

生成AI知財プリンシプル・コード：透明性による信頼構築の新たな枠組み

2026年8月公表。法的な拘束力を持たない「ソフトロー」として、「透明性」を武器に権利者・利用者・事業者間の情報格差を解消し、市場の信頼を促す新たな枠組み。

原則1：モデルと学習データの概要開示



学習データの取得方法や知財保護体制など、公衆向けの概要説明を求めます。

原則2：権利者からの個別照会への回答



権利者が権利行使を準備する場合、特定URLの学習有無について合理的な範囲で回答します。

原則3：AI利用者による類似性確認



利用者が侵害リスクを評価できるよう、特定資料との同一性について回答機会を設けます。

実務的な意義と特徴



「コンプライ・オア・エクスプレイン」の採用
遵守状況を政府に届け出、一発化されることで市場の信頼性を可視化します。



営業秘密と安全性の保護（安全弁）
機密情報の強制開示は求めず、事業者の過大な負担を避ける配慮が明記されています。



商取引における「事実上の標準」へ
法的義務はなくとも、政府調達や企業間契約の判断基準になる可能性があります。

地域別制度比較

日本



制度の性格：ソフトロー（任意）
主な特徴：個別照会メカニズム（原則2・3）が独白。

EU



制度の性格：ハードロー（義務）
主な特徴：AI法に基づき、学習データの詳細な概要公開を法定義務化。

米国



制度の性格：分散型（州法など）
主な特徴：カリフォルニア州などで高レベルな概要開示が義務化。