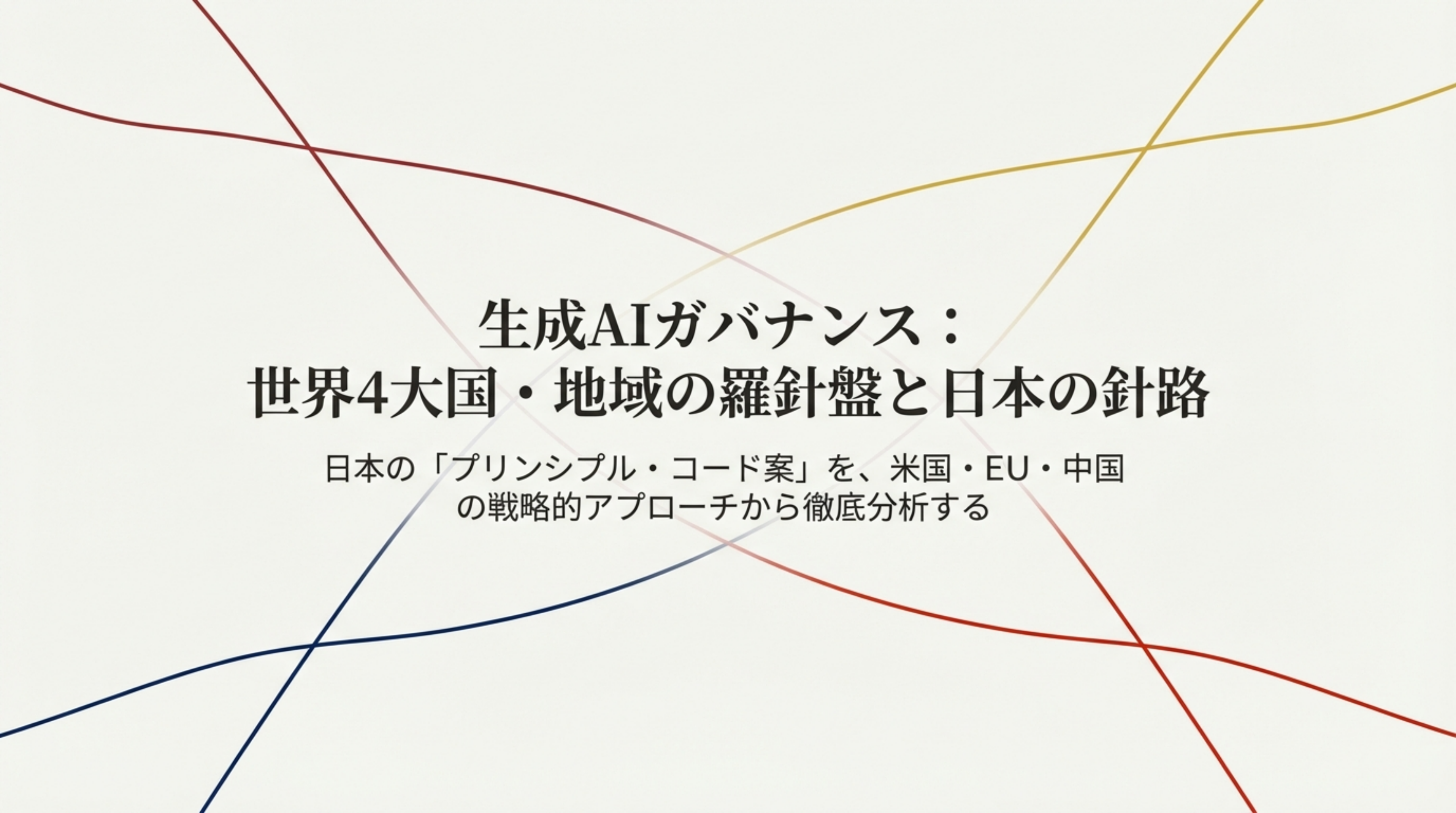




生成AIガバナンス： 世界4大国・地域の羅針盤と日本の針路

日本の「プリンシプル・コード案」を、米国・EU・中国
の戦略的アプローチから徹底分析する

全世界共通の課題：生成AIがもたらす「機会」と「脅威」の奔流

生成AI技術は急速に進化し、経済成長と社会変革の起爆剤となる一方、知的財産権の侵害、偽情報の拡散、ブラックボックス化による説明責任の欠如といった深刻なリスクを顕在化させている。



機会 (Opportunity)

- 生産性の飛躍的向上
- 新たな創造性の解放
- 科学技術研究の加速



脅威 (Threat)

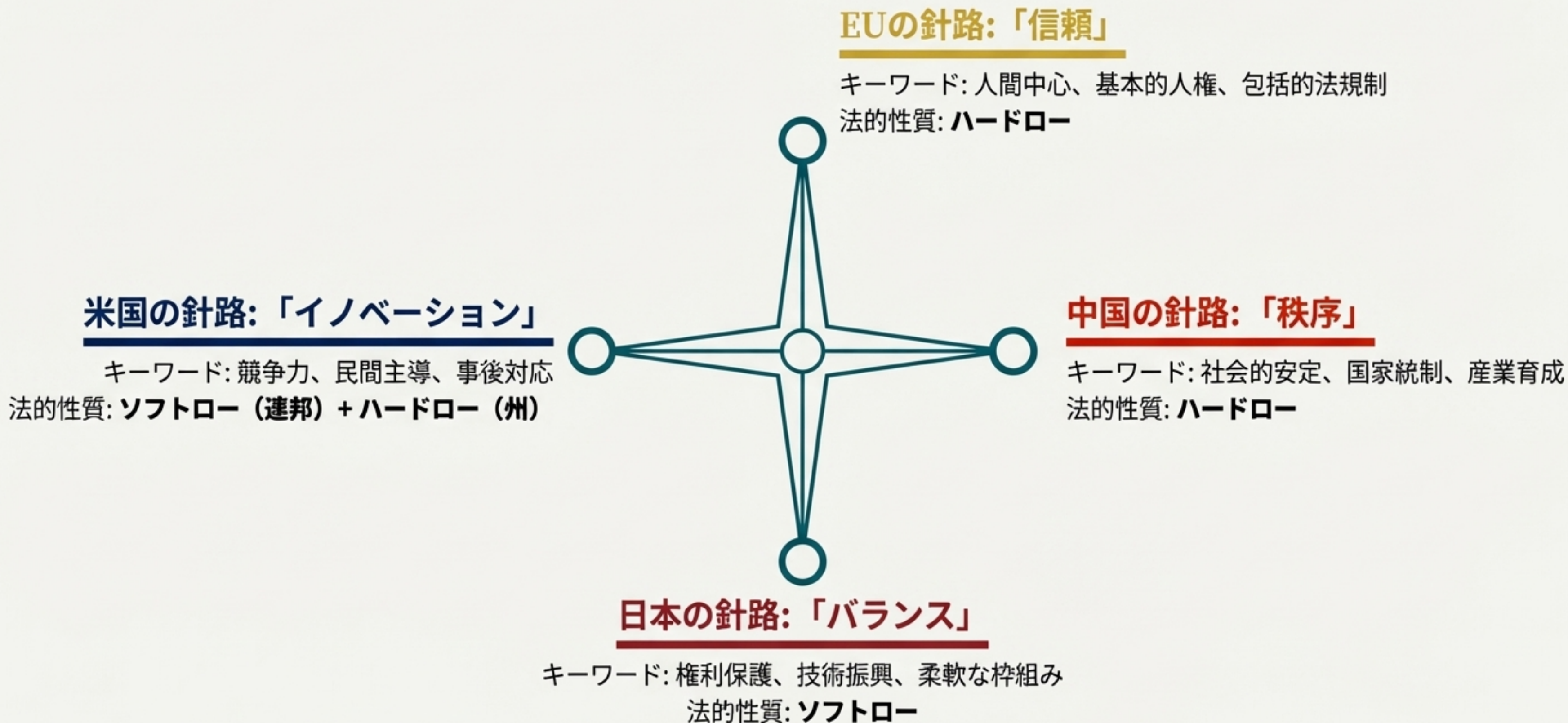
- 著作権など知的財産の侵害
- ディープフェイク等による偽情報の拡散
- プライバシー・データセキュリティの懸念



各国の対応 (Global Response)

この状況下で、主要国は自国の価値観と国益に基づき、独自のルール形成を急いでいる。これは技術覇権のみならず、次世代のデジタル社会における規範を巡る競争でもある。

4つの羅針盤：各国・地域が選択したAIガバナンスの基本方針



EUの道：法による「信頼」の構築と、AI規制のグローバルスタンダード化戦略

羅針盤

人間中心・信頼あるAI (Human-Centric & Trustworthy AI)

基本的人権、民主的価値の尊重を政策の核に据える。

法制度

ハードロー (AI法 / AI Act)

違反時には全世界売上高の最大6%の巨額制裁金。GDPRと同等の執行力。

Key Policy Mechanisms



リスクベース・アプローチ

AIを4段階のリスクに分類。高リスクAI（医療、交通等）に厳しい要件を課す。



訓練データの透明性義務

生成AI（基盤モデル）提供者に対し、著作権保護された学習データに関する「十分に詳細な要約」の公開を義務化。



著作権保護の徹底

権利者による利用拒否（opt-out）の尊重を義務付け。違法なデータ収集や著作権侵害的な出力の防止措置を要求。



生成コンテンツの表示義務

ユーザーがAIと対話していることの通知や、ディープフェイク等への明示的ラベル付けを義務化。

米国の道：イノベーションを最優先する連邦と、規制を先行する州の二重構造

羅針盤	技術競争力・イノベーション（Technological Competitiveness & Innovation） AIのメリットを最大化し、中国に対する技術的優位を維持することが国家目標。
法制度	分散的アプローチ（Decentralized Approach） 連邦レベル: 法的拘束力のない大統領令、ガイドライン、企業の自主的コミットメントが中心。 州・司法レベル: 州法と訴訟が事実上のルールを形成。

Key Policy Mechanisms



企業の自主的コミットメント

大手7社がホワイトハウスと合意。AI生成コンテンツへのウォーターマーク導入等を約束するも、強制力なし。



司法によるルール形成

著作物の無断学習を巡る集団訴訟が多発。裁判所の判断が事実上の規制として機能。



カリフォルニア州の先行

AI透明性法: 生成コンテンツであることの明示的開示を義務化。
トレーニングデータ開示法: 学習データセットの開示を義務化（2026年施行予定）。
「AIのせい」という開発者の免責を禁じる法も成立。

中国の道：国家統制による「秩序」の維持と、戦略的産業育成の両立

羅針盤

社会的安定・国家安全 (Social Stability & National Security)

技術発展は、社会主義体制の維持と国家の管理権限強化に従属する。

法制度

ハードロー (生成AIサービス管理暫行办法)

サイバー空間管理局(CAC)が管轄。事前審査・認可制で、違反には業務停止等の行政処分。

Key Policy Mechanisms



厳格なコンテンツ規制

生成内容は「社会主義核心価値観」を体现しなければならず、国家政権転覆や虚偽情報に繋がるコンテンツは禁止。



開発者の全面的責任

サービス提供者は、内容のフィルタリング、ユーザーの実名認証、当局への協力義務など、AIの出力に全責任を負う。



データ・アルゴリズムの管理

学習データは「合法的ソース」から入手し、著作権を侵害しないことを義務付け。アルゴリズムは当局への登録・説明が必要。



生成物のマーキング

ディープフェイク等にはAI生成であることの明示（透かし等）が義務付けられている。

日本の針路：権利保護と技術振興の「バランス」を模索するソフトロー戦略

羅針盤

バランス・柔軟性 (Balance & Flexibility)

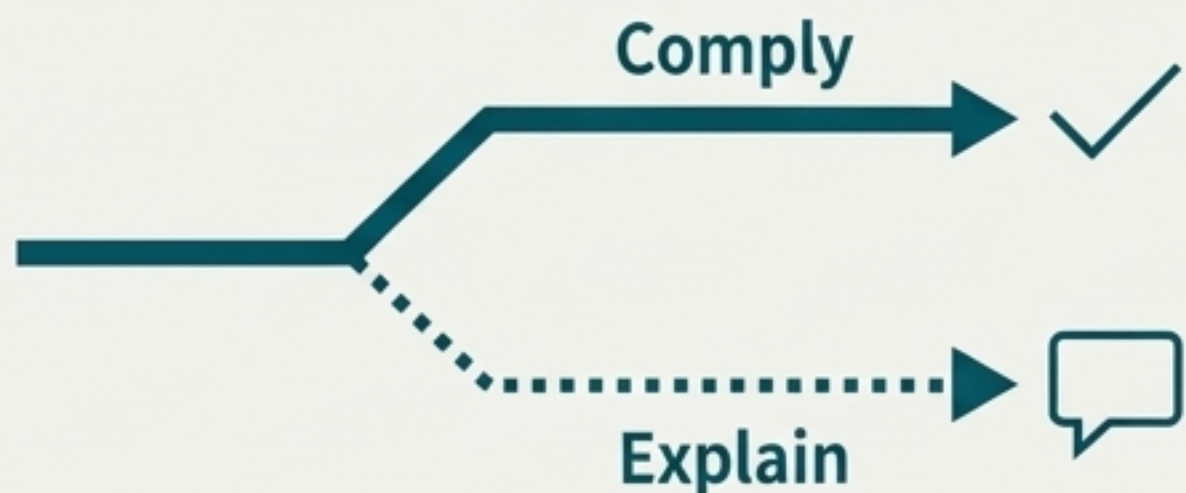
過度な規制によるイノベーション阻害を避けつつ、
権利者・利用者の安心を確保する。

法制度

ソフトロー (プリンシプル・コード案)

法的拘束力のない自主的行動指針。政府は
インセンティブ付与（補助金等）で遵守を促進。

Key Policy Mechanism: 「コンプライ・オア・エクスプレイン」方式 (Comply or Explain)



- 原則を**遵守**するか、しない場合はその**理由を社会に説明**する責任を事業者に課す。
- 一律の強制ではなく、事業者の創意工夫と社会との対話を促す。
- 専門家は「事実上の遵守圧力が働く」と指摘。

日本の原則コード案：その中核をなす「3つの透明性原則」

コード案は、生成AIモデルの開発者・提供者を対象に、主に3つの情報開示を求めている。



原則1：モデル・学習データの透明性確保

内容	事業者はウェブサイト等で「使用モデル名」「学習プロセスの概要」「学習データの種類」等を公開する。
目的	ブラックボックス化しがちな開発過程の透明性を向上させる。



原則2：権利者への情報開示

内容	著作権者等が正当な理由（訴訟準備等）で照会した場合、事業者は当該著作物が学習データに含まれるか回答する義務を負う。
目的	権利者が無断利用を確認し、法的措置を取るための情報基盤を整備する。



原則3：利用者への情報開示

内容	ユーザーが生成物と類似した既存コンテンツを発見した場合、その元データが学習に含まれていたか確認できる仕組みを求める。
課題	専門家からは「技術的・経済的負担が大きく、 スタートアップには困難 」との懸念も指摘されている。

例外規定*：オープンソース (OSS) 利用で詳細開示が困難な場合、OSS利用の事実とライセンス情報の開示で代替可能とし、スタートアップ等に配慮している。

グローバルAIガバナンスの全体像：4つのアプローチの戦略的比較

観点	日本 (原則コード案)	米国	欧州連合 (EU)	中国
法的性質	ソフトロー (遵守・説明責任)	連邦法なし (自主規制＋州法・訴訟)	ハードロー (AI法、高額制裁金)	ハードロー (行政規則、許認可制)
重視価値観	バランス 権利保護と技術振興	イノベーション 技術競争力	人権 人間中心・信頼	秩序 国家安全・社会統制
透明性措置	学習データ開示 モデル概要、データ種別別公開	生成物識別 ウォーターマーク (州法でデータ開示も)	学習データ要約公開 法的義務	当局への説明 アルゴリズム提出
IP・データ 対応	照会に応じ開示 権利者からの問合せに対応	司法判断 訴訟でフェアユースか判断	Opt-out尊重義務 無断収集を禁止	侵害データ使用禁止 合法ソースを義務化
開発者責任	説明責任 社会への説明義務	事後責任 訴訟による損害賠償	法的責任 違反時の罰金・損害賠償	全面的責任 内容検閲・是正義務

焦点比較①「透明性」：何を、誰のために明らかにするのか？

各国・地域が求める「透明性」の中身は根本的に異なる。

モデルA: 「何で学習したか」の透明性 (Input Transparency)

主導

日本・EU

アプローチ

訓練データやモデルの構築過程に関する情報開示を重視。

目的

権利者が自身の著作物の利用状況を把握し、バイアスの原因を追跡し、開発者の説明責任を問うことを可能にする。

キーワード

著作権保護、アカウントビリティ、ブラックボックス解明

モデルB: 「AIが作ったか」の透明性 (Output Transparency)

主導

米国・中国

アプローチ

生成されたコンテンツ自体に、AI製であることを示す識別子（ウォーターマーク、ラベル等）を付与することを重視。

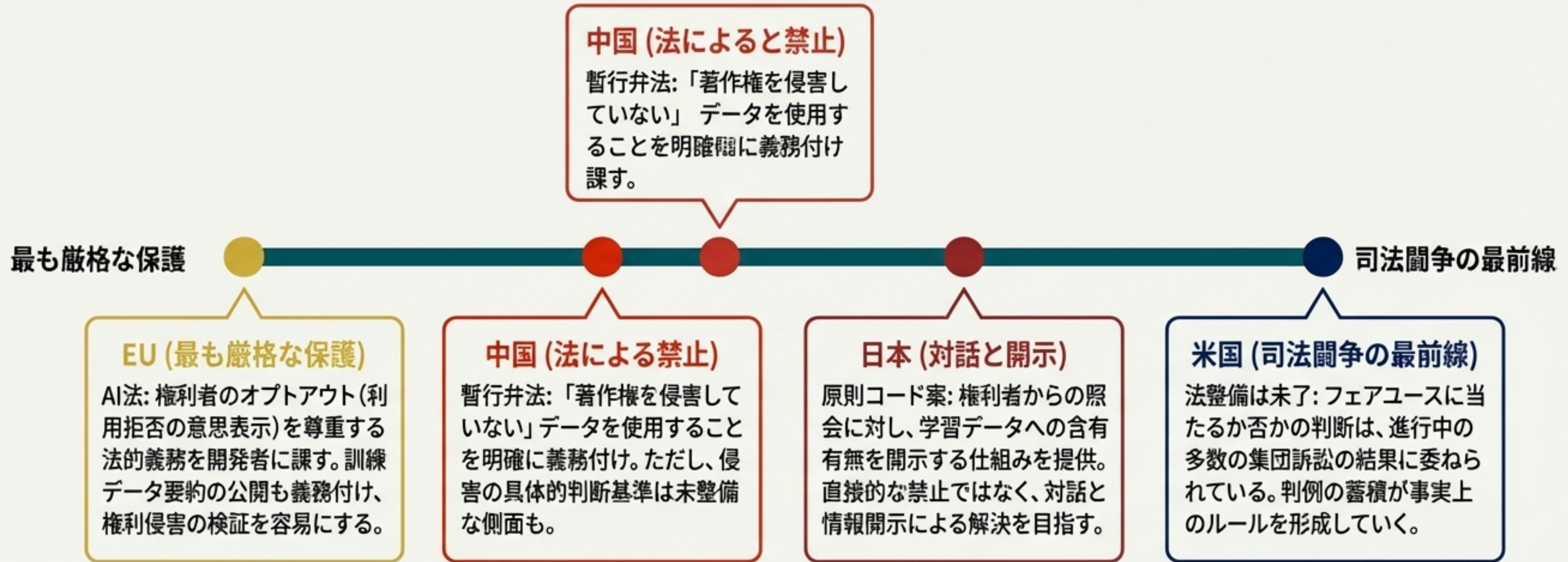
目的

ユーザーが人間による創作物とAI生成物を区別できるようにし、偽情報やディープフェイクによる混乱を防ぐ。

キーワード

偽情報対策、消費者保護、コンテンツの真贋判別

焦点比較②「知的財産」：訓練データの無断利用にどう向き合うか？



日本の針路の戦略的評価：柔軟な「ソフトロー」が持つ機会と課題

日本の「コンプライ・オア・エクスプレイン」方式は、EUの厳格な事前規制と米国の事後対応型モデルの中間に位置する、意図的な戦略的選択である。

+ 機会 (Opportunities / Pros)

機動性：急速に進化する技術に対し、法改正を経ずに柔軟かつ迅速に原則を示せる。

イノベーション促進：過度な規制負担を避け、特にスタートアップの参入と成長を阻害しにくい。

対話の文化醸成：事業者と社会（権利者、利用者）との対話を促し、自主的な最適解の発見を期待できる。

国際的ブリッジ役：欧米中の中間的立場として、国際的なルール形成で橋渡し役を担える可能性がある。

— 課題 (Challenges / Cons)

実効性の担保：法的拘束力がなく、遵守しない事業者への直接的な制裁手段がない。

「説明」の形骸化リスク：理由の説明が不十分でも、社会的なプレッシャーが働かなければ実質的な野放しになる恐れ。

技術的実現性：原則3（利用者への情報開示）など、一部要求の実現には高い技術的・コスト的ハードルが存在する。

今後の日本のAIガバナンスを進化させるための4つの論点

プリンシプル・コード案を起点とし、国内外の動向を踏まえ、日本の政策をより実効的なものへと発展させるために、以下の4つの視点が不可欠となる。

1



法的拘束力の段階的強化

まずはソフトローで産業を育みつつ、違反が横行する場合にはハードロー化も視野に入れた、規制の組み合わせ（スマート・レギュレーション）を検討する。

2



透明性と知財保護の実効性向上

権利者照会システム（原則2）の実用的な運用ルールを策定する。EUのように、権利者が検証可能な情報粒度を確保する方策を探る。

3



開発者責任とイノベーションの調和

スタートアップへの配慮は維持しつつ、米加州の例のように「AIのせい」で責任逃れできない法的整理や、損害発生時の責任分界点を明確化する。

4



国際連携と標準化への積極的関与

日本の中間的立場を活かし、グローバルな原則策定や規制の相互運用性確保に向け、欧米との連携を強化し、国際議論を主導する。