

日本DPO協会 第36回個人情報保護セミナー

「AI規制の国内外の最新動向と DeepSeek騒動が個人情報保護に与える影響」

アンダーソン・毛利・友常法律事務所外国法共同事業
パートナー弁護士 中崎 尚

2025年3月

本日のトピック

- I. 国際的なAIガバナンスの動向
- II. EUと欧州各国のAI規制の動向
- III. 米国（連邦・各州）・カナダのAI規制の動向
- IV. 東アジア（日中韓）のAI規制の動向
- V. DeepSeekと各国データ保護当局の動向
- VI. AI（LLM）に対する、各国のデータ保護当局の姿勢
- VII. （ご参考）顔認識技術とAI規制

I – 1. 主な国際協力の動向

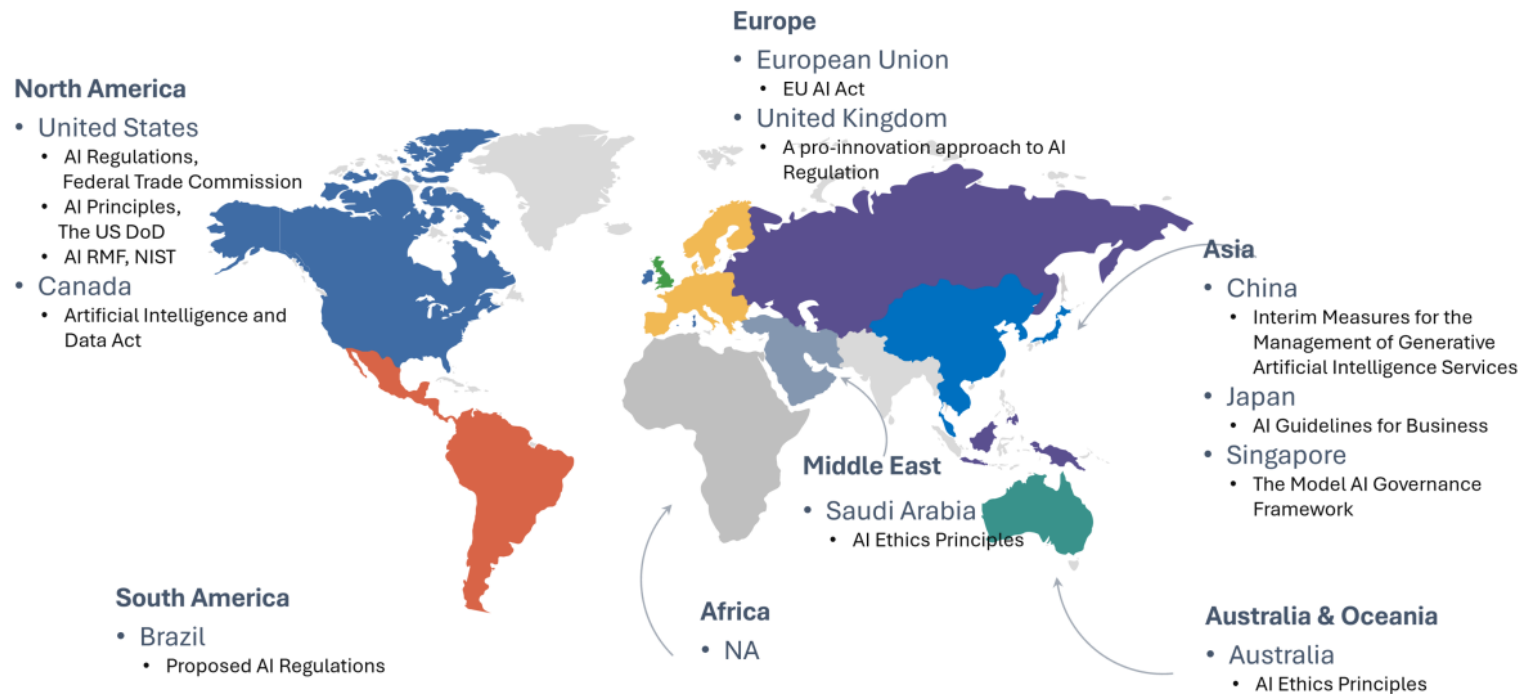
- 2024年9月：欧州評議会「人工知能枠組み条約」の批准がスタート。
- 2025年2月：同条約に日本とカナダが署名。
- 2025年2月：AIアクション・サミットがパリで開催。
- 2025年2月：同サミットにて「人と地球のための持続可能で包括的なAIに関する宣言」を採択。米国・英国は署名せず。
- 2025年2月：5カ国のデータ保護当局の責任者が、AIとデータ保護の基本ルールを定めた宣言に署名。

I – 1. 主な国際協力の動向

- 2025年2月：10カ国が「公益のための人工知能に関するパリ憲章」を採択。
- 2025年2月：OECD、「スクレイピングされたデータでAIを訓練する際の知的財産権に関する報告書」を公表。
- 2025年2月：OECD、「AIのためのデータへのアクセスと共有の強化に関する報告書」を公表。
- 2025年2月：OECD、グローバルなAIインシデント報告のフレームワーク構築を提唱。

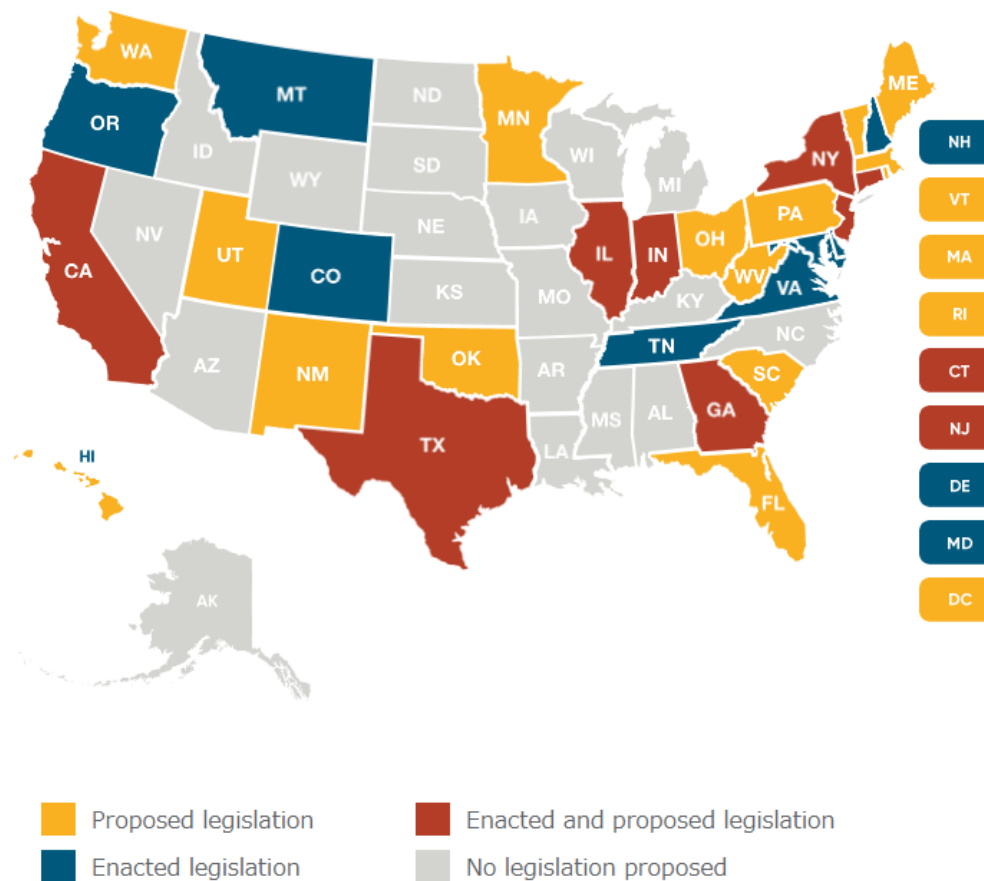
I – 2. 世界各国の動向

2024年12月時点の世界各国のAI規制の導入状況 (<https://www.linkedin.com/pulse/ai-regulations-globally-same-goal-different-paths-dr-irina-steenbeek-1cipe/>)



I – 2. 世界各国の動向

2024年2月時点の米国各州のAI規制の導入状況（Bryan Cave Leighton Paisner LLP 「US STATE-BY-STATE AI LEGISLATION SNAPSHOT」）



I - 2. 世界各国の動向

- 米国（連邦法レベル）：バイデン時代の大統領令（2023年10月、2024年12月）のもと連邦政府による対応が進められてきたが、トランプ大統領の就任直後に一部が撤回された。
- 米国（州法レベル）：カリフォルニア州ほか多数の州で制定されている。
- 欧州（EU加盟国）：2026年8月のEU AI規則の全面施行を目指して、加盟国各国の国内法整備が進行中。ルクセンブルクは2025年1月に制定した。
- 英国（UK）：2025年1月、労働党政権がイノベーション重視の政策を示す。2025年3月、AI法案が第一読会を通過。

I – 2. 世界各国の動向

- 韓国：2025年1月、世界で2番目とされる包括的なAI規制法（AI Basic Act）が制定された。2026年1月に施行予定。
- 中国：EU AI Actに先駆けて、2023年8月より、生成人工知能サービス管理暫定弁法が施行され、2024年には同法に基づくウルトラマン事件で、AI提供事業者が著作権侵害の責任を問われた。
- 日本：2025年2月末、内閣が国会に「人工知能関連技術の研究開発及び活用の推進に関する法律案」を提出した。
- ブラジル：EU AI規則 と類似した法案が、上院を通過した。

I – 3. 国際的な潮流の転換？

- AI分野における国際競争力の強化のためのAI規制の見直し
- 2024年9月、欧州中央銀行（ECB）前総裁マリオ・ドラギ氏による、EUの競争力強化に向けた提言をまとめた報告書「ドラギレポート」が公表される
 - ・ 企業サステナビリティ報告指令（CSRD）の簡素化
 - ・ サステナブル金融開示規則（SFDR）の簡素化
- 2025年2月、AI責任指令案、e Privacy規則案の廃案
- 2025年2月、AIアクション・サミットにおける数々の事象（マクロン仏大統領発言、ヴァンス米副大統領発言、米国・英国の署名拒否、英国AISIの名称変更等）

I – 3. 国際的な潮流の転換？

- 「AI責任指令」が求められていた理由
- 製造物責任指令改正ではAIシステムを含むソフトウェアを適用対象としているものの、AIシステムはその複雑性やブラックボックス問題により、その欠陥や因果関係を立証することが特に困難かつ高額であるという特有の問題状況が指摘されてきたところである。また、製造物責任指令改正案では、救済対象が、安全性が不十分であることから生じた損害に限定されるという限界がある。このため、AI責任指令案を新たに定め、より充実した被害者救済を図ることを目指すものとされていたが……。

Ⅱ－１． EUにおけるAI規制の全体像

- 「AIに関する整合的規則の制定及び関連法令の改正に関する欧州議会及び理事会による規則」 (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) (EU AI規則) ⇒ 2024年8月、発効
- 新・製造物責任指令 ⇒ 2024年12月、発効
- AI責任指令案 ⇒ 2025年2月、e Privacy規則案とともに廃案に
- データ法制 (GDPR、Data Act、Data Governance Actほか)

Ⅱ－２． EU AI規則の適用スケジュールと全体像

■ EU AI規則の発効（entry into force）及び適用（application）のスケジュール

2024年7月12日	官報への掲載（公布）	条項	罰則
2024年8月1日	EU AI法の発効	なし	
2025年2月2日	禁止AIの義務の適用開始 AIリテラシー	1編及び2編	35,000,000ユーロ or 前会計年度の全世界に おける年間総売上高の 7%以下
2025年5月2日	一般目的（GP）AI に関する Codes of Practice		
2025年8月2日	一般目的（GP）AI の義務の適用開始 AIガバナンス・ルールの適用開始 罰則の適用開始	3 編 4章、5編、 7編、12編（101 条を除く）及び 78条	15,000,000ユーロ or 前会計年度の全世界に おける年間総売上高の 3%以下
2026年2月2日	post-market monitoring（市場投入・使用開始後のモニタリング）の実装行為（implementing act）の開始		
2026年8月2日	高リスクAIシステム（付属書III）の義務の適用開始 透明性原則の適用開始	EU AI法的全編	同上
2027年8月2日	高リスクAIシステム（付属書 I）の義務の適用開始	6条1項	同上

Ⅱ－２． EU AI規則の適用スケジュールと全体像

- Guidelines on prohibited artificial intelligence practices established by Regulation (EU) 2024/1689 (AI Act) (2025年2月4日公表)
- Guidelines on the definition of an artificial intelligence system established by Regulation (EU) 2024/1689 (AI Act) (2025年2月6日公表)
- General-Purpose AI Code of Practice (2025年3月に草案第3版が公表)

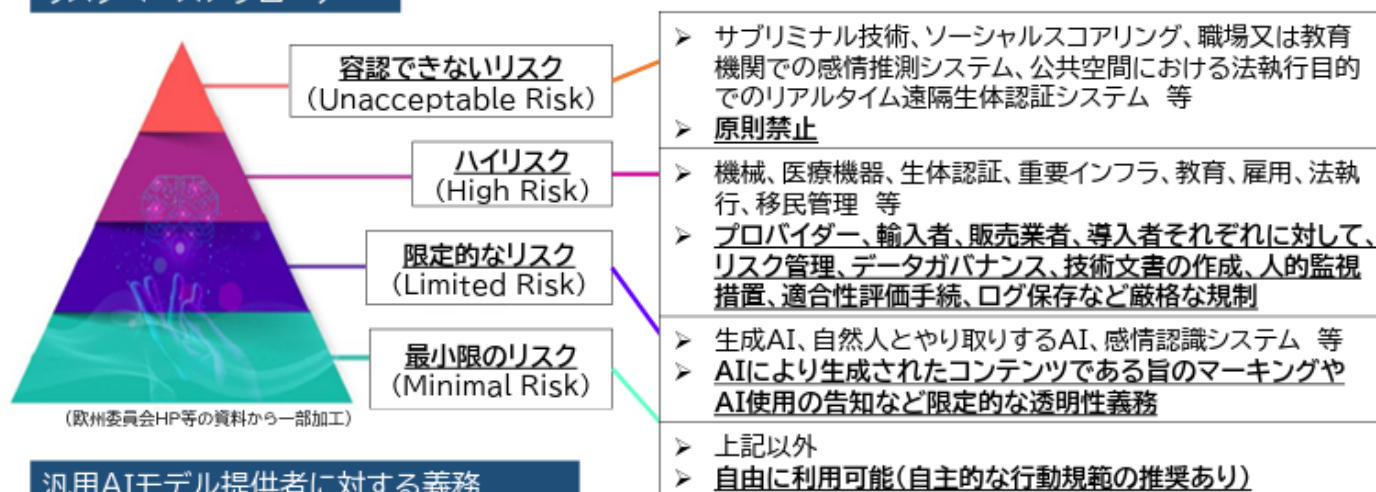
II - 2. EU AI規則の適用スケジュールと全体像

EU AI規則の概要（欧州連合日本政府代表部 2024年9月）より抜粋。

リスクベースアプローチ

- AI規則では、**リスクベースアプローチ**を採用し、4つのリスクレベルを設け、各々のリスクに応じた規制を規定。それに加え、汎用AIに関する規制あり。

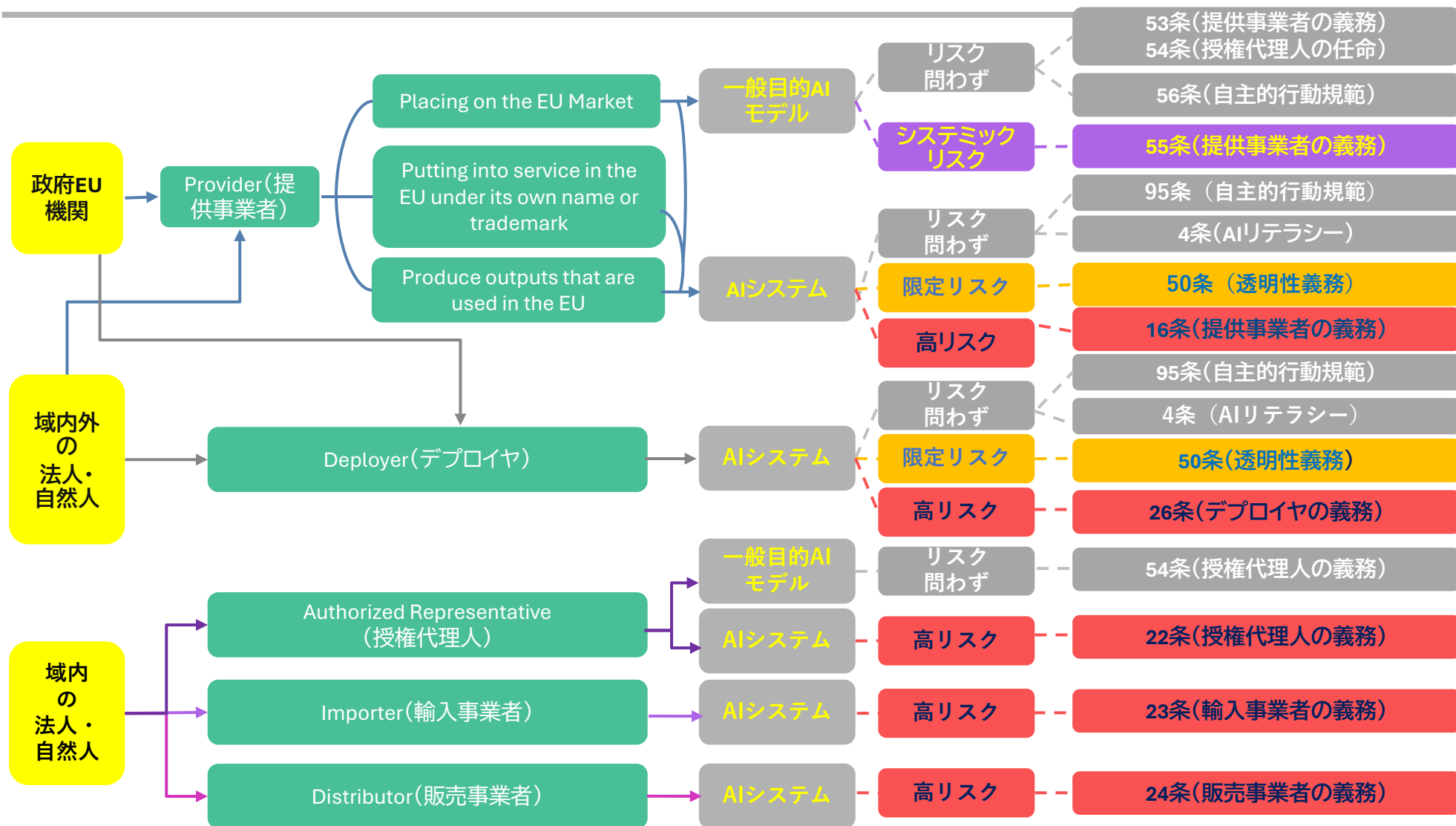
リスクベースアプローチ



汎用AIモデル提供者に対する義務

汎用AIモデル一般	システムリスクを有する汎用AIモデル
- 技術文書の作成及び更新 - 汎用AIモデルをAIシステムに統合する提供者向けの情報・文書の作成、更新及び提供 - 著作権法を遵守するためのポリシーの実行 - 汎用AIモデルの学習に使用したコンテンツに関する十分に詳細な要約の作成及び公開 - 域内代理人の指名	(左記に加えて) - モデル評価の実施 - EUレベルでのシステムリスクの評価及び軽減 - 深刻なインシデント及びそれに対する是正措置のAIオフィスへの報告 - 適切なレベルのサイバーセキュリティ保護

II - 2. EU AI規則の適用スケジュールと全体像



Ⅱ－３． EU AI規則の適用範囲とガイドライン

AI規則が適用される「AIシステム」とは

- 当初の幅広な定義から、OECDの狭い定義に近づける修正がなされた。
- 三者協議を経て、再び定義の拡大が試みられ、最終的に「**さまざまなレベルの自律性で動作するように設計されており**、明示または黙示の目標のために、物理的または仮想の環境に影響を与え得る予測、コンテンツ、推奨、もしくは決定などの出力を生成する方法を、**受信した入力から推測する機械ベースのシステム**」と定義された。

Ⅱ－３． EU AI規則の適用範囲とガイドライン

AI規則の定める義務を負う主体

■ 「提供事業者（provider）」

有償・無償を問わず、自己の名称または商標でこれを市場に置きたまたは実用化する目的のために、**AIシステムを開発**しまたはAIシステムを開発させる自然人または法人、公的機関、専門機関、もしくは他の組織。EU域内に拠点を有するか否かは問わない。

■ 「デプロイヤー（deployer）」

AIシステムが個人的な非専門的活動の過程において使用される場合を除き、その権限の下で**AIシステムを使用**する自然人、法人、公的機関、専門機関、もしくはその他の組織。

Ⅱ－３． EU AI規則の適用範囲とガイドライン

Guidelines on prohibited artificial intelligence practices

- AIシステムの「placing on the market」「putting into service」「use」といった主要な用語を定義し、「provider」と「deployer」の役割を説明。
- 第5条に基づく具体的な禁止事項についても概説しており、顔画像の無差別スクレイピングや感情認識など、これらの禁止事項が適用される基準についても詳しく説明している。
- さらに、ガイドラインでは、禁止されたAIの使用について、基本的人権への影響評価（FRIA）を実施するよう指示している。

Ⅱ－３． EU AI規則の適用範囲とガイドライン

Guidelines on the definition of an artificial intelligence system

- AIシステムの特徴をなす7つの要素が詳細に説明されている。すなわち、機械ベース、自律性の程度、適応性、明示的または暗黙的な目的、推論能力、環境に影響を与える出力、環境との相互作用である。この定義には、導入前と導入後の両方の段階が含まれているが、すべての要素が両方の段階で必要とされるわけではないと明記されている。
- また、ガイドラインでは、複雑な推論能力を持たないシステムや、完全に人間の介入に依存するシステムは除外されることも概説されている。

II – 4. General-Purpose AI Code of Practice

EU公式サイト (<https://digital-strategy.ec.europa.eu/en/policies/ai-code-practice>) より抜粋。



Ⅱ－4． General-Purpose AI Code of Practice

一般目的AIモデルの提供事業者の義務

項目	義務の内容
①技術文書の当局への提供 (第53条第1項(a))	AI Office及び国内管轄当局に要求に応じて提供するために、モデルの技術文書を作成し、最新の状態に保つこと。
②提供事業者への情報提供 (第53条第1項(b))	一般目的AIモデルをAIシステムに統合しようとしているAIシステム提供事業者に対して、情報及び技術文書を作成・更新・提供すること。
③著作権の尊重(第53条第1項(c))	特に、指令(EU)2019/790(デジタル単一市場における著作権指令)の第4条(3)に従って表明された権利の留保を特定し、尊重するために、最先端技術によるものを含め、EUの著作権法を尊重する方針を導入すること。
④学習用コンテンツの情報公開(第53条第1項(d))	AI Officeが提供するテンプレートに従って、一般目的AIモデルの学習に使用されるコンテンツに関する十分に詳細な要約を作成し、公開すること。
⑤当局への協力(第53条第3項)	欧州委員会及び国内管轄当局が本規則に基づく権限及び権限を行使する際に、必要に応じて協力すること。

Ⅱ－４． General-Purpose AI Code of Practice

システミックリスクを伴う一般目的AIモデルの提供事業者の義務

項目	義務の内容
①モデル評価の実施(第55条第1項(a)号)	システミック・リスクを特定し、軽減することを目的としたモデルの敵対的テストの実施及び文書化を含む、最新技術を反映した標準化されたプロトコル及びツールに従って、モデル評価を実施すること。
②システミック・リスクの評価・軽減(第55条第1項(b)号)	システミック・リスクを伴う一般目的AIモデルの開発、市場投入または使用から生じる可能性のある、その発生源を含む、EUレベルで起こりうるシステミック・リスクを評価および軽減すること。
③重大インシデント情報の把握・報告(第55条第1項(c)号)	重大インシデント及びそれに対処するための可能な是正措置に関する関連情報を把握し、文書化し、AI Office及び必要に応じて国内管轄当局に過度な遅延なく報告すること。
④サイバーセキュリティの確保(第55条第1項(d)号)	システミック・リスクを有する一般目的AIモデル及び当該モデルの物理的インフラストラクチャについて、適切なレベルのサイバーセキュリティ保護を確保すること。

II – 4. General-Purpose AI Code of Practice

日本企業への影響

- 直接的な影響：既存の一般目的AIモデルを調整・ファインチューニングした事業者が、新しいモデルのプロバイダーとして位置付けられる可能性がある。※General-Purpose AI Models in the AI Act – Questions & Answers
- 間接的な影響：一般目的AIモデルの開発受託をする場合、取引先から、開示義務の対象となる情報の提供や文書作成について協力を要請されることが想定される。※GDPRのProcessorと位置づけられ、協力を求められるのと類似の場面。

Ⅱ－４． General-Purpose AI Code of Practice

第2版から第3版への主要な変更点

- AIモデルの透明性、著作権、リスク評価、リスク緩和に焦点を当てた4つのコアワーキンググループを維持しながら、2024年12月の第2版を大幅に再構築した。
- 注目すべき変更点としては、より明確な義務規定を優先し、多数の主要業績評価指標（KPI）が削除されたことが挙げられる。
- 第3版では、下流のAIシステム開発者と共有すべき情報と規制当局が保有すべき情報を明確に区別する標準テンプレートを作成する包括的なモデル文書のひな形が導入された。

Ⅱ－４． General-Purpose AI Code of Practice

- 第3版では、「システムリスク分類」が改良され、すべてのプロバイダーが評価すべき4つの主要リスク、すなわち「サイバー攻撃能力」、「化学・生物兵器リスク」、「有害な操作」、「制御の喪失」が定義されている。

Ⅱ－5． UKのAI法案

UK Artificial Intelligence (Regulation) Billの概要

- AI規制の監督、規制当局間の整合性の確保、ギャップ分析の実施、法規制の見直し、規制の有効性の監視を担当するAI Authorityを設立する。
- AI規制の原則を定め、安全性、透明性、公平性、説明責任、救済を強調し、企業に対してこれらの原則を順守し、倫理的なAI利用を確保するためのAI officerの任命を義務付けている。
- AIを、認知能力に近似した知覚、解釈、意思決定を行う技術と定義し、これにはLLMモデルのような生成AIも含まれる。また、この法案は、罰則・罰金を課すことも認めている。

Ⅱ－6． スイス、AI条約の批准に前向きな姿勢を示す

- 2025年2月12日、スイス連邦参事会は、欧州評議会の「人工知能に関する枠組み条約」（AI条約）の批准に前向きな見解を含む報告書を公表した。3つの目標（イノベーションの中心としてのスイスの強化、基本的人権の保護、AIに対する国民の信頼の向上）を達成するために、相互補完の3つの規制アプローチとして、①テーマ別及び分野別の規制活動の継続、②AI条約の批准（最小限度またはより広範な実施）、③EU AI規則との整合化を提案する。
- 同報告書では、2026年末までに、特に透明性、データ保護、非差別、監督の分野における必要な法的措置を定義するための協議草案法案を準備する予定である、としている。

Ⅲ－１． 米国（連邦レベル）の規制動向

- 米国では、2023年10月30日に当時のバイデン大統領が出した「人工知能（AI）の安心、安全で信頼できる開発と利用に関する大統領令」に基づいて、連邦政府の各機関が規則・ガイドラインの策定やレポートの公表が進められてきた。
- 2025年1月23日、トランプ大統領が「人工知能（AI）に対する規制緩和を指示する大統領令」を発表し、米国のAIイノベーションの障壁となっている既存のAI政策を無効とする方針を示し、同日に公表したファクトシートでは、バイデン時代の大統領令はイノベーションを阻害するものであるとして、撤回する方針を示した。具体的には、AIの安全性評価や、公平性と公民権に関するガイダンス、AIが労働市場に与える影響に関する調査の在り方への影響が懸念されている。

Ⅲ－１． 米国（連邦レベル）の規制動向

- 労働者保護や企業規制を重視し、AI分野でもビッグテックに対して厳しい目線を注いできた連邦取引委員会（FTC）のトップも、政権交代に伴い、2025年1月に交代している。新たなトップは、民主党政権下の積極的な企業規制に反対する姿勢を示しており、AIを含む競争政策の分野において、路線変更が見込まれている。
- 2025年2月6日、米国ホワイトハウスは、新たなAI戦略について一般からの意見を募集すると発表した。Fair Useに関して、大手IT事業者とハリウッドスターが対立する意見を提出したことが報じられている。

（ご参考） Thomson Reuters v. Ross Intelligence事件（2025年2月11日、デラウェア地区連邦地方裁判所の略式判決。AIの学習データ利用についてFair Useの抗弁を認めず、著作権侵害を認定した。）

Ⅲ－２． 米国NISTの動向

- 米国商務省の下部組織であるNIST（National Institute of Standards and Technology、米国国立標準技術研究所）は、「AIリスク管理フレームワーク」（AI Risk Management Framework、AI RMF）を2023年1月に公表した。
- AI RMFは、法的な拘束力はないが、実務では多くの企業が準拠している。日本国内では、AI事業者ガイドラインが近い位置づけであり、日本のAISIのウェブサイトで、両者の比較が公表されている。
- 2023年10月の大統領令を受けて、2024年7月に、生成AI に特有のリスクを特定し、これを管理するための指針を示す「AIリスク管理フレームワーク：生成AIプロファイル」を公表している。

Ⅲ－２． 米国NISTの動向

- 2025年2月、トランプ政権は、NISTの497名の試用期間中の職員を解雇する計画であることが報じられた。これにより米国のAI研究・開発能力が大きく損なわれる可能性がある。NISTはCHIPS法と科学法の実施において中心的な役割を果たしており、AIに不可欠な国内半導体生産を促進している。この解雇は、米国のAI分野における世界的リーダーシップを維持するというトランプ政権の目標を損なう恐れがある。
- 米国では、政府内でのAI人材不足が長年の課題となっており、多くの専門家が高額報酬の民間セクターに引き抜かれてきた。バイデン政権下ではこのギャップを埋めるため、「AI人材サージ」計画により2024年から2025年にかけて500人以上のAI専門家を採用することを目指していたが、大規模解雇によりこの取り組みが頓挫する可能性がある。

Ⅲ－３． 米国（州レベル）の規制動向

- もともと、州法レベルでは、雇用場面における AI 技術の規制——①アルゴリズムの規制、②自動化された意思決定の規則（主に雇用分野）を中心に、③基本的権利、④何らかの検討体制を整備するための取り組みの観点からのルール整備が進んでいた。
- イリノイ州：2020年1月より、AI 面接技術を使用している場合は、事前通知と当該技術の説明義務を課し、同意を得ることを義務づけた。
- メリーランド州：2020年10月より、採用面接において顔識別技術を使用する場合には、応募者に書面による同意を必要とする規制が設けられていた。

Ⅲ－３． 米国（州レベル）の規制動向

- ニューヨーク州：2021年11月、雇用決定における AI の使用を規制する法律が可決され、2023年1月1日より施行されている。同法では「自動化された雇用決定ツール」を、「機械学習、統計モデリング、データ分析又は AI から派生した、スコアを含む単純化された出力を発行する任意の計算プロセス」として広く定義し、当該ツールを使用する場合は、1年ごとにバイアス監査をしなければならない、雇用決定に使用する場合は、応募者に対しその旨通知する義務を課す。
- ユタ州：2024年3月、生成 AI を使用する事業者に透明性義務を課す AI Policy Act を制定し、同年5月より施行している。個人との対話を生成AIに実施させる場合、求めがあれば、対話しているのが生成AIであることの開示義務を負う。免許・認証が必要なサービスの提供に際して、生成AIが使用される場合には、その旨の明示を求める。

Ⅲ－３． 米国（州レベル）の規制動向

- コロラド州：2024年5月に制定されたコロラド州AI法（Colorado AI Act）は、米国初の包括的なAI規制法と言われており、2026年2月から施行予定。
- カリフォルニア州：2024年9月、当時、最先端モデルに対する包括的な規制として世界的に注目を集めていたSB1047は、議会を通過するも、知事が署名を拒否し廃案に。その一方で、同じタイミングで、複数のAI規制法が成立している。特に注目されるのが、AI透明化法（AI Transparency Act、Senate Bill No. 942）と、生成AI学習データの透明化に関する法律（Assembly Bill No. 2013）であり、いずれも2026年1月から施行される予定。

Ⅲ－３． 米国（州レベル）の規制動向

- ヴァージニア州：2025年2月に、高リスクAI規制法案（Virginia's House Bill 2094, the High-Risk Artificial Intelligence Developer and Deployer Act）が州議会を通過するも、翌月3月に、共和党出身の知事が拒否権を行使。
- テキサス州：テキサス州選出のジョバンニ・カプリグリオーネ（共和党）下院議員が、同州におけるAIイノベーションを厳しく規制する法案「テキサス責任あるAIガバナンス法（TRAIGA）」を大幅に改訂したバージョンを提出。当初のバージョンはヴァージニア州の法案と非常に類似していたが、新しいバージョン（House Bill 149）では、以前の法案の最も厳格な要素が排除されている。

Ⅲ－４． カナダの規制動向

- カナダにおけるAIの開発と利用に関する規制枠組みを確立すると同時に、プライバシー保護の大幅な改革を導入することを目的として、C-27法案の策定が計画され、2022年にその一部として「AI及びデータ法（AIDA）」法案が国会に提出された。
- AIDAでは、①AIシステムの国際間および及び各州間の取引規制、②個人に重大な損害を与える、または利益を害するおそれのある行為の禁止等を目指していた。
- 2025年に委員会において廃案となり、事実上、法案全体が廃案となった。このため、カナダにおけるAIガバナンスの将来について懸念が生じている。

Ⅲ－４． カナダの規制動向

- 2025年3月、カナダ政府は、複数の政府省庁にわたるAI開発とガバナンスに関する複数のイニシアティブを発表した。
- 3月4日、カナダ財務省事務局は、連邦公共サービスにおける同国初のAI戦略を発表した。この戦略では、AI専門センターを設立し、政府業務におけるAI利用におけるシステムセキュリティ、人材開発、透明性に重点を置いている。
- 3月6日、フランソワ＝フィリップ・シャンパーニュ革新担当大臣は、新たに「安全でセキュアなAI諮問グループ」を設置するとともに、再編された「AI諮問委員会」を発表した。さらに、6つの組織がカナダの「高度生成AIシステムの責任ある開発と管理に関する自主行動規範」（2023年9月）に署名し、現在では46の組織が署名している。

Ⅳ－１． 韓国の規制動向

- 2025年1月21日、「AIの開発および信頼基盤の確立に関する基本法」（AI基本法）が公布された。1年間の準備期間を経て、2026年1月22日に施行予定である。この期間中に、高影響AIの具体的な種類と範囲を定義する下位法令や分野別ガイドラインが策定・施行される予定。
- 目的：市民の権利と尊厳を保護し、生活の質を向上させ、AIの健全な発展と信頼基盤の確立に必要な基本事項を規定することで、国家競争力を強化すること。EUとは異なるアプローチを採用する。
- 対象：AIを開発・提供する「AI開発事業者」及び、開発者が提供するAIを使用してAI製品やサービスを提供する「AI活用事業者」に分類されている。

IV－1． 韓国の規制動向

■ ガバナンス

- ①AIの発展と信頼基盤造成のための推進体系の構築
- ②競争力強化のためのAI基本計画の策定
- ③AI Safety Institute の運営

■ AI産業の育成支援

- ①AI研究開発・標準化・学習用データ施策の策定・AIの導入および活用
に対する政府支援
- ②AI集積団地の指定
- ③専門人材の確保
- ④中小企業のための特別支援

Ⅳ－１． 韓国の規制動向

■ 規制

- ①高影響AI・生成型AIに対する安全・信頼基盤の構築
- ②透明性・安全性確保義務の規定
- ③事業者の責務規定

■ 高影響AI：人命、身体 of 安全、基本的権利に重大な影響を及ぼす、またはリスクをもたらす可能性のある11の特定分野で使用されるAIシステム

■ 生成型AI：入力データの構造や特性を模倣して、テキスト、音声、画像、動画などのさまざまな出力を生成するAIシステム

Ⅳ－１． 韓国の規制動向

- 透明性義務：高影響AIや生成型AIを使用した製品やサービスを提供する際に、事前にユーザーに通知することや、生成AIの出力にラベルを付ける。
- リスクマネジメント：学習のための累積計算能力が一定の閾値を超えるAIシステムは、AIライフサイクル全体でリスクの特定、評価、軽減措置を実施し、AI関連の安全インシデントを監視・対応するリスク管理システムを確立する必要がある。
- 高影響AI事業者は、大統領令で定められたリスク管理計画、AI出力の説明方法と基準、ユーザー保護計画、高影響AIの人間による監視、安全性と信頼性措置の文書化、国家AI委員会が審議・決定したその他の事項などの安全性と信頼性措置を実施する必要がある。

IV－2． 日本の規制動向

- 2024年4月、自民党が「AIホワイトペーパー」を公表し、最小限度ではあるがEUのような包括的な法規制の導入を提言した。
- 2024年8月より、AI制度研究会において、日本でもAIの法規制を導入すべきか検討が行われ、2024年12月には、同研究会の「中間取りまとめ（案）」のパブリックコメントの募集が実施された。2025年2月には、パブリックコメントを踏まえた「中間取りまとめ」が公表された。
- 2025年2月末には、日本版「AI法」の法案（人工知能関連技術の研究開発及び活用の推進に関する法律案）（「AI法案」）が閣議決定され、衆議院に提出された。今国会での成立を目指して、今後審議が行われる。

IV－2． 日本の規制動向

- 法案の名称の通り、AI技術の研究開発・活用を推進すべく、政府内にAI戦略本部を設置し、政府の施策の基本的な方針としてAI基本計画を策定する旨を定め、その基本的な施策の内容を明確にするなど、大半の条項は国や政府を名宛人としており、いわば「AI基本法」ともいうべき内容となっている。もちろん推進一辺倒というわけではなく、多くの国民がAIに対して不安を感じているという現実を踏まえて、「犯罪への利用、個人情報情報の漏えい、著作権の侵害その他の国民生活の平穏及び国民の権利利益が害される事態を助長するおそれがあること」がないよう、「透明性の確保その他の必要な施策が講じられなければならない」と定めている（1条4項）。

IV－2. 日本の規制動向

- 民間事業者に関しては、EUのAI規則と異なり、国や地方自治体の施策への協力義務を定めるにとどまっており（7条）、協力を拒んでもそれを理由に罰則が適用されることもない。ただ、罰則は定めていないものの、16条後段では「（国は）研究開発機関、活用事業者その他の者に対する指導、助言、情報の提供その他の必要な措置を講ずるものとする」と定めており、民間事業者が協力を拒んだ場合に、当該事業者に対して、国から指導・助言がなされる可能性があることを明確に定めている。報道情報では、さらに踏み込んで、事案によっては、民間事業者名を公表することも検討中と伝えられている。国から非協力的な事業者であるというレッテルを貼られることによる、レピュテーションリスクの影響は計り知れず、罰則がなくとも、事実上協力を強制される状況に陥る可能性がある点は注意が必要である。

IV－2. 日本の規制動向

「人工知能関連技術の研究開発及び活用の推進に関する法律案（AI法案）概要」より抜粋。

法案の概要	目的	国民生活の向上、国民経済の発展
	基本理念	経済社会及び安全保障上重要 → 研究開発力の保持、国際競争力の向上 基礎研究から活用まで総合的・計画的に推進 適正な研究開発・活用のため透明性の確保等 国際協力において主導的役割
	AI戦略本部	本部長：内閣総理大臣 構成員：全閣僚 関係行政機関等に対して必要な協力を求める
	AI基本計画	研究開発・活用の推進のために政府が実施すべき施策の基本的な方針等
	基本的施策	研究開発の推進、施設等の整備・共用の促進 人材確保 教育振興 国際的な規範策定への参画 適正性のための国際規範に即した指針の整備 情報収集、権利利益を侵害する事案の分析・対策検討、調査 事業者・国民への指導・助言・情報提供
	責務	国、地方公共団体、研究開発機関、事業者、国民の責務 関係者間の連携強化 事業者は国等の施策に協力しなければならない
	附則	見直し規定（必要な場合は所要の措置）

IV－2． 日本の規制動向

- 日本版「AI法」の中核的な概念である「人工知能関連技術」について、AI法案では「人工的な方法により人間の認知、推論及び判断に係る知的な能力を代替する機能を実現するために必要な技術並びに入力された情報を当該技術を利用して処理し、その結果を出力する機能を実現するための情報処理システムに関する技術」と定義している（2条）。つまり、①人工的な方法により人間の認知、推論及び判断に係る知的な能力を代替する機能を実現するために必要な技術、あるいは、②入力された情報を当該技術を利用して処理し、その結果を出力する機能を実現するための情報処理システムに関する技術、のいずれかに該当すれば、「人工知能関連技術」に該当することになる。

IV－2． 日本の規制動向

- EUのAI規則の中核的な概念である「AIシステム」の定義（「様々なレベルの自律性で動作するように設計され、導入後に順応性を示す可能性のある、機械ベースのシステムであって、明示的又は暗黙的な目的のために、物理的又は仮想的な環境に影響を与え得る予測、コンテンツ、推奨又は決定等のアウトプットを生成する方法を、受け取ったインプットから推論するもの」）と比べると、ややシンプルであり、広範囲を対象としているものと考えられる。

V – 1. DeepSeekと各国のデータ保護当局の対応

- イタリア、韓国では、データ保護当局が、DeepSeekアプリのダウンロードの一時停止を命令。
- オランダ、ドイツ、米国テキサス州、多数の国のデータ保護当局が、DeepSeekに対する調査を開始。
- 多くの国で、政府機関によるDeepSeekの利用を禁止。

V - 2. 個人情報保護の観点から見たDeepSeekのリスク

- GDPRの観点から見たDeepSeekの問題点
 - ・ データ収集と処理に求められる保護のレベルに達していない。
 - ・ 規約上、ユーザーが入力するプロンプトを制限なく、保存・記録し、分析できる。
 - ・ 規約上、外国政府または第三者の潜在的な関与が想定される。
 - ・ DeepSeekの運営事業者が、EU/EEAの域内に拠点を有していないことから、GDPRの遵守を求めても、実効性に乏しい。
- 日本法の観点から見たDeepSeekの問題点
 - ・ 「学習」のオプトアウトができず、プロンプト入力が、第三者提供として規律される。
 - ・ 越境移転規制

VI. AI（LLM）に対する、欧州のデータ保護当局の姿勢

- 2023年のイタリアのデータ保護当局GaranteによるChatGPTの一時停止騒動もEU各国のデータ保護当局によるOpenAIに対する調査は継続。
- 2024年5月、EDPBが「ChatGPTタスクフォースによる進捗報告書」を公表。
- 2024年12月、イタリアのデータ保護監督当局Garanteは、ChatGPTに関するGDPR違反について、OpenAIに対し、1,500万ユーロの制裁金を課したことを公表した。

VI. AI（LLM）に対する、欧州のデータ保護当局の姿勢

ドイツ ハンブルク州DPAのディスカッションペーパー（2024年7月）

- LLMの処理とデータ保存：従来のデータシステムとは異なり、LLMはトークンとベクトル関係（埋め込み）を処理するが、ハンブルク州DPAは、これらはGDPRにおける個人データの「処理」または「保存」には該当しないと主張している。個人データの処理が行われなければ、GDPRは適用されないため、このような整理ができればAI事業者にとってのメリットは計り知れないのだが……
- トークン化と個人データ：LLMにおけるトークンと埋め込みは、欧州司法裁判所の判例法で個人データとみなされるために必要とされる、個人との直接的な識別可能なリンクを持たない。この側面は、LLMの運用を通じて証明される必要がある。

VI. AI（LLM）に対する、欧州のデータ保護当局の姿勢

ドイツ ハンブルク州DPAのディスカッションペーパー（2024年7月）

- 暗記攻撃：LLMからトレーニングデータを抽出することは可能であるが、これらの攻撃は多くの場合非現実的であり、また法的に問題があるため、現行の法律では個人データの特定が常に可能であるとは限らない。つまり、通常の運用においては、個人データの処理は発生しないはずである。
- パーソナルデータを誤って学習してしまったLLMの利用の合法性：LLMの開発中に個人情報が入って取り扱われたとしても、その結果として生じたモデルの使用が必ずしも違法となるわけではないため、サードパーティのモデルを導入する人々にとっては安心材料となる。

VI. AI（LLM）に対する、欧州のデータ保護当局の姿勢

EDPB「AIモデルにおける個人データ処理 に関する意見」（2024年12月）

- ①いかなる場合、AIモデルが「匿名」と見なすことができるか
- ②いかなる場合、GDPRが求める、AIモデルの開発や利用時の適法性の根拠として「正当な利益」に依拠できるか
- ③依拠可能な適法性の根拠のない「違法に処理された個人データ」を用いて学習したモデルを引き続き利用することができるか

※以下56頁目までの和訳は「EDPBの「AIモデルにおける個人データ処理に関する意見書」について」（小泉雄介）を参照している。

VI. AI（LLM）に対する、欧州のデータ保護当局の姿勢

いかなる場合、AIモデルが「匿名」と見なすことができるか？

- これまで日本やEUにおいては、学習データに個人情報が含まれていても、通常、作成されたAIモデルは個人情報に該当しないとされてきた。
- 一部のAIモデルは、モデルの学習に用いられた個人に関する個人データを提供するように、又は何らかの方法でそのようなデータを利用できるように特別に設計されている。このような場合、そのようなAIモデルには、識別された又は識別可能な自然人に関する情報が本質的に（通常は必然的に）含まれ、したがって個人データの処理を伴う。したがって、タイプのこれらのAIモデルは匿名であるとは見なされない。これは例えば、以下の場合に当てはまる。（29項）

VI. AI（LLM）に対する、欧州のデータ保護当局の姿勢

- (i) 個人の音声録音に基づいて微調整され、その個人の声を模倣する生成モデルの場合。
- (ii) 特定の人物に関する情報を求められたときに学習した個人データを利用して応答するように設計されたモデルの場合。
- その他の場合でも、個人データを含む「学習データセットの情報が、AIモデルのパラメータに吸収」されたままの可能性、すなわち数学的オブジェクトを通じて表現されている可能性がある。AIモデルのパラメータは、元の学習データそのものとは異なるとしても、最終的にはモデルから直接的または間接的に元の学習データである個人データが抽出または取得される可能性がある。このような状況が見られる場合、当該AIモデルは匿名とは言えなくなる。（31項）

VI. AI (LLM) に対する、欧州のデータ保護当局の姿勢

- AIモデルには通常、直接分離またはリンクされる可能性のあるレコードは含まれず、むしろAIモデルの学習データ間の確率的関係を表すパラメータが含まれる。これは、個人データをAIモデルから推論できる可能性があるというのが合理的なシナリオの範疇であることを意味する。データ保護監督機関が特定のAIモデルについて匿名性ありと判断するためには、少なくとも以下に関する十分な証拠が揃っているべきである。（38項）
- (i) 学習データに関連する個人データがAIモデルから抽出できない。（抽出には、クエリインターフェイスをほとんど又は全く使用せずに、AIモデル自体から個人データが導き出される場合が特に含まれる。）
- (ii) AIモデルをクエリしたときに生成される出力が、モデルの学習にデータが用いられた個人に関連しない。

VI. AI（LLM）に対する、欧州のデータ保護当局の姿勢

- 上記の条件を満たしていること（個人識別の可能性がないこと）の評価は、管理者や第三者が個人を識別するために「合理的に使用する可能性のあるすべての手段」を考慮して行われるべきである（41項）。
- この評価を行うために、監督機関は、AIモデルの匿名性を実証するために管理者から提供された文書を確認する必要がある。匿名性の実証方法には、学習に用いられる個人データの収集の防止・制限、個人データの識別可能性の低減性推論・メンバーシップ推論、個人データの抽出の防止、最先端の攻撃に対する耐性（(i) 属(ii) 窃盗(exfiltration)、(iii) 学習データの逆流(regurgitation)、(iv) モデル反転、(v) 再構築攻撃）などの、モデルの開発フェーズ中に管理者が取るアプローチが含まれる。（3.2.2節）

Ⅶ. (ご参考) 顔認識技術とAI規制

- 顔認証技術企業クリアビューAIは、2025年3月21日、米国連邦裁判官シャロン・ジョンソン・コールマンが集団訴訟のプライバシー訴訟における型破りな和解案を承認したことで、重要な法的な節目を迎えた。シャロン・ジョンソン・コールマン判事は、イリノイ州の集団訴訟における異例の和解案を最終承認した。
- この和解案では、原告に現金補償ではなく、約5000万ドル相当の同社株式の23%が与えられる。この取り決めは、以前は州検事総長の連合から激しい反対に直面していた。この訴訟では、クリアビュー社が何十億もの顔画像をインターネットから同意なしに収集し、このデータベースへのアクセスを法執行機関に販売したことが、イリノイ州のバイオメトリックプライバシー法違反にあたると主張された。

Ⅶ. (ご参考) 顔認識技術とAI規制

- 中国は、2025年6月1日に発効する顔認証技術に関する包括的な規制を発表した。これにより、生体認証データの収集に関する世界で最も詳細な法的枠組みが構築されることになる。同日、物議を醸した米国とClearview AIの和解により、この分野における世界的な規制格差が浮き彫りになった。
- 「顔認証技術の応用における安全管理措置」では、クリアビュー社のようなデータ収集方法を事実上禁止する要件が定められている。企業は顔データの処理を行う前に、「自発的かつ明確な個別同意」を得なければならないと定められている。



中崎 尚

takashi.nakazaki_grp@amt-law.com

1998年 東京大学法学部卒
 2001年 弁護士登録
 2008年 米国Columbia University School of Law (LL.M.) 卒
 2008年9月
 -2009年6月 米国Washington D.C.のArnold & Porter法律事務所勤務
 2025年 当事務所パートナー就任

2010年～AIPPI（国際知的財産保護協会）編集委員、2016年～経済産業省「経済産業省・総務省 IoT推進コンソーシアム データ流通推進WG」委員、2018年～経済産業省「AI・データの利用に関する契約ガイドライン検討会」委員、2019年～「エンターテインメント・ローヤーズ・ネットワーク（ELN）」幹事、2020年～経済産業省「AI社会実装ガイド・ワーキンググループ」委員、2022年～内閣府「メタバース上のコンテンツ等をめぐる新たな法的課題への対応に関する官民連携会議」構成員、2022年～経済産業省「AI 原則実践のためのガバナンス・ガイドライン」ワーキンググループ構成員、2023年～経済産業省「AIガバナンスのルールに関する調査研究及び検討会運営」有識者検討委員会委員。国内外の個人情報保護案件・IT・インターネット案件・著作権・クロスボーダー案件等を広く取り扱う。その他、多数の技術分野の事件にも豊富な経験を有する。

【著作物・講演等】

『生成AI法務・ガバナンス』（商事法務 2024年）『Q&Aで学ぶGDPRのリスクと対応策（第2版）』（商事法務 近刊）、『Q&Aで学ぶメタバース・XRビジネスのリスクと対応策』（商事法務 2023年）、『生成AIの出力結果について、AI提供事業者の責任を認めた世界初の裁判例（広州ウルトラマン事件）』（NBL 1264号（2024.4.15））、『生成AIをめぐる米国・中国における近時の裁判状況』（NBL 1258号（2024.1.15））ほか個人情報・メタバース・ビッグデータ・SNS・決済ビジネス・システム開発・AIを中心に多数。

【受賞】

- ・ Who's Who Legal Global Guide 2018-2024: Recommended (IT)
- ・ Who's Who Legal Global Guide 2018-2024: Recommended (Telecoms & Media)
- ・ Best Lawyers 2020 -2024 (Information Technology Law, Japan)

【非営利活動】

- ・ International Bar Association (IBA) Technology Committee, Officer (2018-)
- ・ IAPP Asia Advisory Board (2025-)