

イベントレポート

シンポジウム「日本の AI ガバナンス、世界での役割」

安全性かイノベーションか？独自路線か国際協調か？模索が続く中での日本の AI 政策を考える

2025 年 3 月 14 日、国際大学 GLOCOM はシンポジウム「日本の AI ガバナンス、世界での役割」を開催した。AI ガバナンスに関する制度設計のあるべき像を探る議論が世界的に活発化している。2 つの講演の後、日本の政策の現状、そして国際的な文脈でのガバナンスの最新動向と日本が果たすべき役割を議論する 2 部構成で進行した。

■AI 法案、リスク軽減とイノベーションの両立を目指す

内閣官房 内閣審議官 渡邊昇治氏

第一部の最初に登壇した渡邊昇治氏は、最近の日本の AI 政策動向と 2 月末に国会に提出した法案の要点を紹介した。日本の AI 政策が本格化したのは 2010 年代後半と思われ、人間中心の考え方を掲げて検討を進めてきた。状況を一変させたのが生成 AI である。ブームを受けて、2023 年 5 月の G7 広島サミットを機に立ち上がった「広島 AI プロセス」を含む多くの場面で、ガバナンスの議論が活発に行われた。特に生成 AI をリスクとみるか、イノベーションの機会とみるかで激論が続いているが、最近ではイノベーションあつてのガバナンスという方向感が強いのではないかと。

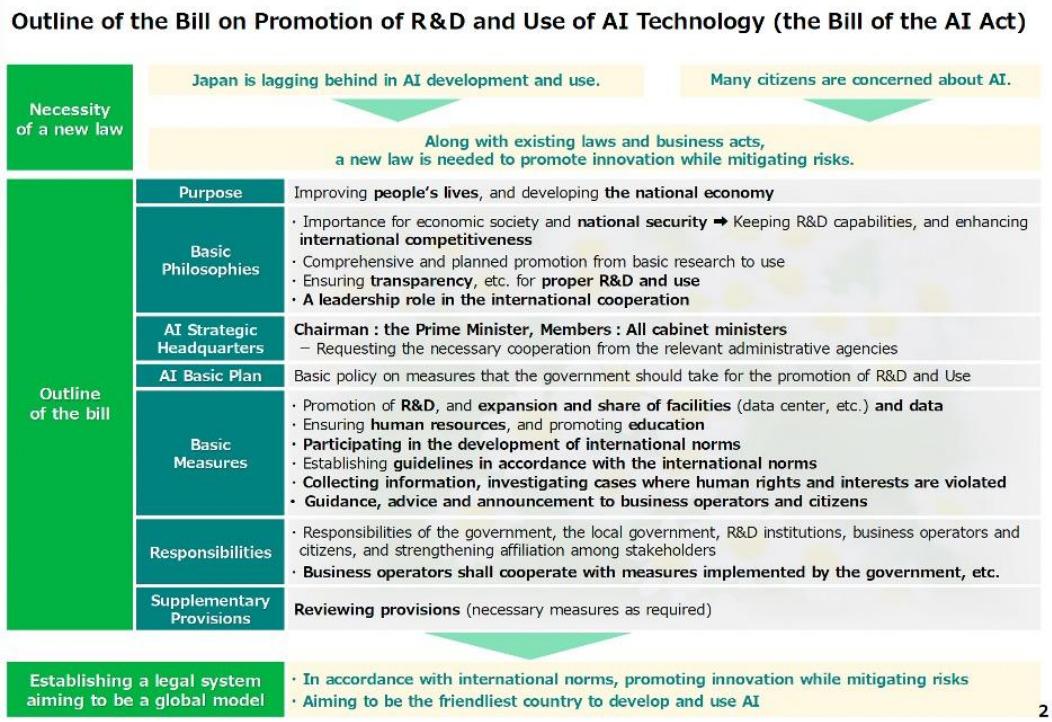
この現状を踏まえ、政府は AI 戦略会議や関係省庁の連携を中心に議論を深め、日本の AI 政策の柱を確立した。その方向性は、リスク対応とイノベーション促進の両方に目配りするものだ。まず、リスクへの対応では、特定領域の AI（重要なインフラや製品の安全性に関わるもの）利用と一般的な AI 利用を分けて考えることができる。前者のリスクについては、業種ごとの業法や既存の製品安全に関する法令が存在している。一方、後者のリスクについては、刑法に代表される既存の法制度による手当がされているほか、各種ガイドラインが作成され、AI Safety Institute も設立されている。次に、イノベーション促進については、AI の活用と研究開発の促進による成長を重視する。先進国の中では遅れている懸念がしばしば指摘されるが、ポテンシャルはあり、人材育成やインフラの整備を官民で取り組んでいく。

そして 2025 年 2 月末、AI 活用と研究開発をより強力に進めていくための法律案として、国会に提出したのが「人工知能関連技術の研究開発及び活用の推進に関する法律案（AI 法案）」である。この法案は「国民生活の向上及び国民経済の健全な発展」を目的に据えた。罰則のない法律とし、規制法ではなく推進法の性格を持たせた。

同法案の要点を 3 つ挙げる。第一に、この法律は PDCA サイクルを作る狙いがある。AI は技術の進化が速いため、常に実態を把握し、必要に応じて計画を見直すようにする意図がある。第二に、国際的な整合性を持たせるため、広島 AI プロセスの国際指針に則す。第三に、自主的な取り組みを尊重する。過剰な規制にならないよう、この法案には罰則はない。AI を監視、検閲に使うつもりはない政府の考え方を法律によって明らかにし、強くアピールする。この法案は AI という技術に対して

の国の考え方を示す初めての機会になる。

図 1：AI 法案の概要



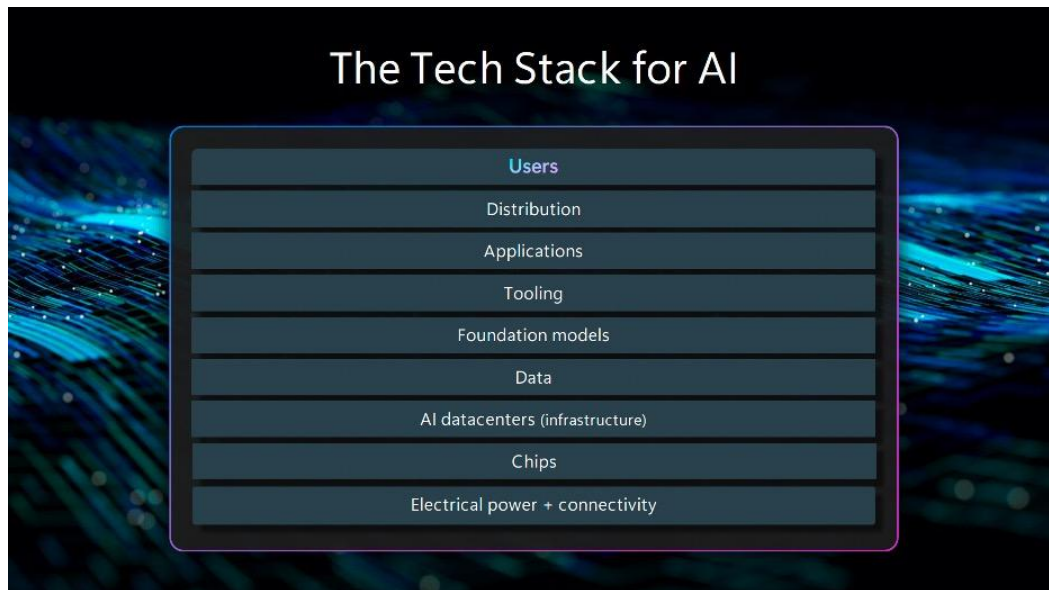
出典：内閣府

■ エージェント AI の登場で汎用技術としての真価を問われる AI

マイクロソフト 公共政策担当アジア地域シニアディレクター Marcus Bartley Johns 氏

続いて Marcus Bartley Johns 氏が登壇し、グローバルの視点から、AI エコノミーの成長で進化するガバナンスの構図について講演を行なった。過去 3 年間を振り返ると、GPT (Generative Pre-trained Transformer)、すなわち大規模言語モデルの能力開発に多くの注目が集まっていた。しかし、これからの AI エコノミーを考える上で重要なのはもう 1 つの GPT (General Purpose Technology)、汎用技術の方だ。汎用技術はスタックに分解でき、その発展の歴史は社会における AI の役割を理解することに役立つ。

図 2：AI の技術スタックの構成要素



出典：マイクロソフト

汎用技術の恩恵が社会全体に行き渡るまでのプロセスは、インフラへの大規模投資から始まる。電気が同様のプロセスを経たように、現在は AI インフラへの投資が世界中で行われている。マイクロソフトが 2024 年からの 2 年間で 29 億ドルを日本での事業に投資すると発表したことはその好例だ。また、AI 技術スタックの上方に位置する基盤モデルについても、OpenAI の GPT に代表される LLM に加えて、非常に多様な規模や用途のモデルが存在し、われわれのプラットフォームでも幅広く提供している。NTT の日本語処理能力に優れる「tsuzumi」のような比較的軽量のモデルも登場した。

しかし、これからの AI エコノミーの成長を牽引する主役は、技術スタックの最上位に位置するアプリケーションだ。医療や教育、災害対策のような公共の利益に資するユースケースへのアプリケーションに期待が高まる。エージェンティック AI の活用事例も出てきた。トヨタは 9 つの AI エージェントを実装した「O-Beya (大部屋)」という社内システムを構築し、運用を始めた。ベテランエンジニアの知見の引き継ぎにも有用という。このようなイノベーションが、特定の国に集中することなく、世界各地で試行錯誤が進行していることは見逃せない。AI ガバナンスに関する関心が世界中で高まっているのはそのためだ。

2025 年 2 月にパリで行われた AI Action Summit では、AI ガバナンスの将来を考える上で重要な方向性について議論された。重要な方向性は 4 つある。第一に、報告面の国際協力で透明性を高めることだ。広島 AI プロセスが提供する「報告枠組み」は、国境を超えた情報共有をより速く行うことに貢献するだろう。第二に、リスクについての研究を強化することだ。AI Safety Institute のグローバルネットワークの取り組みは、AI の影響力をよりよく理解することに貢献できる。第三に、これらに関連したオープンソースの AI ツールの選択肢を増やすことだ。AI ツールへの自由なアクセスは、汎用技術としての AI が社会に受け入れられるために欠かせない。最後が広島 AI プロセス・フレンズグループの活動に代表される包摂性を高める取り組みだ。同グループの取り組みは、世界中の大小の様々な規模の国々の意見を取り入れていることを国内外に示している。

■パネルディスカッション：AI ガバナンスの動向、追及すべき価値、日本の役割

渡辺：ここまでの話を振り返ると、日本の AI 法案は、罰則を設けることや安全性を重視してのイノベーションを遅らせるものではありません。持続的なイノベーションを推奨する方向で、政策や制度づくりの舵取りをすると受け止めました。また、AI Action Summit で、イノベーションと規制の関係について議論が交わされたが、日本の考え方はこの法案からも伺えます。

ここからのパネルディスカッションでは、時間の許す限り、用意した質問を基に、自由に意見を述べてもらいたいと思います。まず、パネルディスカッションから参加した 3 名の皆様に、自己紹介と最初の「AI ガバナンス」に関する質問へのコメントをお願いします。

Business Software Alliance アジア太平洋（APAC）政策担当シニアディレクター Jared Ragland 氏

AI セーフティ・インスティテュート 所長/損害保険ジャパン株式会社 執行役員 CDaO 村上明子氏

富士通株式会社人工知能研究所 シニアリサーチマネージャー 中尾悠里氏

モデレーター：渡辺智暁

図 3：ディスカッションアジェンダ

Panel Discussion

1. AIガバナンスの動向

近年のAIガバナンスに関する世界の動きの中で注目しているものを1つか2つ挙げるとしたら何でしょうか？

If you pick one or two, which of the so many recent developments in AI governance around the world do you think are particularly noteworthy (and why)?

2. 追及すべき価値

日本のAIガバナンスの今後を考えていく上で重視すべき価値を挙げるとすると何でしょうか？

What are the values do you consider important for Japan's AI governance going forward?

3. 日本の役割

日本はこれまでG7に端を発する国際協調の推進や米国の半導体輸出制限への協力など、欧州評議会の枠組み条約への署名など、様々な役割を果たしてきました。今後日本が果たすべき役割についてどう考えますか？

What is role Japan should play in global AI governance for the coming years? It has been promoting international cooperation through Hiroshima AI process, been a partner in the U.S.-led trade restriction regime on some semiconductor chips, signed Council of Europe's Framework Convention on AI, and played other roles.

出典：国際大学 GLOCOM

Jared：Business Software Alliance（BSA）は、エンタープライズソフトウェア産業を代表する業界団体で、米国、欧州、ラテンアメリカなどをカバーしています。私はシンガポールを拠点に アジア太平洋地域の市場をカバーしています。現在、あらゆる所で AI ガバナンスとそれを支える技術についての議論が活発に行われています。私たちの活動は政策に焦点を当てているので、その観点からお話

ししたいと思います。

まず、日本を含む先進国のほとんどが、AI ガバナンスに関する議論を 5 年以上続けてきたことを忘れてはならないと思います。決して新しい論点ではない AI ガバナンスですが、世界で初めて包括的なアプローチを採用したのが EU の AI 法です。この法律には、AI の開発を一律に規制するのではなく、リスクベースのアプローチを採用したことに代表される、評価すべき点が多くあります。また、日本政府が、広島 AI プロセスを通じて、グローバルリーダーシップを発揮したことも注目すべき進展です。

アジア太平洋地域では、各国政府がさまざまなアプローチを採用しています。韓国では、AI に関するリスクベースのアプローチを重視する法律が採択されました。日本と豪州でも AI 法制度の議論が進み、リスクベースのアプローチに注目が集まっています。その一方で、シンガポールのように、拘束力のない枠組みやガイダンス指向のアプローチを採用する国もあります。各国の動向は、うまく行ったことといかなかったことの共有に役立つと思います。さらに、AI Safety Institute が立ち上がったことも良いことです。

これからの発展に向けて、規制アプローチにソフトローで拘束力がないものを採用するか、法的拘束力を課すのか、最終的に国際的な整合性を保てるかどうかに関心があります。

村上：AI Safety Institute (AISI) は、2023 年 11 月に行われた AI Safety Summit 2023 をきっかけに、AI Safety を考えなければならぬと、英国が率先して最初の組織が立ち上がりました。この動きに米国と日本が続ぎ、日本では 2024 年 2 月に設立されました。この数年で AI を取り巻く環境は大きく変わりました。特に、広島 AI プロセスの前に、EU の AI 法ができたことは大きい。韓国で行われた AI Safety Summit 2024 に続いて 2025 年 2 月にパリで行われた 3 回目の Summit は、名称を「AI Action Summit」と改め、安全性だけでなく、イノベーションのためにどんな行動が必要か、より踏み込んだ意見交換が行われたと思います。

AI の安全性に関する動向には、大きく二つの潮流があると思います。一つは、米国で第 2 次トランプ政権が発足し、安全性よりもイノベーションに軸足を置く傾向が明らかになったことです。元々、安全性は非常に広い概念です。特に、National Security、すなわち国家の安全保障に関わることに国の機関は注力するべきという考え方が発展してきたと感じます。最近、英国の AISI は機関名を AI Security Institute に変更しました。ミッションは変わらないものの、軸足が安全保障に移ったことになります。

また、当初は、安全性基準を満たしている認証を発行することを検討していた国が多かった。けれども、今の技術動向のスピードを鑑みると現実的ではない。広島 AI プロセス・フレンズグループが報告枠組みの策定を進めているのを見ても、安全性に配慮し、かつ透明性のある AI を開発していることへのお墨付きを与える方向に変わってきたと思います。最先端の技術動向を理解した上で AI の安全性を語ることが重要になってきました。

中尾：所属する富士通の人工知能研究所には、LLM を開発している部署もあれば、AI ガバナンスのツールを開発している部署もあり、さまざまな研究を行なっています。私個人は 2022 年に『AI と人

間のジレンマ：ヒトと社会を考える AI 時代の技術論』を出版したことの縁で、内閣府の AI 制度研究会のメンバーとしても活動しています。

私が注目している AI ガバナンスの動向は、EU の AI 法の今後の実効性です。トランプ政権の影響が、米国だけでなく世界中に及ぶかもしれない。村上さんの話に出てきた安全保障については、NIST の舵取りを含めて、AI ガバナンスにどう影響するのかが気になるところです。また、私自身は、研究者としてイノベーションを起こすことがミッションなので、世界の AI 研究開発が縮小するのか、それとも方向性が変わるのか、その先行きを注視しています。また、例えば、広島 AI プロセスの報告枠組みに合わせて、人間中心の AI に研究開発の方向性を沿わせていくことも重要だと考えています。

■AI の進化のスピードにガバナンスが追従できるのか？

渡辺：渡邊さんと Marcus さんにも、AI ガバナンスの動向についての意見を聞かせてほしいです。

渡邊：Marcus さんの話に、エージェンティック AI の事例がありました。この技術が社会に浸透することを踏まえ、ガバナンスがどう変わるのか。今は人間の関与がありますが、人間が AI エージェントに委任した結果への責任は誰が負うのでしょうか。委任した人か、それとも AI 開発・提供者か。技術の発展を止めてはいけないし、止められない。その時に備えて、どうすべきかに関心があります。

Marcus：2つの観点があります。1つは AI の適用と行動の両方に焦点を当てていることです。過去3年間を振り返ると、アジアでは「AI イノベーションと適用」と「ガバナンスと行動」の取り組みが並行して行われてきました。日本でも2つが並行して進んでいますが、グローバルな議論をするときは2つを区別するべきだと思います。

もう1つは、実践的な動きがあることです。世界は、基本的フレームワークと原則の設定から、実践的な活動へと焦点が移行しています。すでに紹介された例でも、広島 AI プロセスの報告枠組みを包括的なものにしようとする取り組みは、実践的なものの好例です。AISI のようなグローバルネットワークでの活動もあります。その活動は世界の研究者間の実践的な技術交流につながり、AI の安全性研究に関する科学を前進させられる。これらの動向は将来性のあるものばかりです。

渡辺：他の皆さんは、今のコメントに補足したいことがありますか。

中尾：研究開発をしている立場からすると、新しい技術の登場で、今までできなかったことが次々できるようになりました。エージェンティック AI のような新しい事例の話がありましたが、リスクが顕在化してから対応技術が出てくる一方で、ガバナンスはその技術がない状態で方向性を考えなくてはならない。その状態で規制していいのか。ガバナンス対策なしの AI が広まってもいいのか。この点をどう考えればいいかが重要な論点だと思います。

渡邊：その通りだと思います。技術が変化すれば、ガバナンスも変化する。技術がダイナミックに変化することを前提に、ガバナンスも考えなくてはならないと思います。

村上：言語処理学会に出席したところ、AI の安全性が話題になりましたが、面白いことに、エージェント AI に関しての安全性を議論している人は誰もいなかった。私も研究者なのでよくわかるのですが、AI エージェント自体は古くから研究されてきたテーマなので、今更その安全性を議論するのかわかれている。研究者だけでなく、実務家や政策立案者が一緒に話し合わなくてはならないと思いました。

渡辺：アカデミアには、すでに起こりうる問題の解決策の知恵があるのですか。

村上：いいえ。生成 AI 時代に AI エージェントがどう変わるかという観点での安全性は何も議論していない。それが問題です。

Marcus：非常に重要な問題提起です。エージェント AI は素晴らしいイノベーションですが、規制やガバナンスが、実現したいことが基本原則に基づいているかを確認しなくてはならないと示しています。特定の技術を規制しようと努力しても、その技術に追いつくだけの規制になってしまう。それが避けられないとすると、原則に基づく規制の中で枠組みを定め、その技術の利用に伴う避けるべき実害は何かを明確にする。その方が直接的な規制よりも効果的だと思います。

Jared：ここまでの話から追加したいことが 2 つあります。1 つは、ガバナンスの文脈で過去数年間にできたリスクと機会です。そのリスクとは、AI 関連政策の分断化です。国際的な分断と国内の分断の両方があります。国際的な分断のリスク軽減には、広島 AI プロセスの活動が貢献すると思いますが、国内の分断リスクとは何か。例えば、AI ガバナンスのビジョンを示す国のリーダーシップが弱い場合、セクター別の規制当局が自分たちの管轄内だけで対処しようとするでしょう。個人情報や著作権など、さまざまな問題に対して一貫性のないルールや法的要件を生み出すリスクを内在化させてしまうことになりかねません。

その一方で機会もあります。中央政府がアプローチを考えながら、渡辺さんが話していたように、既存の法律を利用すること、そして既存のデータガバナンス規制が AI エコノミーに適用できるかを検証することです。韓国の個人情報保護委員会では、AI の開発力に不必要な影響を与えないように、個人情報保護の厳しい要件の一部を調整することを検討していると聞きました。これはほんの一例です。私たちが考えなくてはならないのは、国単位でも国際的にも整合性のある規制環境をどう作るかです。

■これからの日本の AI ガバナンスで追求すべき価値原則とは？

渡辺：Marcus さんが指摘した、規制が技術に追いつくためには、価値原則を重視すべきという話に、2 つ目に用意した質問「追求すべき価値」とも関連します。また、Jared さんが分断化リスクには 2 つあると話していました。広島 AI プロセス・フレンズグループにおける日本のリーダーシップについて評価してもらったのですが、今後に向けて期待することはありますか。どちらかを選んで、最後にパネリストの皆さんに一言をお願いして締めくくりにしたいと思います。

渡辺：日本は中立的な役割を担える国だと思います。グローバルに活躍する特徴的な企業が少なくともあって、中立的に振る舞える。上から目線で「こうしなさい」ではなく、下から AI ガバナンス

を支えることができればと思います。もう一つ、「規制」という言葉はあまり使わない方がいいのかもしれない。Marcus さんが話していたように、特定の技術が登場する都度、規制するのはほぼ不可能です。ガードレールやフレームワークを明確にした上で、その範囲内で事業者に自主的に活動してもらうことが現実的だと思いました。

Marcus：日本のような民主主義社会では、法律にその国の価値観が反映されているものです。AI 技術を考えるときも、その価値観に根ざした基本原則を当てはめることだと思います。そして、これからの AI の機会を考える時に重要な価値の一つは、包摂性だと思います。電気を例にとると、日常的にアクセスできない人々が世界中に何千億人もいます。AI が汎用技術だとすると、電気と同様の状況になるのは避けなくてはならない。その恩恵が一部の国にとどまらず、できるだけ広く普及させるには、包摂性の価値が重要になると思います。

Jared：日本が国際的なリーダーの役割を果たしていることは、私たち BSA が、日本国内の規制や政策の動向に注目している理由の 1 つでもあります。日本の取り組みは、他の国々のモデルとして大きな影響力を持ちます。特にアジア太平洋地域では、米国が明確な方向性を示さない場合、日本を参考にする。EU の AI 法には、称賛に値するところがたくさんありますが、アジア太平洋地域の国々は、規範的なものではなく、イノベーションの促進と国の競争力を高めることのできるアプローチを求めています。そこに日本が国際的な議論における強力なリーダーシップを発揮できる場所があります。

村上：渡邊さんの話にもあったように、日本という国の立ち位置は非常に面白い。独自の言語を話し、独自の文化を持っています。その意味では、AI の安全性は国ごとの文化的背景の相違が反映されるはずです。一方で、世界に一つの AI 安全性があると信じている人たちがいて、全てに規制をかけようとしている。それが今の混乱の原因ではないかと思います。各国が大事にしている価値観を尊重し合いながら、AI を安全に使うために、何が世界共通の脅威なのか、それを洗い出すためにどんなフレームワークが必要なのかを議論すべきです。広島 AI プロセスをきっかけに、日本はその議論のリーダーになれる。多様な文化を世界に認めてもらう。それが日本のできることだと信じています。

中尾：私が重視すべきだと考える価値原則は 2 つあります。一つがマルチステークホルダーの参加です。LLM は非常に複雑で、専門家ですらもう中身を完全にトレースすることはできない。今までの技術では、予防的観点からアプリケーションに応じたリスクの洗い出しの議論が行われてきました。今後は社会全体で AI という技術自体のリスクを考えなくてはならない。その意味で、研究者や実務家だけでなく、市民参加が不可欠になると思います。もう一つは、文化的な多様性の尊重です。AI の安全性に文化的多様性があるように、LLM にも文化的多様性があります。どこか一国の汎用的な価値観ではなく、その土地ごとの、日本であれば日本の価値観の中の尊重すべき部分を議論すべきだと思います。