

EU AI 法ナビゲート: グローバル企業のための AI Pact および行動規範に関する戦略的分析

Gemini Deep Research

エグゼクティブサマリー

欧州連合(EU)の AI 法(AI Act)は、人工知能(AI)に関する世界初の包括的な水平的規制であり、データプライバシーにおける GDPR のように、信頼できる AI のグローバルスタンダードを確立することを目指しています。本法は、EU 市場に AI システムを提供する、あるいはその出力が EU 域内で利用されるすべての企業に適用される広範な域外適用性を有しており、日本企業を含む非 EU 企業にとって、その遵守は喫緊の経営課題となります。

本報告書は、EU AI 法の複雑な規制体系、特にユーザーの関心事である「行動規範」について、詳細かつ戦略的な分析を提供するものです。この「行動規範」は、主に 2 つの異なるメカニズムを指します。

1. **AI Pact(AI 協定)**: 法的拘束力のない自主的なイニシアチブであり、欧州委員会が企業に対し、法的期限に先駆けて AI 法の要件を早期に実施することを奨励するものです。これは、企業の積極的な姿勢を示し、ブランド価値を高める機会となる一方で、市場からの参加圧力という側面も持ち合わせています。
2. **Codes of Practice(実践規範、COP)**: AI 法自体に定められた正式な共同規制ツールです。特に汎用 AI(GPAI)モデルの提供者が、承認された COP を遵守することで、特定の法的義務への「適合性の推定」を得ることができます。これは、事実上のコンプライアンス遵守の指針となる重要な文書です。

本報告書は、これらの「行動規範」を明確に区別し、その目的、法的地位、戦略的価値を分析します。さらに、AI 法の根幹をなすリスクベースアプローチ(「許容できないリスク」「高リスク」「限定的なリスク」「最小リスク」の 4 分類)を解き明かし、特に最も厳格な義務が課される「高リスク AI」および「システミックリスクを伴う GPAI」の提供者が負う具体的な責任を詳述します。

また、EU の規制中心アプローチを、米国の市場主導・自主規制アプローチ、中国の国家主導・安全保障重視アプローチ、そして日本の産業主導・ソフトローアプローチと比較分析することで、グローバル企業が直面する「コンプライアンスの断片化」という戦略的課題を浮き彫りにします。

最後に、これらの詳細な分析に基づき、当社が取るべき具体的な行動計画を、評価・マッピング(即時)、ガバナンス構築(短期)、実装・文書化(中期)、戦略的エンゲージメント・継続的監視(長期的)の4段階に分けて提言します。本報告書は、当社の経営陣がEU AI法という新たな規制環境を的確にナビゲートし、法的・財務的リスクを軽減し、戦略的機会を特定するための羅針盤となることを目的としています。

セクション I: EU AI 法: AI 規制の新たなグローバルベンチマーク

本セクションでは、AI法の基礎となる法的・概念的枠組みを確立します。これは、「行動規範」が機能する文脈を理解するために不可欠です。

1.1 AIにおける「ブリュッセル効果」: グローバルスタンダードの確立

EU AI法は、AIに関する世界初の包括的かつ水平的な規制として位置づけられています。これは、データプライバシーの分野で一般データ保護規則(GDPR)が果たした役割と同様に、信頼できるAIに関する世界的な基準を設定することを目指すものです¹。この法律の主な目的は、EU市場に出回るAIシステムが安全であり、基本的権利とEUの価値を尊重することを保証すると同時に、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提供することにあります³。

AI法は、EU加盟国による国内法化を必要とせず、全加盟国に直接的かつ統一的に適用される「規則(Regulation)」として制定されました。これにより、AIアプリケーションのための分断されていない単一のEU市場の発展が促進されます⁴。このアプローチは、EUが設定した基準が、グローバルに事業を展開する企業にとって事実上の世界標準となる「ブリュッセル効果」をAIの領域でもたらす可能性を秘めています。

1.2 リスクベースアプローチ: 4つの規制階層の解体

AI法の中核的なアーキテクチャは、AIシステムがもたらす潜在的なリスクのレベルに応じて規制の厳格さを調整する「リスクベースアプローチ」です。システムは以下の4つのカテゴリーに分類されます³。

- **許容できないリスク(Unacceptable Risk):** EUの価値観に反し、基本的権利を侵害するようなAIの利用は原則として禁止されます。具体例としては、個人の潜在意識を操作する技術、脆弱性を悪用するシステム、公的機関によるソーシャルスコアリング、法執行目的での公共空間におけるリアルタイム遠隔生体認証システムの大部分などが含まれます³。これらの禁止措置は、2025年2月から適用が開始されます¹³。

- 高リスク(High Risk):** 規制の主要な対象であり、個人の健康、安全、または基本的権利に重大なリスクをもたらす可能性のある AI システムが該当します。これらは主に 2 つのカテゴリーに分けられます。
 - EU の製品安全法規(例: 医療機器、玩具、自動車)の対象となる製品の安全コンポーネントとして使用される AI システム⁸。
 - 附属書 III(Annex III)にリストアップされた特定の重要分野で使用されるスタンドアロン AI システム。これには、生体認証、重要インフラの管理、教育・職業訓練、雇用・労働者管理、必要不可欠な公的・民間サービスへのアクセス(信用スコアリングや保険料算定を含む)、法執行、移民・国境管理、司法・民主的プロセスなどが含まれます⁴。
- 限定的なリスク(Limited Risk / Specific Transparency Risk):** 特定の透明性義務が課される AI システムです。これには、利用者が AI と対話していることを開示する必要があるチャットボット、利用者に通知する必要がある感情認識システムや生体認証分類システム、そして人工的に生成されたコンテンツであることを明示的にラベル付けする必要があるディープフェイクなどが含まれます³。
- 最小リスク(Minimal Risk):** 上記のいずれにも該当しないその他の AI システム(例: AI 搭載のビデオゲームやスパムフィルター)です。これらのシステムは基本的に規制の対象外ですが、自主的な行動規範の遵守が奨励されています⁷。

このリスクベースアプローチの運用において、企業のコンプライアンス戦略上、極めて重要な点が存在します。それは、「高リスク」の定義が静的なものではないという事実です。規制の中核をなすのは、高リスクなユースケースを列挙した附属書 III です⁸。欧州委員会は将来、このリストを更新する権限を有しています¹⁷。これは、現在「最小リスク」に分類されている製品が、将来の法改正によって「高リスク」に再分類される可能性があることを意味します。したがって、企業は自社の現行製品ポートフォリオを評価するだけでなく、附属書 III の改正につながりうる政治的・社会的動向を継続的に監視し、コンプライアンス体制を動的に調整し続ける必要があります。コンプライアンスは一回限りのチェックリストではなく、恒久的なリスク監視活動となるのです。

リスクレベル	定義と具体例	主な義務	適用開始時期
許容でき	EU の価値観に反し、基本的権利を侵害する AI。例: ソーシャルスコ	禁止	2025 年 2 月

リスクレベル	定義と具体例	主な義務	適用開始時期
ないリスク	アリング、職場での感情認識、法執行によるリアルタイム遠隔生体認証（一部例外あり）。		
高リスク	健康、安全、基本的権利に重大なリスクをもたらす AI。例：医療機器、採用選考 AI、信用スコアリング、重要インフラの安全部品。	リスク管理、データガバナンス、技術文書作成、人的監視、適合性評価、EU データベースへの登録など、厳格な要件。	2026 年 8 月 / 2027 年 8 月
限定的なリスク	利用者を欺くリスクがある AI。例：チャットボット、ディープフェイク、感情認識システム。	透明性義務（AI との対話であることの通知、生成コンテンツのラベリングなど）。	2026 年 8 月
最小リスク	上記のいずれにも該当しない AI。例：スパムフィルター、AI 搭載ビデオゲーム。	法的義務はなし。自主的な行動規範の遵守を推奨。	N/A

出典：³

1.3 域外適用：日本企業および非 EU 企業への影響

AI 法は、その適用範囲が地理的に限定されていません。提供者の所在地にかかわらず、EU 市場に AI システムを上市 (place on the market) またはサービス提供 (put into service) するすべての提供者（プロバイダー）に適用されます¹。また、EU 域内に所在する AI システムの導入者（デプロイヤー）にも適用されます。

特に日本企業にとって重要なのは、AI システムの「出力 (output)」のみが EU 域内で利用される場合にも、この法律が適用される可能性があるという点です⁴。例えば、日本国内で運用されている AI サービスであっても、その分析結果や生成コンテンツが EU 域内の顧客に提供される場合、その AI システムの提供者は AI 法の規制対象となる可能性があります。この広範な域外適用性は、グローバルにサービスを展開する企業にとって、重大なコンプライアンス上の課題を提起します。

1.4 AI バリューチェーンにおける主要な主体：義務の定義

AI 法は、AI のバリューチェーンに関与する様々な経済事業者 (economic operator) に対して、それぞれの役割に応じた義務を課しています¹。

- **提供者 (Provider)**: AI システムまたは汎用 AI モデルを開発し、自らの名称または商標の下で EU 市場に上市またはサービス提供する自然人または法人を指します¹⁵。提供者は、特に高リスク AI システムに関して、コンプライアンス義務の大部分を負います。
- **導入者 (Deployer / User)**: 個人的・非職業的な活動の過程を除き、自らの権限の下で AI システムを利用する自然人または法人です¹²。高リスク AI システムの導入者は、人的監視の実施やシステムの監視など、特定の義務を負います³。
- **輸入者 (Importer) および 販売者 (Distributor)**: EU 域外から AI システムを EU 市場に上市する者、および市場で利用可能にする者を指します。彼らは、提供者が適切な適合性評価手続きを実施し、必要な文書を作成しているかなどを確認する義務を負います¹。

ここで注意すべきは、「導入者」と「提供者」の区別が流動的であるという点です。第三者から高リスク AI システムを調達した導入者が、そのシステムに自社ブランドを付して提供したり、元の提供者が意図しない目的のためにシステムを大幅に変更したりした場合、その導入者は AI 法の下で「提供者」とみなされ、提供者としての全ての重い法的責任を負うこととなります¹⁹。これは、AI システムの調達、導入、およびブランディング戦略において、隠れたコンプライアンスリスクを生み出す可能性があります。マーケティング部門による安易なホワイトラベルの決定が、会社全体に予期せぬ巨大なコンプライアンス負担を引き起こすことになりかねません。

1.5 EU AI オフィス: ガバナンスと執行の中核機関

AI 法は、欧州委員会内に新たな専門機関として「欧州 AI オフィス (EU AI Office)」を設立します¹。AI オフィスの主な機能は以下の通りです。

- AI 法の効果的な実施の監視と、加盟国間の統一的な適用促進。
- 加盟国の監督当局との連携・調整。
- 標準化プロセスや実践規範 (Codes of Practice) の策定奨励。
- 特に汎用 AI (GPAI) モデルに関する規則の監督と執行¹。

AI オフィスは、EU レベルでの AI ガバナンスの中核を担い、規制の調和と一貫した執行を確保する上で中心的な役割を果たします。

セクション II: 「行動規範」の解説: AI Pact と実践規範

本セクションでは、ユーザーのクエリに直接的に応え、AI 法に関連する 2 つの主要な「行動規範」を分析し、それぞれの明確な役割を明らかにします。

2.1 AI Pact: コンプライアンスへの自主的な架け橋

2.1.1 目的と性質

AI Pact (AI 協定) は、欧州委員会が立ち上げた、法的拘束力のない自主的なイニシアチブです²²。その主な目的は、AI 法の各規定が法的に適用開始となる前に、企業や組織がその準備を進め、AI 法の措置を早期に実施することを奨励し、支援することにあります¹¹。これは、2026 年に予定されている本格的な法施行までの移行期間を埋め、円滑な制度移行を促すための「橋渡し」としての役割を担います。

このイニシアチブは、単なる広報活動以上の戦略的な意味合いを持ちます。欧州委員会は AI Pact を通じて、早期導入者のコミュニティを形成することで、①実際の導入における課題に関する実世界からの情報を収集し、②法的拘束力を持つ前にガイダンスをテスト・改良し、③コンプライアンスを標準化して遅れている企業への圧力をかける「意欲ある有志連合」を創出することができます²⁴。これは、規制当局が規制対象者から学ぶことで、2026 年の本格施行のリスクを低減させる、マクロレベルでのアジャイル・ガバナンスの一形態と言えます。

2.1.2 構造 (2 つの柱)

AI Pact は、2 つの主要な柱で構成されています²⁴。

- **第 1 の柱: ネットワークと知識共有:** AI オフィスがウェビナーやイベントを主催し、AI 法の目的についての共通理解を深め、関係者が準備を進めるための支援を行います。これにより、ベストプラクティスを共有するための実践的なコミュニティが形成されます。
- **第 2 の柱: 企業の誓約の促進と伝達:** これが AI Pact の中核です。企業は、「エンゲージメント宣言 (declarations of engagement)」という形で自主的なコミットメント (誓約) を行い、コンプライアンスをどのように先取りしているかを公に表明することができます。2024 年 9 月時点で、既に 200 社以上の企業がこの誓約に署名しています²²。

2.1.3 誓約: 中核となるコミットメント

企業が AI Pact で行う誓約には、基本的なものと、より野心的な追加のものが含まれます。

- **基本的な誓約:** 以下の 3 つが中核となります²²。
 1. 組織内での AI 導入を促進し、将来的な AI 法遵守に向けた **AI ガバナンス戦略**の採択。

2. AI 法の下で高リスクに分類される可能性のある AI システムの特定とマッピング。
 3. 倫理的で責任ある AI 開発を保証するための、従業員の AI に対する認識とリテラシーの促進。
- **追加の誓約:** 企業はさらに、高リスク AI システムに求められる要件に関連する、より具体的な誓約を行うことができます。これには、人的監視、リスク軽減、ディープフェイクのような AI 生成コンテンツの透明性確保などが含まれます²⁴。

2.1.4 戦略的計算:エンゲージメントのビジネスケース

AI Pact への参加は任意ですが、企業にとっては戦略的な判断が求められます。

- **メリット:** AI Pact への参加は、企業が積極的かつ責任ある AI 開発に取り組んでいることを示すシグナルとなり、ブランドの評判を高め、顧客や投資家からの信頼を構築する上で有利に働く可能性があります²⁵。また、規制が完全に固まる前に、自社のソリューションをテストし、幅広いコミュニティと知識を共有することで、業界のベストプラクティスの形成に影響を与える機会も得られます²⁴。これは、規制の潮流に先んじるための戦略的手段と見なすことができます²⁷。
- **リスクと圧力:** 自主的な取り組みであるにもかかわらず、特に業界の主要企業にとっては、AI Pact に参加しないことが評判上のリスクとなるような、市場からの強い圧力が存在する可能性があります¹¹。また、早期にコミットメントを行うことは、最終的な規則や標準が確定する前にコンプライアンスのためのリソースを割り当てることを意味し、将来的にプロセスを再構築する必要が生じるリスクも伴います¹¹。

2.2 実践規範(Codes of Practice): 正式な適合性推定への道

2.2.1 法的根拠と機能

AI Pact とは対照的に、実践規範(Codes of Practice, COP)は、AI 法自体に定められた正式な共同規制の手段です。承認された COP を遵守することにより、提供者は特定の法的義務に対する「適合性の推定(presumption of conformity)」を享受することができます。これは特に、汎用 AI(GPAI)モデルの提供者にとって、コンプライアンスを証明するための重要な手段となります¹⁰。

2.2.2 共同策定プロセス

AI オフィスは、産業界、市民社会、学界、その他の利害関係者と協力して、COP の策定を奨励・促進する役割を担います⁸。このプロセスには、草案の作成、パブリックコメントの募集、そして最終的な欧州委員会による承認が含まれます。この共同策定

プロセスは、規制が技術の最先端と乖離しないようにするための重要なメカニズムです。

2.2.3 GPAI への焦点

COP は、特に汎用 AI (GPAI) モデルの提供者にとって極めて重要です。AI 法の条文は高レベルの原則を定めるに留まりますが、COP はモデル評価、リスク管理、透明性確保といった義務を実際にどのように履行するかについての、具体的かつ最先端の技術的詳細を規定することになります¹⁰。AI 法は、提供者が準備期間を確保できるよう、これらの COP を法の施行後 9 ヶ月以内 (2025 年 5 月頃まで) に準備が整うよう求めています¹²。

GPAI モデルの提供者にとって、COP は AI 法の条文そのものよりも重要な文書となるでしょう。なぜなら、COP は「モデル評価」や「リスク軽減」といった抽象的な法的義務を、具体的なエンジニアリング要件や文書化プロセスに変換する運用上の翻訳機となるからです。COP の草案作成プロセス²¹を見れば、そこには「6 ヶ月ごとの再評価」や「訓練データの詳細」といった具体的な要求事項が含まれることがわかります。このプロセスに参加する企業は、自社および競合他社が準拠すべき「最先端技術 (state of the art)」の基準形成において、大きな影響力を持つことができます。一方、参加しない企業は、他社によって設定された基準に適応せざるを得ないリスクを負うことになります。

2.2.4 標準との違い

COP は、CEN-CENELEC のような欧州標準化機関によって策定される正式な整合規格 (harmonised standards) とは区別されます。COP は、GPAI のような新興分野において、より迅速かつ柔軟に策定されることを意図しています。一方で、整合規格は、高リスク AI システムのより確立された技術要件 (例: サイバーセキュリティ、正確性) をカバーすることが想定されています⁷。

特徴	AI Pact (AI 協定)	実践規範 (Codes of Practice, COP)
目的	AI 法の早期かつ自主的な準備・実施の促進	法的義務遵守の正式な証明
法的地位	法的拘束力のない自主的な誓約	「適合性の推定」を生む正式な共同規制ツール
主な対象	全ての AI 関係者、特に「先進企業」	主に GPAI モデル提供者、高リスク AI 提供者

特徴	AI Pact (AI 協定)	実践規範 (Codes of Practice, COP)
主な内容	高レベルのコミットメント (ガバナンス、マッピング、リテラシー)	特定の法的義務を満たすための詳細かつ技術的な方法論
管轄機関	欧州委員会 / AI オフィス (促進者として)	AI オフィス (調整・承認者として)
タイムライン	現在から AI 法適用開始まで活動	法施行後に策定、初版は 2025 年 5 月頃までに完成予定
戦略的価値	評判構築、知識共有、ベストプラクティスへの影響力行使	法的コンプライアンスとリスク軽減のための直接的なツール

出典: ⁸

セクション III: グローバル企業の義務と戦略的必須事項

本セクションでは、法的枠組みを、企業が準備しなければならない具体的な義務へと落とし込み、特に規制が厳しい分野に焦点を当てて解説します。

3.1 高リスク AI システム: コンプライアンス要件の詳細

高リスク AI システムの提供者と導入者は、AI 法の最も厳格な要件の対象となります。これらの義務はシステムのライフサイクル全体にわたって適用され、一回限りの認証プロセスではなく、継続的な運用活動が求められます。

- リスク管理システム:** 提供者は、AI システムのライフサイクル全体を通じて、リスク管理システムを確立、導入、文書化し、維持しなければなりません ⁹。これには、既知および合理的に予見可能なリスクを特定・分析し、それらを排除または最小化するための措置を講じることが含まれます。
- データとデータガバナンス:** AI システムの訓練、検証、テストに使用されるデータセットは、意図された目的に対して関連性があり、十分に代表的で、可能な限りエラーがなく、完全でなければなりません。偏見を助長する可能性のあるデータについては、その検出、防止、および軽減のための措置が求められます。厳格なデータガバナンスと管理慣行の確立が不可欠です ⁷。

- **技術文書と記録保持:** 提供者は、システムを市場に投入する前に、AI 法への適合性を証明するための詳細な技術文書を作成し、常に最新の状態に保つ必要があります⁵。また、システムは、その運用のトレーサビリティを確保するために、イベントを自動的に記録(ロギング)する機能を備えていなければなりません⁷。
- **透明性と利用者への情報提供:** 導入者が AI システムを理解し、適切に使用できるよう、提供者はシステムの意図された目的、能力、限界、性能に関する明確かつ包括的な使用説明書を提供しなければなりません¹⁰。
- **人的監視:** システムは、人間による効果的な監視が可能となるように設計・開発される必要があります。これには、システムの運用を監視、判断、介入するための適切なヒューマン・マシン・インターフェースツールが含まれます³。導入者は、この監視が十分な訓練を受け、権限を与えられた担当者によって行われることを保証する責任があります³。
- **正確性、堅牢性、サイバーセキュリティ:** システムは、そのライフサイクルを通じて一貫した性能を発揮し、適切なレベルの正確性、堅牢性、サイバーセキュリティを達成するように設計されなければなりません。これには、エラーや不整合に対する回復力や、モデルポイズニングのような悪意のある第三者による攻撃に対する防御策が含まれます⁵。
- **適合性評価と登録:** 高リスク AI システムは、市場に投入される前に適合性評価手続きを経る必要があります⁴。多くの場合、これは提供者による自己評価が可能ですが、特定の極めて重要なシステム(例:遠隔生体認証)については、第三者機関による評価が義務付けられます⁸。全ての高リスク AI システムは、EU が管理する公開データベースに登録されなければなりません⁵。
- **基本的権利影響評価(FRIA):** 高リスク AI システムを導入する者は、そのシステムの利用が基本的権利に与える影響を、利用開始前に評価することが求められます³。

3.2 汎用 AI(GPAI)モデル:規制の階層をナビゲートする

AI 法は、特定の目的を持たず、多様な下流タスクに適応可能な GPAI モデル(生成 AI の基盤モデルが典型例)に対して、新たな規制の層を導入しました。

3.2.1 全ての GPAI 提供者の基本義務

全ての GPAI モデルの提供者は、以下の義務を負います。

- AI オフィスの要請に応じて、モデルの訓練プロセスや評価結果を含む詳細な技術文書を維持し、提供すること¹⁰。

- 自社の GPAI モデルを AI システムに統合する下流の提供者が、自らの義務を果たすために必要な情報と文書を提供すること¹⁰。
- EU の著作権法を遵守するためのポリシーを策定し、実行すること¹⁰。
- モデルの訓練に使用されたコンテンツに関する「十分に詳細な」概要を作成し、公開すること¹⁰。

これらの義務は、AI のサプライチェーンにおける透明性を確保することを目的としています。GPAI 提供者から下流の AI システム提供者への情報提供義務は、AI のサプライチェーンがコンプライアンスの連鎖であることを示しています。下流の企業は、自社の高リスク AI システムの適合性評価を完了するために、上流の GPAI 提供者からの情報に完全に依存することになります。このため、サプライチェーンにおけるデューデリジェンスは、単なるビジネス上のベストプラクティスから、法的な必須要件へと格上げされます。

3.2.2 「システミックリスク」の閾値: ハイインパクトモデルへの追加義務

GPAI モデルの中でも特に強力で広範な影響を及ぼす可能性のあるモデルは、「システミックリスク」を持つと見なされ、さらに厳しい義務が課されます。

- **定義:** モデルの訓練に使用された累積計算量が 1025 FLOPs (浮動小数点演算回数) の閾値を超える GPAI モデルは、システミックリスクを持つと推定されます¹²。また、欧州委員会は、ユーザー数や影響の大きさなど他の基準に基づいて、モデルをシステミックリスク指定することもできます⁸。
- **追加義務:** システミックリスクを持つモデルの提供者は、基本義務に加えて、以下のことを行う必要があります。
 - 敵対的テスト(レッドチーミング)を含む、最先端のモデル評価を実施すること¹⁰。
 - EU レベルでの潜在的なシステミックリスクを評価し、軽減するための措置を講じること⁵。
 - 重大なインシデントを追跡、文書化し、AI オフィスや関連する国内当局に遅滞なく報告すること⁵。
 - モデルとその物理インフラに対する、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¹⁰。

3.2.3 施行タイムライン

GPAI モデルに関する義務は、法の発効から 12 ヶ月後、すなわち 2025 年 8 月から適用が開始されます¹²。

3.3 透明性義務: 限定リスクシステムと AI 生成コンテンツ

限定的なリスクを持つと分類された AI システムには、主に透明性に関する義務が課されます。

- チャットボットのように人間と対話することを意図したシステムの提供者および導入者は、個人が AI システムと対話していることを通知されるようにしなければなりません³。
 - 感情認識または生体認証分類システムの導入者は、そのシステムにさらされる個人にその旨を通知する義務があります³。
 - AI によって生成または操作された画像、音声、動画、テキストコンテンツ(いわゆるディープフェイクを含む)は、それが人工的なものであることを示すラベルを、機械可読な形式で付与しなければなりません。特に、実在の人物、場所、出来事を模倣したディープフェイクコンテンツには、その人工的な性質を明確に開示する義務があります³。
-

セクション IV: グローバルな規制のモザイク: 比較分析

本セクションでは、EU のアプローチを他の主要なグローバルプレイヤーと比較し、多国籍企業が航海しなければならない断片化された規制の状況を明らかにすることで、重要な文脈を提供します。

4.1 欧州連合: 法による規制、権利中心

- **アプローチ:** 包括的、水平的、かつ法的拘束力のある「ハードロー」による枠組み⁴。
- **哲学:** 基本的権利、安全、そして EU の価値観の保護を中核に据えています³。「信頼できる AI」を厳格な規制を通じて創出することを目指します。
- **メカニズム:** リスクベースの分類(禁止、高リスク要件、透明性義務など)と、中央執行機関(AI オフィス)による監督が特徴です¹。

4.2 米国: 市場主導、自主的なフレームワーク

- **アプローチ:** 業界の自主規制を奨励する、「ソフトロー」的、非セクター的なアプローチ³⁰。
- **哲学:** 事前規制よりも、イノベーションを促進しつつリスクを管理することに重点を置いています。企業の自己規律と透明性が強調されます³¹。
- **メカニズム:** 米国国立標準技術研究所(NIST)が策定した「AI リスク管理フレームワーク(AI RMF)」が主要なツールです。これは、組織が AI リスクを特定、評価、管理するための自主的なガイドであり、法的拘束力はありませんが、契約や将来の規制に組み込まれる可能性があります³²。ガバナンス、マ

ッピング、測定、管理という4つの機能を通じて、リスク管理をAIライフサイクルに統合することを推奨しています³²。

4.3 中国:国家主導、安全保障重視、個別指令

- **アプローチ:** 2030年までに世界のAIリーダーになるという国家計画と、特定のAIアプリケーションを対象とした個別的かつ拘束力のある規制の組み合わせです³⁶。EUのような単一の統一法ではありません。
- **哲学:** 技術的リーダーシップの目標と並行して、国家による統制、社会の安定、国家安全保障の懸念が強く反映されています³¹。
- **メカニズム:** 国務院が全体戦略を設定し³⁶、国家インターネット情報弁公室(CAC)が生成AI、ディープシンセシス、推薦アルゴリズムに関する個別規制の実施を主導します³⁸。コンテンツのラベリングとトレーサビリティに重点が置かれています³⁶。包括的なAI法は計画段階にあります。が、制定には数年を要する可能性があります³⁹。

4.4 日本:産業主導、「ソフトロー」ガバナンス

- **アプローチ:** 政府による非拘束的なガイドラインに基づき、企業の自主的なガバナンスを促進する「ソフトロー」アプローチを採用しています³⁰。
- **哲学:** 規範的なハードローを避け、アジャイルで柔軟なガバナンスを通じて、イノベーションの促進と国際的な相互運用性を優先します³⁰。
- **メカニズム:** 経済産業省(METI)と総務省(MIC)が発行した「AI事業者ガイドライン」が中心的な文書です⁴³。このガイドラインは、ゴールベース、リスクベースであり、政府による直接的な強制ではなく、民間契約を通じて実施されることが意図されています⁴⁴。日本は、G7広島AIプロセスに見られるように、ガバナンスに関する国際協力を積極的に推進しています³⁰。

これらの規制哲学の衝突は、「コンプライアンスの断片化」という深刻な課題を生み出しています。EU(権利ベースのハードロー)、米国(市場ベースのソフトロー)、中国(国家ベースの安全保障法)は、根本的に異なる3つのパラダイムを代表しています。当社のような多国籍企業にとって、これは単一のAI製品であっても、これら主要市場にアクセスするためには、3つの異なる方法で設計、文書化、統治する必要があります³¹。この断片化は、法務、技術、運用上の莫大なオーバーヘッドを生み出し、主要な戦略的課題となります。グローバル企業は、単一の普遍的なAIガバナンスフレームワークを導入するのではなく、これらの相反する法的要件に適應できるモジュール式のフレームワークを構築する必要があります。

比較項目	欧州連合(EU)	米国	中国	日本
コアアプローチ	権利中心の規制	市場主導のイノベーション	国家主導の安全保障	産業主導の相互運用性
法的性質	ハードロー(規則)	ソフトロー(自主的)	ハードロー(個別指令)	ソフトロー(自主的ガイドライン)
主要なツール	AI 法、GDPR	NIST AI RMF、AI 権利章典	生成 AI 管理弁法、サイバーセキュリティ法等	AI 事業者ガイドライン、広島 AI プロセス
執行哲学	中央集権的(AI オフィス)および各国の監督当局	産業界の自主規制、FTC による不正行為への執行	中央集権的な国家統制(CAC)	企業主導、契約による執行

出典:¹

セクション V:ステークホルダーの言説と批判的視点:論争の領域を航海する

本セクションでは、法の条文を超えて、それを取り巻く継続的な議論を分析し、その潜在的な弱点や論争点についての多角的な理解を提供します。

5.1 市民社会の懸念(EDRi、BEUC):基本的権利と抜け穴

- 不十分な基本的権利の保護:** European Digital Rights (EDRi) や Access Now のような市民社会団体は、AI 法が「人権のゴールドスタンダード」を確立するには至っていないと主張しています²⁸。彼らは、多くの有害なシステムが完全な禁止ではなく、透明性義務の対象となるに留まっている点を批判しています¹⁷。

- **広範な免除と抜け穴:** 最大の批判の一つは、「国家安全保障」目的で開発・利用される AI システムに対する広範な免除規定です。これは、法の保護を骨抜きにする重大な抜け穴と見なされています²⁸。この免除は、当社が EU の政府機関や法執行機関と契約する際に、深刻な法的・評判上のリスクを生み出します。当社が供給する「民間用」高リスク AI システムが、顧客である政府によって「国家安全保障」目的で展開された場合、法のセーフガードの対象外となる可能性があり、その結果、当社は管理不能な権利侵害利用に関連付けられるリスクを負うことになります。
- **禁止規定の弱さ:** 予測的取締りや遠隔生体認証(RBI)のような有害な慣行に対する禁止規定は、法執行機関向けの例外が多すぎ、範囲が狭すぎると批判されています²⁸。「リアルタイム」と「事後(post)」の RBI の区別は人為的であり、どちらも同様に有害であると指摘されています²⁸。
- **効果的な救済措置の欠如:** AI によって損害を受けた個人に対する救済メカニズムが弱いと見なされています。特に、公益団体が被害者に代わって監督当局に苦情を申し立てる権利などの規定が欠けている点が問題視されています²⁸。
- **有害技術の輸出リスク:** AI 法は、EU 域内で禁止されている AI システムを EU 拠点の企業が開発し、域外に輸出することを妨げていません。これは、EU 製の技術によって EU 域外国の人々の権利が侵害されるリスクを生み出します²⁸。
- **消費者保護(BEUC):** 欧州消費者機構(BEUC)は、不公正な AI システムが消費者に与えるリスク、効果的な法執行の必要性、そして国際的な貿易協定が EU の基準を損なう可能性について警鐘を鳴らしています⁴⁷。また、AI オフィスに助言する科学専門家パネルの独立性と利益相反に関する厳格な規則を求めています⁴⁹。

5.2 学界と産業界の分析: イノベーション対規制の議論

- **自己評価の問題:** 多くの高リスクシステムに対するコンプライアンスが、提供者による自己評価に依存しており、義務的な第三者検証がないことが、法の執行力を弱める可能性があるという重要な批判があります¹⁷。
- **リスクの定義と線引き:** 学界からは、物質的、非物質的、社会的なものを含む全ての潜在的損害を、法的なリスクの枠組みで捉えることの難しさが指摘されています。GPAI に対する「システムックリスク」という概念はこの問題に対処しようとする試みですが、損害の線引きをさらに複雑にしています⁵⁰。
- **規制の将来性:** AI 法の急速な技術トレンドへの対応能力について懸念が示されています。高リスクリストを更新するためのより柔軟なメカニズムや、国内機関と EU 機関間の情報フローの改善などが提案されています¹⁷。

- ・ **イノベーションへの負担:** 産業界からは、特に高リスクシステムや GPAI に対する厳格な要件がイノベーションを阻害し、中小企業に不釣り合いな負担を強いる可能性があるとの懸念が頻繁に表明されています²。

これらの議論から明らかなように、AI 法のテキストがほぼ確定した今、真の戦いの場は、法律そのものから、それを具体化する標準化と実践規範(COP)の策定プロセスへと移行しています²¹。BEUC のような市民社会団体は、既に COP の草案や専門家パネルの規則に対する意見提出を開始しています⁴⁹。これは、企業にとって、議会へのロビー活動の時代は終わり、影響力を行使するための新たな重要な舞台が、AI オフィスや標準化団体が運営する技術委員会やステークホルダー協議の場になったことを意味します。

セクション VI: AI ガバナンスとコンプライアンスのための戦略的提言

本最終セクションでは、これまでの分析を統合し、当社が取るべき段階的な戦略的ロードマップを提示します。

6.1 フェーズ 1: 評価とマッピング(即時行動)

- ・ **全社的な AI インベントリの作成:** 全事業部門で開発・導入されている全ての AI システム(サードパーティ製システムを含む)を網羅したカタログを作成します。
- ・ **予備的なリスク分類の実施:** AI 法の基準⁶を用い、インベントリ内の各システムを 4 つのリスク階層に予備的に分類します。特に附属書 III にリストされたユースケース⁴に細心の注意を払います。
- ・ **AI サプライチェーンのマッピング:** 全ての上流提供者(GPAI モデル、データセット提供者など)と下流導入者を特定し、法的な依存関係の連鎖を可視化します。これは、セクション III で指摘したサプライチェーンリスクを管理するために不可欠です。

6.2 フェーズ 2: ガバナンスアーキテクチャの構築(短期的優先事項)

- ・ **部門横断的な AI ガバナンス組織の設立:** 法務、コンプライアンス、研究開発、サイバーセキュリティ、および各事業部門の代表者からなる委員会を設立し、AI 戦略とリスクを統括します。この組織は、経営層の強力な後援を得る必要があります。
- ・ **AI リスクの全社的なリスク管理(ERM)への統合:** AI リスクをサイロ化させてはなりません。AI リスク評価の結果を、サイバーセキュリティやデータプライバシー

ーに関する既存のプロセスと連携させ、当社の全社的リスク管理(ERM)フレームワークに統合します⁸。

- **社内 AI 倫理・コンプライアンスポリシーの策定:** AI 法の原則(人間中心、公平性、透明性など)に基づき、全ての AI 開発・調達の指針となる、拘束力のある社内ポリシーを策定します。その際、日本の「AI 事業者ガイドライン」⁴⁴ を原則の出発点として参照します。

6.3 フェーズ 3: 実装と文書化(中期的プロジェクト)

- **コンプライアンス準拠インフラの構築:** 高リスクと特定されたシステムについて、要求される技術的・手続き的インフラ(リスク管理システム、データガバナンスプロトコル、ロギングメカニズム、人的監視インターフェース)の構築を開始します¹⁰。
- **技術文書の標準化:** 高リスクシステムおよび GPAI モデルに要求される広範な技術文書を生成・維持するためのテンプレートとプロセスを確立します⁵。
- **イノベーション支援ツールの活用:** 利用可能な場合、「AI 規制サンドボックス」を活用し、規制当局の監督下にある管理された環境で高リスクシステムをテストすることを検討します。これにより、新規アプリケーションのリスクを低減できる可能性があります²。

6.4 フェーズ 4: 戦略的エンゲージメントと継続的監視(恒久的活動)

- **AI Pact への参加に関する戦略的決定:** AI Pact への参加がもたらす評判上の利益と、それに伴うリソースコストを慎重に評価します¹¹。参加しないという決定を下す場合でも、それは意識的な選択であるべきであり、準備されたコミュニケーション戦略を伴う必要があります。
- **標準化および COP 策定への関与:** 当社の事業に関連する整合規格や実践規範(COP)の策定プロセスを積極的に監視し、可能な限り参加します。これは、将来のコンプライアンス義務を形成するための極めて重要な機会です。
- **継続的な規制監視の実施:** AI を取り巻く環境は静的ではありません。AI 法(特に附属書 III)の変更、AI オフィスからの新たなガイダンス、そして執行事例を継続的に監視するプロセスを確立し、当社のガバナンスフレームワークが常に最新かつ効果的であることを保証します。