

国際 AI 規制関連文書マップ

第 1 版
(Revision 1.0.0)

2026 年 3 月 31 日

国立研究開発法人産業技術総合研究所

インテリジェントプラットフォーム研究部門
テクニカルレポート IPRI-TR-2026-04

前書き

本マップの作成は、国立研究開発法人新エネルギー・産業技術総合開発機構 (NEDO) の 2025 年度委託業務「AI の安全性確保に関する研究開発・検証等の推進事業／AI セーフティ強化に関する研究開発」(P25006)の一環として行った。

本マップは、AI のセーフティその他の品質の実現に取り組もうとする方々に参考情報として提供するものである。規範的な意味を持つものではない。

This document is distributed on an AS IS BASIS, WITHOUT WARRANTIES OF CONDITIONS OF ANY KIND, either express or implied.

ライセンス表示 (Copyright and Licensing)



本マップの著作権は国立研究開発法人産業技術総合研究所が保有します。国立研究開発法人産業技術総合研究所は、本マップの利用を、クリエイティブコモンズライセンス CC BY 4.0 (<https://creativecommons.org/licenses/by/4.0/>) の下で許可します。

Copyright © 2026 by the National Institute of Advanced Industrial Science and Technology. This document is licensed under the Creative Commons CC BY 4.0 License. To view a copy of the license, visit <https://creativecommons.org/licenses/by/4.0/>.

1 はじめに

近年、人工知能の社会実装が急速に進展する中で、AI の安全性や信頼性、社会的影響に対する関心が高まっている。これに伴い、各国・各地域において、法規制、ガイドライン、評価枠組みなど、AI に関する多様な文書が発行されてきた。これらの文書はいずれも、AI の利活用を一律に制限することを目的とするものではなく、リスクの低減と価値の創出を両立させるための考え方や枠組みを示している。

一方で、こうした AI 規制関連文書の数や分量は増加しており、個々の文書を通読したうえで相互の関係性や位置づけを把握することは容易ではない。文書ごとに用いられる用語、整理の観点、抽象度には差があり、ある文書で重視されている論点が、別の文書では別の形で表現されている場合もある。このため、AI の安全性や品質の確保に実務として取り組もうとする立場からは、個別文書の内容理解に加えて、それらを横断的に眺めるための視点が求められている。

本書は、そのような状況を踏まえ、AI 規制関連文書を横並びに概観するための整理を提供することを目的とする。本書が目指すのは、特定の文書や制度の解釈を示すことではなく、複数の文書に共通して見られる考え方や構造を整理し、それぞれの文書がどのような位置にあるのかを把握するための参照点を示すことである。

この目的を達成するため、本書では、個別文書の詳細な解説や比較評価を行うのではなく、文書間に共通する構造や整理軸に着目する。具体的には、AI 規制関連文書が何を目標として掲げているのか、どのような対象に対して対策を講じようとしているのか、また、それらを運用するための枠組みをどのように位置づけているのかといった観点から整理を行う。

本書の構成は、こうした考え方に基づいて設計されている。まず、AI 規制関連文書に共通する基本的な構造を整理し、その後、その構造を用いて複数の文書を横断的に概観する。さらに、生成 AI を対象とする文書については、AI 全般を対象とする文書とは異なる整理上の論点が現れるため、それを独立して整理する。

本書は、規範的な立場から遵守すべき内容を示すものではなく、AI 規制関連文書の全体像を把握するための参考情報として提供されるものである。読者が自らの立場や目的に応じて個別文書を読み解く際の補助となることを、本書の狙いとする。

本書が対象とする AI 規制関連文書

本書が対象とする文書は、表 1 に一覧として示すとおり、日本、米国、欧州において発行された、AI の利用、提供、開発に関わる規制またはガイドライン、評価枠組みである。これらの文書は、法的拘束力の有無や形式の違いはあるものの、いずれも AI の安全性、信頼性、またはそれらを含む品質の確保を目的としている点で共通している。

表 1 本書で取り上げる AI 規制関連文書

Formal name	Short name for this paper	Issuer	Version	Date	URL
AI Guidelines for Business Ver1.1 (AI事業者ガイドライン (第1.1版))	AI Guidelines for Business	Japan MIC and METI	1.1	2025/3/28	https://www.meti.go.jp/shingikai/mono_info_service/ai_shakai_jisso/20240419_report.html
Guide to Evaluation Perspective on AI Safety (Version 1.10) (AIセーフティ評価観点ガイド (第1.10版))	Guide to Evaluation Perspective	Japan AISI	1.10	2025/3/28	https://aisi.go.jp/output/output_framework/guide_to_evaluation_perspective_on_ai_safety/
Artificial Intelligence Risk Management Framework (AI RMF 1.0) (NIST-AI-600)	AI RMF	US NIST	1.0	2023/1	https://www.nist.gov/itl/ai-risk-management-framework
AI RMF PLAYBOOK	Playbook	US NIST	-	2023/1	https://airc.nist.gov/airmf-resources/playbook/
Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile (NIST-AI-600-1)	GenAI Profile	US NIST	-	2024/7	https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.600-1.pdf
2024/1689 REGULATION (EU) 2024/1689 OF THE EUROPEAN PARLIAMENT AND OF THE COUNCIL of 13 June 2024 laying down harmonised rules on artificial intelligence and amending Regulations (EC) No 300/2008, (EU) No 167/2013, (EU) No 168/2013, (EU) 2018/858, (EU) 2018/1139 and (EU) 2019/2144 and Directives 2014/90/EU, (EU) 2016/797 and (EU) 2020/1828 (Artificial Intelligence Act)	AI Act	EU	-	2024/6/13	https://eur-lex.europa.eu/legal-content/EN/TXT/PDF/?uri=OJ:L_202401689
Code of Practice for General-Purpose AI Models (Transparency Chapter; Copyright Chapter; Safety and Security Chapter)	GPAI CoP	EU	-	2025/7/10	https://digital-strategy.ec.europa.eu/en/policies/contents-code-gpai
Guideline for Designing Trustworthy Artificial Intelligence – AI Assessment Catalog	AI Assessment Catalog	Germany Fraunhofer IAIS	-	2023/2	https://www.iais.fraunhofer.de/en/publications/studies/2023/ai-assessment-catalog.html
White Paper - Developing trustworthy AI applications with foundation models	White Paper for Foundation Models	Germany Fraunhofer IAIS	-	2024/1	https://arxiv.org/abs/2405.04937

対象文書は、大きく二つの系統に分けられる。

第一は、AI 全般、あるいは予測 AI を含む幅広い AI システムを対象とした文書である。これらの文書は、特定の技術類型に限定されず、AI 一般に共通する品質やリスク、あるいはマネジメントの考え方を扱っている。

第二は、生成 AI を主な対象とした文書である。これらの文書は、大規模言語モデル等に代表される生成 AI の特性を前提として、固有のリスクや対応の考え方を整理している。

対象文書の選定にあたっては、発行主体の地域的な偏りを避けること、AI の開発者、提供者、利用者といった複数の立場を想定していること、あるいは国際的に参照される構造や枠組みを持つことなどを考慮している。一方で、発行時期や改訂頻度、実務における普及度合いなどについては、本書の選定基準とはしていない。

表 1 に示した文書群は、本書の執筆時点で入手可能であった公開文書であり、今後の改訂や新規文書の発行によって、AI 規制を取り巻く状況が変化する可能性がある。本書は、特定時点における文書群を前提とした整理であることに留意されたい。

なお、本書では、これらの文書に示された要求事項や考え方について、遵守の要否や妥当性を評価する立場は取らない。本書の位置づけは、あくまで参照情報として、AI 規制関連文書の全体像を把握するための整理を提供することにある。

2 AI 規制関連文書に共通する基本構造

本章では、本書全体で用いる整理の前提として、AI 規制関連文書に共通して見られる基本的な構造を整理する。

本書の目的は個別文書の内容を詳細に解説することではなく、複数の文書を横並びに眺めることで全体像を把握することにある。そのためには、文書間で共通に用いられている思考の枠組みをあらかじめ明確にしておく必要がある。

AI に関する規制やガイドラインは、法的拘束力の有無、対象とする技術の範囲、想定する読者などが異なっている。しかし、それらの違いにもかかわらず、多くの文書には共通する構造が確認できる。本章では、その構造を抽象化し、後続の整理の基盤とする。

2.1 管理指標と介入領域という基本整理

AI 規制関連文書の多くは、達成を促す目標と、その目標を実現するための対応対象という二つの要素から構成されている。

本書では、前者を管理指標、後者を介入領域と呼ぶ。

管理指標とは、AI システムの開発、提供、利用にあたって望ましい状態として示される品質や、低減すべきリスクの方向性を指すものである。文書ごとに表現や抽象度は異なるが、何を実現すべきか、どのような点に配慮すべきかという目標が示されている点は共通している。

介入領域とは、管理指標の達成に向けて具体的な対応を講じる対象となる組織的、人的、あるいはシステム上の要素を指す。多くの文書では、管理指標が示された後に、それを実現するための考え方や取組の方向性が述べられるが、その対応対象の想定や具体性には幅がある。

管理指標はほぼすべての文書で明示されている一方、介入領域については詳細な整理がされていない文書もあれば、特定の領域に重点を置いて記述している文書もある。この点は、文書を横断的に整理する際の重要な観点となる。

2.2 管理指標の主な類型

AI 規制関連文書で示される管理指標は多様であるが、複数の文書を整理すると、共通して取り上げられる論点の類型が存在する。本書では、その中でも特に多くの文書に共通して確認できるものとして、五つの類型を整理する。

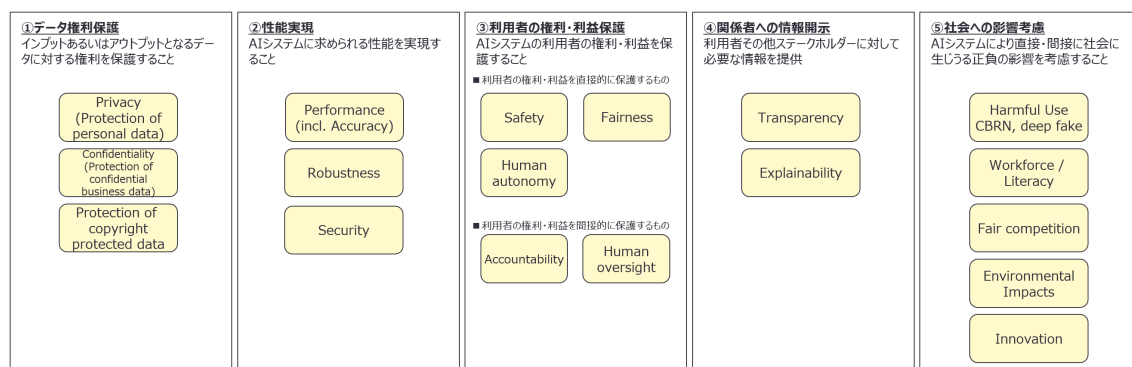


図 1. 管理指標の類型

第一に、データ権利保護である。これは、AI システムの入力データおよび出力データに関する権利を保護することを指す。個人情報、著作物、企業機密といったデータの取扱いが含まれ、学習データだけでなく、生成されたアウトプットに関する権利の扱いが論点となる場合もある。

第二に、性能および頑健性である。AI システムが想定された目的に照らして一定の性能を発揮すること、また、環境の変化や入力の変動に対して安定した動作を保つことが求められる。AI システム固有の性質として、性能指標の設定や評価が容易でない点が指摘されることも多い。

第三に、利用者の権利および利益の保護である。ここでいう利用者とは、契約関係の有無にかかわらず、AI システムを直接利用する者や、その出力や挙動によって直接の影響を受ける者を含む。AI システムの利用によって、不当な不利益や権利侵害が生じないように配慮することが求められる。

第四に、透明性および情報開示である。AI システムは、その判断過程や動作原理が外部から把握しにくい場合がある。そのため、利用者や関係者が AI システムを適切に理解し、安心して利用するための情報提供が管理指標として示される。

第五に、社会的影響への配慮である。AI システムの影響は、個々の利用者にとどまらず、社会全体に及ぶ場合がある。イノベーションへの寄与、悪用の防止、環境への影響、雇用や公正競争への影響などが、この類型に含まれる。

これらの管理指標は、AI システムであるがゆえに特に考慮する必要がある品質やリスクを整理したものであり、情報システム一般に共通するすべての品質属性を網羅するものではない。

2.3 介入領域の整理と本書での扱い

管理指標の達成に向けた対応は、単一の対象に限定されるものではなく、複数の介入領域にまたがって行われる。本書では、文書間の比較を行うための便宜的な整理として、介入領域をいくつかの要素に分けて整理する。

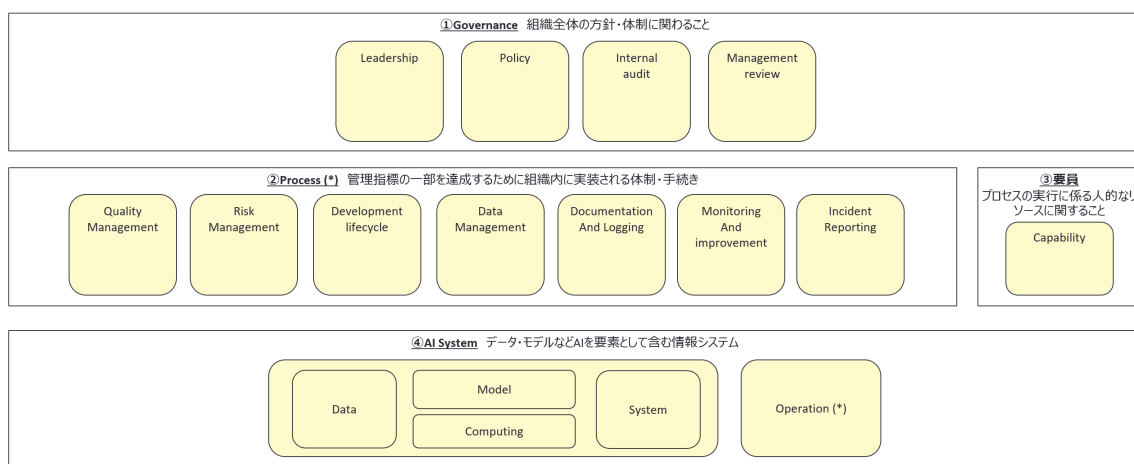


図 2. 介入領域の整理

一つ目は、ガバナンスである。これは、組織全体の方針、意思決定体制、責任分担など、AI に関する取組を支える統治に関わる要素を指す。

二つ目は、プロセスである。管理指標の一部を達成するために、組織内で定められる手続きや業務プロセスが該当する。

三つ目は、要員である。プロセスを実行し、AI システムの開発や運用、管理に関与する人的リソースに関する要素を指す。

四つ目は、Operation である。本書では、利用者に対して AI システムをサービスとして提供するために必要となる人間による運用のうち、必要性が高い部分を Operation として整理する。利用申請の受付や利用規約の提供など、サービス提供の前提となる運用が含まれる。

五つ目は、System である。ここでは、AI を要素として含む情報システムのうち、AI 固有のデータ、モデル、計算資源以外に、情報システムとして一般に求められる構成要素を指す。

これらの介入領域は、すべての文書において同一の形で示されているわけではない。また、文書によっては、これらを明確に区別せずに記述している場合もある。本書における整理は、文書間の比較を可能にするための枠組みとして採用するものであり、唯一の分類体系であることを意図するものではない。

3 AI 全般対象文書の横断的マップ

本章では、AI 全般を対象とする主要な AI 規制関連文書について、本書で定義した管理指標および介入領域の枠組みを用いて整理する。ここでの目的は、個々の文書の内容を詳細に説明することではなく、複数の文書を同一の観点から眺めた場合に、どのような構造的な違いが見えてくるかを確認することにある。

AI 全般を対象とする文書には、法令、ガイドライン、評価カタログなど、性格の異なるものが含まれている。本章では、それらを同列に扱い、それぞれが管理指標をどのように設定し、介入領域についてどの程度踏み込んでいるかを整理する。

3.1 本章における比較の視点

本章で用いる比較の視点は、主に三つである。

第一に、管理指標がどのように示されているかである。具体的な品質属性やリスク類型として整理されているのか、あるいはより抽象的な原則や価値として示されているのかという点を確認する。

第二に、介入領域への踏み込み方である。管理指標の達成に向けて、どのような組織的、人的、システムの対応が想定されているか、また、その想定がどの程度具体的に記述されているかに着目する。

第三に、マネジメントプロセスの位置づけである。管理指標や介入領域を実際に運用していくためのプロセス、すなわちリスク評価、見直し、改善といった活動が、文書の中でどのように扱われているかを確認する。

なお、AI ライフサイクルや対象システムの文脈といった観点も、文書によっては言及されているが、本章では参考的な位置づけとし、比較の主軸とはしない。

3.2 各文書の位置づけ

ここでは、AI 全般を対象とする代表的な文書について、本章の比較視点に基づき、その位置づけを整理する。整理は文書ごとの特徴を把握するためのものであり、優劣や成熟度を評価することを意図するものではない。

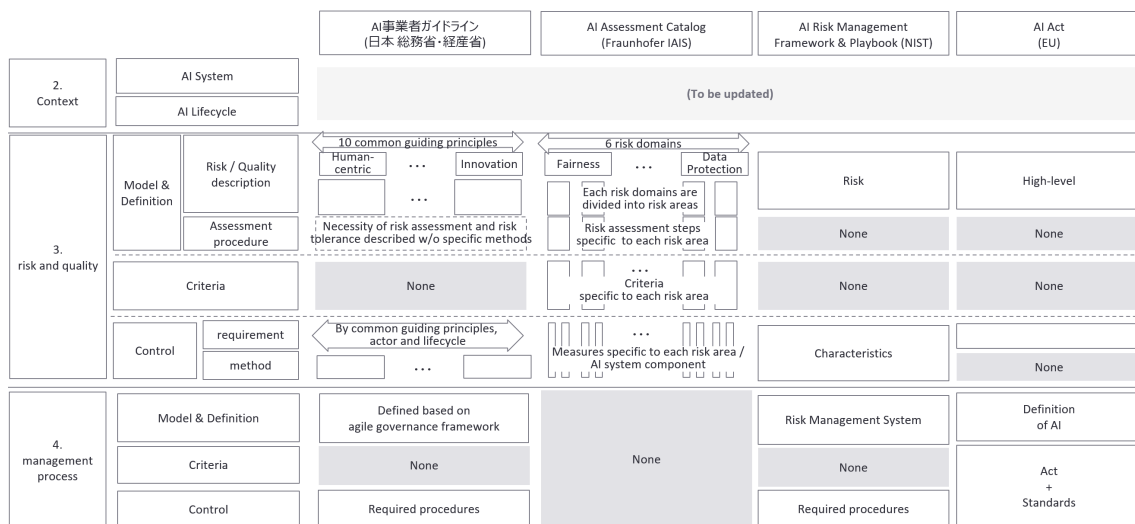


図 3. AI 全般を対象とする規制関連文書

日本の AI 事業者ガイドラインは、AI に関わる主体を開発者、提供者、利用者といった立場に整理したうえで、共通の指針を示す構成をとっている。管理指標は、原則や指針の形で明示されている一方、それらを実現するための具体的な評価手法や実装方法については、本文では限定的に扱われている。マネジメントプロセスについては、全体像としての考え方が示されている。

Fraunhofer による AI Assessment Catalog は、AI システムの信頼性を評価することに主眼を置いた文書である。管理指標は、信頼性に関する複数の次元として体系的に整理されており、それぞれについて評価基準や測定の観点が示されている。一方、マネジメントプロセスについては、評価を中心とした簡潔な想定が置かれている。

米国における NIST AI Risk Management Framework および関連文書は、AI リスクマネジメントの枠組みそのものを中心的な主題としている。管理指標は、信頼できる AI システムが備えるべき特性として整理されており、それを実現するための取組がマネジメントプロセスとして体系化されている。文書全体を通じて、プロセスの構造が明示されている点が特徴である。

欧州の AI 法は、AI に関する規制を法的に定めた文書であり、他のガイドライン型文書とは性格を異にする。管理指標に相当する要件は、特定のリスク区分に応じた義務と

して示されている。介入領域やマネジメントプロセスの詳細については、法律本文では抽象的な記述にとどまり、具体化は関連する整合標準等に委ねられている。

3.3 横断比較から見える構造的な整理

これらの文書を横断的に整理すると、いくつかの構造的な特徴が確認できる。

第一に、管理指標に相当するものは、すべての文書において何らかの形で示されている点である。表現や粒度は異なるものの、何を達成すべきかという方向性が欠けている文書は見られない。

第二に、管理指標に関する評価方法や評価尺度まで踏み込んでいるかどうかには差がある。一部の文書では、評価基準や測定の観点が明示されている一方、他の文書では、それらが読者の判断や後続文書に委ねられている。

第三に、マネジメントプロセスの扱い方にも違いが見られる。プロセスそのものを文書の中心に据えて体系化しているもの、評価行為を中核に据えて簡潔なプロセスを想定しているもの、原則や指針を示すにとどめているもの、法的枠組みとして外部の制度や標準との連携を前提としているものなど、文書ごとの立場が反映されている。

本章では、これらの違いを個別に評価するのではなく、AI 全般を対象とする規制関連文書が、管理指標、介入領域、およびマネジメントプロセスをどのように構成しているかという観点から整理した。

4 生成 AI 対象文書に現れる新たな分岐

本章では、生成 AI を主な対象とする AI 規制関連文書を取り上げ、それらがどのような観点から安全性や品質を捉えようとしているかを整理する。

ここでの目的は、生成 AI に関する個別のリスクや対策を詳細に論じることではなく、AI 全般を対象とする文書と同じ整理枠組みを適用した際に、どのような点で整理が難しくなるのかを明らかにすることにある。

生成 AI を対象とする文書群は、発行時期に近いものが多い一方で、その対象や考え方には幅がある。本章では、その幅を無理に統合するのではなく、どのような分岐が存在しているかを整理する。

4.1 生成 AI 文書の多様性と比較の難しさ

生成 AI を対象とする文書は、AI 全般を対象とする文書と比べて、対象システムの想定や論点設定が一樣ではない。

同じ生成 AI という用語を用いていても、文書によって、主な対象が利用システムである場合もあれば、モデルそのものを想定している場合もある。

		GenAI QM GL (AIST)	GL for Business (Japan MIC, METI)	Safety Guide (Japan AISI)	EU AIA & GPAI CoP (EU)	AI Assessment Catalog and White Paper (Fraunhofer IAS)	NIST AI RMF & GenAI Profile (US NIST)
Target	Model / System	Software systems with generative AI models	Advanced AI systems (incl. FM, GenAI)	LLM	GPAI	AI systems with FM	Generative AI
	Actor	Provider	Model Developer / Provider	Provider	Model Developer / Provider	Provider	All actors (Not distinguished)
Risk / Quality		Sufficiency to requirements; Reliability; Interactiveness (and other seven criteria)	Human-centric Safety Privacy protection (and other seven criteria)	Control of toxic output; Prevention of Misinformation, Disinformation and Manipulation; Fairness and inclusion (and other seven perspectives)	Four risk levels are specified	Same as AI Assessment Catalog (six dimensions). Specific risks are provided	Same as NIST AI RMF. Specific risks are provided
Control		Exemplary controls provided by component / quality (Components: overall system, LLM selection and learning, prompt and relevant components)	Eleven actions are provided (not only for individual AI system, but for community)	No control (only assessment perspectives)	Some measures to be taken by the developer / provider are specified (Details to be provided as standards)	The steps supplementary to the steps defined in AI Assessment Catalog are provided.	Suggested actions to manage GAI risks are provided on top of the suggested actions defined in AI RMF

図 4. 生成 AI を対象とする規制関連文書

利用システムを主な対象とする文書では、生成 AI モデルを構成要素の一つとして扱い、特定の用途における品質や安全性の確保が主な関心となる。一方で、モデル自体を主な対象とする文書では、特定の用途を前提としない形で、モデルの能力や特性に内在するリスクが論点として取り上げられる。

このように、対象の取り方が異なるため、管理指標や介入領域の整理だけでは、文書間の考え方の違いを十分に説明できない場合が生じる。本章では、その点を踏まえ、生成 AI 文書の整理において追加的な視点が必要となる理由を確認する。

4.2 新たな整理軸としての用途起点と能力起点

生成 AI を対象とする文書を整理するにあたり、本書では、管理指標および介入領域に加えて、もう一つの整理軸を導入する。

それは、安全性や品質を考える際の出発点がどこに置かれているかという観点である。

本書では、この観点を用途起点と能力起点という二つの起点として整理する。

用途起点とは、生成 AI が組み込まれた特定の利用システムや利用文脈を前提とし、その用途に即して必要な品質やリスク対応を考える立場である。この立場では、生成 AI モデル単独の性質よりも、システム全体としての振る舞いが重視される。

能力起点とは、生成 AI モデルそのものが持つ汎用的な能力や特性に着目し、それによって生じうるリスクを問題にする立場である。この立場では、特定の用途に限定せず、広く利用されることによって生じる影響が論点となる。

この二つの起点は、どちらが正しいという関係にあるものではなく、文書ごとにどちらを主に採用しているかが異なっている。

4.3 用途起点で整理される生成 AI 文書

用途起点の立場をとる生成 AI 文書では、生成 AI モデルを部品として用いるシステムを想定し、そのシステムが特定の目的を果たすにあたって満たすべき品質や安全性が整理されている。

この種の文書では、生成 AI モデルそのものについて直接的な制御ができない場合を前提とし、モデルを組み込んだシステム全体として、どのような管理や評価が可能かが論じられる。管理指標は、システムの用途に紐づいた形で設定され、介入領域も、システム構成や運用に重点が置かれることが多い。

また、リスクの考え方も、特定の利用場面において、どのような影響が生じうるかという観点から整理される。このため、同じ生成 AI モデルであっても、用途が異なれば、考慮すべきリスクや対応は異なるという前提が置かれている。

4.4 能力起点で整理される生成 AI 文書

能力起点の立場をとる生成 AI 文書では、生成 AI モデルが持つ汎用的な能力そのものに着目し、その能力がもたらしうるリスクが整理される。

ここでは、モデルの規模や性能、自己学習能力といった特性が、どのような影響を及ぼしうるかが論点となる。

この種の文書では、特定の利用システムや用途を前提としないため、管理指標は、透明性、セーフティ、セキュリティといった横断的な性質として示されることが多い。介入領域についても、モデルの開発や提供に関わる主体の行為や、情報開示のあり方が中心となる。

また、能力起点の整理では、モデルが幅広い用途に利用される可能性を前提としているため、個別用途では把握しにくい影響についても考慮することが求められている。

4.5 両アプローチの間にある整理上の課題

用途起点と能力起点という二つの整理は、それぞれ異なる前提と関心に基づいており、単純に一方を他方に置き換えることはできない。用途起点では、具体的なリスク評価や対策が想定しやすい一方で、用途に依存しない影響を整理することが難しい場合がある。

一方、能力起点では、生成 AI モデルの特性に起因する広範な影響を捉えようとするが、特定の利用場面におけるリスク評価との接続が課題となる場合がある。

生成 AI を対象とする規制関連文書の整理においては、このような出発点の違いが存在すること自体を認識したうえで、文書の位置づけを把握する必要がある。本章では、生成 AI 文書に見られる分岐と、その整理上の課題を確認した。

5 国際 AI 規制関連文書マップから得られる示唆

本章では、前章までの整理を踏まえ、本書が提示した国際 AI 規制関連文書マップから得られる示唆を整理する。ここで示す示唆は、特定の制度や文書の優劣を論じるものではなく、複数の文書を俯瞰的に整理した結果として確認できる観点をまとめたものである。

5.1 共通して確認できる前提

まず、対象とした AI 規制関連文書全体を通じて、いくつかの共通した前提が確認できる。

第一に、AI の利活用に伴うリスクの存在が広く前提とされている点である。文書ごとに用語や整理の仕方は異なるが、AI システムが社会や利用者に影響を及ぼするという認識は共有されている。

第二に、リスクの低減と価値創出を対立概念として扱っていない点である。多くの文書では、AI による便益を前提としつつ、その実現のためにどのような配慮や対応が必要かという形で整理がなされている。

第三に、AI を単体の技術要素としてではなく、組織的、社会的な文脈の中で扱う必要があるという認識である。管理指標に加えて、ガバナンスやプロセスといった介入領域が繰り返し言及されていることは、その表れと位置づけられる。

5.2 文書ごとの重心の違い

一方で、AI 規制関連文書の整理を通じて、文書ごとに重心の置き方が異なっていることも明らかになる。

ある文書では、達成すべき管理指標の整理が中心となり、別の文書では、それらをどのように運用するかというマネジメントプロセスに重点が置かれている。また、法的枠組みとして義務を定める文書と、ガイドラインとして考え方を示す文書とでは、介入領域への踏み込み方にも差が見られる。

生成 AI を対象とする文書については、用途起点と能力起点という異なる整理の出発点が存在し、それぞれが異なる論点を浮かび上がらせている。この違いは、生成 AI という技術の性質に起因するものでもあり、単一の整理軸で全体を説明しきることが難しい状況を示している。

こうした重心の違いは、文書間の不整合を意味するものではなく、文書が想定する対象や役割の違いを反映したものとして理解することができる。

5.3 実務的な読み取りに向けた観点

本書の整理から得られる示唆の一つは、AI 規制関連文書を読む際に、個別の記述をそのまま自組織に当てはめるのではなく、文書がどの位置づけで書かれているかを意識することの重要性である。管理指標を示す文書なのか、介入領域やプロセスに重点を置く文書なのかによって、読み取り方は異なる。

また、AI 全般を対象とする文書と、生成 AI を対象とする文書では、前提としているリスクの捉え方や整理軸が異なる。その違いを理解せずに文書を横断的に比較すると、要求事項の多さや抽象度の違いがそのまま混乱につながる可能性がある。

本書が提示するマップは、こうした文書の位置関係を整理し、どの文書がどの論点に応答しようとしているのかを把握するための参照点を提供することを意図している。

5.4 本書の位置づけと今後の活用

本書は、AI 規制関連文書に関する最終的な解釈や判断を示すものではない。また、特定の文書への適合や遵守を促すものでもない。本書の役割は、複数の文書を俯瞰的に整理し、読者が自らの立場や目的に応じて文書を読み解くための補助的な枠組みを提供することにある。

AI を取り巻く規制やガイドラインは今後も変化し続けることが想定される。本書で示した整理は、特定時点の文書群を前提としたものであり、今後の改訂や新たな文書の出現によって見直しが必要となる可能性がある。その意味で、本書の整理も固定的なものではなく、状況に応じて更新されることを前提とした参照資料として位置づけられる。

本書が、AI の安全性や品質の確保に取り組む際の出発点として、文書全体を見渡すための視点を提供することを期待する。