

2026 年 1 月 26 日

「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）」 に関する意見

総論

【意見内容のカテゴリー:3.その他】

Business Software Alliance（ビジネス・ソフトウェア・アライアンス、以下 BSA）¹は、「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）」²（以下、コード案）に関して意見する機会を得られたことに感謝します。BSA は、エンタープライズソフトウェア業界を代表するグローバルな業界団体であり、会員企業は、人工知能（AI）、クラウドストレージサービス、サイバーセキュリティ・ソリューション、量子コンピューティング等で他の企業の事業活動を支えています。BSA の会員は、AI システムやサービスの開発、カスタマイズ、インテグレーション（統合）、導入、そして AI システムおよびアプリケーションの開発に活用されるツールにおいて、世界をリードする事業者であり、デジタルトランスフォーメーションを促進する AI テクノロジーの大きな可能性と、AI の責任ある利用を最も効果的に支援する政策について、独自の知見を有しています。

¹ Business Software Alliance (<https://bsa.or.jp/> (日本語)、www.bsa.org (英語))は、アメリカ、ヨーロッパ、アジアの 20 を超える市場で活動し、あらゆる分野の産業また一般消費者がイノベーションの恩恵を受けられるよう、テクノロジーに対する信頼を構築する政策を推奨しています。

BSA の会員には以下の企業が含まれます：Alteryx, Amadeus, Amazon Web Services, Asana, Atlassian, Autodesk, Avalara, Bentley Systems, Box, Cisco, Cloudflare, Cohere, Cohesity, Dassault Systemes, Databricks, Docusign, Dropbox, Elastic, EY, Graphisoft, HubSpot, IBM, Informatica, Kyndryl, MathWorks, Microsoft, Notion, Okta, OpenAI, Oracle, PagerDuty, Palo Alto Networks, PTC, Rubrik, Salesforce, SAP, ServiceNow, Shopify Inc., Siemens Industry Software Inc., Trend Micro, TriNet, Veeam, Workday, Zendesk, aZoom Communications Inc.

² <https://public-comment.e-gov.go.jp/pcm/download?seqNo=0000305362>

日本だけでなく世界的にも、AI と知的財産権の関係は機微かつ進化し続ける課題であると我々は認識しています。クリエイティブ産業は日本経済を支える重要な分野であり、日本の文化的アイデンティティにも大きく寄与しています。権利保護を継続しつつ、ソフトウェア開発者を含む、権利者、アーティスト、コンテンツクリエイターを活性化させるような、AI の責任ある利用を促進する知的財産政策を、協働的に検討することが重要です。同時に「世界で最も AI を開発・活用しやすい国」を目指す日本において、生産性を向上させ、さまざまな分野でイノベーションを生み出し、国際競争力を強化するために AI を活用する力を、こうした政策が意図せず妨げてしまわないことも重要です。

提言

[意見内容のカテゴリー: 2. この文書が示す原則及び例外]

AI の透明性措置において法的確実性とイノベーション促進に資する運用を確保すること

BSA の会員企業は、イノベーションの推進と知的財産の保護に強くコミットしており、透明性の促進の重要性を認識しています。透明性の目的は、AI システムのエンドユーザーがモデルの能力と限界を理解し、モデル利用時に情報に基づいた選択を行えるようにすることにあるべきです。しかし、コード案の広範な適用範囲と厳格なアプローチには、いくつかの懸念が生じます。

コード案では、AI 事業者の受入れを任意と位置付けつつも、受入れた企業は開示情報を掲載した自社サイトのリンクと共に公開リストに掲載され、受けなかった企業はその理由を説明（エクस्पレイン）することが求められるとしています。さらに、政府は AI 事業者の取組の状況を評価し、政府が実施・運用する各種の事業や制度において一定のインセンティブを設ける意向を示しています。このような構造は、コード案が自主的な枠組であるとしながらも、事実上の義務を生み出してしまいうおそれがあります。

さらにコード案は、生成 AI 事業者に対して自社の情報（営業秘密を含む）の強制的な開示を求めるものではないと明記しているにも関わらず、各 AI 事業者の取組の蓄積により開示される情報が標準化されることを期待しているとも記しています。実質的に、このようなアプローチは企業活動に重大な制約を課す可能性があり、AI 開発促進における日本の著作権法の有効性を損なう恐れがあります。

日本の著作権法やその他の規制において、AI モデル開発者や AI サービス提供者に AI 開発・学習に関する情報の開示を義務付ける法的要件は存在しないことを明確にすることが重要です。特に著作権法第 30 条の 4 は、大量のデータの抽出、比較、分類、統計的处理を含む、表現内容の享受を目的としない情報解析のために必要な範囲において、著作物の複製その他の利用を明示的に認めています。この例外規定は、AI、機械学習、データサイエンスといった分野におけるイノベーションが、取引コストや法的な不確実性によって阻害されないようにするという、日本の立法府の政策的判断を反映したものです。

AI 事業者は、適用法令の遵守を確実にすれば、追加的な要件や自主的措置を課されないことを当然期待しています。しかし、コード案で採用されている「コンプライ・オア・エクスプレイン」のアプローチは、AI 事業者にとって法的不確実性を生じさせるおそれがあり、AI イノベーションを促進するという日本の政策方針とも整合しません。

提言： 知的財産戦略本部（以下、知財戦略本部）は、原則 1 を修正し、「生成 AI の開発・学習等も含めたデータの活用に関しては、他者の知的財産権を侵害しないこと」といった文言や、権利者の救済を確保するため、AI 事業者に窓口を整備することを求める条項など、著作権物を用いた学習自体が著作権侵害の執行措置の対象となる可能性があることを示唆する条項を削除するか、これらの条項を修正し、合法的な学習には適用されないことを明記すべきです。

また知財戦略本部は、コード案が法的拘束力のない、参照ガイドラインとして明確に位置づけられていることを明記すべきです。

過度に広範な適用範囲・公的開示措置に対する懸念

コード案は、EU AI Act およびコーポレートガバナンスの分野におけるスチュワードシップ・コード等の取組を参考にしていると説明しているものの、その適用範囲は著しく広範です。コード案は、EU における、2025 年 7 月の汎用 AI の行動規範（General-Purpose AI Code-of Practice）が基盤モデル開発者やシステミックリスクを呈する主体に限定されていたのとは異なり、公衆に AI 製品・サービスを提供する全ての生成 AI 開発者・提供者に適用されます。さらに、規制当局や下流の AI 提供者に対する透明性ではなく、公衆に対する透明性を重視しており、これは EU のリスクベース・階層的アプローチから逸脱しています。

AI 事業者は、自社の学習手法や AI 開発プロセスの構築に多大なリソースを投じており、これらを機密性が高く、商業的に重要な営業秘密とみなしています。AI モデルの開発に関する情報は極めて機微性が高く、これを公開すれば AI 事業者の正当な利益が大きく損なわれるおそれがあります。さらに、技術革新のスピードが非常に速く、AI 開発の状況が常に変化していることを踏まえると、継続的または頻繁な情報開示義務は、不正確、不完全、または最新ではない情報が意図せず公にされる結果を招くおそれがあります。

AI 事業者に対し、詳細な学習手法や学習データに関する極めて粒度の高い情報（例えば URL の一覧など）といった、AI の学習に関する具体的な情報の開示を求めることは、透明性の確保に必要な範囲を超えており、権利者やその他のステークホルダーの正当な利益にも資するものではありません。

AI モデルは学習データの複製を保存したり、取り出すのではなく、学習データ全体にわたるパターンを学習し、特徴を一般化します。特定の生成結果を、一つの情報源や URL に遡ることは技術的に不可能です。AI モデルは統計的パターンを分析しているに過ぎないため、AI 学習（入力側の分析）は著作権侵害を判断する上で、技術的にも法的にも関連性はありません。出力に類似性が見られる場合、それは学習に用いられた特定の著作物に由来するのではなく、多様な学習データに繰り返し現れる著名なキャラクターなど、一般的によく見られる特徴に起因する可能性が高いといえま

す。さらに、入力側の分析が法的に無関係なのは、AI の学習（すなわち情報解析）が、先述したとおり、著作物の表現内容の享受に該当しないためです。

著作権侵害の有無を判断する上で、開示は技術的にも法的にも関連性がなく、商業的に機微な情報を露出させるリスクを伴い、研究開発への投資を損ない、さらには AI モデル開発者が日本でモデルを提供する意欲を削ぐおそれがあります。また、海賊版サイトからのコンテンツ利用を回避する取組については支持しますが、いかなる義務も、実行可能なものであり、かつオンライン上の情報を検閲することのないよう配慮されるべきです。AI 開発者が、学習において海賊版サイトのコンテンツを使用しないための措置を既に講じていることを認識することが重要です。海賊版サイトのコンテンツ利用を避けるための自主的なベストプラクティスを促進することは、有効な取組であり、特に、日本の著作権法に違反して、著作権で保護された作品の無断複製物を、継続的かつ反復的に、商業目的で公衆に提供していることが確認されているウェブサイトについては、その意義が大きいと考えます。

また、著作権は著作者が保有する私的権利であり、不正競争防止法の下で保護される利益も同様に、特定の個人や法人に帰属する私的権利です。これらの権利に関する係争は、関係当事者間で解決されるべきものです。

政府は、AI 開発者に対して専有情報や機密情報の開示をを迫ることによって、生じ得る係争で一方の当事者に実質的な優位性を与えるような介入を控えるべきです。AI 開発者と権利者は、商用ライセンス交渉を含む、直接協議のための既存の仕組みをすでに有しています。

提言：知財戦略本部は、透明性に関する措置が概要レベルの、非機密情報に限定されることを明確にし、営業秘密や契約上の機密情報、または機密性の高い事業・技術情報の開示は求められないことを、コード案において明示すべきです。

また、知財戦略本部は、関係当事者の一方に有利または不利となり得る追加的な開示権限や手続きをコード案において新たに設けることを避けるべきです。

透明性義務におけるサイバーセキュリティと国家安全保障の確保

さらに、開示要件の詳細を策定するにあたっては、深刻な国家安全保障上の影響が生じ得ることを十分に考慮するよう、強く政府に求めます。ネットワーク防御のために、どのように AI システムが用いられ、どのように学習されているかといった情報の開示を求めることは、意図せずして、サイバー攻撃者が当該防御を回避するための手掛かりを与える結果となり、ひいてはネットワークおよび情報システムの根本的な安全性を損なうおそれがあります。

原則 1 における「モデルのトレーニングプロセスの内容」や「パラメータの設定」といった技術情報の開示要件は、サイバーセキュリティ強化を目的とする AI に適用される場合、重大なリスクを伴います。このような詳細情報は、攻撃者に防御アーキテクチャの「設計図」を提供する結果となり得るためです。例えば、脅威行為者がサイバーセキュリティ AI モデルがどのようにマルウェア

を検知するよう調整されているかを理解した場合、検知を回避するように攻撃手法を変更することが可能になります。したがって、サイバーセキュリティを支える AI システムについては、詳細なアーキテクチャ情報の開示義務から免除することが特に重要です。

さらに、法的措置のためにデータアクセスの提供を AI 事業者を求める原則 2 に関しては、サイバーセキュリティ向け AI の学習に用いられるデータには、機微な脅威インテリジェンスや専有的なサンプルが含まれることを認識することが重要です。第三者が「法的措置を準備している」という理由のみでこうしたデータへのアクセスを認めることは、機微なセキュリティインテリジェンスへのアクセスを狙う悪意ある主体による濫用的な、情報漁りを招くおそれがあります。このように高度に機微なデータへのアクセスについては、より厳格な法的要件が課されるべきです。

同様に、学習データの URL の提供を AI 事業者を求める原則 3 についても、重大なセキュリティ上の懸念が生じます。サイバーセキュリティ企業は、ウェブフィルターやファイアウォールを学習させるために大量の URL データを使用しています。特定の URL が学習に使用されたかどうかを開示することは、当該インフラがセキュリティシステムによって検知されていることを、意図せずして悪意ある主体に示す結果となり、防御を回避できるようインフラを変更する手がかりを、そのような主体に与えてしまう可能性があります。

提言：上記の提言にあるように、知財戦略本部は原則 1～3 を改訂し、サイバーセキュリティを可能にする AI システムについては、防御ロジックが明らかとなり得る内部のパラメータ設定や学習プロセスの開示が求められないようにするとともに、開示内容を当該システムの概要レベルの機能に限定すべきです。また、濫用を防止する観点から、セキュリティシステムに関連する学習データへのアクセスについては、裁判所命令等、より高い基準を設けるとともに、脅威インテリジェンスの情報源や手法を保護するため、サイバーセキュリティモデルの学習に使用された特定の URL を開示することが求められないようにすべきです。

不相応かつ実行不可能な措置の導入を避けること

また、コード案の原則 2 および原則 3 では、EU AI Act では想定されていない追加的な遵守が期待されていることを指摘したいと思います。これらの原則は、法的措置を講じている、または講じる準備をしている当事者や、生成 AI の利用者からの開示要求に対して、AI 事業者が対応する措置を求めています。両原則はいずれも、学習および検証に用いられたデータに関する情報の提供を AI 事業者に求めており、その対象には、ウェブクロールや第三者から取得した非公開データセット、公開データセット、さらには合成データも含まれています。原則 2 は、実際の、または想定される法的手続の文脈で適用される一方、原則 3 は、生成または検証データに、自らの生成物と同一または類似のコンテンツが含まれているかどうかについて情報を求める利用者に対して適用されます。これらの原則は総合すると、不相応かつ実行不可能な措置の大幅な拡大となっています。

原則 2 および原則 3 における義務の問題点は、AI モデルの学習データ（すなわち入力）と生成 AI モデルの出力とを不適切に結び付け、その結果として、著作権侵害の申し立てにおいて本来重視さ

れるべき「出力側」の分析ではなく、「入力側」の分析に焦点を当てている点にあります。コード案においては、出力に関する問題が入力に関する対応によって解決されるかのように扱っていますが、実際には、前述のとおり、入力データは関連性を有しておらず、仮に侵害の疑いのある出力が存在する場合には、現行の著作権法の下で既に救済手段が用意されています。

大規模な AI の学習コーパスには、膨大な種類のデジタル情報源から取得された数十億単位の個別データ要素が含まれる場合があります。前述のとおり、こうしたデータをトークン化し、学習のために分析する過程で行われる計算的情報解析は、著作物の表現内容の享受を伴うものではありません。むしろ、当該分析は、トークン化されたデータセット全体を対象として、確率、相関関係、傾向、その他のパターンに関する数学的計算を行うものです。この分析が理解しようとするのは、データセット全体に分布する統計的なパターン（特定のトークンが他のトークンとどのような関係にあるか）に限られます。こうした統計的パターンそれ自体は、著作権法によって保護される表現内容には該当しません。また、著作権による保護は、これらの計算的・統計的分析プロセスの対象となる事実、アイデア、または数学的概念に及ぶものではありません。

これらの事実に加えて、日本には既に先見性のある計算的情報解析に関する例外規定が存在していることを踏まえると、著作権侵害への懸念を判断するために学習データの詳細な開示を求めること、または期待することは、実務上、非現実的であるだけでなく、法的にも関連性を欠くものです。

原則 2 において、単に法的措置を「検討している」段階にある当事者に対しての情報開示を AI 事業者に求めることは、本来裁判所の管轄に属する複雑な法的主張の判断について、確立された司法手続を事実上迂回することとなり、問題のある前例を生むおそれがあります。開示要求が認められる主体の範囲が広く、要求のハードルが低く設定されていることにより、AI 開発者は、推測的かつ濫用的な開示要求にさらされる高いリスクを負うことになります。その結果、要求者が機密情報へのアクセスを要求できるようになり、営業秘密の開示、また、研究開発から人的・時間的なリソースを奪うことにつながるおそれがあります。日本の現行の民事訴訟法にはすでに訴訟の文脈で関連情報を取得するための仕組みが存在します。追加措置を検討する前に、現行の司法手続で正当な開示ニーズに十分に対応できるかを検証することが不可欠です。

さらに、原則 3 に基づく義務を含めることにより、AI によって生成された出力に関する著作権侵害を判断するために、AI 利用者が特定の学習データの開示を求めることが可能となり、コード案の当該部分における技術的および法的な実現可能性の低さを一層深刻化させています。

前述のとおり、AI の学習入力は、著作権侵害を判断する上で、技術的にも法的にも関連性を有していません。したがって、こうした技術的に実現不可能な要件は、著作権法上の問題を判断するために不要なものです。AI の出力が著作権で保護された作品と一定の類似性を有する場合であっても、著作権侵害に関する責任は、出力側の分析のみに基づいて判断することが可能です。

これらの理由等から、AI の出力に焦点を当てることが極めて望ましいと考えます。特定の出力によって著作権者の権利が侵害された場合には、権利者に十分かつ実効的な救済が与えられるべきであることに、我々は強く賛同します。この原則は、AI システムを用いて生成された出力にも、その他

の方法で生成された出力にも、等しく適用されるべきものです。したがって、侵害に関する判断については、現行の著作権法で十分に対応可能であると考えます。

実際、利用者が AI によって生成された出力がオンライン上のコンテンツと類似していることに気付いた場合、当該利用者は、類似性や当該出力が生成された状況を判断するために必要な関連情報、すなわち、自ら入力したプロンプトや反復的な創作過程を含め、既に関連情報にアクセスできる立場にあります。

著作物と類似するコンテンツが著作権侵害に該当するか否かという法的判断は、複雑であり、事実関係に大きく依存します。現行の著作権法の法理が示すとおり、類似性があるという事実のみをもって、意図的な複製を立証することにはなりません。たとえば、多くの場合、芸術的な画像などの作品は、複数のプロンプトやその他の手法を用いて反復的に生成され、AI ツールによって構築・創作されることがあります。いずれにしても、AI と著作権に関する議論において適切な焦点となるのは、入力側の分析ではなく、こうした AI の出力側の分析です。出力側の分析において検討されるべき主な論点には、(1) ある著作物と生成 AI の出力との類似の程度、(2) 独自に創作されたことを示す証拠の有無、(3) 当該出力を生成した主体、(4) 生成に至った状況、(5) 生成の意図が挙げられます。これらはいずれも、現行の著作権法の枠組みに照らして、AI の出力を検討することにより十分に対応可能です。

最後に、AI の出力に基づく解決に焦点を当てることは、著作権侵害となり得る AI 生成による出力に対処するための有効な道筋を示します。多くの AI 開発者は、既に、出力フィルタリング、メタプロンプト、契約条項といった技術的・契約的措置を講じ、AI ツールが保護された作品と実質的に類似する複製を生成してしまうというわずかなリスクを低減しています。これらの措置は、侵害出力を防止するという目的に適切に合致しています。

原則 2 および原則 3 が入力側の開示要件に重点を置いていることにより、権利者および彼らの利益に実質的なメリットがほとんどないにもかかわらず、日本の AI 開発目標を不要に損ない、また、損害や多額の費用を伴うような、根拠を欠いた訴訟が助長されることを我々は懸念しています。

こうした理由から、原則 2 と原則 3 に基づいて AI 事業者 URL や情報源特定の照会に対応する措置を課すことには、合理的根拠も実務的な妥当性もありません。これらの原則は、モデル学習の技術の実態を反映しておらず、軽率な訴訟を誘発し、AI の入力と出力の区分においても、誤った側面に焦点を当てています。法的、事実に、技術的観点のいずれに照らしても、AI の出力ではなく入力に焦点を当てることは、著作権上の懸念に対して実効可能かつ効果的な解決策とはなり得ません。

提言：知財戦略本部は、コード案の全体的な構造および枠組みを再検討し、原則 2 および原則 3 をコード案から完全に削除すべきです。

リスクベースで国際的整合性のある AI 政策立案の再確認

コード案を含む、AI 政策の策定にあたっては、AI 事業者や権利者をはじめとする影響を受けるステークホルダーとの十分かつ実質的な協議を行うことを政府に強く求めます。これは、日本が G7 広島 AI プロセス（HAIP）を通じて主導してきた、国際的に認められた AI ガバナンスのアプローチと整合するものです。残念ながら、コード案は、HAIP の基盤となる国際的に整合したリスクベースの原則から逸脱したアプローチを採用しています。

著作物を含み得る AI 学習に関して、著作権者が自らの意向を表明できる新たな技術的解決策を開発するために、著作権者と AI 開発者の間で進められている国際的な標準化の取組があります。我々はコード案がこの取組を損なったり、これに矛盾する恐れがあることを懸念しています。検索エンジン向けの現行の「クロール禁止（do not crawl）」ツールと同様に、現在、ウェブサイトが学習目的に使用されることを権利者が望まないことを示す自動化ツールに関して、活発な業界横断的な協議が行われています。こうした議論は、Internet Engineering Task Force（IETF）、World Wide Web Consortium（W3C）、Coalition for Content Provenance and Authenticity（C2PA）などの組織で進められています。これらの議論には、クリエイティブ産業、テクノロジー産業、また、市民社会、学術界の代表者が参加しています。

提言：知財戦略本部は、日本の AI イノベーションにおけるリーダーシップを不意に損なう、あるいは現在進行中の国際標準化の取組と日本を不要に対立させかねないような、AI 事業者に対する自主的措置や期待を課すことを控えるべきです。日本が既存の政策目標を堅持することを我々は奨めます。AI に関するリーダーシップの目標や既存の規制・国際標準化の枠組みと整合しない、不要な、あるいは介入的な措置を回避することで、日本は引き続き AI 分野のグローバルリーダーとしての地位を維持することができます。

コード遵守を公共調達の基準として用いることのリスク

前述のとおり、コード案の「(4)その他の事項」では、政府が AI 事業者の公表内容や具体的取組の状況等を評価し、政府が実施・運用する各種の事業や制度等において、一定のインセンティブを設けることが想定されており、コード遵守が公共調達に影響を及ぼし得ることを示唆しています。しかし、コード案は日本法上の明確な法的根拠を欠いているため、こうした措置は、適用される国際貿易協定上、日本が負っている物品およびサービスの市場アクセス義務との不整合を生じさせるおそれがあります。これには、WTO 政府調達協定（GPA）において日本が負っている責務、また、経済連携協定（EPA）及び地域貿易協定（RTA）における付随的責務が含まれます。

AI 利用者による責任ある利用の促進

生成 AI システムは、幅広い用途に使用できるツールです。汎用目的のツールを提供する AI 開発者は、利用者が多様な作業を行えるようにしています。もし特定の利用者が権利者が懸念する行為に及んでいる場合、適切な対応は、基盤技術そのものを制限することではなく、そうした利用者によ

る AI ツールの責任ある適法な利用を促すことに焦点を当てるべきです。とりわけ前述のとおり、多くの AI モデルにはすでに、既存の著作物と実質的に類似する出力が生成されるリスクを低減するための、出力フィルタリングやプロンプト入力制限といった安全措置が組み込まれています。したがって、AI システムの適切な利用を促進する施策を政府が優先することを求めます。

提言：知財戦略本部は、出力に関する利用者の責任を明記した、文化庁からのガイドライン³に沿って、技術的安全策の活用、利用者リテラシーと認知向上の強化を通じて、AI システムの適切かつ責任ある利用促進への確固たる取組を再確認すべきです。

結論

[意見内容のカテゴリー:3.その他]

BSA の意見をご検討頂けることに感謝します。知財戦略本部が、影響を受けるステークホルダーとの議論を継続し、今回の意見募集後に公表される修正コード案に意見を述べられるように、2回目の意見募集を実施することを強く推奨します。利用者、権利者、その他のステークホルダーを適切に保護しながら、AI イノベーションを推進するために、引き続き政府と協働できることを我々は期待しています。

³ 文化庁「AIと著作権に関する考え方について」

https://www.bunka.go.jp/seisaku/bunkashingikai/chosakuken/pdf/94037901_01.pdf