

経済情報： AIの経済・金融への影響と 各国の規制や支援の枠組み

2024年9月

三菱UFJ銀行 経営企画部 経済調査室

【目次】

1. AI発展の経緯	3
2. マクロ経済への影響	4
3-1. 金融セクターへの影響（AI活用の機会や課題）	5
3-2. 金融セクターへの影響（生成AIの活用）	6
3-3. 金融セクターへの影響（生成AIのリスクや課題）	7
4-1. AIに関する規制動向	8
4-2. OECD AI原則の概要	9
4-3. 広島プロセス国際指針概要	10
4-4. EU	11
4-5. 米国・日本・中国	12
5. インプリケーション	13
■主要参考資料	

1. AI発展の経緯 - AI、生成AI、AGIへ -

- AIとは人間のような知能を必要とするタスクを実行するコンピュータシステムを指す幅広い用語。
- AIのルーツは1950年代後半まで遡ることができるが、1990年代の機械学習（ML）分野の進歩が現在の生成AIの礎を築いた。（機械学習とはデータからパターンを検出し、それを予測や意思決定に役立てるために設計された技術の総称）
- 2010年代には深層学習（DL）の開発とNVIDIAのGPU処理能力の向上が大きな飛躍の契機となった。深層学習はニューラルネットワーク^(注)を使用し、顔認証や音声アシスタントなどの日常的なアプリケーションを支えている。
- さらに2020年に入り大規模言語モデル（LLM）の進歩により、自然言語を用いたプロンプトによりテキスト、画像、音楽などを生成できる生成AIが登場した。
- 今後は、人間に近い思考回路や感情を持つAGI（汎用人工知能/Artificial General Intelligence）の開発も期待されている。

（注）ニューラルネットワーク：人間の脳の働きを模倣したコンピュータシステム。多層構造で構成され、層を深くすることで、より複雑な問題を解決できるようになった。

AI（人工知能）の発展の経緯

1990年代	2000年代	2010年代	2020年以降
● 機械学習の台頭 機械学習のアルゴリズムが発展し画像認証、音声認識、自然言語処理などの分野で成果を上げ始める。 • 1993年 Microsoftは開発者が音声認識機能をアプリケーションに組み込むためのAPIを発表し、音声コマンドや音声入力をサポート。 • 1997年 IBMが開発したチェス専用のスーパーコンピューター（ディープブルー）はチェスの世界チャンピオンを撃破。 • 1990年代後半にはメールサービスプロバイダーによるスパムメールフィルタリング技術や、Amazonの商品推奨サービスが開始された。	● インターネットの普及 スマートフォン、ソーシャルネットワーキングサービス（SNS）が登場し、大量のデータ収集・分析が容易になり、機械学習の発展が加速。 • 2007年初代iPhoneが発売されスマートフォンが急速に普及。 • ウェブページの重要度を決定するGoogleのPageRankアルゴリズムにより検索エンジンの精度が飛躍的に向上。 • TwitterやFacebookやといったSNSが相次いで登場し、データ収集と分析が大規模に行われるようになる。 • デジタルカメラの顔認証技術、自動車の衝突回避システム、音楽ストリーミングサービス等の提供開始。	● 深層学習の登場 画像認識、音声認識、自然言語処理などの分野で機械学習を超える成果を上げ始める。 • 音声認識技術の進化により自然言語理解の精度が大幅に向上し、AppleのSiri、Googleアシスタント、Amazon Alexaといったサービスが続々と登場した。 • Google翻訳はニューラル機械翻訳を導入し、翻訳の精度が飛躍的に向上した。 • スマートフォンの顔認証や写真補正機能の搭載。 • 車両の自動運転技術の進展。 • 医療診断システムの進化。	● 大規模言語モデルの進歩 言語理解、生成、翻訳などの能力が飛躍的に向上。 • 2022年 GPT-3の登場に続き、各社から様々なサービスが提供されている。高品質で自然な文章の作成のほか、音声での応答が可能となり、チャットボットやバーチャルアシスタント等のビジネスでの利用のほか、教育、エンターテインメント等の様々な用途で活用されている。 • 今後、人間のように汎用的な知能を持ち、複雑な概念や感情をも理解し、より幅広いタスクに対応できるAIGの開発も期待されている。

（資料）各種資料より国際通貨研究所作成

2. マクロ経済への影響 - 国際通貨基金(IMF)、国際決済銀行(BIS) -

- 国際通貨基金(IMF)は、世界の雇用の40%がAIの高い影響に晒されており、特に先進国ではその割合は60%に達している。また、IMFはAIの採用が世界の中期的成長率を+0.1%~+0.8%高める可能性を指摘(2024年4月世界経済見通し)。生産性向上は途上国より先進国においてより早期に顕現化するとの見方。
- IMF、国際決済銀行(BIS)とも、AIが経済や社会に及ぼす影響は予見しがたいとの前提ではあるが、主に高い認知能力を必要とする職業を中心に生産性を高め、企業の投資、家計の消費にも好影響を与える可能性があるとしている。
- 一方、AIの導入により、経済、社会の不平等が発生するといったマイナスの側面についても言及している。

	国際通貨基金(IMF)	国際決済銀行(BIS)
影響度	- 世界の雇用の約40%がAIの高い影響に晒されている。 - 特に先進国では、その割合は60%に達する。	—
影響を受ける対象	先進国、認知集約型の職業、高学歴労働者 - AIの普及は先進国で進んでおり、特に認知集約型の職業においてその影響が顕著。 - 大学教育を受けた若手労働者がAIの恩恵を受けやすい一方で、高齢労働者は新技術に適応しにくい傾向がある。	高い認知機能を必要とする職業 - 労働者や企業の生産性を高める。 - イノベーションを促進し、将来の生産性を間接的に高める。 - イノベーションの大半は高い認知能力を必要とする職業で生まれる。
プラスの側面	生産性の向上と労働者の所得水準の上昇 - AIは長期的に経済を変革し、特に先進国でその効果が早期に現れると予測されている。 - AIはルーチン作業から創造的な仕事まで幅広い活動に適用でき、生産性を向上させる可能性がある。 - 生産性の向上が十分に大きければ、多くの労働者の所得水準が上昇する可能性がある。	生産性の向上、総供給の増加、総需要の増加 - AIは経済の生産能力を拡大し、総供給を増加させる。 - 生産性の向上は企業の投資の変化を通じて総需要にも影響を与える。 - また、①消費者の探索能力の向上や企業のマーケティング能力の向上による消費需要の喚起、②労働需要を高め賃金の上昇による消費需要の拡大を通じて、家計消費にも影響を与える。
マイナスの側面	労働所得、資本所得の不平等 - AIと高所得労働者の補完性が強い場合、労働所得の不平等が増加する可能性がある。 - AI資本への投資による資本所得の増加により、富の不平等が拡大するリスクがある。	労働者の失業、経済的不平等の拡大 - AIの効果は仕事や職種により異なり、AIの導入により一部の労働者は失業し、賃金の低下、消費需要の低下をもたらす可能性がある。 - AIの恩恵を受けられる人とそうでない人との間の経済的不平等に影響を与える可能性がある。

(出所) IMF, “[Artificial Intelligence and the Future of Work](#)” 14 Jan 2024 ,BIS, “[Annual Economic Report 2024 Chapter III](#)” 24 Jun 2024, BIS,

3-1. 金融セクターへの影響（AI活用の機会や課題） - 国際決済銀行（BIS） -

- BISによれば、金融セクターは、多様な情報を理解し、処理し、判断をくだす認知的タスクの割合が高く、データ集約的な性質を持つことから、AIの影響を大きく受けるとされる産業分野のひとつである。
- 構造化されていないデータを巨大な計算能力を活用し構造化することに優れているAIの機能を活用することにより、流動性管理、KYC、クレジットリスク分析、リスク評価、アルゴリズム取引といった業務の効率を高め、コスト削減につながる事が期待されている。
- 一方、AIの活用が進むと、サイバー攻撃への懸念の高まり、膨大なデータを取り扱うことによるプライバシーや機密性確保の問題、アルゴリズムによる差別や偏見の問題、群衆行動や市場のボラティリティ増幅など、金融セクターに新たな課題をもたらす可能性も指摘されている。

金融セクターにおけるAIの機会、課題

	決済	貸出	保険	アセットマネジメント
金融セクター特有の機会	- 流動性管理：過去の取引や市場情報等の膨大なデータを活用した需要予測精度の向上。 - KYC、AML/CFTプロセスの改善：迅速なデータ処理と不正行為検出能力の向上。	- クレジットリスク分析：銀行口座取引も含めた非構造化データの活用。 - 金融包摂：クレジットスコアが低くとも質の高い借手を検出する能力を高める。	リスク評価、プライシング、請求処理：画像や映像を自動的に解析し、自然災害による物的損害を評価したり、保険請求が実際の損害と一致しているかどうかを判断。	ポートフォリオの配分、アルゴリズム取引、ロボ・アドバイザー：伝統的な構造化データにはない情報も含めた大量のデータを活用し、より迅速に正確な情報を投資家に提供。
金融セクター特有の課題	- 流動性危機：アルゴリズムの協調による予期せぬ価格変動や、データの偏りによる不適切な流動性管理。 - 巧妙な詐欺やサイバー攻撃：フィッシングメール、マルウェアの生成。	- アルゴリズムによる差別：AIに学習させたデータに偏りや不正確さがあると、一部のグループをサービスの対象から除外するリスク。 - プライバシーへの懸念：増大するデータを扱う際のデータプライバシーと機密性の確保。		- 個人の利益のためのゼロサム競争：市場参加者は技術やインフラへの多額の投資を行うが、市場全体の利益を向上させるわけではなくゼロサムゲームに陥る。 - 群衆行動、アルゴリズム協調
金融安定化に関わる課題	- 群衆行動：投資家は他の投資家の行動を模倣する傾向があり、市場動向に対する過剰反応、バブルの形成、急激な価格変動等を引き起こす可能性。 - ネットワークの相互接続性とプロシクリシティ（景気循環増幅効果）：金融機関や市場が相互に結びついているため、一つの金融機関の問題が他の金融機関に波及しやすく、景気変動に対して金融システムが過剰に反応するリスク。 - 単一障害点：システム全体が一つの要素に依存している場合、障害が発生するとその影響がシステム全体に及ぶリスク。 - 代表性のないデータの短いサンプルに基づく誤った決定：誤ったデータに基づいて意思決定を行うと、誤った判断が下され、金融市場の安定性が損なわれるリスク。 - 実体経済からの波及効果：金融機関の健全性への悪影響が、金融市場全体に波及するリスク。			

（出所）BIS, “[Annual Economic Report 2024 Chapter III](#)” 30 Jun 2024

3-2. 金融セクターへの影響（生成AIの活用） - 経済協力開発機構（OECD） -

- 従来のAIは主にパターン識別、分類、予測のために利用されてきたが、生成AIは人間が生成したものと見分けがつかないコンテンツを作成することができるため、金融セクターのバックオフィス（オペレーション）やミドルオフィス（コンプライアンスやリスク管理）の自動化を中心に、業務の効率化や生産性向上の手段として利用されている。
- 一方、生成AIには様々なリスクが指摘されており（詳細次項）、金融の安定性確保、利用者保護、リスク管理等の観点から、顧客と直接やり取りするフロント業務やアプリケーション開発への活用は途上である。

金融セクターにおける生成AIのユースケース

生産性向上／ミドル・バック業務		価値創造／フロント業務	
法令順守のために必要な報告書等の生成 企業内のデータセットに基づいた報告書等の作成プロセス自動化、効率化。	融資業務 借り手の信用力を評価するための情報管理やデータ処理の効率化、信用履歴の乏しい借り手への金融包摂の促進。	新商品開発 市場のトレンドや顧客ニーズの分析による革新的な商品の開発。	
AML/CFTプロセスと不正検知 与えられたデータセットの中から、通常の行動に適合しない異常値を自動的に特定し、AML/CFTプロセスと不正検知の両方を改善。	合成データの生成 シナリオ分析のための金融市場のシミュレーションデータや、AIのモデルのテストのためのデータ作成。	ターゲットマーケティング 特定のターゲットグループに最適なマーケティング戦略を立案、マーケティング効果の最大化、顧客エンゲージメントの向上	
資産運用会社や機関投資家のリスク管理支援・強化 非構造化データの利用によるセンチメント分析の強化やパターン認識による取引後の追加的な洞察の効率化。	コーディング ソフトウェア等の開発において、新しいコードの生成、スクリプトのバグの解消、コーディングエラーの解決策の提供。	カスタマーオンボーディング 顧客のオンボーディングプロセスを自動化、効率化により、顧客のスムーズなサービス利用と、企業側のコスト削減。	
ESGデータの処理や分析 構造化されていないESG関連データの処理や分析。	その他 人事管理（人事評価の要約や生成）、翻訳（契約書や報告書などの翻訳や要約）。	カスタマーサポート チャットボットやバーチャルアシスタント等を通じた顧客サポート機能の提供による、顧客対応の迅速化と顧客満足度の向上。	
		アセットアロケーション 市場データのほか非構造化データのリアルタイム分析により、投資ポートフォリオの最適化。	
		トレーディング戦略 アルゴリズムの進化によるオーダー金額や期間の最適化、流動性管理の強化、他のトレーダーの行動予測能力の強化。	

（出所）OECD, “[Generative artificial intelligence in Finance](#)” 15 Dec 2023

3-3. 金融セクターへの影響（生成AIのリスクや課題） - 経済協力開発機構（OECD） -

- 生成AIは大量のデータ学習を前提に構築されていることから、学習するデータの適切性、プライバシーや知的財産権保護に関わる事項に加え、同一のAIモデルを使用することによる金融安定化へ及ぼす影響や、エネルギー消費の問題まで、様々なリスクや課題が指摘されている。

生成AI活用のリスクや課題	
公平性	アルゴリズムによる意思決定における偏りや区別の可能性は機械学習モデルの初期からの確立された概念だが、AI駆動モデルも、質の低い不十分なデータセットを使用することで、偏った内容や公平でない結果を生成する可能性がある。
説明能力の欠如	生成AIモデルは複雑で不透明、何千億ものパラメータに依存しているため、生成結果の説明が難しい可能性がある。
データ関連リスク	生成AIは、非常に大規模なデータセットをベースに構築されており、そのデータの品質、適合性、信頼性の点で欠点がある場合、健全なデータガバナンスの課題を増加させる可能性がある。また、生成AIが学習する大量の非構造化データの中にはプライバシーに関わるデータや知的財産権で保護されたデータを含んでる可能性が高い。
堅牢性、信頼性、市場操作リスク	生成AIモデルの出力の信頼性や精度が低いと投資アドバイス等のユースケースにおいてリスクが高まる。 悪意ある行為者による、株式やその他投資に関わる虚偽の情報を提供することで、生成AIの能力を利用して投資家や消費者に欺瞞的なアドバイスを提供させるリスクも生じる。
ガバナンス	組織に生成AIの基盤モデルを構築、拡張、統合するための適切なスキルや知識が不足していたり、効果的な監督を行うためのガバナンスが十分でない可能性がある。また、生成AIに関連するサービスやインフラを社外の第三者に依存する場合には、更なるガバナンスの問題が生じる。
金融の安定性 競争関連リスク	多数の金融関係者が同じAIモデルを採用することで、群集心理や一方通行の市場が発生する可能性があり、その結果、特にストレス時にシステムの流動性と安定性に対するリスクが高まる可能性がある。 ビジネス参入障壁から、生成AIの提供が少数の先行するプロバイダーに限定され、競争原理が働かないリスク。
雇用や環境への影響	ミドル、バックオフィスの多種多様な業務を自動化する可能性を秘めており、将来的に雇用上の課題が生じる可能性がある。また、生成AIのトレーニングには大量のデータ学習を必要とし、同時に大量のエネルギーを消費する。

（出所）OECD, “[Generative artificial intelligence in Finance](#)” 15 Dec 2023

4-1. AIに関する規制動向 - 技術の急速な進歩に対しEU以外の法整備は途上 -

■ AIに関する国際的取り組み

- OECDは2019年5月「AI原則」を公表。AIに関する最初の国際的な原則とされ、信頼できるAIの責任あるスチュワードシップのための5つの原則と、国際政策と国際協調のための5つの原則を定めている。
- G7は2023年10月に「先進的AIシステムを開発する組織のための広島プロセス国際指針、行動規範」を公表。OECDのAI原則を踏まえ、国際的なルール作りの議論を促進するため11項目の指針について首脳会議で合意。

■ AIに関する各国の動向

- ユーロ圏では、リスクベースアプローチによるAIに関する包括的な法規制を整備し、民間のイノベーションにも配慮。
- 米国、日本、中国は、包括的な法整備には至らないものの、ガイドラインや関連法令を示しつつ、民間のイノベーションを重視。

AIに関する規制動向							
	2019	2020	2021	2022	2023	2024	2025
OECD G7	OECD AI Principles (AI原則) 2019年5月				G7 広島AIプロセス 国際指針、行動規範 2023年12月		
	AI白書 欧州データ戦略 2020年2月		EU AI法案 2021年4 月	AI法 2024年5月			
	大統領令 - AI分野における米国の主導権維持 2019年2月			AI権利章典の青写真(ホワイトハウス)- AI設計、使用、開発の原則 2022年10月 安全・安心・信頼できるAIに関する大統領令 2023年10月			
	AIガイドライン(2017年/総務省) AI利活用ガイドライン(2019年/総務省) AI原則実践のためのガバナンスガイドライン(2021年/経産省)			(統合) AI事業者ガイドライン(経産省) 2024年4月			
	新一代人工知能ガバナンス準則 2019年6月		データセキュリティ法 個人情報保護法 2021年	インターネットコメントサービス管理規定 2022年 生成人工知能サービス管理暫行弁法 2023年7 月			

(資料) 各種資料より国際通貨研究所作成

4-2. OECD AI原則の概要

- OECDのAI原則(AI Principles)は、AIの開発と利用に関する最初の国際的なガイドラインであり、AIの信頼性を確保し、政策立案者に効果的なAI政策の推進を支援するためのガイドラインとして、加盟国を中心に47の国と地域^(注)がこの原則に賛同している。

(注)47の内訳は、OECD加盟38カ国、非加盟8カ国、EU。中国、ロシアは入っていない。

OECD AI Principles	
信頼できるAIのステュワードシップのための原則	国際政策と国際協力のための原則
1. 包摂的な成長、持続可能な開発および幸福 ステークホルダーは人々と地球にとって有益な結果を追求することにより、信頼できるAIの管理責任を積極的に果たすべきである。	1. AIの研究開発への投資 信頼できるAIのイノベーションを促進するため、研究開発やオープンサイエンスの分野で各国政府は長期的な公共投資を検討し、かつ民間投資を推奨すべきである。
2. 法の支配、人権並びに公平性及びプライバシーを含む民主主義的価値の尊重 AIアクターは、AIシステムのライフサイクルの全体を通じて、法の支配、人権並びに民主主義的及び人間中心の価値観を尊重すべきである。	2. 包括的なAIを推進するためのエコシステムの整備 各国政府は信頼できるAIのための包括出来でダイナミック、持続可能で相互運用可能なデジタル・エコシステムの開発及びそれへのアクセスを促進すべきである。
3. 透明性及び説明可能性 AIアクターはAIシステムに関する透明性及び責任ある開示に積極的に関与すべきである。これを達成するためにAIアクターは状況に適した形で、かつ技術の水準を踏まえ、意味ある情報提供を行うべきである。	3. AIを推進するための相互運用可能なガバナンス及び政策環境の形成 各国政府は信頼できるAIシステムの研究・開発段階から、展開・稼働の段階への移行を支援するための機動的な政策環境整備を促進すべきである。また、イノベーションと競争を奨励するため、AIシステムに適用される政策や規制の枠組みとその評価メカニズムについて見直し、かつ状況に順応させるべきである。
4. 頑健性、セキュリティ及び安全性 AIシステムはそのライフサイクル全体にわたって頑健で、セキュア、かつ安全であるべきである。	4. 人材育成及び労働市場の変化への備え 各国政府は仕事の世界と社会全体の変化に備えるために、ステークホルダーと緊密に協議すべきである。また、AIの普及がもたらす労働市場の変化が労働者にとって公平なものであるよう、様々な措置を講じるべきである。
5. 説明責任 AIアクターは、その役割と状況に基づき、かつ技術の水準を踏まえた形で、AIシステムが適正に機能していること及び上記の原則を尊重していることについて、説明責任を果たすべきである。	5. 信頼できるAIのための国際協力 各国政府は本原則の推進、および信頼できるAIのステュワードシップの進展のために積極的に協力すべきである。AIナレッジの共有を促進するための協働や国際的な技術水準の開発等を推進すべきである。

(注) AIシステム：実環境又は仮想環境に影響を及ぼす、予測、コンテンツ、推奨、又は決定などの出力をどのように生成するか、受け取った入力内容から推論する機械ベースのシステム。

AIアクター：AIシステムのライフサイクルにおいて能動的な役割を果たす者であり、AIの展開又は稼働を行う組織や個人が含まれる。

AIシステムのライフサイクル：通常、計画と設計、データの収集と処理、評価、検証、稼働とモニタリング、稼働停止と終了など、いくつかの段階から構成される。

ステークホルダー：直接的なものであるか、間接的なものであるかを問わず、AIシステムに関与するか、又はAIシステムから影響を受ける組織及び個人のすべてが含まれる。

(出所) OECD, “[AI Principles](#)”, 22 May 2019 (updated 3 May 2024)

4-3. 広島プロセス国際指針概要

- G7各国はOECDのAI原則を基に、AIシステムの倫理的な開発と利用を促進するための議論を開始し、2023年12月広島サミットにおいて「広島プロセス国際指針、行動規範」を発表した。
- 広島プロセス国際指針は「安全、安心、信頼できるAIを世界に普及させることを目的とし、最先端の基盤モデル及び生成AIシステムを含む。最も高度なAIシステムを開発・利用する組織のための指針」であり、「AI技術によってもたらされる利益を捉え、リスクと課題に対処することを補助することを意図している」としている。
- 2024年6月に開催されたプーリア（イタリア）サミットにおいても、広島プロセス国際指針、行動規範の実践や、G7以外の国・地域とも協力しながら、AIガバナンスに関する取り組みを進める重要性が共有された。

広島プロセス国際指針

1. AIライフサイクル全体にわたるリスクを特定、評価、軽減するために高度なAIシステムの開発全体を通じて、その導入前及び市場投入前も含め、適切な措置を講じる	7. 技術的に可能な場合は、電子透かしやその他の技術等、ユーザーAIが生成したコンテンツを識別できるようにするため、信頼できるコンテンツ認証及び来歴のメカニズムを開発し導入する
2. 市場投入を含む導入後、脆弱性、及び必要に応じて悪用されたインシデントやパターンを特定し、緩和する	8. 社会的、安全、セキュリティ上のリスクを軽減するための研究を優先し、効果的な軽減策への投資を優先する
3. 高度なAIシステムの能力、限界、適切・不適切な使用領域を公表し、十分な透明性の確保を支援することで、アカウンタビリティの向上に貢献する	9. 世界の最大の課題、特に気候危機、世界保健、教育等（ただしこれらに限定されない）に対処するため、高度なAIシステムの開発を優先する
4. 産業界、政府、市民社会、学界を含む、高度なAIシステムを開発する組織間での責任ある情報共有とインシデントの報告に向けて取り組む	10. 国際的な技術規格の開発を推進し、適切な場合にはその採用を推進する
5. 特に高度なAIシステム開発者に向けた、個人情報保護方針及び緩和を含む、リスクベースのアプローチに基づくAIガバナンス及びリスク管理方針を策定し、実施し、開示する	11. 適切なデータインプット対策を実施し、個人データ及び知的財産を保護する
6. AIのライフサイクル全体にわたり、物理的セキュリティ、サイバーセキュリティ、内部脅威に対する安全対策を含む、強固なセキュリティ管理に投資し、実施する	あわせて、国際指針に基づいて、最も高度なAIシステムを開発する組織による行動のための自主的な手引きを提供する行動規範も公開されている。

（出所）総務省、広島AIプロセス <https://www.soumu.go.jp/hiroshimaaiprocess/>（2024年9月10日閲覧）
外務省、G7 プーリア・サミット https://www.mofa.go.jp/mofaj/ecm/ec/pageit_000001_00745.html（2024年9月10日閲覧）

4-4. EU - 強い法律規制と、開発・利用を促進する支援策 -

- EUは2020年2月「AI白書」「欧州データ戦略」を公表し、安全に利用できるAIの世界的リーダーを目指す目標を掲げ、2021年4月「AI白書」のフォローアップ文章として「EU AI法案」を公表。
- 2024年5月、世界初の包括的なAI規制法となる「AI法」が成立。リスクに応じてAIを4つのレベルに分類し、禁止事項、要求事項、罰則等を定めている。AI法は段階的に適用され、24か月後の2026年5月より全面的な適用が開始される予定。
- EU域内で利用されるAIシステムおよび汎用目的型AIモデルと、その開発者や使用者等が対象となるため、EU域内に拠点のない日本企業であっても、EU域内でAIシステムを活用したサービス提供を行っている場合等は適用対象となる。
- AI規制を進めるだけではなく、イノベーションを後押ししていくための各種支援策も設けられている。

リスクレベル	主な対象
許容できないリスク	人間にとって脅威とみなされ禁止されるAIシステム ・ 人々や特定の脆弱なグループに対する認知的行動操作 ・ ソーシャルスコアリング ・ 生体認証による人物の識別と分類 ・ 顔認証などのリアルタイムモニタリングおよびリモート生体認証システム
ハイレスク	安全性や基本的権利に悪影響を及ぼすAIシステムで、市場に投入される前、ライフサイクル全体を通じて評価される ・ EUの製品安全法*の対象となる製品(玩具、航空機、自動車、医療機器、エレベーター等)に使用されるAIシステム ・ EUデータベース**に登録する必要がある特定の分野に該当するAIシステム
特定の透明性が必要なリスク	出力した内容が間違っただ認識等を生む可能性があるAIシステムには、コンテンツがAIによって生成されたことを開示することや、学習に使用されたデータの概要を公開すること等が義務付けられる。 ・ 汎用AIシステム(OpenAI等)
最小リスク	上記以外のAIシステム

支援策
■ 2024年1月欧州委員会より公表された、信頼できる人工知能の開発において、新興企業と中小企業を支援するための一連の施策であるAIイノベーションパッケージの主な内容は以下のとおり。 ・ AIファクトリー設立のための「Euro HPC規制」の改正によるスーパーコンピュータの開発促進 ・ 欧州委員会内にAIオフィスの設立による、AI法の実施と施工の監督 ・ ホライズン・ヨーロッパやデジタルヨーロッパプログラムを通じて、2027年までに約40億ユーロの追加的な財政支援により公的、民間投資を実施 ・ AI人材プールを強化するため育成、トレーニング、スキル習得活動 ・ ベンチャーキャピタルや株式支援等を通じた、AIの新興企業やスケールアップ企業への官民投資 ・ 欧州共通データ空間の開発と展開を加速し、AIモデルのトレーニングと改善のためAIコミュニティに提供 ・ 「GenAI4EU」イニシアティブによる、14の産業エコシステムと公共部門における新しいユースケースとアプリケーションの開発をサポート ■ また、欧州委員会はAIのイノベーション促進のため、いくつかの加盟国と共同で、2つの欧州デジタルインフラコンソーシアムを設立した。 ・ ALT-EDIC: ヨーロッパ言語による大規模言語モデルの開発サポート ・ Citi VERSE: 最先端のAIツールを適用してスマートコミュニティ向けのローカルデジタルツインを開発し、交通管理から廃棄物管理まで、都市のプロセスのシミュレーションと最適化を支援

(注)*【EU 製品安全法】EUにおけるすべての消費者製品が安全であることを保障するための規制。食品以外のすべての製品および販売チャネルに適用される。

**【EUデータベース】EUにおける製品の安全性と追跡性を確保するため構築されている、医療機器、化学部室、化粧品、食料品などの各種データベース、

(資料)各種資料より国際通貨研究所作成

4-5. 米国・日本・中国 - 日米はガイドラインを示しつつイノベーション重視 -

- 米国、日本についてはEUのような法整備には至らないものの、ガイドラインを整備し、民間主導のイノベーションを重視。
- 中国については、AIを包括的に規制する法規制はないが、既存のサイバーセキュリティ関連法令に加え、アルゴリズム規制、倫理規定等を組み合わせる、独自のAI規制を行っている。

主なガイドライン・法令等

	米国	日本	中国
主なガイドライン 法令等	「安全・安心・信頼できるAIに関する大統領令」 2023年10月	「AI事業者ガイドライン」(総務省、経済産業省) 2024年4月	「生成人工知能サービス管理暫行弁法」 2023年7月
目的	AIの無責任な使用による詐欺や差別、国家安全保障上のリスクを軽減し、AIがもたらす利益を享受する	安全で信頼性の高いAIの利用促進を目的とし、特に、生成AIの普及に対応し、AI技術がもたらす社会的リスクを低減しつつ、イノベーションを促進するための枠組みを提供する	AIシステムの透明性を確保するためプロバイダーに対して、届出制度を設け、アルゴリズムを利用していることを明示すること等を義務付けている。
内容	1. AIの安全性やセキュリティの新たな基準 2. 国民のプライバシー保護 3. 公平性と公民権の推進 4. 消費者、患者、学生の権利を守る 5. 労働者の支援 6. イノベーションと競争の促進 7. 国際的なリーダーシップの強化 8. 政府による責任ある効果的なAI利用の確保	1. 基本理念 2. 共通指針 3. 高度なAIシステムに関する指針 4. AI開発者、提供者、利用者に関する事項	サービス提供を行う事業者は所定の手続に従い、自社の組織に関する情報に加え、AI製品サービスへのアクセス方法、アルゴリズムの分類、学習用データソースに関する情報等を政府に届出。情報は中国インターネット情報弁公室(CAC)が一般公開している。
その他	2019年2月の大統領令ではAI技術開発投資への積極的な投資方針が謳われ、2022年10月の「AI権利章典の青写真」ではAI開発にあたり考慮すべき原則がまとめられている。 米国において、連邦レベルで法律が成立する見通しはたっていないが、カリフォルニア州のAI規制法案が2024年8月上下院を通過した。 2024年9月末までに州知事が法案に署名するか否かが注目されている。	「AI事業者ガイドライン」は、それまで各省庁が定めた、AIガイドライン(2017年/総務省)、AI活用ガイドライン(2019年/総務省)、AI原則実践のためのガバナンスガイドライン(2021年/経産省)を、AIのリスクへの対応やAIの最適な利用に向け統合したもの。 2024年4月政府の有識者会議「AI戦略会議」が開催され、AI法規制の議論を開始。	2019年6月の「新一代人工知能ガバナンス準則」においては、責任ある人工知能の発展のため8項目のガバナンスを定義。 そのほか、2021年データセキュリティ法、個人情報保護法、2022年インターネットコメントサービス管理規定など、様々な法令や規制を制定し、独自のAI規制を行っている。

(資料)各種資料より国際通貨研究所作成

5. インプリケーション

- 自然言語により操作可能な生成AIの登場により、AIの活用はこれまでにない大きな広がりをみせている。AIの活用により、主に認知的タスクの割合が高い、高スキルの業務を中心に、生産性を高めることで総供給や総需要への好影響が期待されている。
- 一方、AIの活用が一部の労働者の失業に繋がる懸念や、労働所得や資本所得の不平等が生じるといったマイナスの影響も指摘されており、現時点ではAIが経済や社会に及ぼす影響は予見しがたい。
- 金融セクターにおいても、バックオフィス・ミドルオフィスを中心に業務効率化や生産性向上の手段としてAIの活用が進んでいるが、膨大なデータを学習させるため、金融の安定性の確保、プライバシーや著作権保護等の観点から、フロント業務への活用は途上。また、大量のエネルギー消費による環境面への影響も指摘されている。
- こうしたリスクや課題に対し、OECDのAI原則を皮切りに、国際的なガイドラインや行動規範が制定され、国・地域によって対応は異なるものの、AIの規制や活用支援に向けた法制度やガイドラインの整備も進められている。
- 今後、AIを活用したイノベーションの効果を享受できるかどうかは、民間企業や労働者がこのテクノロジーをどう活用し適応するかや、各国政府による法規制、ガイドライン、支援策の制定、活用に伴い生じる経済的不平等への対策にかかっている。

以上

主要参考資料

- OECD, “ [AI Principles](#) ”, 22 May 2019 (updated 3 May 2024)
- OECD, “ [Generative artificial intelligence in Finance](#) ”, 15 Dec 2023
- IMF, “ [The Macroeconomics of AI](#) ”, Dec 2023
- IMF, “ [Artificial Intelligence and the Future of Work](#) ”, 14 Jan 2024
- BIS, “ [Annual Economic Report 2024 Chapter III](#) ”, 30 Jun 2024
- McKinsey & Company, 「 [生成AIがもたらす潜在的な経済効果](#) 」, 2023年6月
- European Parliament, “ [EU AI Act : first regulation on artificial intelligence](#) ”, 18 Jul 2024
- European Commission, “ [Commission launches AI Innovation package](#) ”, 24 Jan 2024
- The White House, “ [Executive order on the safe, secure, and trustworthy development and use of AI](#) ”, 30 Oct 2023
- 総務省, 「令和6年版 情報通信白書」
- 経済産業省, 「AI事業者ガイドライン(第1.0版)」, 2024年4月19日
- 三菱総合研究所, 「生成AIコラム」<https://dx.mri.co.jp/generative-ai/column/risks-03/> (2024年9月10日閲覧)
- KPMG, 「各国のAI規制動向と欧州連合AI規制案について」, 2023年10月17日
- PwC, 「中国におけるAI関連規制」, 2023年7月19日
- PwC, 「米国における「人工知能の安全・安心・信頼できる開発と利用に関する大統領令」の解説」, 2024年1月23日
- PwC, 「欧州(EU)AI規制法」の解説, 2024年6月21日

当資料は情報提供のみを目的として作成されたものであり、何らかの行動を勧誘するものではありません。
ご利用に関しては、すべてお客様御自身でご判断下さいますよう、宜しくお願い申し上げます。当資料は信頼できる
と思われる情報に基づいて作成されていますが、当行はその正確性を保証するものではありません。
内容は予告なしに変更することがありますので、予めご了承下さい。また、当資料は著作物であり、著作権法により
保護されております。全文または一部を転載する場合は出所を明記して下さい。

会社名：株式会社三菱UFJ銀行 経済調査室
〒100-8388 東京都千代田区丸の内2-7-1

執筆：国際通貨研究所 経済調査部 宮川 真一
照会先：三菱UFJ銀行 経済調査室 長谷川 裕輝 e-mail: yuuki_4_hasegawa@mufg.jp