

第3回 AIと著作権文化庁「AIと著作権に関する考え方について」 及びジュリスト No.1599 号の概説 (青山綜合法律事務所弁護士 内藤 篤)

はじめに

内藤篤：2024年3月に文化庁が「AIと著作権に関する考え方について」（以下「考え方」といいます）を公表し、それに呼応する形で、同年7月に、わが国における著名な法学雑誌である「ジュリスト」（2024年6月・第1599号）において、この「考え方」に関する論考（以下「ジュリスト論考」といいます）が掲載されました。以下、文化庁「考え方」及び当該論考の概要とそれについての更なる考え方をご披露させていただくというのが本日の趣向です。

1. 文化庁「考え方」の概観

内藤篤：文化庁の「考え方」に対して、「ジュリスト」では四つのポイントに分けて論考が載せられています。1. 開発・学習段階、2. 生成・利用段階、3. 生成物の著作物性、4. その他の論点（主に、裁判においてAI利用に関する著作権侵害が争われた場合の要件事実面等）となっております。

1. 開発・学習段階

...ジュリストp62~67 澤田将史弁護士解説箇所

2. 生成・利用段階

...ジュリストp68~73 中川達也弁護士解説箇所

3. 生成物の著作物性についてAIの技術的な背景について

...ジュリストp74~79 出井甫弁護士解説箇所

4. その他の論点について

※ジュリストp80~85 高部元裁判官解説箇所は主張立証責任に関する論稿のため割愛

◆ ELSIとの関係では主に1・2が論点

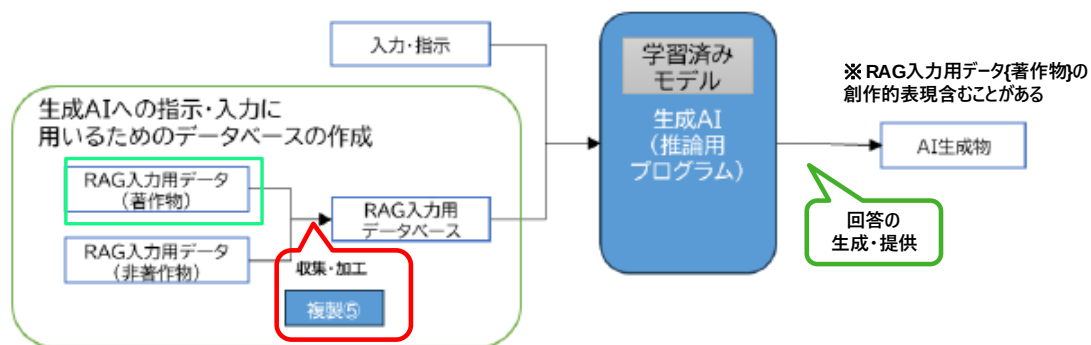
内藤篤：本プロジェクトにとって重要なのは前の二つ、開発・学習段階と生成・利用段階ですので、これらの論点に絞りつつ、文化庁の「考え方」とこのジュリスト論考とを概観していきたいと思います。

2. 開発・学習段階の論点

内藤篤：まずは開発・学習段階の話ですが、ジュリスト論考の筆者である澤田将史弁護士が冒頭で取り上げているのは、文化庁の「考え方」の18～19ページにかけて、ポンチ絵とともに説明されている部分です。

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

以下に検索拡張生成 (Retrieval-Augmented Generation : RAG) 等において生成 AI への指示・入力に用いるためのデータベースの作成についてのポンチ絵を示します。この他、文化庁の「考え方」には、AI 学習用データセット構築のための学習データの収集、それから基盤モデル作成に向けた事前学習、そして既存の学習済みモデルに対する追加的な学習に関する三つのポンチ絵が挙げられています。



内藤篤：私としては、この三つないし四つのデータ収集なり学習というものが、それぞれ独立なのか、それとも一種の発展段階を表しているのかもよく分からない状況でした。

しかし、本プロジェクトにおける勉強会で国立がん研究センター研究所の小林和馬先生が挙げておられた汎用コーパス、医療コーパス、指示コーパスという三つのものが、このうちの三つ、AI 学習用データセット構築のための学習用データの収集、基盤モデ

ル作成に向けた事前学習、既存の学習済みモデルに対する追加的な学習なのだろうと想像すると、非常に明快です。

小林和馬先生からのご質問は、指示コーパスに著作権法第 30 条の 4 の適用があるのかというものでしたが、この指示コーパスでのデータの取り込み方には、他の二つ（AI 学習用データセット構築のための学習用データの収集、基盤モデル作成に向けた事前学習）とはやや質的に異なるものがあることが示唆されていると思われました。

【小林和馬先生の質問】

小林和馬：先日の勉強会では、今取り上げていただいている RAG は対象外とさせていただいています。この RAG を使う前の、いわゆる開発段階におけるコーパスごとの特性に応じた著作権法第 30 条 4 の適用範囲を教えてくださいという質問の趣旨は、いま先生方に汲み取っていただけているかと思います。



2

Q1. LLMの学習過程のうち、非享受目的といえる範囲について

- 汎用コーパスから構築された汎用ベースモデルに対して、以下の多段階の学習・評価の過程が想定される。
- ここで、著作権法30条の4「情報解析に関する権利制限規定」（情報解析規定）の適応範囲を教えてください。

汎用コーパス 特定の分野に特化せず、Web等から恣意的に収集したテキストの集合。	汎用ベースモデル 特定のタスクに依存しない、一般的な言語パターンが学習されたモデル。	文章における規則性や特徴、事実関係など、著作権で保護されない要素のみを学習する目的とされるか？
医療コーパス 医療分野に特化したテキストの集合。Web上の公開テキスト、書籍、病院データ等が想定。	医療ベースモデル¹⁾ 一般的な医学知識や医学分野に特徴的な言語パターンが学習されたモデル。	
指示コーパス ユーザーの指示に応じて、適切な応答を行うように調整するための、指示と応答のペアからなるデータセット。	アライメントモデル ユーザーに提供され、実際に利用されることが想定されるモデル。	著作物に含まれるコンテンツの全部あるいは一部を、モデルが出力できるような目的性を持って学習している場合、 非享受利用と考えるか？

医療コーパスに含まれるそれぞれのテキストの利用条件に応じて、医療ベースモデルは増強のモデル・パターンに分類することも想定されている。尚、今回は医療拡張生成（RAG）における論点はカバーしていない。

内藤篤：RAG は、データ自体を自分の中に取り込むのではなく、外部データベースから関連情報を引っ張ってきて、質問やプロンプトに対する回答を生成する仕組みという事です。ただ、RAG への入力・指示を行うためのデータベースというのは、事前

に作る必要になっているようです。

✓ **検索拡張生成（RAG）**
 学習データ以外の追加のデータベース（下図「RAG入力用データベース」）を検索して関連情報を取得し、質問やプロンプトに対する回答を生成する手法
 生成した回答において、データベース上の情報（元の新聞記事など）と類似する表現が出力されることがある

✓ **開発・学習段階の使用行為**：データベース上の著作物（黄枠）の創作的表現の出力を目的とする場合、データベース作成の際に当該著作物を複製等する行為（赤枠）は、「考え方」（30条の4適用否定説）によれば「享受目的」併存のため法30条の4適用否定

✓ **生成・利用段階の使用行為**：生成結果にてデータベース上の著作物（黄色枠）の創作的表現を出力する行為（橙枠）は複製又は翻案に該当

内藤篤：では、今回の SIP プロジェクトが行おうとしていることは、この RAG に該当するのでしょうか。例えば、レントゲン写真などを大量に読み込ませて、それを読影

に役立てるような AI というのは、RAG に該当し得るように思われます。結局は、どのような出力形態を志向するかによって結論も違ってくるだろうと考えています。

この点に絡めて、まずは、著作権法第 30 条の 4 と第 47 条の 5 の関係はどうなっているかを、先んじて説明しておきたいと思います。この RAG については、「考え方」において第 47 条の 5 の適用可能性が指摘されています。第 47 条の 5 と第 30 条の 4 との主な差異は、前者が「非享受目的の限定がない」、「情報解析（入口）のみならず結果の提供（出口）まで対象」、「各号の行為との付随性、軽微利用性が要件」の三つであるといえます。

✓ **47条の5の適用可能性の指摘**

- RAGについては「考え方」にて47条の5の適用可能性が示唆されている
※文化審議会著作権分科会（第4回）「[法30条の4と法47条の5の適用例について](#)」も同様
- 47条の5も情報解析に関する権利制限規定だが、30条の4との主な差異は
 - －非享受目的の限定がない
 - －情報解析（入口）のみならず結果の提供（出口）まで対象（2号）
 - －各号の行為との付随性、軽微利用性が要件

✓ **付随性の充足**

- 付随性：①情報解析・結果提供行為、②著作物の利用行為とを区分し、①が主で②が従という関係性が必要
- ジュリスト澤田弁護士解説は、解析結果としてデータベースに記録されたデータへのリンクを提供しつつ、確認の便宜のために当該データの一部をスニペット表示するケースであれば付随性要件を充たすとの言及（**要約の提供等には否定的**）

✓ **軽微利用性の充足**

- 外形的な要素（割合、分量、表示の制度等）を総合的に考慮し、著作物の本来的な市場に影響を与える可能性が典型的に低い利用態様であるかによって判断
→学術論文の要約等を出力する行為は付随性・軽微性の観点からしても許容されない可能性が高い

内藤篤：このジュリスト論考における澤田弁護士の解説では、RAG について「解析結果としてデータベースに記録されたデータへのリンクを提供しつつ、確認の便宜のために当該データの一部をスニペット表示するケースであれば付随性要件を充たす」という指摘がされています。つまり「スニペット表示」のような短い要約文ならば 47 条の 5 でいけるということです。ここでは、著作権法第 30 条の 4 においては、要約文を出力することを目的とする場合は、取り込んだ著作物の創作的表現の再現であり、いわ

ゆる享受目的になってしまうため、同条は適用されないという議論が前提とされているわけです。

しかし、私見としては、何千字、何万字からなる論文を数十字の要約にまとめたものは、論文の著作権侵害にはならないと考えています。

次に、著作権法第 30 条の 4 の「非享受目的」について説明します。

✓ 「非享受目的」

- 法30条の4の1～3号に掲げる場合
- その他当該著作物に表現された思想又は感情を自ら享受し又は他人に享受させることを目的としない場合

✓ 「享受目的」が併存する場合は30条の4の適用なし

文化庁「考え方」における享受目的が併存する場合の具体例

- 既存の学習済みモデルに対する追加的な学習...のうち、意図的に、学習データに含まれる著作物の創作的表現の全部又は一部を出力させることを目的とした追加的な学習を行うため著作物の複製等を行う場合
- 既存のデータベースやインターネット上に掲載されたデータに含まれる著作物の創作的表現の全部又は一部を、生成AI を用いて出力させることを目的として、これに用いるため著作物の内容をベクトルに変換したデータベースを作成する等の著作物の複製等を行う場合

→著作物の表現上の本質的特徴を感得することができる出力を目的とした場合
「享受目的」が併存するとみられる

✓ 「享受目的」の有無は開発・学習段階で判断

生成・利用段階において学習著作物の創作的表現の生成が著しく頻発するという事情は開発・学習段階における「享受目的」の存在を推認する上での一要素

内藤篤：文化庁の「考え方」では、「著作物に表現された思想又は感情を自ら享受し又は他人に享受させることを目的としない」という、非享受目的の場合に、著作権法第 30 条の 4 が適用になるというのが基本的な考え方です。この非享受目的の有無は、開発・学習段階で判断されるのですが、実際にクローリングをしてデータを集めている段階において、そこに享受目的があるかないかということは外から見て分からないわけです。そして、享受目的が非享受目的に併存する場合には、この 30 条の 4 の適用があるかないかということが、割と論点として言われています。

2 開発・学習段階の論点

論点：「享受目的」が併存する場合の30条の4の適用

- ✓ 30条の4適用否定説（文化庁「考え方」）
 - 30条の4の文理からすると専ら非享受目的であることが求められる
 - 享受目的が併存する場合は30条の4の適用余地はなく、47条の5の問題
- ✓ 30条の4適用肯定説（高部、愛知）
 - 30条の4「著作物は、次に掲げる場合（注:1～3号）その他著作物に表現された思想又は感情を自ら享受し又は他人に享受させることを目的としない場合」について法令用語で「その他の」は、前の用語がその後の用語の例示であることを示すものであるから同条1～3号は非享受利用の例示
1～3号の要件を充足するにもかかわらず（享受目的が併存する場合などに）非享受利用に該当しないという解釈は文言解釈として妥当なものではない
 - 但し、肯定説も既存著作物の創作的表現の出力が目的とされた場合に他の要件において考慮することを排斥するものではない
ex) 愛知「学習・推論段階での入力行為それ自体に対する判断に際して、事後的にコンテンツの生成を目的としているという事情は、但書適用の可否において斟酌される」（上野・奥邨「AIと著作権」p. 29）

内藤篤：この文化庁の「考え方」では、著作権法第30条の4の文理からすると、もっぱら非享受目的であることが求められるということになります。即ち、享受目的が併存した場合は、著作権法第30条の4の適用余地はなく、著作権法第47条の5の問題であると言っているのですが、これに対しては有力な反対説がいくつか存在しています。

まず、文言解釈上、著作権法第30条の4には「著作物は次に掲げる場合その他著作物に表現された」と書いてあるが、法令用語で「その他」とある場合は、「前の用語がその後の用語の例示であることを示すもの」だということを根拠にする説です。つまり、この1～3号に列挙されている場合は、もう非享受利用として確定しているというように読むべきなのだと主張されています。

この説に則れば、仮に出力段階での出力物が、データとして取り入れたものの創作的表現の再現をも意図していたとしても、このデータクローリングが著作権法第30条の4本文の適用を受けないということにはならないということになります。

【参加者からの RAG (Retrieval-Augmented Generation) に関する補足】

児玉安司：典型的には、「RAG (Retrieval-Augmented Generation) は、大規模言語モデルによるテキスト生成に、外部情報の検索を組み合わせることで、回答精度を向上させる技術のこと。『検索拡張生成』、『取得拡張生成』などと訳されます。外部情報の検索を組み合わせることで、大規模言語モデルの出力結果を簡単に最新の情報に更新できるようになる効果や、出力結果の根拠が明確になり、事実に基づかない情報を生成する現象（ハルシネーション）を抑制する効果などが期待されています」とされています。

簡潔に幾つかの命題を言われているのですが、情報学のプロの方々、これでいいでしょうか。それだけコメントいただけますか。

原田達也：一般的な大規模言語モデルというのは、ニューラルネットワークの中に学習したデータがそのまま記憶されているわけではなくて、抽象化されたかたちで分散表現されています。この RAG の場合は、我々が何かレポートを書くプロセスに少し似ています。例えば、何か分からないことがあったら Google などで検索しますよね。Google で検索して、これはいいなと思った 10 件の情報があったとしたら、10 件の情報を参照して、それを要約してレポートを書くという操作をすると思うのです。

RAG も本当にそういう感じで、生データが多くあります。生データと抽象化された大規模ランゲージモデルがあります。そうすることによって、大きなランゲージモデル、ニューラルネットワーク自体が知らない情報だとしても、この生データの中に入っていれば回答が可能になるということなのです。その生データ自体の質が高いことが担保されているものであれば、ランゲージモデルが嘘をつく可能性が低くなるということになります。そのため、普通のランゲージモデルの回答では生データが出る可能性は非常に低いのですが、RAG の場合は、もともと生データ自体から引っ張ってきたも

のを直接的、明示的に使うようなプロセスが入ってくるので、生データが直接的に出てくる可能性も高まり、著作権侵害の問題が起こる可能性もあるのではないかと心配している人が多いと思います。

内藤篤：次に著作権法第 30 条の 4 の但し書きの話に移ります。

「著作権者の利益を不当に害する場合」は適用されないというのが、著作権法第 30 条の 4 の但し書きです。「考え方」では、具体例として、「販売されている（又は販売が予定されている）情報解析用データベースを利用する行為」を挙げており、より具体的には次の三つだと言っています。

文化庁「考え方」における「著作権者の利益を不当に害する場合」の具体例 1

1. 大量の情報を容易に情報解析に活用できる形で整理したデータベース（情報解析用データベース）の著作物が販売されている場合に、当該データベースを情報解析目的で複製等する行為
2. インターネット上のウェブサイトで、ユーザーの閲覧に供するため記事等が提供されているのに加え、データベースの著作物から容易に情報解析に活用できる形で整理されたデータを取得できるAPIが有償で提供されている場合において、当該APIを有償で利用することなく、当該ウェブサイトに掲載された記事等のデータから、当該データベースの著作物の創作的表現が認められる一定の情報のまとまりを情報解析目的で複製する行為
3. AI学習のための著作物の複製等を防止する技術的な措置（robots.txt”への記述等）が講じられており、かつ、このような措置が講じられていることや、過去の実績（情報解析に活用できる形で整理したデータベースの著作物の作成実績や、そのライセンス取引に関する実績等）といった事実から、当該ウェブサイト内のデータを含み、情報解析に活用できる形で整理したデータベースの著作物が将来販売される予定があることが推認される場合に、この措置を回避して、クローラにより当該ウェブサイト内に掲載されている多数のデータを収集することにより、AI学習のために当該データベースの著作物の複製等をする行為

【参加者からの質問】

小林和馬：以下のスライドにあるとおり、LLM に限らず、一般的に情報系の世界において、データセットというかたちで、主に研究用途に活用できるようなデータベースが公開されていることがあります。しかし、そのデータベースの利用において、ライセンスが設定されていることが多くあります。

例えば学術研究目的に限るとか、ものによっては商用利用も可能というように、おそらく著作物性という枠をもう少し超えた広い範囲で、そのデータセットを利用するにおいての利用許諾契約というニュアンスで、ライセンスが設定されていることがあります。こちらの著作権法との兼ね合いについて少し教えていただきたいと思ったのが、

こちらのスライドの趣旨になります。また、生のデータセットを使って作ったニューラルネットワークに抽象化されてしまったモデルにも、その前のデータセットの規約、例えばアカデミック利用可能だが商用利用不可というライセンス規約といったものが引き継がれるのでしょうか。

Q6. データセットのライセンスについて

8

- モデルの学習を含む情報解析を目的に構築されたデータセットが多く公開されていますが、そのうちの一部には下記のようにライセンスが課せられています。
- 本邦の著作権において、こうしたデータセットのライセンスはどのように考慮されますでしょうか？
- また、こうしたデータセットを利用して学習したモデル自体のライセンスは、どのように設定されるべきでしょうか？

Corpus	Publisher	Release Time	Size	Public or Not	License
BBT-FinCorpus	Fudan University et al.	2023-2	256 GB	Partial	-
FinCorpus	Du Xiaoman	2023-9	60.36 GB	All	Apache-2.0
FinGLM	Knowledge Atlas et al.	2023-7	69 GB	All	Apache-2.0
Medical-pt	Ming Xu	2023-5	632.78 MB	All	Apache-2.0
Proof-Pile-2	Princeton University et al.	2023-10	55 B Tokens	All	-
PubMed Central	NCBI	2000-2	-	All	PMC Copyright Notice
TigerBot-earning	TigerBot	2023-5	488 MB	All	Apache-2.0
TigerBot-law	TigerBot	2023-5	29.9 MB	All	Apache-2.0
TigerBot-research	TigerBot	2023-5	696 MB	All	Apache-2.0
TransGPT-pt	Beijing Jiaotong University	2023-7	35.8 MB	All	Apache-2.0

データセットのライセンス



OpenLLMなどを用いてデータセットを整形する場合、当該OpenLLMのライセンスも考慮される

最終的に構築されたLLMのライセンスはどう考えるべきか？

内藤篤：ご質問の前提として、公開されてアクセス可能な状態になっているものをクローリングで取り込む場合が想定されている理解です。まず、著作権法第30条の4の観点からすると、そもそも規約の定めの内容に関わらず、まず、勝手に取り込むこと自体は問題ないと思います。文化庁の「考え方」も、そうした規約をいわば無視して取り込みを行っても、そのこと自体では著作権法第30条の4の適用は排除されないという立場です。では、規約が引き継がれるのかということになると、規約は人間と人間の間の約束ということになるので、規約を承認しますというのをクリックして入れ

ば、それはそういった契約関係が成立したことになるのですが、クローラで取り込んでしまったものに関しては、規約という概念が成り立たないと見ていいと思います。

内藤篤：スライドに戻ります。この著作権者の利益を不当に害する場合の例として、先ほど1～3の例が出ていましたが、その他にもあり得るだろうと言われているものとして、以下の①、②があります。

①「情報解析などの非享受目的利用を本来の用途とする著作物（データベース以外）を利用する場合」

このような著作物にはどんなものがあるのかよく分からないのですが、「情報解析の用に供することを目的として創作された画像などの著作物」というのが例示されているとのこと。

②「享受目的利用を本来の用途とする著作物について、情報解析などの非享受目的利用のライセンス提供が行われている場合にこれを利用すること」

そして、「学術論文のデータベースなど元々享受目的を本来の用途とするデータベース著作物が、情報解析用にライセンス販売されている場合にこれを利用する行為」という例が挙げられています。しかし、この②に関しては、学説が割れています。

論点：享受目的利用を本来的な用途とする著作物について非享受目的利用のライセンス提供が行われている場合の30条の4但書き該当性

✓ 限定的該当説

- 非享受目的のライセンスが提供されていることをもって直ちに但書きに該当するわけではないが、非享受目的のライセンス市場が発展し当該著作物について非享受目的の利用が本来的な利用と客観的に評価できるに至った場合には、そのライセンス市場と衝突するような利用は但書きに該当することもあり得る（澤田将史（『著作権法コンメンタール別冊 平成30年・令和2年改正解説』勁草書房、2022年 32頁））。

✓ 非該当説

- 30条の4により著作権侵害を構成しない行為について、余計な紛争・訴訟リスクや手間を避けるという理由などのために、本来は不要なはずのライセンスに応じるという取引慣行が一般化した既成事実それ自体が非侵害行為を侵害行為に転化させる理由とはならない（愛知靖（「AIと著作権」勁草書房、2024年 27頁））
- 対価を収受すべき本来的利用に該当するかは、著作物の性質や一般的な取引の実情によって定まると考えるべき。著作権者が非享受利用にかかるライセンスを事実上行っていたとしても、ただし書適用の根拠にはならない（前田健（「知財とパブリックドメイン第2巻著作権法篇」208頁））
- 本来30条の4により自由に利用できるため、ライセンス市場が形成されているとしても著作権法で保護することにはならない（柿沼太一・<https://storialaw.jp/blog/10806>）

内藤篤：この「非該当説」のほうを読むと、どのような対立状況かが分かりやすいかと思います。具体的に言うと、例えば医学雑誌を DVD にまとめたような類の（享受目的利用の）データベースがあった場合、これを学習用データとして非享受目的で利用する際にライセンスやその対価の支払いが必要かという話です。非該当説であれば、ライセンスやその対価の支払いは不要という話になります。一方で、限定的該当説の方は、市場が競合するのであれば、ライセンスを受けて対価を支払うべきだという説です。

3. 生成・利用段階の著作権侵害の論点

内藤篤：次に生成・利用段階の著作権侵害の話です。生成・利用段階では、①類似性、②依拠性、この2点がある場合に著作権侵害が認められてしまうのです。

【文化庁「考え方」：生成・利用段階の著作権侵害】

①類似性

- ✓ AI生成物と既存の著作物との類似性の判断については、AIを使わずに創作したものについて類似性が争われた既存の判例と同様

②依拠性

「考え方」では次の3つの場合に分けて考えを示している

a. AI利用者が既存著作物を認識していた場合

- ✓ 「考え方」ではこの場合依拠性が認められるとする
但し、AI利用者が既存著作物と同一・類似のものを生成する意図を有していることまで必要かについては議論有り

例) ミッキーを知る人が「何かキャラを出して」とプロンプトを入れた結果ミッキーに類似するキャラクターが生成された場合
生成意図不要とする説に立てば上記も依拠性が認められるが、生成意図を必要とする説によれば上記事例は依拠性が否定される。但し、必要説でも
・多数回生成を試行しその中から同一・類似物を選択した場合
・同一・類似物を生成するために詳細な指示を出した場合
等には依拠性が認められるとの立場もあり（愛知）。

内藤篤：似ている、似ていないというのは一般的な著作権侵害の考え方と同じですが、依拠という点においては、従来は、後行する作品を作る人が先行する作品を知っていたか知っていなかったかが、基本的には問題となってきました。しかし、AIに関しては、開発段階で元の著作物にアクセスがあれば、基本的には依拠ありと推認されるというところが、今回の「考え方」での目新しいポイントになります。

ただ、反論の余地はあります。「学習に用いられた著作物と創作的表現が共通した生成物が出力されないよう出力段階においてフィルタリングを行う措置が取られている場合」や、同じようなことですが、「生成 AI 自体の仕組み等に基づき、学習に用いられた著作物の創作的表現が生成・利用段階において生成されないことが合理的に説明可

能な場合」は依拠性の推認を覆すことができる可能性あると言われてしています。しかし、AI の利用者には分かりません。開発側において、そのような説明が可能なものなのかどうかというところが、この生成・利用段階において、安全・安心な AI という意味では、求められることになります。

4. 準拠法

内藤篤：最後の論点は準拠法です。ジュリスト論考では触れられていませんが、「考え方」には準拠法の話が出てきます。「損害賠償請求」と「差止請求」とで分けていて、差止請求は利用行為に関する保護国法が準拠法です。つまり、日本で裁判が行われたら日本国の著作権法です。損害賠償に関しては結果発生地の法律であると言っていて、「利用行為に関する準拠法は個別・具体的な事案に応じて」裁判所で判断される、と述べられています。

そして、開発・学習段階については、「AI 学習データの収集を行うプログラムが我が国内に所在するサーバー上で稼働していて、この収集に伴う既存の著作物等の複製を行っていること」を理由として、日本法が適用になる可能性が高いのではないかと述べられています。

【文化庁「考え方」：日本の著作権法が適用される範囲】

- ✓ **損害賠償請求**：結果発生地（法の適用に関する通則法第17条）
著作権侵害の結果発生地が日本国内と評価される場合、損害賠償請求（及びその著作権侵害に関する権利制限規定の適用）については、結果発生地法＝日本の著作権法が適用
- ✓ **差止請求**：利用行為に関する保護国法（ベルヌ条約第5条第2項第3文）
利用行為が行われた地が日本国内と評価される場合は、当該利用行為に伴う差止請求（及びその著作権侵害に関する権利制限規定の適用）については、保護国法＝日本の著作権法が適用
- ✓ 利用行為に関する準拠法決定は個別・具体的な事案に応じて、裁判所において判断されるが、次のような事情は日本の著作権法が適用される可能性を高める一要素になる。
 - ・ 生成AIの開発・学習段階において、AI学習データ収集を行うプログラムが我が国内に所在するサーバー上で稼働しており、この収集に伴う既存の著作物等の複製を行っていること
 - ・ 生成AIの生成・利用段階において、生成AIが我が国内に所在するサーバー上で稼働しており、既存の著作物等を含む生成物を生成していること
 - ・ 生成AIの生成・利用段階において、インターネット上で提供される生成AIサービスが、我が国内のユーザーに向けて既存の著作物等を含む生成物を公衆送信していること
- ✓ 「考え方」によれば結果発生地の法律が準拠法となるが、上記記載のようにAI学習のための学習データ収集を行うプログラムが我が国内に所在するサーバー上で稼働していること等をもって、日本の著作権法準拠と評価し得る有用な事実となるかは疑問

内藤篤：しかし、これは少しアプローチ的に違うのではないかというのが私の見方です。まず、インターネットが絡んだ際の結果発生地主義というのは、様々な要素を考慮に入れざるを得ないところはあるのです。従来は、ウェブサイトが日本語で書かれているとか、日本の顧客を主たるターゲットにしているというような、様々な理屈で日本に「結果発生地」を持ってくるという話がよくありました。しかし、クローリングに関しては、それは当てはまりません。

また、サーバー所在地主義という主張もあるのですが、これは、従来からあまり説得力がないと言われていた考え方でもあります。そのため下手をすると、例えばアメリカの著作物がクローリングの対象になったということがある場合、この考え方では結果発生地はアメリカだということになりかねない問題であると思います。