガバナンス整備を両立するための有効な政策を立案すべきであろう。
- もっとも、政策の立案をゴールにするのではなく、AI技術が発展途上である点を十分に鑑み、AIを活用する企業の課題や状況を継続的に確認するとともに、柔軟に変更できるような制度設計を考慮しておくことも、日本のAI産業の発展において重要な要素となろう。



総括



- AIの技術革新や対話型の生成AIの台頭によって、AIの社会実装がグローバルに加速しているが、米国などのデジタル先進国の状況と比較すると、国内でのAI活用が大きく進展しているとは言い難い。
- 日本企業がAI導入を進める際の障壁として挙げられるのは、サイバーセキュリティや、AIを利用する際のリスク管理、プライバシー保護等に関する一連のガバナンス施策の立案と推進である。AIは多様な場面での活用が期待される反面、安全性や信頼性、法規制への準拠状況など、AIが世の中に与える影響を多角的に検証することも併せて求められる。
- AI活用を積極的に進める企業の先進事例からは、AIガバナンスの強化において、特定の部署が単一の施策を行うのではなく、組織が部門横断的に戦略を立案し、透明性の確保や法規制への順守、AI管理組織の構築など、さまざまな施策を有機的に繋げて、全社的な活動として発展させていく重要性が示唆されている。
- 日本企業がAI競争力を維持しながら、グローバルスタンダードにも対応できる柔軟な管理・運用体制を構築するためには、欧米のAI政策の動向やAIを活用する企業の取り組みを継続的に検証するとともに、企業のAI施策をサポートする法規制や、AIリスクを評価するための明確な基準を整備・立案していくことも必要になるだろう。