

中間とりまとめ

2025 年 2 月 4 日

AI 戦略会議・AI 制度研究会

目次

概要	1
I. はじめに	2
II. 制度の基本的な考え方	5
1. 近年の AI の発展	5
2. 関係主体	5
(1) 主な主体	5
(2) 国外事業者	6
3. イノベーション促進とリスクへの対応の両立	6
(1) イノベーションの促進	6
① 研究開発への支援	6
② 事業者による利用	7
(2) 法令の適用とソフトローの活用	8
(3) リスクへの対応	10
4. 国際協調の推進	11
(1) AI ガバナンスの形成	11
(2) 国際整合性・相互運用性の確保	12
III. 具体的な制度・施策の方向性	13
1. 全般的な事項	13
(1) 政府の司令塔機能の強化、戦略の策定	13
(2) 安全性の向上等	14
① AI ライフサイクル全体を通じた透明性と適正性の確保	14
② 国内外の組織が実践する安全性評価と認証に関する戦略的な促進	15
③ 重大インシデント等に関する政府による調査と情報発信	16
2. 政府による利用等	17
(1) 政府調達	17
(2) 政府等による利用	17
3. 生命・身体の安全、システミック・リスク、国の安全保障等に関わるもの	19
IV. おわりに	19

概要

AI戦略会議・AI制度研究会 中間とりまとめ 概要

2024年7月以降、AI制度研究会¹⁾を計7回開催。計15の研究者、事業者等からのヒアリングを含む議論、パブリックコメントを経て中間とりまとめを作成。

背景

- AIは我が国の発展に大きく寄与する可能性がある一方、**様々なリスクが顕在化**。
- AIに対する不安の声が多く、諸外国と比べても**開発・活用が進んでいない**との指摘。
- ▶ AIの透明性など、**適正性を確保し、AIの開発・活用を進める**必要がある。

基本的な考え方

■ イノベーション促進とリスク対応の両立

(Ⅱ.3.)

- 研究開発支援、人材育成、データや計算資源の整備などイノベーションの促進
- 法令とガイドライン等の適切な組合せ
- OECD原則、広島プロセス国際指針等の共通的な指針等と個別の既存法令の活用



■ 国際協調 (Ⅱ.4.)

- AIガバナンスの形成に向けて議論をリード
- 国際整合性・相互運用性の確保



具体的な制度・施策の方向性

■ 全般的な事項 (Ⅲ.1.)

「世界で最もAIを開発・活用しやすい国」を目指す

● 政府の司令塔機能の強化、戦略の策定

- ・ 全体を俯瞰する**司令塔機能強化**
- ・ AIの安全・安心な研究開発・活用のための**戦略（基本計画）の策定**

速やかな法制度化が必要
世界のモデルになるような制度

● 安全性の向上等

- ・ **国による指針（広島AIプロセス準拠）の整備、事業者による協力**
- ・ **国による調査・情報収集、事業者・国民への指導・助言、情報提供等**



■ 政府等による利用 (Ⅲ.2.)

- 適正なAI政府調達・利用 等

■ 基盤サービス等における利用 (Ⅲ.3.)

- 各業法等による対応 等

1) 官房長官が議長、全閣僚が構成員となっている「統合イノベーション戦略推進会議」の下に「AI戦略会議」を設置。その下に「AI制度研究会」を設置。

2) 上記の政策を講じた上で、今後のリスク対応のため引き続き制度の検討を実施すべき。

I.はじめに

2022 年秋以降、生成 AI の性能は飛躍的に向上し、豊富な情報を処理する AI による、自然な会話、プログラム、精巧な動画等の出力が可能となっている。このような AI は、これまで人が行っていた作業を代替し、又は人が行っていた以上の成果を創出することが可能であり、人間が携わるあらゆる分野における活用が想定され、今後、産業や国民生活の様々な分野において効率性や利便性を大きく向上させるほか、新しい科学的発見や人間の創造性を補う役割も期待されており、国民生活の向上、国民経済の発展に大きく寄与する可能性がある。

他方で、AI による偽サイトや合成音声が詐欺等に使用される犯罪の巧妙化、偽・誤情報の作成に AI が使用され、拡散されることによる情報操作等多様なリスクが顕在化しつつある。また、AI にはデュアルユース技術の側面もあり、CBRN（Chemical [化学]・Biological [生物]・Radiological [放射性物質]・Nuclear [核]を用いた兵器等）の開発やサイバー攻撃等に AI が使用される安全保障上のリスクも指摘されている。

このような中、EU においては、2024 年 8 月、AI に関する包括的な規制である AI Act が発効した。AI Act は、4 段階のリスクに応じたアプローチを採用し、人間の安全や基本的権利を脅かす AI についてはその市場投入や使用を禁止し、人間の健康・安全や民主主義・法の支配に重大な害を及ぼす恐れのある AI についてはハイリスクな AI システムとして市場投入前に影響評価・適合性評価を行う義務を課すといった規制を導入している。また、汎用 AI モデルの提供者には、技術文書の作成や学習データの開示等の透明性の義務を課し、学習の計算量が 10^{25} FLOPs を超える等システムック・リスクを伴う汎用 AI モデルの提供者にはモデル評価の実施やインシデントに関する報告義務を上乗せで課すといった規制を設けている。

米国においては、2023 年 7 月以降、AI のリスク管理に関する情報共有や AI システムがもたらす可能性のある社会的リスクに関する研究、サイバーセキュリティへの投資等を実施する旨のボランティア・コミットメントを米国の大手 AI 開発企業が発表した。また、2023 年 10 月、安全保障上のリスクへの対応のため、国防生産法に基づき、学習の計算量が 10^{26} FLOPs を超える潜在的なデュアルユース基盤モデルを開発する米国の事業者に対し、モデルの訓練、開発又は製造に関する活動内容等の情報を継続的に政府に提出するよう指示する内容を含む大統領令を発出した。なお、本大統領令は 2025 年 1 月 20 日に撤回され、同年 1 月 23 日に、AI に関する新たな大統領令が発出されており、2023 年 10 月 30 日の大統領令の下でとられた政策や規制等の措置の見直しと、行動計画の策定が指示されている。その他、カリフォルニア州においては、2024 年 9 月、AI が生成したコンテンツの透明性を高める州法 SB942 と、AI の学習に使用したデータを開示する州法 AB2013 が制定された。

我が国が議長国を務めた 2023 年の G7 では、生成 AI に関する国際的なルールの検討のため、「広

島 AI プロセス¹」を立ち上げ、安全、安心で信頼できる AI の実現に向け、「全ての AI 関係者向けの広島プロセス国際指針」及び「高度な AI システムを開発する組織向けの広島プロセス国際行動規範」（以下「広島プロセス国際指針等」という。）を策定した。その後、我が国は、G7 を超えたアウトリーチとして、広島 AI プロセス・フレンズグループを立ち上げ、広島 AI プロセスの精神に賛同する国々の拡大を図っているところである。また、2024 年には、G7 議長国のイタリアが広島 AI プロセスを引き継ぎ、国際行動規範の履行状況の報告枠組みについて議論が行われている。このほか国連、欧州評議会、OECD 等の多国間の枠組みにおいても、AI ガバナンスに関する議論が活発に行われている。

国内においては、2023 年 5 月以降、「AI に関する暫定的な論点整理」（2023 年 5 月 26 日 AI 戦略会議）や「AI 制度に関する考え方」について」（2024 年 5 月 AI 戦略チーム）をとりまとめ、AI に関する論点整理を行い、AI 制度の考え方を示したほか、「統合イノベーション戦略 2024」（2024 年 6 月 4 日閣議決定）では、AI 分野の競争力強化と安全・安心の確保を目的とした戦略を策定している。さらに、AI の安全・安心な活用の促進のため、「AI 事業者ガイドライン（第 1.01 版）」（2024 年 11 月 22 日 総務省 経済産業省）にて、事業活動における AI 開発、提供、利用にあたっての考え方を示す等、関係府省庁が連携して検討し、対応している。

AI に関する意識調査の結果（図 1 参照）によると、日本では、「現在の規則や法律で AI を安全に利用できる」と思う回答者は 13%と低く、77%の人が「AI には規制が必要」と考えている。また、「品質の不安定さ」、「プロセスのブラックボックス化」等についてリスクを感じているほか、政府に求めることとして、「AI の悪用や犯罪に対する法的対策の強化」が挙げられている。AI のリスクへの対応については、欧米を中心とする各国においては AI に関する法制度の議論や検討が進んでいる一方で、我が国においては、ガイドライン等のソフトローによる対応が中心であり、AI に特化した法制度の検討は行われていない状況である。

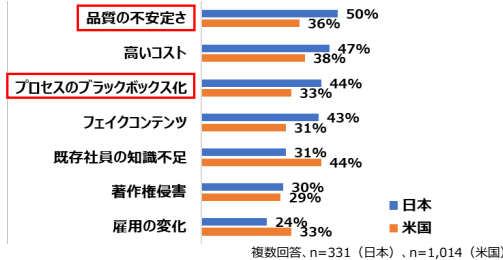
上記の状況を踏まえ、2024 年 7 月、AI 戦略会議の下、AI 制度研究会が設置され、事業者、有識者、自治体を含む様々な関係者からヒアリングを行い、法制度の要否を含む、AI 制度のあり方について検討を行った。検討にあたっては、広島 AI プロセスの精神に基づき、また、「リスク対応とイノベーション促進の両立」、「技術・ビジネスの変化の速さに対応できる柔軟な制度の設計」、「国際的な相互運用性」及び「政府による AI の適正な調達と利用」の 4 つを基本原則として、議論を行った。本とりまとめは、当該ヒアリングや議論等を踏まえた検討結果をとりまとめたものである。なお、AI の開発・利用の過程における著作権を含む知的財産権法等の扱いは、関係省庁で検討が行われたほか、安全保障の観点での AI の活用については、安全保障関係省庁を中心に別途検討が進められている。

¹ 2023 年 5 月に開催された G7 広島サミットの結果を踏まえ、その急速な発展と普及が国際社会全体の重要な課題となっている生成 AI について議論するために、2023 年 5 月に立ち上がったもの。

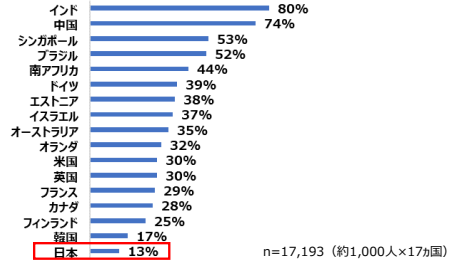
AIのリスクや安全性に関する意識調査（国際比較）

- 日本は米国と比較して、「品質の不安定さ」「プロセスのブラックボックス化」「フェイクコンテンツ」に不安を感じている企業が多く、AIリスクへのガバナンスの取組をする企業が少ない
- 日本では、「現在の規則や法律でAIを安全に利用できる」と思う回答者は13%と調査対象国の中で最も低く、『AIには規制が必要』と考える回答者は77%と調査国の中で中位に位置する

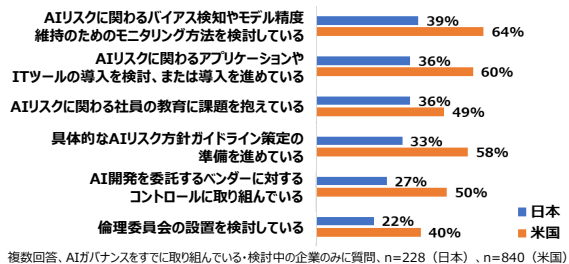
生成AI活用に関して感じるリスク



現在の規則や法律でAIを安全に利用できると思う

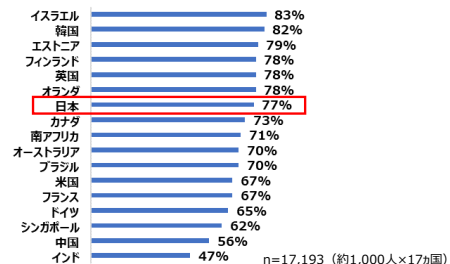


AIリスクへのガバナンス施策の取組状況



出典：PwC「2023年AI予測」を基に内閣府で作成

AIには規制が必要だと思う

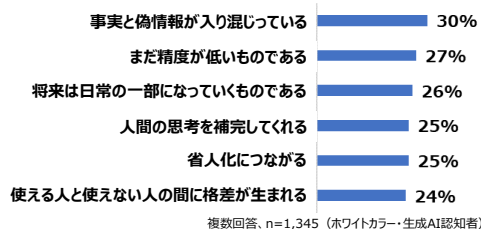


出典：KPMG「Trust in Artificial Intelligence: A global study」を基に内閣府で作成

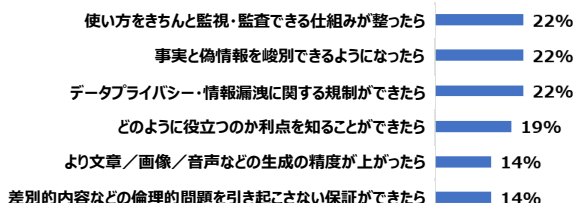
生成AIに関する意識調査（日本）

- 生成AIのイメージは「事実と偽情報が入り混じっている」という回答が30%と多く、どのようになれば生成AIを使いたいと思うかは「監視・監査できる仕組み」「事実と偽情報の峻別」「プライバシー等に関する規制」という回答が22%と多い
- 企業に求めたいことは「顧客データの安全性とプライバシー保護の強化」が66%と一番多かったのに対し、政府には「AIの悪用や犯罪に対する法的対策の強化」が66%と一番多い

生成AIのイメージ

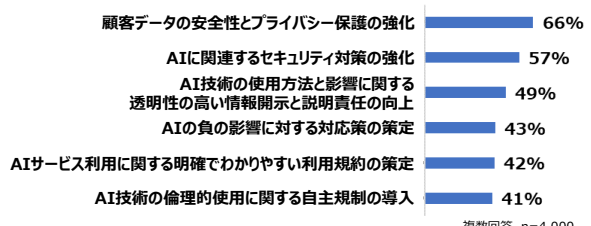


どのようになれば生成AIを使いたいと思うか

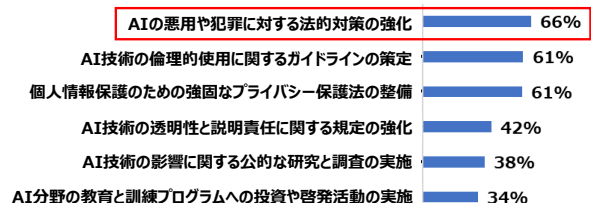


出典：クロス・マーケティング「生成AIに関する調査（2023年）利用実態・意識編」を基に内閣府で作成

企業に求めたいこと



政府に求めたいこと



出典：国際大学グローバル・コミュニケーション・センター「Innovation Nippon 2024 生成AIと日本」を基に内閣府で作成

図1 生成AIに関する意識調査

II.制度の基本的な考え方

1.近年の AI の発展

AI の分類にあたっては、様々な方法が考えられるが、例えば、特化型 AI と汎用型 AI に分類できる。特化型 AI は、音声認識、画像認識、自動運転等の特定のタスクを処理することに特化した AI である。汎用型 AI は、特化型 AI よりも大量のデータで学習され、高い汎用性を示し、様々なタスクを処理できる AI であり、近年、その可能性から大きな注目を集めている生成 AI²は一般的に汎用型 AI に属される。汎用型 AI は、一般的に、学習データやモデルのパラメータ数が多くなれば性能が向上すると考えられていたが、最近では学習データ等の規模によらず性能が高いものも登場している。将来的には、様々なタスクを人間と同等のレベルで実現できる能力を持つ AGI（Artificial General Intelligence）が登場するという意見もある。このように、AI の技術は近年目覚ましい発展を遂げている。なお、「AI」、「事業者」等 AI に関する用語の定義については、国際的にも議論がなされており、今後とも変化すると考えられることから、これらの議論も踏まえ検討することが重要である。

2.関係主体

（1）主な主体

本とりまとめでは、データ収集、モデルの開発等を行い、AI システム（AI モデルを含むものとする）を開発し、最終的に AI サービスを利用するまでのライフサイクルにおいて、主に AI 開発者、AI 提供者、AI 利用者の3つの主体が存在するものとして、述べることとする（図2 参照）。

まず、AI 開発者は、データ収集やモデル学習、そのほかモデルのシステム基盤構築や入出力機能等の開発を行う者とする。次に AI 提供者は、既存または新規システムに AI を組み込み、サービスに利用可能な状態で AI システムを提供する者、または AI の組み込みから AI サービスの提供まで実施する者とする。最後に AI 利用者は、他者が実装した AI システムをサービスに組み込み、AI サービスとして利用する者、または提供されている AI サービスを利用する者とする。

なお、上記3つの主体のほか、学習データを提供する者、データセンター等の AI に必要な資源を提供する者、研究を行う者等、様々な関係者が存在することにも留意する必要がある。

² 「AI 事業者ガイドライン（第1.01版）」では、「文章、画像、プログラム等を生成できる AI モデルにもとづく AI の総称」と定義している。

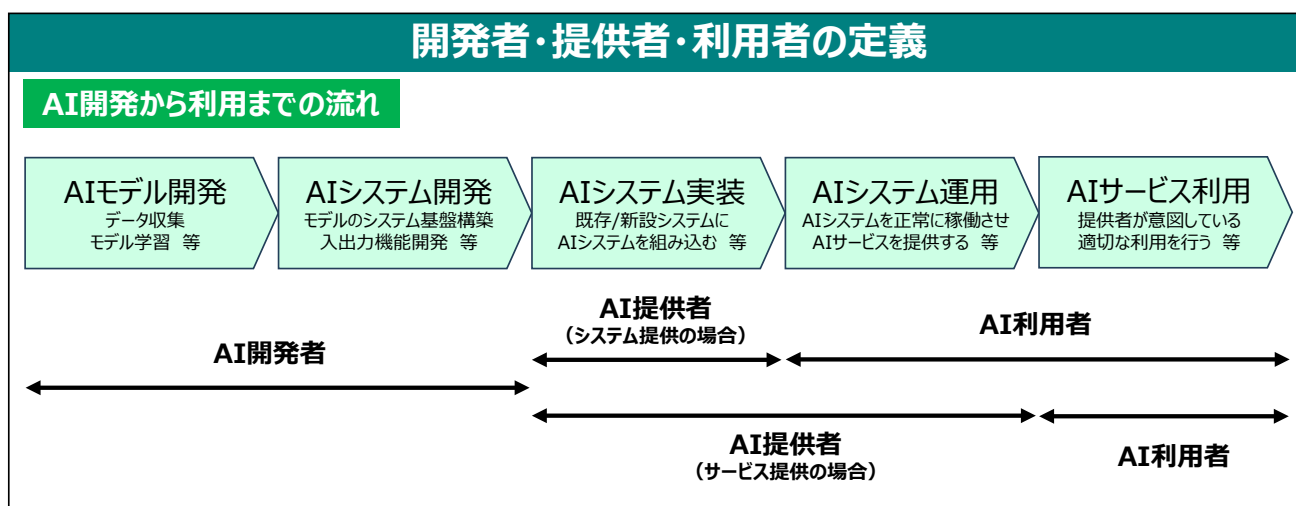


図 2 各主体の位置づけ

(2) 国外事業者

国内で利用される生成 AI の多くは国外事業者が提供しており、AI に係る制度の検討にあたって、国外事業者をその対象から一律に除くことは適当ではない。インターネットを通じて、我が国に在住する人々が日常的に国外事業者の提供する AI サービスに容易にアクセスし、利用できる状況にあること、仮に、事業者に対して何らかの義務付けをするような制度について国外事業者のみを対象外とする場合、国内事業者が一方的に不利益を被る可能性もあること等を踏まえると、国外事業者についても、国内事業者と同じく制度の対象とすべきである。この際、地理的な要因等からコンプライアンスの協力を得られにくい国外事業者に対しても制度の実効性を確保するため、国外事業者も明確に対象とするルール化を検討すべきである。実務においては、国外事業者の日本支社、日本における代表者等が存在する場合には、当該者を通じて対応を求めること等も考えられる。

3.イノベーション促進とリスクへの対応の両立

AI については、その開発・利用方法等によっては、様々なリスクを生じさせ得る一方で、国民生活の向上、国民経済の発展に大きく寄与する可能性がある。また、世界で最も AI を開発・活用しやすい国を実現するため、AI のイノベーションの促進とリスクへの対応を両立させることが重要である。

(1) イノベーションの促進

① 研究開発への支援

AI の研究開発には大量の学習データやこれを取り扱うことのできる大規模な情報処理、情報通信、データ保管等の施設や設備が必要となる。スタートアップ企業を含め、多様な主体がこのような施設や設備を活用することによりイノベーション促進が図られるようにする観点から、政府による整備を進めることが重要である。また、研究開発を行う人材については、国際

的に獲得競争が激化していることもあり、我が国における AI 人材の不足は深刻化しており、人材育成も積極的に推進していくことが重要である。

政府等では、既に質の高い日本語データ等を整備・拡充し我が国企業等に適切な形で提供する取組を行うとともに、データセンターの整備や、電力の確保に向けた検討のほか、AI 半導体の支援を進めている。また、次世代 AI 人材育成プログラム³といった人材育成事業も実施している。引き続き、様々な側面から研究開発支援を実施すべきである。

加えて、例えば大規模言語モデルの開発に関する基礎的な研究開発においてブレイクスルーがあったことによって、生成 AI が開発され、急速な活用の拡大につながったことに見られるように、AI は、基礎研究と活用の拡大が密接な関連性を有しており、基礎研究の振興への配慮が必要である。

現在、国立研究開発法人理化学研究所では AI 基盤モデルを科学研究に活用する AI for Science の研究に取り組んでおり、科学研究の手法や研究そのものに大きな変革をもたらす可能性がある。また、2024 年 4 月には米国アルゴンヌ国立研究所と AI for Science で連携することが日米首脳共同声明に盛り込まれ、国際的な連携も行われている。このような他分野の発展に寄与する研究開発の活動を継続して取り組むことが重要である。

②事業者による利用

我が国の競争力を向上させるためには、①の研究開発の支援だけではなく、AI 技術が広く社会実装され、国内事業者にて利用されることが重要である。例えば、政府の研究機関や大学等の研究成果を社会還元・技術移転する取組を推進することが、事業者による新たなビジネス参入や各業界の活性化となり得ると考えられる。

また、我が国においては、事業者自らが AI システムを開発し、AI サービスを消費者等に提供するケースがよく見られるが、今後、他社が開発した AI システムを利用し、AI サービスを消費者等に提供する事業者⁴が少なからず増えると考えられるところ、そのような事業者が円滑に国内外に AI サービスを提供できる環境を政府が整備することも重要である。具体的には、市場を広げるために安全性の高い AI システムを普及させるべく、後述する、国際整合性の確保や安全性評価や認証の実施が有効である。例えば、認証を取得した者に対するインセンティブを付与し、認証取得事業者の数を増やすことで、多くの国民がより安心して AI を利用することができ、市場が拡大し、イノベーションが促進されることが考えられる。また、国内事業者に

³ 「国家戦略分野の若手研究者及び博士後期課程学生の育成事業 次世代 AI 人材育成プログラム」は、国立研究開発法人科学技術振興機構において、次世代 AI 分野（AI 分野及び AI 分野における新興・融合領域）に資する研究開発に取り組もうとする若手研究者及び博士後期課程学生に対して支援を行う。

⁴ 株式会社ベネッセコーポレーションにおいては、他社が開発した AI システムを利用し、顧客に対し、「自由研究お助け AI」や蓄積した知見やデータを活用したサービスを展開している（AI 制度研究会（第 2 回）資料 2 参照。）。

よる AI 市場への新規参入を促すため、事業者が AI の基礎知識等を学べる環境を整備することも重要である。

また、ロボット、医療、防災等の分野における AI の活用や、アジア諸国との連携など国際連携・国際貢献も重視していくべきである。

(2) 法令の適用とソフトローの活用

AI のもたらし得るリスクの例と関係する主要法令は、図 3 のとおりである。国内においても、性的な動画コンテンツの出演者の顔を芸能人の顔にすり替えインターネット上に公開した事件（刑法・名誉毀損罪と著作権法違反）、生成 AI を用いてコンピュータウイルスを作成した事件（刑法・不正指令電磁的記録に関する罪）、生成 AI を用いて女性のキャラクターを作って現金をだましとった事件（刑法・詐欺罪）などが実際に発生している。リスクに対しては既存の法令で一定の対応がなされていることを前提に、更なる制度の検討を行う必要がある。

リスクへの対応にあたっては、法令による対応とガイドライン等のソフトローによる対応がある。

我が国では、現時点においては、AI に起因するリスクや問題の対処にあたって、各分野の所管府省庁が法令やソフトローにより対応しているところである。例えば、2023 年 6 月には、個人情報保護委員会が個人情報取扱事業者や行政機関等⁵に対し、生成 AI サービスの利用に際しての個人情報の取扱いに関する注意点を示しつつ、個人情報の保護に関する法律（平成 15 年法律第 57 号）に従って、個人情報を適正に取り扱うよう注意喚起を行った。また、2024 年 3 月、文化庁の文化審議会著作権分科会法制度小委員会は、生成 AI が膨大なデータを学習し、コンテンツを生成する際に著作権の侵害が生じるのではないかという懸念に対し、AI と現行の著作権法の関係についての解釈に当たっての一定の考え方を示したほか、2024 年 5 月には、内閣府の AI 時代の知的財産権検討会が、AI と知的財産権等との関係について、法的ルールの考え方の整理とともに、AI 技術の進歩と知的財産権の適切な保護が両立するエコシステムの実現に向けて、法・技術・契約の手段の組合せにより各主体が対応する必要性について、考え方を示した。その他、冒頭で述べたとおり、2024 年 4 月には総務省及び経済産業省が「AI 事業者ガイドライン（第 1.0 版）」を公表⁶し、法の支配、人権、民主主義、多様性及び公平公正な社会を尊重するよう AI システム・サービスを開発・提供・利用すべきである旨等を示している。

法令に基づく罰則がある場合には、公的機関が何かしらの強制力を発動することが可能であり、規律の実効性の確保が得られやすいという利点がある一方で、規制を行った分野の発展を阻害する可能性があるほか、国民の権利利益に影響を及ぼす規制が明確である必要があることに鑑み、

⁵ 行政機関、地方公共団体の機関（議会を除く。）、独立行政法人等（個人情報保護法別表第 2 に掲げる法人を除く。）及び地方独立行政法人（地方独立行政法人法第 21 条第 1 号に掲げる業務を主たる目的とするもの又は同条第 2 号若しくは第 3 号に掲げる業務を目的とするものを除く。）をいう。

⁶ 2024 年 11 月には時点更新等を施した第 1.01 版を公表している。

AIのもたらし得るリスクの例	具体事例・想定ケース	主要法令等
AIへの秘密情報の入力	外部のAIサービスに企業の秘密情報を入力し情報が漏洩	不正競争防止法、民法（契約） ※ 2024年2月「秘密情報の保護ハンドブック」においてAI利用時の留意点を整理（経済産業省）
AIの開発・学習及び生成・利用の過程での他者の著作権の侵害	特定の漫画・アニメのキャラクター等のイラストに類似した画像を生成する目的での学習や、そうしたイラストに類似する画像の生成・利用	著作権法 ※ 2024年3月「AIと著作権に関する考え方について」を公表し、解釈を明確化（文化庁）
AIの開発・利用の過程での他者の産業財産権の侵害	他者の登録商標を学習して、登録商標と同一又は類似の商標を作成し、その指定商品・役務と同一又は類似の商品・役務について使用	意匠法、商標法 ※ 2024年5月「AI時代の知的財産権検討会中間とりまとめ」を公表し、法的ルールの考え方を整理（AI時代の知的財産権検討会）
AIの開発・利用の過程でのプライバシー侵害・個人情報保護違反	本人の同意なしに個人情報を含むデータをAI学習に利用	憲法（プライバシー権、パブリシティ権）、個人情報保護法 ※ 2024年6月「個人情報保護法 いわゆる3年ごとの見直しに係る検討の中間整理」を公表し、AI利用時の論点を整理（個人情報保護委員会）
AI搭載製品の誤作動	自動運転車が誤作動により生命・身体の安全に影響	道路運送車両法、薬機法、労働安全衛生法、民法（不法行為等）、製造物責任法、自動車損害賠償保障法、国家賠償法
ディープフェイク（AIで合成した肖像・声等の悪用）	本人の同意なしに個人の画像をポルノその他の性的な画像に合成し拡散する行為や、AIにより有名人・知人になりすました音声通話による詐欺	民法（人格権・不法行為）、刑法（脅迫罪、名誉毀損罪、わいせつ物頒布等罪、詐欺罪、偽計業務妨害罪等）、児童ポルノ禁止法、情報流通プラットフォーム対処法（権利侵害情報）
バイアス（差別・偏見）の助長	不適切なAIによる採用や退職に関する判断、人事評価の実施	ヘイトスピーチ解消法、労働関係法令、民法、個人情報保護法、障害者差別解消法、部落差別解消法
偽・誤情報による情報操作	立候補者に関する偽情報をAIで作成し、SNS等で拡散し選挙を妨害	民法（人格権・不法行為）、刑法（名誉毀損罪）、行政法規、公職選挙法、情報流通プラットフォーム対処法（権利侵害情報）
国民の権利利益の侵害	AIの誤った判断で個人が行政サービスを受けられない等不利益を被る可能性	憲法（適正手続）、行政手続法
ウィルスの作成等のサイバー攻撃	生成AIを悪用しコンピュータウイルスを作成	刑法（不正指令電磁的記録に関する罪）、不正アクセス禁止法
ハルシネーション（AIが虚偽の情報を作成）	生成AIが虚偽の情報を作成し利用者を誤解させる	民法（不法行為、契約）
環境負荷の増大	AI開発過程での電力需要等の増大に伴うCO2排出量増大	地球温暖化対策の推進に関する法律
人間とAIの負の相互作用	AIとの対話にのめり込んだ人が人生に悲観して自殺	自殺対策基本法
AGIが制御不能になる懸念	人間がAGIを制御不能になり社会混乱を引き起こす可能性	なし

※ AI事業者ガイドライン（第1.01版）において、上記リスクの複数について記載があり、10個の共通の指針（人間中心、安全性、公平性、プライバシー確保、セキュリティ確保、透明性、アカウントビリティ、教育・リテラシー、公正競争の確保、イノベーションの促進）の下にその対応の在り方について示している。

図3 AIのもたらし得るリスクの例に関する整理

その範囲を検討するには一定の時間を要するため、柔軟性に欠けるといった欠点がある。その他、罰則を伴わない法令であっても、法令に事業者の義務や責務が明記されること自体によって国内外の事業者に対し規律を働かせ、一定の実効性を確保することが可能である。

他方、ガイドライン等のソフトローには、国際情勢や最新技術の動向に合わせた迅速かつ柔軟な対応が可能であり、イノベーションに与える負の影響が少ないという利点がある。一方で、ソフトローでは事業者等の自主的な対応に頼らざるを得ない可能性がある。

一般的に、我が国の企業等は法令遵守の意識が高いとされており、新たな規制が制定された場合、当該規制の遵守を意識するあまり、新たな研究開発やサービスの開発・展開を必要以上に躊躇する可能性がある。技術の発展やサービスの変化が急速な AI の分野において、過度な規制により研究開発やサービスの開発・展開を抑制させてしまうことは、将来にわたって我が国の国際競争力を損なう危険性をはらむものである。そのため、新たな制度を検討・導入するにあたってはイノベーションに与える影響を十分に留意する必要がある。

上記の観点から、イノベーション促進とリスクへの対応の両立を確保するため、法令とガイドライン等のソフトローを適切に組み合わせ、基本的には、事業者の自主性を尊重し、法令による規制は事業者の自主的な努力による対応が期待できないものに限定して対応していくべきである。

また、既存の個別の法令の存在する領域においては、AI が各領域で様々な用途で利用され始めており、権利利益の保護の必要性が生じる場面も AI の用途に応じて異なることから、まずは当該法令の枠組みを活用しつつ対応すべきである。その上で、そのような領域以外に関して新たな制度を創設する場合も含め、仮に法律上の規制による対応を行う場合には、事業者の活動にもたらす影響の大きさを考慮しつつ、(3) で後述する AI のもたらすリスクを踏まえた上で、真に守る必要のある権利利益を保護するために必要な適用内容とすべきである。その際、政府と事業者との役割分担を意識した上で、何が規制の対象となり、事業者の活動はどこまで許容されているのかといった線引きを明確化することが重要である。また、「規制はその目的を達成するために、特定の種類の技術の使用を強制したり、優遇したりすべきではない」という規制の技術中立性の原則も踏まえた検討も重要である。なお、AI の安全性に関する正当な研究を行うために不適切な AI を試作するケース等における規制の適用については、その要否も含め検討を行う必要がある。また、広く事業者一般を対象とする制度を検討する際には、スタートアップ企業も含め、どのような規模の事業者であっても対応可能なものとなるよう、制度への対応に伴う事業者の負担を考慮する必要がある。

(3) リスクへの対応

リスクへの対応にあたっては、AI 関係者が守るべき共通的な内容を明確にするとともに、AI を特定の領域において利用する際の個別の基準を設け、対応することが有効である。

現時点においては、国際的な枠組みとして、OECD「人工知能 (AI) に関する理事会勧告」(OECD AI 原則)、広島プロセス国際指針等が存在しており、これらを踏まえた国内向けの規範として、様々な事業活動において AI の開発・提供・利用を担う全ての者を対象とした「AI 事業者ガイドライン」が公表されている等、広く一般に守るべき事項は形成されている。イノベーション促進

との両立を確保しつつ、より適切にリスクに対応するためには、「AI 事業者ガイドライン」を技術の進展等に合わせて内容を更新していくとともに、各主体が適切に遵守するように、普及啓発等を進めつつ、必要に応じて対応することが重要である。適切なリスクへの対応のためには、ガイドラインに沿って開発者、提供者等の各主体の役割を明らかにしたうえで、責任を明確化する必要がある。同時に、開発者と提供者、提供者と利用者といった各関係者の間で、必要な情報共有を行い、密に連携していく必要がある。

AI を特定の領域に利用する場合は、目的、利用方法等を考慮し個別に対応することが重要である。例えば、国民生活や経済活動の基盤となるインフラやサービス等（以下「基盤サービス等」という。）や製品安全に関する AI については、各業所管府省庁により既存法等を中心とした対応がなされる。

このほかにも、AI の急速な発展に伴い今後新たに顕在化するリスクについても各分野の内容に応じて適切に対応する必要がある。特に、人の生命、身体、財産といった人間の基本的な権利利益や社会の安全、我が国の安全保障に対して実際に重大な問題を生じさせる、あるいは生じさせる可能性の高い AI に対しては、そのリスクの内容や当該リスクの社会的な影響の重大性に応じて規律の必要性の有無を検討すべきである。

そのためには、政府と事業者とで連携をしつつ実際に起きている事例等を踏まえ、AI のモデルや用途が様々存在する中で、開発、提供、利用といった AI のライフサイクルの各場面において顕在化する可能性のあるリスクとは、いかなる種類の AI モデルのどのような性質に起因するリスクであって、誰にどのような影響を与えるものかといった要素を分析する必要がある。その前提として、まずは AI の開発、利用等に関する実態を調査・分析し、社会全体で認識を共有した上で必要な対応を適時適切に行うことが重要である。例えば、不適切な AI による求職者の選別や AI による消費者の混乱に対する懸念があり、政府は実態の把握に努めるとともに技術的対応も含め必要な対策を検討することも重要である。また、政府による実態の把握のため、各主体に協力の要請を行う際は、広島 AI プロセス等でも確認された、法の支配、適正手続き、民主的責任行政、及び基本的人権、特にプライバシー権などの基本原理を遵守すべきであり、政府の恣意的権限行使を抑止し、事業者等の予見可能性低下や萎縮効果を生じさせないよう対応する必要がある。

なお、諸外国においては学習の計算量といった AI の規模や利用者数により規制等を設けているが、規模に依存しない高性能な AI が開発されていること等を踏まえ、どのような要素を考慮すべきか検討が必要である。

4. 国際協調の推進

（１）AI ガバナンスの形成

AI ガバナンスについては多国間の枠組において活発な議論がなされている。

「I.はじめに」で述べたとおり、G7 においては、2023 年 5 月、G7 広島サミットを受け、生成 AI に関する国際的なルール検討のため「広島 AI プロセス」を立ち上げ、同年 12 月には、広島 AI プロセス国際指針等を含む「広島 AI プロセス包括的政策枠組み」が G7 首脳により承認され

た。その後、我が国は、安全、安心で信頼できる AI の実現に向け、様々な国際会議等の場において、広島 AI プロセスで掲げた精神を発信している。また、G7 を超えたアウトリーチとして、2024 年 5 月には、当該精神に賛同する 49 の国・地域の賛同を得て、「広島 AI プロセス・フレンズグループ」を立ち上げ、支持の拡大を図っている。国際的な AI ガバナンスの形成は、今後の AI の発展の方向性を定めるものであり、国益にも資するため、様々な国際会議等において、引き続き、広島 AI プロセスの考え方にに基づき、議論をリードしていくべきである。また、我が国が各国のモデルとなるような AI 制度を構築し、世界に向けて発信していくべきである。

その他、国連においては、2024 年 9 月に未来サミットの成果文書の附属文書として「グローバル・デジタル・コンパクト」が採択され、AI のリスクや機会の評価を通じた科学的理解を促進するための AI に関する国際科学パネルの設置や、各国政府及び関連するステークホルダーが参加する AI ガバナンスに関するグローバル対話の開始等が盛り込まれたほか、欧州評議会において「人工知能と人権、民主主義及び法の支配に関する欧州評議会枠組条約」⁷が成立し、現在、EU に加え、米国、英国等の 10 か国が署名している状況である。OECD においては、2019 年に公表した、包摂的な成長、持続可能な開発及び幸福、人間中心の価値観及び公平性等からなる AI 原則を、急速な発展を遂げる生成 AI による偽・誤情報のリスク等に対応させるため、2024 年 5 月に改定したところである。人間中心の考え方に立ち、「責任ある AI」の開発・利用を実現するため設立された国際的な官民連携組織である GPAI (Global Partnership on Artificial Intelligence) においては、東京で GPAI サミット 2022 を開催し、人間中心の価値に基づく AI の活用促進、AI の違法かつ無責任な使用への反対、持続可能で強靱かつ平和な社会への貢献等について各国で合意したほか、結果を閣僚宣言として初めて公表した。2024 年 7 月には OECD との統合パートナーシップを経て GPAI 参加国は 44 に達するとともに、アジア地域初の GPAI 専門家支援センターが東京に設置され、広島 AI プロセスが推進する生成 AI の安全性評価に関するプロジェクト等を支援している。このほかにも、国際的に様々な議論がなされているところ、AI に係る制度・施策の実施にあたっては、これら国際的な枠組等において合意あるいは認められた取決や考え方を踏まえ対応すべきである。

(2) 国際整合性・相互運用性の確保

先に述べたとおり、我が国の国民はインターネットを通じて様々な国の事業者等の AI サービスを利用することができ、また、我が国の AI サービスは世界各国の人々が利用することができる。このような状況において、満たすべき安全性等に係る国際的な規範と我が国において適用される規範の相互運用性が確保されている場合、我が国の事業者が円滑に海外市場に進出できるほか、我が国の国民が全世界の AI サービスにアクセスすることが可能となるため、国際整合性・相互運用性の確保は重要である。

⁷ 正式名称は「The Council of Europe Framework Convention on Artificial Intelligence and Human Rights, Democracy, and the Rule of Law」であり、「人工知能と人権、民主主義及び法の支配に関する欧州評議会枠組条約」は和訳名の仮称。

このため、広島 AI プロセスの重要性はもちろんのこと、ISO、IEC 等における国際規格を定める標準化活動は、将来の世界的な市場や事業活動の発展にとって重要であり、幅広い関係者による議論を行い長期的な視野を持ち、積極的に進めていく必要がある。

AI の安全性に関する評価手法や基準の検討・推進を行うための機関である AISI（AI セーフティ・インスティテュート）は、国内外の AI 安全性の知見のハブとして、国内外の関係機関とのネットワークの構築のほか、安全性評価のためのガイダンスの作成等を進めている。2023 年 11 月に英国にて開催された AI 安全性サミットを契機に、AI の安全性に係る技術的な側面からのガバナンスに関する議論が進み、2024 年 2 月、英国、米国に続く形で、我が国においても AISI が設置されたところ、上記の国際整合性・相互運用性の確保のため、AISI による取組を進めていくことが重要である。

III. 具体的な制度・施策の方向性

1. 全般的な事項

生成 AI のような AI は汎用性が高く、様々な分野で利用されており、リスクへの対応も様々である。個人情報や著作権の取扱い、偽・誤情報への対処といったリスクへの対応にあたっては、既存法等を中心とする対応が前提であるが、AI については横断的な対応が必要なケースもあるため、全体を俯瞰する政府の司令塔機能の強化、戦略の策定、また、安全性の向上のため、透明性や適正性の確保等が求められており、必要に応じて制度整備することが適当である。

（1）政府の司令塔機能の強化、戦略の策定

AI は、利用分野や用途の広がり、汎用型 AI の登場等により、研究開発から活用に至るまでの期間が短い場合も存在し、その間の各段階における取組がほぼ同時並行的に行われ得るものである。このため、研究開発から活用に至るまでに介在する多様な主体や過程における取組が互いに密接に関連し、一体的・横断的に行われる必要があり、研究開発から経済社会における活用までの一体的な施策を推進する政府の司令塔機能を強化すべきである。

AI は、政府・自治体での活用を含め、国民生活の向上のための様々な場面での利用だけでなく、犯罪への悪用の懸念もあるほか、デュアルユース技術の側面も持つため、司令塔機能の強化に際しては、広く関係府省庁が参加する政策推進体制を整備する必要がある。

また、総合的な施策の推進にあたっては、司令塔が戦略あるいは基本計画（以下「戦略」という。）を策定する必要がある。AI については、安全・安心の確保が AI の活用の促進、イノベーションの促進、安全保障リスクへの対応、犯罪防止等にとって重要であることから、AI の安全・安心な研究開発、活用の促進等に資する戦略とすべきである。国際的な協調を図りつつ、イノベーションの促進とリスクへの対応の両立を図るために政府全体で取り組むことが必要となる施策を当該戦略に盛り込むべきである。

上記については、AI の司令塔機能の強化や、司令塔による関係行政機関に対し協力を求めることができる等の権限を明確化するため、法定化すべきである。

(2) 安全性の向上等

AI の安全性を向上させるためには、研究開発から活用までのライフサイクルにおいて、少なくとも透明性や適正性を確保していく必要がある。また、事業者が自主的に取り組む安全性評価や第三者による認証などを活用することも一つの有効な手段となると考えられる。さらに、政府が、進化の著しい AI の技術や利用動向等の実態を調査して情報提供を行うとともに、必要に応じて、関係各主体に対応を求めていくべきである。

これらの実施にあたっては、事業者を含む関係各主体からの情報共有や協力が必須であることから、官民が協調して取り組むことが必要である。

① AI ライフサイクル全体を通じた透明性と適正性の確保

AI の研究開発から活用に至るまでには、例えば、モデル構築のため膨大な量の学習データを使用し、その後、チューニング等を経て、開発者から提供者に AI システムが渡り、その後、提供者がさらに追加学習をしたうえで、利用者に AI サービスを提供する場合、開発者が開発した AI システムと提供者が利用者に AI サービスを提供する際の AI システムとではその能力は変化し、リスク評価の結果も変化している可能性がある。利用者の側では必ずしも開発時点でのリスクへの対応について把握することができない中で、リスクに対応するにあたっては、このようなリスクに係る必要な情報を関係者に適切に共有しなければ、誤った認識で AI システムを提供者が利用者に提供し、あるいは利用者が不適正に AI サービスを利用し、リスクが顕在化する可能性がある。

このため、AI の安全・安心な研究開発や活用には、開発者－提供者間、提供者－利用者間において必要な情報を共有する透明性を確保すべきである。他方で、かかる透明性の確保のための措置が事業者の事業運営に過度な負担や広範すぎる情報開示とならないようにするため、また、研究開発段階においては未だ実際の利用に供されていないこと、及び研究開発に関する情報は企業の機密にかかわることが多いことから、情報の共有は真に必要な範囲に留めることが重要である。

適正性に関しては、広島 AI プロセスで合意された「全ての AI 関係者向けの広島プロセス国際指針」において、高度な AI システムの開発時、市場投入前後におけるリスクやインシデントの特定と対応、信頼でき責任ある利用の促進等を AI 関係者に求めたほか、各国の AISI、ISO 等においても様々な議論がなされており、適正な研究開発や活用を進めていく必要がある。

適正性の確保にあたっては、広島 AI プロセス等の国際的な規範の趣旨を踏まえた指針を政府が整備などを行い、事業者に対し各種規範等に対する自主的な対応を促していくことが適当である。当該指針については、例えば、不適切な出力を減らすための工夫を AI 開発者等が講じること、AI の概要やリスク等に関する情報を AI 開発者等が開示・共有すること、訓練用データセットの適切な透明性も支援されるべきであること等について検討することが重要である。

また、透明性と適正性の確保については、調査等により政府が事業者の状況等を把握し、その結果を踏まえて既存の法令等に基づく対応を含む必要なサポートを講じるべきである。政府による事業者の状況等の把握や必要なサポートについては、事業者の協力なしでは成り立たな

いため、国内外の事業者による情報提供等の協力を求められるように、法制度による対応が適当である。

偽・誤情報対策等として、他国において、AI を利用したサービスであることを顧客に通知することや、AI で作成したコンテンツに AI 生成であると表示することを、一部義務付けている例があり注視する必要がある。また、透明性や適正性確保のため、電子透かし⁸や来歴証明⁹等による技術的な対応も重要である。関連して、2024 年 9 月には、総務省のデジタル空間における情報流通の健全性確保の在り方に関する検討会において、偽・誤情報等について、AI が生成した情報であるか否かを判断する技術や、情報コンテンツ、その発信者の信頼性等を確保する技術等の研究開発や社会実装を進めていくことが重要であるとし、政府に対する具体的な方策を提言している。

②国内外の組織が実践する安全性評価と認証に関する戦略的な促進

AI の安全・安心な活用の促進にあたっては、安全性評価や認証制度の実施も 1 つの有効な手段である。安全性の評価や認証制度は、AI システムに関する評価・認証と、AI を利用する組織のガバナンス等に関する評価・認証とに大別される。

AI システムの安全性の評価について、リスクを理解する AI の開発者、提供者や利用者は、組織内外の専門家チームや評価用のツールなどを使って、基本的には自らリスク評価を行い、対処していくべきである。将来的に有用な第三者認証が確立されるならば、AI 開発者や提供者がかかる第三者認証を取得することにより、一般国民を含め多くの利用者やこれまで AI サービスを扱っていなかった提供者が AI の安全性を評価することも考えられる。提供者や利用者は、第三者認証の有無によって、安全性の高い AI 事業者やその AI システムを認識し、選択することが可能となる。

一方、AI を利用する組織のガバナンスに関しては、利用者が自ら体制を構築して評価する場合と、第三者認証制度を活用する場合が考えられる。第三者認証制度は、AI の安全・安心な活用を促進し、我が国の AI 産業の活性化に寄与する方策の一つとなると考えられる。

なお、AI システム、AI を利用する組織のガバナンスの認証については、ISO 等で検討がなされている状況である。

安全性については電気用品安全法（昭和 36 年法律第 234 号）、医薬品、医療機器等の品質、有効性及び安全性の確保等に関する法律（昭和 35 年法律第 145 号）といった個別法により規律が既に導入されているほか、「AI 事業者ガイドライン」による対応もなされている。また、各国の AISI のコミュニティ、ISO、IEC 等においても議論が行われている状況である。

国際的な基準や規格を作る活動は、将来の世界的な市場において AI システムの相互運用性を確保するために重要であり、積極的かつ戦略的に対応する必要がある。

⁸ 電子透かしは AI が生成したことを示す識別情報をコンテンツに埋め込む技術で、例えば「SynthID」等の技術が存在する。

⁹ 来歴証明は作成者等の来歴情報を検証可能な形でコンテンツに付与する技術で、例えば「Originator Profile」「C2PA」等の技術が存在する。

また、国内における制度整備は、国際的な規範を踏まえ、かつ、制度の実効性も考慮し対応すべきである。AI の評価や認証を実施する場合には、利用者や利用目的に従ってレベルを設けることや、一定の安全性を確認するための利用者の負担が軽減する仕組みや評価・認証を実施する機関を認定する仕組みを構築できれば、より効果的で持続可能な制度となると考えられる。ただし、この仕組みを構築する際は、AISI や ISO 等の活動を前提にしつつ、どのような主体を巻き込み、どのような基準で評価を行っていくのか、詳細な検討が必要である。なお、AI 安全性の確保のため、AISI には、関係省庁、関係機関と連携し、調査、分析、整理、情報発信などに引き続き取り組み、司令塔となる組織を支援することが期待される。

③重大インシデント等に関する政府による調査と情報発信

上述のとおり、AI は近年急速な発展を遂げており、様々なリスクが増大している。このような中、AI のリスクに対処し、政府として適切な施策を実施するためには、技術及び事業活動の双方の側面から時々刻々と変化する AI の開発、提供、利用等に関する実態をまず政府において情報収集・把握し、事業者において AI が効果的かつ適正に利用されるとともに、広く国民が AI の研究開発や活用の促進に対する理解と関心を深められるよう、企業秘密等に配慮しつつも説明責任を果たせるように、必要な範囲で国民に情報提供することが適当である。中でも、多くの国民が日々利用するような AI モデルについては、政府がサプライチェーン・リスク対策を含む AI の安全性や透明性等に関する情報収集を行い、国民に対して広く情報提供されることで、利用者は安全性の高い AI 事業者やその AI システムを認識・選択しやすくなる。また、基盤サービス等における AI 導入の実態等に関しては、政府による情報収集が重要である。

また、AI の利用に起因する重大な事故が実際に生じてしまった場合、政府としては、その拡大又は再発の防止を図るとともに、AI を開発・提供する事業者による再発防止策等について注意喚起を行っていく必要がある。すなわち、国内で利用される AI について、国民の権利利益を侵害するなどの重大な問題が生じた場合、あるいは生じる可能性が高いことが検知された場合において、その原因等に関する事実究明を行い、必要に応じて関係者に対する指導・助言を行い、得られた情報の国民に対する周知を図るべきである。なお、事故が生じているといえるか否かについては、上記の情報収集・把握を通じて政府に蓄積された事例や知見をもとに判断していくことが重要である。

この調査や情報発信は事業者の協力なしでは成り立たないため、国内外の事業者による情報提供等の協力を求められるように、法制度による対応が適当である。

2.政府による利用等

我が国における、個人及び企業による AI の利用率は、他国と比較すると著しく低迷している状況である¹⁰。AI は国民生活や経済活動の発展の基盤として、その利用の重要性が増していくことが見込まれる中、このような状況を放置すれば、我が国の国際競争力が損なわれるおそれがある。このため、まずは政府が率先して AI を利用し、国民による活用を促進することが考えられる。

(1) 政府調達

政府が AI を利用する際の基本的な考え方を示すことは、政府調達への参入を検討する AI 開発者、AI 提供者に安全性の向上等のための自主的な取組を促す観点からも有用である。

政府が情報システム等を調達する際は、「IT 調達に係る国等の物品等又は役務の調達方針及び調達手続に関する申合せ」（2018 年 12 月 10 日 関係省庁申合せ）に基づき、サプライチェーン・リスクに対応することが必要であると判断されるものについては、内閣サイバーセキュリティセンター及びデジタル庁と協議のうえ、必要な措置を講じることとしており、AI についても、同申合せの対象となる。

一方で、現時点において、AI 調達に特化したガイドライン等はないところ、AI には前述のとおり様々なリスクが存在することから、これに対応するため、政府が AI を調達する際に参考となる AI に特化した政府調達ガイドライン等の整備や既存の AI に関するガイドライン等の深掘りなどを行うことが重要である。

また、このような政府調達に関するガイドライン等を整備することで、AI を利用する事業者が何らかの AI システム・サービスを利用する際の参考にもなり、安全性の高い AI の普及促進に貢献できると考えられる。

なお、ガイドライン等の整備にあたっては、AI のリスクの軽減や品質の確保を図ると同時に、透明性の確保を求める場合は利用シーンに応じた企業の負担を考慮し、多くの事業者による容易かつ迅速な参加が可能となるように配慮すべきである。

(2) 政府等による利用

政府が AI を利用することにより、行政サービスや業務等の質・効率を向上させることができるほか、政府がユースケースやその有用性を示し、具体事例、留意点等を周知することにより、AI の活用を促進し、国内の AI 市場の発展等にも貢献できるため、政府が率先して AI を利用し

¹⁰ 総務省「令和 6 年版情報通信白書」によれば、我が国において生成 AI を利用している個人は 9.1%にとどまり（中国（56.3%）、米国（46.3%）、英国（39.8%）、ドイツ（34.6%））、また、企業向けのアンケートでは、生成 AI を業務で利用している割合は 46.8%となっている（米国（84.7%）、中国（84.4%）、ドイツ（72.7%））状況にある。

ていくことは重要である。ただし、国民の権利利益に重大な影響を及ぼしかねないものについては、AI の出力結果を自動的に採用することのリスク¹¹を踏まえ、慎重に取り組むべきである。

地方自治体についても、行政サービスの大きな部分を占めており、国民生活への影響も大きい。ため、AI の利用を推進し、行政サービスや業務等の質・効率を向上させていくことが重要である。また、地方自治体における AI の利用の推進にあたっては、各自治体の先進的な取組¹²を含む利用事例を参考にするほか、各自治体の地域課題に応じた AI の利用も重要である。

政府等による AI サービスの利用については「ChatGPT 等の生成 AI の業務利用に関する申合せ（第 2 版）」（2023 年 9 月 15 日 デジタル社会推進会議幹事会申合せ）にて、不特定多数の利用者に対して提供する、定型約款や規約等への同意のみで利用可能となるクラウドサービスでは、要機密情報の取扱いは不可とし、また、個別契約等に基づく生成 AI の利用を検討する場合には、政府の AI 戦略チームの承認を得たうえで、適切なリスク分析を行い、一部の機密性 2 情報まで取り扱うことが可能としている。政府が AI を利用するにあたっては、機密情報の取扱い等には注意が必要であり、上記のような申合せ等を遵守し、また、諸外国における取組¹³も参考にしつつ、必要に応じてアップデートしていくことが重要である。

その他、国民による活用を促進するため、リスクリングやリテラシー向上のための施策を実施していくことも重要である。

¹¹ 海外におけるリスク発現事案としては、米国における失業保険に関する事例（ミシガン州の失業保険庁が請求者の詐欺を検出するために使用した統合データ自動化システムは、2013 年から 2 年間で 93% のエラー率を記録し、2 万人の州民を詐欺で誤って告発した。）や、教師の失職につながった事例（ヒューストン独立学区は、教師が生徒の学力成長に及ぼす影響を推定する付加価値モデル（VAM）を導入し、2007 年から 2016 年までに教師の契約更新の拒否や解雇に利用したため、多くの教師が VAM の使用を止めるための司法救済を求めた。）が存在する。

¹² 神戸市においては、2024 年 3 月、一定のルール下での AI の効果的かつ安全な活用を目的として AI 条例を制定するほか、文章要約、アイデア出し、プログラミングコードの生成等 AI の活用を進めている（AI 制度研究会（第 2 回）資料 3 参照）。

また、東京都でも、ガイドラインの制定・改定のほか、2024 年 12 月には東京都 AI 戦略会議も開催するなど取組を進めている。

¹³ 諸外国においては政府機関向けの AI 利用ガイドラインや訓練プログラムを提供している。例えば、2023 年 9 月にカナダ政府は連邦政府職員向けの生成 AI 利用ガイドを公表した。また、2024 年 3 月に米国カリフォルニア州政府が、州政府機関による生成 AI 製品購入のためのガイドラインを公表し、州政府職員のための生成 AI 調達に関する訓練プログラムの提供も開始した。2024 年 5 月に成立した EU AI 法においても「AI オフィスが講じる措置」の一つとして、AI システムに関連する公共調達手続のベストプラクティスの評価・促進が記載されている。

3.生命・身体の安全、システミック・リスク、国の安全保障等に関わるもの

医療機器、自動運転車、基盤サービス等、特に国民生活や社会活動に与える影響が大きい生命・身体の安全やシステミック・リスク¹⁴に関わるものについては特に注力して対応する必要があると考えられるが、業界毎に各業所管省庁が既存の業法に基づき対応し、また、追加的な対応の必要性の有無を判断するため、AI技術の発展、利用状況について随時業界と対話している状況である。

現時点においては、引き続き、各業所管省庁が既存の法令あるいはガイドライン等の体系の下で対応すべきであるが、今後、新たなリスクが顕在化し、既存の枠組で対応できない場合には、政府は、関連する枠組の解釈を明確化したうえで、制度の見直しあるいは新たな制度の整備等を含めて検討すべきである。システミック・リスクについては、将来的に、複数のAIシステムが連動する大規模なAIシステム群が社会システムを支える状況となる可能性があり、その際、当該AIシステム群が予期せぬ挙動をした場合、社会全体に大きな混乱をもたらす可能性があるため、適切に対処することが重要である。

また、CBRN等の開発やサイバー攻撃等へのAIの利用といった国の安全保障に関わるリスクについては、我が国の安全保障を確保するという観点から、関係省庁において、必要な対応をさらに検討していく必要がある。

IV.おわりに

諸外国においては、AIに関する制度整備が進められているなか、以上を踏まえると、我が国においては、イノベーション促進とリスクへの対応の両立を図るため、広く一般的に使われるAIを対象とする指針を政府が整備などを行い、透明性・適正性の確保が事業者主導で進むよう促しつつも、生命・身体の安全や安全保障の確保は当然の前提として、AIの開発、利用等の実態を含めた様々なリスクへの対応状況を政府が調査・把握し、重大な問題が生じた場合、あるいは生じる可能性が高い場合には、既存の法令等に基づく対応や、必要なサポートを講じるべきである。特に、AIは世界中で急速な技術革新が起きており、政府は迅速な状況の把握と対応が可能となるよう体制を整える必要がある。政府による調達・利用や、基盤サービス等については、業界との対話を継続しつつ、まずは各業法、ガイドライン等で対応すべきである。このように官民協調によるリスクガバナンスを確立していくことが重要である。

上記の政府による指針の整備・対応やAIに関する実態の調査・把握にあたっては、事業者による自主的な対応も重要であるが、実効性を確保することが必要であるため、事業者の活動にもたらす影響等を考慮しつつ、法制度により実施すべきである。なお、これらの法制度による対応にあたっては、広島AIプロセス等でも確認された、法の支配、適正手続き、民主的責任行政、及び基本的人権、特にプライバシー権などの基本原理を遵守し、法制度がイノベーション促進の阻害とならないように留意する必要がある。

¹⁴ 特定のシステムの不全が関連するシステムにも波及し、広範囲に深刻な影響を及ぼすリスク。

政府に対しては、本とりまとめを踏まえ、世界で最も AI を開発・活用しやすく、かつ、人間中心の価値に基づく、他国のモデルとなるような AI に係る法制度を含む制度整備を速やかに実施していくことを期待する。

AI 戦略会議 構成員名簿

座 長	松尾 豊	東京大学大学院工学系研究科 教授
	江間 有沙	東京大学国際高等研究所東京カレッジ 准教授
	岡田 淳	森・濱田松本法律事務所 弁護士
	川原 圭博	東京大学大学院工学系研究科 教授
	北野 宏明	株式会社ソニーリサーチ 代表取締役 プレジデント
	佐渡島庸平	株式会社コルク 代表取締役社長
	田中 邦裕	さくらインターネット株式会社 代表取締役社長
	山口 真一	国際大学グローバル・コミュニケーション・センター 准教授

AI 制度研究会 構成員名簿

座長	松尾 豊	東京大学大学院工学系研究科 教授
座長代理	村上 明子	独立行政法人情報処理推進機構 A I セーフティ・インスティテュート 所長
	生貝 直人	一橋大学大学院法学研究科 教授
	岡田隆太郎	一般社団法人日本ディープラーニング協会 専務理事
	岡本浩一郎	一般社団法人ソフトウェア協会 副会長／株式会社リアルソリューションズ 代表取締役社長
	柿沼 由佳	公益社団法人全国消費生活相談員協会消費者教育研究所 副所長
	工藤 郁子	大阪大学社会技術共創研究センター 特任准教授
	殿村 桂司	長島・大野・常松法律事務所 弁護士
	中尾 悠里	富士通株式会社富士通研究所人工知能研究所 シニアリサーチマネージャー
	永沼 美保	一般社団法人日本経済団体連合会デジタルエコノミー推進委員会 国際戦略WG 主査／日本電気株式会社 品質・エンジニアリング推進部門 主席プロフェッショナル
	原山 優子	東北大学 名誉教授／GPAI 東京専門家支援センター長
	平野 晋	中央大学国際情報学部 教授・学部長
	福岡真之介	西村あさひ法律事務所・外国法共同事業 弁護士
	松原実穂子	日本電信電話株式会社 チーフ・サイバーセキュリティ・ストラテジスト

AI 制度研究会 開催実績

2024 年 8 月 2 日	AI 戦略会議（第 11 回）・AI 制度研究会（第 1 回）合同会議 ・ AI 政策の現状と制度課題について
2024 年 8 月 23 日	AI 制度研究会（第 2 回） ・ AI のリスクと制度的対応について（ヒアリング） （以下、ヒアリング対象者） ・ 東京大学 大学院法学政治学研究科 宍戸 常寿 教授 ・ 株式会社ベネッセホールディングス ・ 神戸市 ・ 人工知能研究開発ネットワーク（AI JAPAN） ・ 桃尾・松尾・難波法律事務所 松尾 剛行 弁護士
2024 年 9 月 10 日	AI 制度研究会（第 3 回） ・ AI のリスクと制度的対応について（ヒアリング） （以下、ヒアリング対象者） ・ 株式会社 ABEJA ・ 株式会社 Preferred Networks ・ ヤマト運輸 株式会社 ・ 一般社団法人 情報処理学会 ISO/IEC SC42WG1 国内委員会
2024 年 9 月 12 日	AI 制度研究会（第 4 回） ・ AI のリスクと制度的対応について（ヒアリング） （以下、ヒアリング対象者） ・ Google ・ 日本マイクロソフト株式会社 ・ Meta（Facebook Japan） ・ 中央大学 国際情報学部 須藤 修 教授 ・ 日本電信電話株式会社 ・ 株式会社三井住友銀行
2024 年 12 月 26 日	AI 制度研究会（第 5 回） ・ 中間とりまとめ（案）について
2024 年 12 月 26 日	AI 戦略会議（第 12 回）・AI 制度研究会（第 6 回） 合同会議 ・ 中間とりまとめ（案）について
2025 年 2 月 4 日	AI 戦略会議（第 13 回）・AI 制度研究会（第 7 回） 合同会議 ・ 中間とりまとめ（案）について