

生成AI原則コード案（日本）と米・欧・中の政策比較報告書

はじめに（要約）

2025年12月に日本政府が公表した「生成AIの適切な利活用等に向けた知的財産の保護及び透明性に関するプリンシプル・コード（仮称）（案）」は、生成AI開発企業に対し透明性の確保と知的財産（IP）保護のための自主的な行動指針を示すものです¹。本報告書では、この日本の原則コード案と、**米国**、**欧州連合（EU）**、**中国**における類似の政策・原則を比較分析します。各国・地域で重視される価値観や規制アプローチ、生成AI技術の推進と規制のバランス、政策の強制力（ソフトロー／ハードロー）の違いに焦点を当てます。

要約ポイント:- 日本:「コンプライ・オア・エクスプレイン（遵守せよ、さもなくば説明せよ）」方式のソフトロー。生成AI事業者に対しモデルや学習データの概要公開、著作権者への情報開示請求への対応など透明性と権利保護を促す²。EUのAI規制動向を参考にしつつ、日本独自の柔軟な手法を採用³。- **米国:** **連邦レベルでは統一規制なし**で、企業の自主的コミットメントや大統領令による指針が中心。**透明性**ではウォーターマーク（透かし）等によるAI生成物識別を企業が自主導入する動き⁴⁵。**著作権**は訴訟に委ねられており、AI企業による無断学習が争点の集団訴訟が進行中⁶。ただしカリフォルニア州など一部州で**AI透明性法**（生成物の明示表示義務）や**トレーニングデータ開示法**（学習に使った著作物の開示義務）など**ハードロー**が成立・施行予定で、連邦を先行する動きがあります⁷⁸。- **EU:** **AI法（AI Act）**による包括的規制（ハードロー）が整備。高リスクAIの規制に加え、**生成AI（基盤モデル）提供者に対し学習データの「十分に詳細な要約」公開を義務化**⁹し、著作権保護されたデータ利用の透明性確保を図っています。さらに**著作権保護策**として、違法なデータ収集の禁止や権利者意思表示（robots.txt等）の尊重、著作権侵害的な出力の防止措置を求めるなど¹⁰、**開発者責任**を明確化。政策目標は「人間中心の信頼できるAI」の実現であり、**基本的人権の尊重とイノベーション推進の両立**を図っています。- **中国:** **生成AIサービス管理暂行办法**（2023年8月施行）など行政規則による厳格な**ハードロー**規制。**学習データは合法かつ著作権侵害のないものを使用すべきとされ**¹¹、AI生成コンテンツには社会主義核心価値観の体现や不法・有害情報の抑制といった**内容規制**も課しています¹²。提供者は**アルゴリズム登録・当局への説明責任**を負い、違法コンテンツが出た場合の迅速な是正義務やユーザ実名登録制など**開発者責任**が厳格に問われます。**技術発展の推進**も掲げられますが、政治的・社会的安定を最優先する特徴があります。

以下、各国・地域の状況を詳述した上で、最後に比較分析をまとめます。

日本：生成AI原則コード案の内容と特徴

日本政府（内閣府知的財産戦略推進事務局）が2025年12月に公表した「生成AIの適切な利活用等に向けた知的財産の保護及び透明性に関するプリンシプル・コード（案）」は、**生成AIモデルの開発者・提供者**を対象にした倫理的行動規範です¹。この指針案の目的は、**AI技術の発展と知的財産権の保護の両立**によって、権利者・利用者双方に安心な利用環境を確保することにあります¹³。主な内容と特徴は以下の通りです。

- ・**対象事業者と「コンプライ・オア・エクスプレイン」方式:** 原則コード案は、日本国内外を問わず日本向けに生成AIサービスを提供する開発者・提供者を対象とし、「**遵守せよ、さもなくば説明せよ**」（comply or explain）の手法を採用しています¹⁴³。すなわち、**提示された原則を実施するか、実施しない場合はその理由を社会に説明することが求められます**。この柔軟なアプローチにより、一

律の強制ではなく事業者自らの創意工夫と社会との対話を促す狙いです¹⁵。専門家からは、この方式により事実上の遵守圧力が働く可能性が指摘されています¹⁶。

- **原則1：透明性確保と知財保護の措置:** モデルや学習データに関する情報公開が柱です。具体的には、各事業者が自社ウェブサイト等で「使用モデルの名称」「学習プロセスの概要」「学習データの種類（例：ウェブクロールデータの有無、非公開データセットの利用有無）」「責任体制の明確化」といった事項を公開するよう求めています¹。これにより、従来ブラックボックスになりがちだった学習データやモデル構築過程の透明性向上が期待されています¹⁷。
- **原則2：権利者（法的正当性を持つ者）への情報開示:** 著作権者など権利を有する人が、自身の著作物が特定の生成AIの学習データに含まれているかどうかを問い合わせた場合、一定の条件下で事業者は回答する義務を負います¹⁸。例えば訴訟準備など正当な法的手続きを行う者からの照会に対し、学習データへの当該著作物の含有有無を開示するという仕組みです¹⁸。これは権利者が自身の作品の無断利用を確認し、必要に応じて法的措置を取るための情報基盤を整えるものです。
- **原則3：生成AI利用者への情報開示:** 生成AIの利用者（コンテンツ生成者）にも一定の透明性を提供します¹⁹。具体的には、ユーザーが生成AIで作成した成果物と類似した既存コンテンツを発見した場合、その既存コンテンツが当該生成AIの学習データに含まれていたかどうかを確認できる仕組みを求めています¹⁹。これは、自分の生成物がAIのトレーニングデータに由来する「デジャヴュ」的な類似性を持つ場合に、その元ネタの有無を検証できるようにするものです。もっとも、この実現には巨大なデータベース検索等が必要な可能性があり、技術的・経済的負担への懸念も専門家から指摘されています（スタートアップには困難との声）²⁰。
- **OSS利用時の例外:** 上記の情報開示が難しい場合の例外規定もあります。生成AIの開発・学習にオープンソースソフトウェア（OSS）を利用している場合で詳細開示が困難なときは、OSSを利用している事実やライセンス情報を開示することで代替可能としています²¹。これにより、オープンソース活用企業への過度な負担に配慮しています。
- **政策的位置づけ:** 本コード案は現時点で法的拘束力を持たないソフトローですが、政府はインセンティブ付与（補助金や認証制度等）による遵守促進も検討しています¹⁶。EUのAI規則からの示唆を受けつつ、日本の実情に合わせた任意指針であり、ハードローによる厳格規制ではなく業界自主基準による柔軟な対応を目指す点が特徴です³。この背景には、過度な規制でイノベーションを阻害しないよう配慮しつつも、放任すればクリエイターの権利侵害や不透明性への社会不安が高まるとの判断があります¹⁷。まさに「権利保護」と「技術育成」のバランスを模索するアプローチといえます。

米国：自主的コミットメントと州法による分散的アプローチ

米国では、2025年末時点で連邦レベルの包括的AI規制法は存在せず、政策対応は大統領令、行政機関のガイドライン、業界の自主規制、そして州法といった多層的かつ分散的な形で展開されています。

- **連邦政府の動向:** バイデン政権はAIのリスクと恩恵に関する包括的対応を取っており、2023年10月の大統領令（EO 14110）「安全で信頼できるAIの開発・利用の促進」では、各省庁に対し技術標準の策定、AI生成コンテンツの識別技術（ウォーターマーク等）の開発奨励、プライバシー配慮などを指示しました²²²³。また米国商務省標準技術局（NIST）はAIリスク管理フレームワークを公表し、透明性・説明責任・公平性などの原則を推奨しています。ただしこれらはいずれも法的強制力は持たず、あくまで指針に留まります。ホワイトハウスは2022年に「AI権利章典（Blueprint for an AI Bill of Rights）」として安全・有効なシステム、アルゴリズム差別防止、データプライバシー、通知と説明、人間の介在など5原則を掲げましたが、これも法ではなく原則宣言です。
- **企業の自主的コミットメント:** 2023年7月、大手AI企業7社（OpenAI、Google、Meta等）はホワイトハウスとの会合で自主的な安全・透明性確保のコミットメントを発表しました²⁴。具体的には、AIモデルの安全性検証（レッドチーミング）、AI生成コンテンツであることの明示（ウォーターマーク等）、サイバーセキュリティや情報共有への投資など8項目に合意しています²⁴⁵。中でも「AI生成物の識別」については、コンテンツに目に見えない透かし情報を埋め込む技術の開発が挙げられました⁵。これらコミットメントは消費者保護や社会的信頼の観点から歓迎されましたが、法的拘束

力や違反時の制裁が無い「実効性に疑問」との批判もあります⁴。実際「耳障りの良い誓約だが強制力がなく具体策も曖昧」と指摘され、**実際には企業の任意努力に委ねられているのが現状です**⁴。

・著作権と学習データ：訴訟と法案の動き

- ・司法判断: 米国では、**生成AIの学習用にインターネット上の著作物を無断利用することの是非**が大きな争点となっており、**裁判による判断**が進行中です。2024年には著名アーティストらがStable DiffusionやMidjourney等の開発企業を相手取り、**自分たちの作品が無断で学習データに使われたとして集団訴訟**を提起しました²⁵。予備的な段階ながら、裁判所は「**著作物が無断利用されている場合、開発企業に法的責任を問える可能性がある**」との見解を示し、損害賠償やサービス差止めにつながり得るリスクを明確化しました⁶。また2024年末には大手出版社らがOpenAIを提訴し、新聞・雑誌記事が無断でChatGPTの訓練に使われたと主張するなど、**訴訟は出版業界にも拡大**しています²⁶。このように**米国では裁判を通じて著作権とAIの境界線を定めようとする動きが顕著**であり、判例の積み重ねが事実上の規制役割を果たしつつあります。
- ・連邦法案: 立法府でも対応策が模索されています。2024年には下院議員アダム・シフ氏が「**AI訓練に使用した著作物の透明性を義務付ける法案**」を提出し、ハリウッドの映画・音楽業界団体が支持を表明しました²⁷。この法案は、**AI企業が製品提供開始後30日以内に、訓練データとして使用した著作権保護作品の詳細な要約を米国著作権局に提出することを義務付ける内容**で、不遵守には\$5,000の罰金を科すとしています⁸。成立すればOpenAI等は**モデル開発に利用した画像・映像・文章データの開示を迫られ**、クリエイター側は自作品の無断利用を把握し**訴訟の証拠**を得ることが可能となると期待されています⁷。しかし、本法案は連邦レベルで成立のハードルが高く、成立時点でも**訓練データ利用の適法性（フェアユースか侵害か）が未確定**であることから議論が続いています²⁸。加えて、AI企業側は**学習データ開示は企業秘密の侵害**につながるとして強く反対しており、立法実現には時間を要する見通しです²⁹。
- ・州法による先行: 連邦法成立が不透明な中、**州レベルでは先進的な規制が動き出しました**。特にテック産業の中心地である**カリフォルニア州**は2024年に複数の関連法を制定しています。代表的なのが「**AI透明性法（SB 942）**」で、**月間100万以上のユーザーが利用する生成AIシステム提供者に対し、自社AIで生成・編集されたコンテンツであることを判別できる公開ツールの提供と、生成コンテンツに関する明示的な開示を義務付け**ました³⁰³¹。具体的には、**消費者が画像・映像・音声コンテンツの真偽を判断できる仕組み（例：AI生成である旨の表示や検出ツール）**を提供しなければなりません³⁰。この法律は2026年1月施行予定で、米国内外の大規模生成AIプロバイダにも契約等を通じ影響を及ぼすと指摘されています³¹³²。さらに同州では「**トレーニングデータ開示法（AB 2013）**」も2024年9月に成立し、2026年施行予定です。この法律は**大規模生成AI開発者に対し、学習に用いたデータセットの開示を義務付けるもの**と報じられています（著作権保護作品のリスト提出等を含むとされる）。また**開発者責任**に関しては、カリフォルニア州法AB 316で「**AIが自律的に起こした損害だ**」という言い逃れを禁止し、人が関与したシステムの設計・運用上の責任を開発者・提供者が負うことを明確化する法律も成立しました³³。このように米国では、**革新的なAI技術の恩恵を守りつつ透明性と責任を確保するため、州レベルが実験場となり連邦の遅れを補完している状況**です³⁴。
- ・価値観とバランス: 米国のアプローチは、「**AI技術競争でリードし国際的優位を保つ**」ことと「**社会的リスクに対処し信頼を醸成する**」ことの両立を図る点に特徴があります³⁵³⁶。政府高官も「AIのメリットを享受しつつリスクを管理するため、官民・学術・市民社会を巻き込んだ全社会的取組が必要」と強調しています³⁷。ただ連邦政府は**イノベーション促進への配慮からハードな規制は慎重**で、代わりに標準策定や市場原則による誘導に重きを置いています³⁸。そのため現状では**企業の自主性に依存する部分が大きく、「ルールなきまま民間の暴走を許していないか」との懸念と、「下手に規制して技術覇権を中国に譲り渡すな」との懸念がせめぎ合っています**。一方、表現の自由が重視される米国では**AI生成コンテンツの内容規制**については中国ほど強くなく、政治的・芸術的表現の制限は慎重です。また欧州ほどの網羅的規制も未整備であり、**技術発展を優先しつつ問題発生時には事後的に対処する**という、比較的可成りな規制哲学が見て取れます。

欧州連合（EU）：AI法による包括規制と倫理原則の法制化

欧州連合は、**人権・民主主義とイノベーションの両立**を掲げ、「**信頼できるAI**」実現のため世界で最も先進的な法整備を進めています。その中心が**AI法（Artificial Intelligence Act）**で、2024年末に成立し2025～2026年に段階的に施行される見込みです⁹。加えて、製品安全や責任に関する法改正（**AI責任補完指令案**や**製品責任指令改正**）も進行中で、EUはハードローでAI開発・利用を包括的に規律しようとしています。

- ・**リスク分類と規制強度**：AI法は全AIシステムを**リスクに基づき4段階**に分類し、高リスク分野（医療、交通、雇用等）のAIには厳しい要件を課す一方、リスクの低い汎用チャットボット等には基本的な透明性義務のみとする**リスクベース・アプローチ**を取ります³⁹。生成AI（特に大規模汎用AI、いわゆる**基盤モデル**）は、汎用性ゆえに多用途に転用可能であることから独自の規制枠組みが適用されます⁴⁰。一般目的AI（GPAI）提供者はAI法第53条に基づき、以下のような追加義務を負います。
- ・**技術文書と情報提供（透明性）**：モデルの技術文書を作成・維持し、**当局やダウンストリームの利用企業へ必要情報を提供**する義務⁴¹⁴²。提供すべき情報カテゴリは法律で定められ、モデルの特性・用途、トレーニング手法、データの出所、性能・制限事項等が含まれます。EUはこれを具体化するため**モデル文書フォーム**等の標準様式を策定し、提供者が当局やビジネスユーザーに適時情報開示できるよう支援しています⁴³。また**AI生成コンテンツのラベリング**も義務化されており、ユーザーがAIと相互作用していると認識できるよう通知すること、特にディープフェイクなど他者を欺く恐れのあるコンテンツは明示的な開示が求められます⁴⁴。
- ・**著作権・データの扱い（IP保護と透明性）**：EUは生成AIに**訓練データの透明性と著作権尊重**を強く求めています。AI法には、「**著作権で保護された学習データの内容に関する詳細な要約を公開すること**」が盛り込まれ⁹、提供者は**自社モデルがどのようなデータで訓練されたかを第三者にも分かる形で明らかにする法的義務**を負います。この「**トレーニングデータ要約（Training Data Summary）**」には、使用した大規模公開データセット名、ライセンス取得済みデータや第三者提供データの種類など、**データの一般的特性を一覧できる情報が含まれます**⁴⁵。個々の作品名全てを列挙する必要はなく、企業の機密（トレードシークレット）に配慮して**集計的・概要的な開示**とされました⁴⁵。この透明性義務により、権利者は**自作品が含まれている可能性を把握しやすくなり**、不適切な無断利用が疑われる場合に企業へ問い合わせたり法的措置を検討したりする足掛かりとなります⁴⁶。さらにAI法第53条(c)では**著作権法遵守ポリシーの策定**を提供者に求め⁴⁷、EUの著作権指令（DSM指令）で権利者に認められた**テキスト・データマイニングの権利留保（opt-out）**を尊重する措置を講じるよう義務付けています¹⁰。具体的には、**ウェブクロールで収集するデータは合法アクセス可能なものに限定し、技術的保護手段を迂回しないこと**、さらに**サイト運営者がrobots.txt等でデータ利用を禁止している場合は収集しないこと**を求めています¹⁰。またモデル出力についても**著作権侵害となるようなコンテンツ生成を防ぐための適切な措置**を講じ、利用規約にも侵害利用の禁止を明記すること、そして**権利者からの苦情申立て窓口を設置すること**まで規定されました⁴⁸。⁴⁹ EUはこれらを具体化した「**GPAI行動規範（Code of Practice）**」も2025年に策定し、OpenAIやGoogleなど主要企業が署名しています⁵⁰⁵¹。署名企業はAI法上の義務遵守が推定されるというインセンティブがあり、事実上**業界標準**としてグローバルに影響を与えています⁵⁰。
- ・**リスク管理と説明責任**：提供者はモデルの**リスクアセスメント**を行い、例えば有害な出力（ヘイトや誤情報等）の防止策を講じる義務があります。高リスク用途でなくとも、**差別や偏見の回避策、安全な動作の検証**が求められます。さらに当局に対しては**モデルやデータについて説明・協力義務**が規定され、各国の監督機関が必要に応じてアルゴリズムやデータに関する情報提出を命じることが出来ます⁵²⁵³。これはブラックボックス化への対抗策であり、**開発者に高い説明責任（アカウンタビリティ）**を負わせています。
- ・**強制力と制裁**：AI法は規則（Regulation）として加盟国に直接適用され、**違反企業には売上高の最大6%の制裁金を科すことが可能です**（GDPRと同程度の厳格さ）。つまり**ハードロー**としての拘束力が

極めて強いのがEUの特徴です。例えば学習データ要約を怠ったり虚偽報告したりすれば巨額の罰金リスクがあり、**法遵守がグローバルAI企業のビジネス継続に直結**します。EU域内でサービス提供する限り、米国企業であってもこれを無視できません。その意味で、日本のような自主原則とは一線を画す**執行力のある規制**と言えます。

- ・**価値観とバランス**： EUは「**倫理的で人間中心のAI**」を標榜しており、**基本的人権・民主的価値の尊重**が政策の核です。GDPRで個人データ保護を厳格化したのと同様、AIによる差別・プライバシー侵害・安全性リスクに対しても事前に歯止めを設けています⁵⁴⁵⁵。この背景には欧州の歴史的文脈（監視社会化や人権侵害への警戒）があり、技術の社会的影響を慎重にコントロールする傾向があります。他方で、EUは近年デジタル分野で米中に遅れを取っているとの反省から、**イノベーション促進策**も並行して打ち出しています。例えば「**ヨーロッパAI協定**」で研究開発投資を呼びかけ、**中小企業への規制負担軽減**（オープンソースモデルは一部適用除外）も検討しました。しかし規制厳格化への産業界の懸念は強く、「**EU発のAI企業が生まれにくくなる**」「**スタートアップが他地域に流出する**」といった批判もあります。それでもEUは、自らの規制が「**デジタル時代の信頼の基盤**」を築き、**長期的には企業競争力の強化につながると主張**しています。実際、欧州基準はGDPRのように世界標準化する可能性があり、AI法も各国の政策形成に影響を与えています（日本の原則コード案もEU動向を強く意識しています³）。

中国：規制による秩序維持と国家主導のAI発展

中国はAIを国家戦略の重要分野と位置付け、**政府主導の強力な規制と産業振興策**の両面で対応しています。生成AIについては、2023年8月施行の「**生成式人工知能サービス管理暫行弁法**」が世界に先駆けた包括ルールとなりました⁵⁶。この規則はサイバー空間管理局（CAC）が中心となり策定した**行政規制（ハードロー）**で、生成AIモデルの提供・利用に関する細かな義務を定めています⁵⁷。主な内容は次の通りです。

- ・**トレーニングデータ規制（知的財産とデータ法順守）**：生成AI提供者および支援者は、**事前学習やファインチューニング等のデータ処理を法律に則って行わねばならない**とされます⁵⁸。具体的には、「**使用するデータおよび基盤モデルは合法的なソースから入手したものであること**」「**第三者の知的財産権を侵害していないこと**」が明示的に義務付けられました⁵⁹。つまり著作権で保護された**コンテンツを許諾なく含めることは禁止**との趣旨ですが、どこまでが侵害かの解釈は残されています⁶⁰（中国著作権法には教育研究目的の合理使用規定があり、個人研究目的のデータ利用は許容され得る等⁶¹）。さらに**個人情報**が含まれる場合は**本人同意取得等の個人情報保護法（PIPL）上の要件を満たす必要**があり、**データ安全法やインターネット安全法**など関連法令にも適合しなければなりません⁶²。加えて、最終版規則では「**訓練データの品質向上と真実性・正確性・客観性・多様性の確保に効果的な措置を講じる**」ことも求められました⁶³⁶⁴。これは草案段階の「**確実に真実・正確・客観・多様であることを保証する**」より文言が緩和されたものの、**虚偽・偏りのあるデータでモデルを訓練しないよう努める義務**と読めます⁶⁵。
- ・**内容規制と検閲**：中国の特徴として、生成AIが生み出す**コンテンツの思想・言論面の規制**があります。規則第4条は「**生成AIによる内容はコア社会主義価値観を体現し、国家政権転覆や民族分裂の扇動、暴力・ポルノ・虚偽情報の拡散など違法有害情報を含んではならない**」と定めています¹²。つまり**政治的に不穏当な発言やデマ、過激表現は禁止**であり、既存のサイバーセキュリティ法やネット情報内容規定の違法情報リストが準用されます⁶⁶。サービス提供者はこれらを順守すべく**AIモデル出力にフィルタリングを掛け、違反があれば迅速に修正・削除する義務**があります。特に**ディープフェイク等、人を欺く可能性のある合成メディア**に対しては、**目に見えるかたちでの表示（manifest disclosure）**や**透かし埋め込み（latent disclosure）**によってAI生成であることを明示することが、別途の「**ディープシンセシス規定**」（2023年1月施行）で義務付けられています⁶⁷⁶⁸。生成AI暫行弁法自体にも、第11条で「**虚偽・有害な情報の生成防止措置**」や「**識別技術の提供**」が求められており、**利用者や社会に対する透明性確保策**として生成コンテンツへの**マーカー埋込**などが事実上期待されています。

・**開発者責任と当局への説明義務:** 中国ではサービス提供者が全責任を負う設計になっています。規則は提供者に対し、**内容安全管理制度の確立、セキュリティ担当者の設置、利用者の実名認証とログ記録、違法情報の通報窓口設置**などを義務付けました^{69 70}。ユーザが生成AIを悪用した場合でも当局は提供企業を直接指導・処罰でき、企業は「AIが勝手にやった」と言い訳できません³³。また第19条では**監督当局からの検査・資料提出要請への協力義務**が規定され、**訓練データの出所・規模・種類、データラベリングルール、アルゴリズム原理等**について説明を求められたら応じねばなりません⁷¹。このように中国は政府による事前・事後の介入権限が非常に強く、開発者は常に当局の監視下で責任を負います。

・**法的強制力:** 暫行弁法は行政規則であり違反時には是正命令や業務停止・罰金等の行政処分が可能で、さらに2023年は**インターネット情報サービス深度合成規定やアルゴリズム推薦管理規定**も施行されており、これらと合わせた**統合的なAIガバナンス体制**が構築されています⁷²。企業は提供サービス毎に当局へ届け出・審査を受け、認可されなければ公開提供できません。また当局は**モデルの安全評価報告**の提出やソースコード開示を要求する権限も持つとされ、一元的管理体制です。

・**価値観と政策目的:** 中国のAI規制は「**社会の安定維持**」と「**産業発展**」の両立を謳いますが、実際には**政治的統制と情報検閲**が色濃く反映されています。**核心価値観**の強調や不穩情報の禁止は、AIが体制批判やデマ拡散に利用されることへの強い警戒を示しています⁷³。他方で第3条では「**独立自主的イノベーションと国際協力推進**」を掲げており、基礎技術の国産化や有用なAI応用の奨励もうたわれています⁷⁴。実際、規制草案は非常に厳格でしたが最終版では**中小企業負担に配慮して一部要件（事前ライセンス制等）を緩和する修正**がなされました⁷⁵。これは**米国の輸出規制下でも中国独自のAI産業を育成**する必要性から、民間活力を殺がめ調整が図られたとされています⁷⁶。しかし全体として、中国のAI規制は**国家安全・社会秩序を最優先**しており、企業活動や表現の自由よりも**政府の管理権限**が重視されます。この点は民主主義国のアプローチとの大きな違いです。

比較分析：各国・地域の特徴の対比

以上を踏まえ、日本・米国・EU・中国のアプローチの共通点と相違点を**主要な観点ごとに整理**します。

・**法的枠組みの性質（ソフトロー vs ハードロー）:** 日本のプリンシプル・コード案は**業界自主ルール**としてのソフトローであり、**遵守は努力義務**に留まります。一方、EUと中国は**法的強制力を持つ規制（ハードロー）**を導入し、違反に対する制裁も用意されています。米国は連邦レベルではハードローが無く（大統領令やガイドラインは非強制）、**州法や民事訴訟**による事後救済が実質的役割を果たしています。この違いは各国の法文化とAI産業競争力への考え方の差を反映しています。

・**重視する価値観・目的:** 日本は**知的財産権の保護とイノベーション促進のバランス**を重視し、「**透明性による信頼確保**」を軸に据えています¹⁷。米国は**技術革新と経済競争力の維持**に強い関心を持ちつつ、**差別やセキュリティなど具体的リスクへの対処**を図る姿勢です（人権や消費者保護も重視しますが、全体としてビジネス促進のトーンが強め）。EUは**基本的人権の尊重と消費者・労働者等の保護**が最優先で、その上で「**人間中心の技術発展**」を目指します。中国は**社会主義体制の維持・強化**が最上位目標で、コンテンツ規制や国家安全保障を貫徹しつつ、自国企業を世界トップに押し上げる産業政策的関心を持ちます。

・**透明性確保のアプローチ:** どの国・地域も生成AIの**透明性向上**を掲げていますが、中身は異なります。**日本とEU**は主に**訓練データやモデル情報の開示**に焦点を当て、開発過程を白日の下に晒すことで権利侵害の抑止や説明責任の履行を狙います^{2 9}。**米国はアウトプットの識別（ウォーターマーク等）**を重視し、AI生成コンテンツと人間の創作物を混同しないようにする方向です⁵。またカリフォルニア州法で**出所情報（プロビナンス）表示義務**を課すなど、生成物レベルでの透明性に力点があります⁶⁷。**中国は対政府・対社会**での透明性確保に特徴があり、企業が当局にモデル内部情報を提示する義務や、ユーザに対しては生成物へのマーカー埋込を要求しています（ただし一般ユーザに

訓練データの詳細を公開する仕組みはなく、透明性は限定的です）。要するに、日本・EUは「**何で学習したか**」を透明にし、米国・中国は「**AIが作ったものだと分かるように**」透明にする違いが見られます。

・**知的財産・訓練データの扱い**： 日本・EU・中国はいずれも他者の著作物の無断学習への懸念を共有し、対応策を講じています。日本の原則2・3やEUのAI法は**権利者への通知・開示や許諾の尊重**を制度化しようとしています⁷⁷¹⁰。特にEUは法的義務として**権利留保（opt-out）の尊重やデータ要約公開**を課し、違法なデータ収集を禁じました¹⁰。中国も「**著作権侵害データを使うな**」と規定していますが¹¹、実効面では**政府検閲による違法コンテンツ排除**に力点があり、著作権侵害の具体的判定基準は曖昧です⁶⁰。米国は現状**フェアユース（一部許容）か侵害か**の結論が出ておらず、裁判と権利者の交渉に委ねられています。しかし、出版社や映画業界が積極的に声を上げたことで、**訓練データの無断利用に対する反発が強まりつつあり**、上述の法案提出や州法（AB 2013）など**透明性から権利保護につなげる試み**が現れています²⁷⁸。

・**開発者責任・説明責任**： 中国は開発者責任を最も厳格に定め、**AIの出力にも全面的に責任を負わせる**体制です（不適切出力への迅速対処義務、利用者の不法行為も含め監督責任を負う等）。EUも**高リスクAI提供者の厳格責任や違反時の損害賠償責任**を今後法制化予定で、既存の製品責任法制をAI時代にアップデートしています。日本のコード案は法的責任までは踏み込んでいませんが、**社会に対する説明責任を果たすこと（説明なき不遵守は許されない）**という仕組みで道義的・実質的な責任を担保しようとしています³¹⁷。米国ではセクション230（プラットフォーム免責）との絡みなど法律上の論点がありますが、カリフォルニア州AB316のように「**AIのせい**」で逃げられないようにする法改正も出始めました³³。総じて、**中国・EUは開発段階からの事前責任確保（安全策実装や届け出）**を課し、**米国は被害発生後の救済（裁判での責任追及）**が中心、**日本はその中間で事前透明性による牽制**という色合いです。

・**技術推進と規制のバランス**： 日本は規制がソフトローである点にも象徴されるように、**技術推進への配慮**があります。政府もスタートアップ支援やAI活用推進策と両輪で本コード案を位置づけており、「**ルール形成で信頼を醸成しつつ企業の自主的イノベーションを促す**」スタンスです¹⁷。米国は言うまでもなく民間主導の技術革新が原動力で、政府規制は必要最小限にとどめています。ただし安全・倫理面の問題が顕在化すれば素早く（場合によっては州レベルで）対応するという**問題発生後追型**です。EUは「**信頼があってこそそのイノベーション**」という理念から先回り規制を敷いていますが、そのために企業にコスト負担が生じ、短期的には萎縮効果も懸念されています。それでもEUは**規制と投資の両面戦略**（例えばAIサンドボックス制度で革新的実証を支援）で克服を図っています。**中国は規制による秩序維持を重視しつつ、国家予算を投入した研究開発や政府需要による市場形成で技術を強力に推進する**という**国家資本主義的モデル**です。規制自体も「**良質なAIを育てるための必要条件**」と位置づけられており、統制と振興が不可分になっています。

以上の点を表形式でまとめると次のようになります。

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（倫理指針）。 遵守は努力義務だが「説明責任」で実効性担保 ³ 。	連邦レベルで統一法なし。 大統領令・ガイドライン（非強制）と州法・訴訟の組合せ。	ハードロー（AI法）。 直接適用の規則で罰金等制裁あり。	ハードロー（行政規則）。 政府当局による許認可・監督と処罰。

観点	日本（原則コード案）	米国	欧州連合（EU）	中国
重視価値観	権利保護と技術振興の両立。 透明性と安全・安心な利用環境の確保。	技術競争力・イノベーション最優先。 リスク管理（安全・公正・プライバシー）も重視。	基本的人権・消費者保護最優先。 「人間中心・信頼あるAI」で社会受容性確保。	国家安全・社会秩序最優先。 体制維持と国産技術の飛躍的発展。
透明性措置	訓練データやモデル情報の開示。 モデル名・データ種別等を公開 ⁷⁸ 。 権利者・利用者への問い合わせ窓口設置 ⁷⁷ 。	生成物の識別。 水印埋込や出所表示ツール提供 ³¹ 。 州法で訓練データ開示義務（2026年～）も導入予定。	訓練データ要約の公開義務 ⁹ 。 AI生成であることを表示義務（深層偽造対策） ⁴⁴ 。	当局への説明（アルゴリズム・データ提出）。 AI生成物への透かし埋込要求。 一般公開の訓練データ開示は無し。
IP・データ対応	問合せにより学習データへの著作物含有情報を開示 ⁷⁷ 。 OSS使用時はライセンス情報開示で代替可 ²¹ 。	法整備途上。 裁判で適法性判断（無断学習は訴訟係争中 ⁶ ）。 著作権局がAIアウトプットの著作物性指針策定 ⁷⁹ 。	権利留保の尊重義務（無断収集禁止） ¹⁰ 。 訓練データに何を使ったか開示義務 ⁴⁶ 。 苦情処理制度の整備。	違法・侵害データの使用禁止 ¹¹ （ただしフェアユース類似例外あり）。 データの真実性・多様性確保義務 ⁶³ 。 個人情報同意必要。
開発者責任	コンプライアンス説明で社会に説明責任 ³ 。 強制力は無いが事実上遵守圧力。	基本は事後責任（被害時に訴訟で賠償）。 CA州法でAIの自律性悪用の免責禁止 ³³ 。	法的責任を明確化（違反罰金・損害リスク）。 高リスクAIは認証・当局報告義務。 AI責任法制で欠陥時の推定責任導入予定。	全面的責任負担（当局監督下）。 違法出力は企業に是正義務。 利用者不正も連帯責任。 許可なくサービス提供不可。
規制と推進のバランス	過度な縛りを避け業界自主性尊重。 インセンティブで遵守誘導、スタートアップ配慮。	イノベーション優先で規制は最小限。 問題発生時は個別対処・州主導。	事前規制で信頼醸成が前提。 規制負担大も国際標準化で欧州優位狙う。 補助金やサンドボックスで企業支援。	統制と振興を両立（規制は産業発展のため必要と位置付け）。 国家投資と市場保護で国産AI育成。

おわりに（政策提言への示唆）

日本の原則コード案は、各国の潮流を踏まえつつ独自のソフトロー戦略を取った点で注目に値します。欧州型の透明性義務を参考にしながらも、米国型の柔軟さを取り入れた「コンプライ・オア・エクスプレイン」は、日本らしい折衷策と言えます³。これは現時点で急速に発展する生成AI分野において機動的に原則を示し業界の自主的対応を促すメリットがあります。もっとも、実効性確保という課題も残ります。事実上の強制力が働くとの指摘もありますが¹⁶、法的裏付けがないため遵守しない企業への直接的制裁は困難です。また、原則3に見るように技術的実現性に疑問符の付く要求もあり²⁰、将来的には実効性のある措置への改訂や必要に応じたハードロー化も視野に入れるべきでしょう。

各国比較から浮かび上がる示唆として、日本が今後政策を充実させる際には以下の点を検討すると有益です。

- ① **法的拘束力の段階的強化:** 最初はソフトローで産業萎縮を防ぎつつも、違反や不正が横行する場合にはハードロー化する準備を整えておく。例えばEUのように自主コードへの署名企業には優遇（セーフハーバー）を与え、未署名企業には個別監督を厳格化する等、自主規制と法規制を組み合わせた仕組みも考えられます⁵⁰。
- ② **透明性と知財保護の実効策:** 訓練データ開示については、日本版も「概要的情報の開示」から始めていますが⁷⁸、EU同様に権利者が検証可能な情報粒度を確保することが重要です。権利者照会システム（原則2）が有効に機能するよう運用ルールを詰め、必要なら米国の事例にならいデータベース設置も検討すべきです⁸。またクリエイターへの適正な還元については、欧州で進むライセンス交渉の動きを参考に、日本でも業界間対話の場を設けることが望まれます⁸⁰。
- ③ **開発者責任とイノベーションの調和:** スタートアップなど新興企業への過度な負担は避けつつ、基本的な説明責任や安全配慮義務は明確にしておく必要があります。米国カリフォルニア州のように「AIのせい」で責任逃れできない規定³³や、中国のようなリスクある用途への注意義務は、日本でも製造物責任や不法行為法理の中で整理しておくことが重要でしょう。将来的にAIが公共インフラ化することを見据え、事故・被害時の責任分界点を事前に議論しておくことが、健全なイノベーションの土台となります。
- ④ **国際連携と標準への関与:** 日本の動きは各国の中間に位置するため、橋渡し役として国際的なルール形成に寄与できます。欧州とは信頼性・人権の価値観を共有しつつ、米国とは実務的解決策で協調し、将来のグローバル原則や相互運用可能な規制枠組みづくりに積極的に関与すべきです。特に生成AIは国境を越えるため、日本国内で完結する制度にも限界があり、各国政策の調和やデータ流通ルールの国際標準化への働きかけが必要でしょう。

本報告の比較分析から明らかなように、生成AIのガバナンスには多角的視点が求められます。日本は今回のプリンシプル・コード案を起点に、国内外の知見を取り入れつつ柔軟かつ実効的な政策発展を図ることが期待されます。権利者・開発者・利用者の三者が歩み寄り、安全で創造的なAI時代を築くため、引き続きエビデンスに基づく議論と制度設計を深めていく必要があります。

¹ ² ³ ¹³ ¹⁴ ¹⁵ ¹⁶ ¹⁷ ¹⁸ ¹⁹ ²⁰ ²¹ ⁷⁷ ⁷⁸ 内閣府が生成AIの透明性と知財保護に関する「プリンシプル・コード」を公表。パブコメを1月26日まで募集中！

[https://www.aicu.jp/post/naikakufu-generative-ai-principle?](https://www.aicu.jp/post/naikakufu-generative-ai-principle?srsltid=AfmBOooPLu2vf-9pRr2WmPHuqxl-8rRH-2ltQ1sFN9_lxSudjT4Rh0DD)
[srsltid=AfmBOooPLu2vf-9pRr2WmPHuqxl-8rRH-2ltQ1sFN9_lxSudjT4Rh0DD](https://www.aicu.jp/post/naikakufu-generative-ai-principle?srsltid=AfmBOooPLu2vf-9pRr2WmPHuqxl-8rRH-2ltQ1sFN9_lxSudjT4Rh0DD)

⁴ ⁵ ²⁴ Voluntary Commitments from Leading Artificial Intelligence Companies on July 21, 2023 - Harvard Law Review

<https://harvardlawreview.org/print/vol-137/voluntary-commitments-from-leading-artificial-intelligence-companies-on-july-21-2023/>

⁶ ⁹ ²⁵ ²⁶ ⁴⁶ ⁷⁹ ⁸⁰ 【2025年最新版】生成AI技術と著作権問題を徹底解説 | ギグワークスクロスアイティ株式会社

<https://gigxit.co.jp/blog/blog-20273/>

⁷ ⁸ ²⁷ ²⁸ ²⁹ ハリウッドの組合がAI透明性法案を支持、クリエイターへの発言力を与えるか - THR Japan

<https://hollywoodreporter.jp/news/38362/>

¹⁰ ⁴⁰ ⁴¹ ⁴² ⁴³ ⁴⁵ ⁴⁷ ⁴⁸ ⁴⁹ ⁵⁰ ⁵¹ ⁵² ⁵³ Taking the EU AI Act to Practice Decoding the GPAI Code of Practice and the Training Data Summary Te - Bird & Bird

<https://www.twobirds.com/en/insights/2025/taking-the-eu-ai-act-to-practice-decoding-the-gpai-code-of-practice-and-the-training-data-summary-te>

11 58 59 62 63 64 65 69 70 71 China's New AI Regulations

<https://www.lw.com/admin/upload/SiteAttachments/Chinas-New-AI-Regulations.pdf>

12 54 55 56 57 60 61 66 72 73 74 Overview of Draft Measures on Generative AI

<https://www.chinalawtranslate.com/en/overview-of-draft-measures-on-generative-ai/>

22 23 35 36 37 38 President Biden Signs Sweeping Artificial Intelligence Executive Order | Insights |

Sidley Austin LLP

<https://www.sidley.com/en/insights/newsupdates/2023/11/president-biden-signs-sweeping-artificial-intelligence-executive-order>

30 31 32 カリフォルニア州、AIコンテンツの開示を義務付けるAI透明性法を制定 | Jones Day

<https://www.jonesday.com/ja/insights/2024/10/california-enacts-ai-transparency-law-requiring-disclosures-for-ai-content>

33 67 68 米カリフォルニア州、AIの安全性・透明性・プライバシー保護強化への複数の法案を制定(米国) | ビジネス短信 — ジェトロの海外ニュース - ジェトロ

<https://www.jetro.go.jp/biznews/2025/10/ea97d110406cf2b1.html>

34 州のAI規制は連邦政府の政策を先行できるか? - Gamma Law

[https://gammalaw.com/ja/](https://gammalaw.com/ja/%E5%B7%9E%E3%81%AE%E8%A6%8F%E5%88%B6%E3%81%AF%E9%80%A3%E9%82%A6%E6%94%BF%E5%BA%9C%E3%81%AE%E6%94%B)

39 High-level summary of the AI Act | EU Artificial Intelligence Act

<https://artificialintelligenceact.eu/high-level-summary/>

44 Article 50: Transparency Obligations for Providers and Deployers of ...

<https://artificialintelligenceact.eu/article/50/>

75 76 China Finalizes Its First Administrative Measures Governing ...

<https://www.pillsburylaw.com/en/news-and-insights/china-generative-ai-measures.html>