

ブリュッセル効果への対応：日本企業は EU-AI 法にどう備えるべきか 4

日時	2025 年 5 月 12 日（月）17:30-18:30
会場	Zoom ウェビナー
主催	東京大学国際高等研究所東京カレッジ 東京大学未来ビジョン研究センター 東京大学次世代知能科学研究センター
後援	日本 AI セーフティ・インスティテュート 大阪大学 社会技術共創研究センター 日本ディープラーニング協会

はじめに

EU-AI 法は、2024 年 5 月に成立し、2024 年 8 月 1 日に発効した、世界初の AI に関する包括的な規制法である¹。AI 法の規定は 2030 年 12 月 31 日までに段階的に施行され、日本の企業や組織もこの影響を受ける可能性がある。

東京大学では、2024 年 12 月 11 日にウェビナーイベント「ブリュッセル効果への対応：日本企業は EU-AI 法にどう備えるべきか」²を開催し、EU-AI 法に加え、それに関連して策定作業が進められている汎用目的 AI（general-purpose AI、以下「GPAI」）モデルに関する行動規範（Code of Practice、以下「CoP」）の第一草案について、専門家が解説・議論を行った。また、翌年 2025 年 1 月 15 日の第 2 回目³では、第二草案について、3 月 19 日の第 3 回⁴では、第三草案について解説・議論が行われた。そして第 4 回となる今回のイベントは、EU-AI オフィスによるガイドラインのパブリック・コンサルテーション（公開アンケート）の募集⁵を受けて開催され、その公開アンケートの概要と日本企業が留意すべき対応の要点などについて解説・議論があった。

本レポートでは、当日の様子を簡単に紹介するとともに、視聴者から寄せられた質問とそれに対する登壇者の回答を改めて整理して示すこととする。

¹ EU-AI 法は、EU 法における規則（regulation）であり、加盟国に直接適用される。

² 第 1 回イベント：https://www.tc.u-tokyo.ac.jp/ailec_event/13586/

³ 第 2 回イベント：https://www.tc.u-tokyo.ac.jp/ailec_event/13797/

⁴ 第 3 回イベント：https://www.tc.u-tokyo.ac.jp/ailec_event/14275/

1. 当日の様子

開会のあいさつ

はじめに、総務省の飯田陽一氏から開会の挨拶があった。

飯田氏は、EU-AI 法、とくに CoP に関する議論が佳境に入っている中で、マルチステークホルダー・アプローチに基づき、専門家、産業界、市民社会など幅広い関係者の意見を取り入れながら進められている点に言及した。

その上で、現状の EU-AI 法の規制が重たくなりすぎているのではないかという懸念が各方面から聞かれると述べ、自身も中東地域の政府職員を対象とした JICA の研修において AI ガバナンスについて講義した際、「EU-AI 法は過剰規制であると思うか？」との質問を受けた経験を紹介した。その際、明言は避けつつも、個人的には過剰な要素が多分に含まれていると感じる旨を共有したところ、相手は大きく首肯していたという。

さらに、先週行われた G7 の AI ガバナンス議論にも触れ、広島プロセスの推進が中心議題となっていること、そして EU もこれに積極的な姿勢を示していることを報告した。その一方で、EU-AI 法と広島プロセスとの相互運用性（Interoperability）に疑問を呈し、関係者間で相互運用性の解釈や立場の違いがある可能性を示唆した。

最後に、過剰な規制は EU にとっても望ましいものではないと強調し、本イベントでの議論が EU 内の議論にも反映されることを期待していると締めくくった。



飯田氏

論点提供：公開アンケートの概要と日本企業が留意すべき対応の要点

続いて、公開アンケートの概要と日本企業が留意すべき対応の要点について、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と実積寿也・中央大学教授それぞれから解説があった。

まず工藤特任准教授⁶は、EU-AI 法に関するこれまでの経緯と、新たに意見公募が始まったガイドライン案について説明した。

⁶ 講演資料:

https://www.tc.u-tokyo.ac.jp/blog/wp-content/uploads/2025/02/20250512_merged.pdf#page=2

冒頭で工藤特任准教授は、EU-AI 法はリスクベースアプローチを取るハードローであり、リスクの高さに応じた 4 つのリスク類型が整備されており、域外適用があるので日本企業も無関係ではなく、違反時には高額な罰金があることなどについては、報道などを通じてよく知られていると説明した。一方、本イベントではその 4 つのリスク類型ではなく、「GPAI モデル」に関する別枠の規制を取り上げることが主眼としていると述べた。GPAI モデルとは、生成 AI や大規模言語モデル (LLM) のように、広く応用可能で複数のアプリケーションや API に接続される AI を指し、大量のデータで学習されたものを含むとしながらも、それに限定されない定義がなされていると説明した。

また、工藤特任准教授は、研究開発やプロトタイピングのために使われる AI モデルは規制から原則除外されるため、研究者にとっては安心材料となると述べた。

次に、GPAI モデルのうち「システミックリスク」があるかどうかに応じて、義務が分かれることを説明した。システミックリスクとは、化学・生物兵器等の開発障壁の低下、人間の制御不能な AI の出現、有害な差別や偽情報の拡大などを指す。当該リスクがない場合は、技術文書の提供や著作権遵守ポリシーの策定と実施などが求められるが、当該リスクがあると判断される場合は、事前アセスメントやインシデント発生時の当局報告など、より重い義務が課される。また、システミックリスクの判断には FLOP（浮動小数点演算数）による累積計算量が用いられ、**10²⁵ を超える**場合はリスクがあると推定されることも述べた。さらに、有識者との協議を経て EU 当局が個別指定する可能性もあるという。

これまでのウェビナーでは、GPAI モデルに関する「CoP」の策定に向けて、有識者としてインプットを行ってきたと工藤特任准教授は振り返った。当初この CoP は、マルチステークホルダー協議を経て 2025 年 5 月までに策定される予定だったが、イベント開催日までには最終化されてなかった。他方、突如として 2025 年 4 月 22 日に新たに別のガイドライン案が公表され、5 月 22 日まで意見が公募されることになったという。今回のガイドライン案は誰でも意見提出可能である点が従来の CoP と異なるとして、積極的な意見提出を呼びかけた。

さらに、日本企業への影響についても言及した。GPAI モデルの調整やファインチューニングを行った主体も、条件によっては新たなプロバイダーと見なされる可能性があり、日本企業に直接影響を及ぼす点が課題とされた。日本の専門委員からは、この点に関する明確な解釈を EU 側に要望してきたが、CoP の範囲外とされ、明確な扱いは未定だった。今回の新ガイドライン案でこの点が整理されたことは、日本側の要望が反映された成果であり、良いニュースだと述べた。

そのうえで工藤特任准教授は、EU-AI 法の条文に基づいて、CoP では技術文書のテンプレートや安全性・セキュリティフレームワークの整備が行われてきたことを確認した上で、今回の新ガイドライン案では、「汎用」の定義や、どこまでが特化型でどこまでが汎用型かといった境界、（システミックリスクのない場合も含む）汎用モデルの推定基準として 10²² という閾値の提案、調整・ファインチューニングした主体の取り扱い方針など、日本企業が特に注目すべき内容が盛り込まれていると説明した。

また、こうした動きの背景について、工藤特任准教授はあくまで個人的な推察としながらも、EU では現在、国際競争力強化を目的とした AI 規制の見直しが進んでいることに言及した。ドラギレポート⁷による規制の簡素化の提言や、アメリカ政府からの圧力、さらにはフランスのマクロン大統領の規制緩和と要求など、複雑な政治的背景が影響している可能性を指摘した。その一方で、EU-AI 法は既に発効し

⁷ ドラギレポート:

https://commission.europa.eu/topics/eu-competitiveness/draghi-report_en

ているため、条文の大幅な変更は困難であり、ガイドラインでどこまで柔軟に対応できるかが問われていると述べた。

最後に、今回のガイドライン案には FLOP の計算式や、NVIDIA 製チップ（例：A100）を使った場合の具体的な計算方法など、非常に詳細な技術的記載が含まれていると紹介した。そして、今回は、法学や経済学の専門家よりも、ビジネスやエンジニアリングの観点からのインプットが重要となるため、ぜひ多くの方に具体的な改善提案を含めた意見（例：～とある箇所は…という趣旨だと理解したが、XX の観点からは〇〇とすべき）を提出してもらいたいと締めくくった。

次に、実積教授⁸は、工藤特任准教授と同じく GPAI のガイドライン案に関して補足説明を行った。

まず、EU-AI 法における生成 AI、特に LLM や GPAI に関する規制が 2025 年 8 月 2 日に施行予定であると述べた。その施行開始に向け、昨年からの CoP の議論が進められてきたが、1 月のバンス副大統領の演説などを契機に状況が大きく変化し、4 月にはスケジュールが大幅に延期されたうえで、ガイドライン案への公開アンケート募集が始まった経緯を説明した。

今後のスケジュールについては、当初 5 月上旬に CoP が発表される予定であったが、それが延期され、今回のガイドラインと併せて来月中には確定させ、8 月 2 日の施行に間に合わせる方針であると説明した。そのため、公開アンケート提出の機会が延長された一方で、施行までの準備期間は短縮されているとも指摘した。

ガイドライン案は CoP を補完する非法的文書であり、法的拘束力はないと述べた。しかしながら、ガイドラインに従うことによって EU-AI 法の義務を果たしていると主張する根拠になり得ることから、多くの企業や国にとっては、ガイドラインおよび CoP への準拠がコスト面でも有利に働くとは指摘した。

また、EU-AI オフィスはこれまで、CoP を EU-AI 法の実施基準と位置づける一方で、EU-AI 法そのものの解釈は CoP の対象外とされてきたが、おそらくは米国との交渉などの影響もあり、今回のガイドラインでは EU-AI 法の解釈にも踏み込む内容が含まれていると述べた。特に、ファインチューニングの程度によって既存モデルが新たなモデルと見なされるかどうかなど、曖昧な点に対する解釈が示されたことを紹介した。

さらに、今回の公開アンケートでは誰でも意見提出が可能であるとしたうえで、意見募集ページの紹介を行った。主な論点としては、GPAI モデルの定義、Downstream Modifier がどの程度の変更で新たな提供者と見なされるか、モデルの市場投入時期の判断基準、オープンソースモデルにおけるビジネスモデルの許容範囲など、実務的な内容が中心であると説明した。

最後に実積教授は、こうした実務的課題については学術的観点からの断定が難しいため、実際にサービスを提供しビジネスを構築している関係者からの意見提出が強く望まれていると述べ、説明を締めくくった。

パネルディスカッションと Q&A

1. 参加企業 3 社からの論点提供

続いて、パネルディスカッションに移る前に、本イベントに参加する 3 社より話題提供があった。

⁸ 講演資料:

https://www.tc.u-tokyo.ac.jp/blog/wp-content/uploads/2025/02/20250512_merged.pdf#page=12

最初に、日本電気株式会社の吉永和弘氏⁹より、ガイドラインに関する注意点3点について紹介があった。

まず導入として、今回のガイドラインでも使用されている指標「FLOP」と「FLOPS」の違いについて説明した。FLOPは浮動小数点演算の総量を示し、自動車で言えばオドメーター（積算距離計）に相当する。一方でFLOPSは1秒間あたりの浮動小数点演算の回数であり、同じく自動車で言えばスピードメーター（速度計）のような指標である。EU-AI法では、GPAIモデルのトレーニング量の指標として前者の**FLOP**を使用しているが、国内外の解説記事では後者のFLOPSと混同されることも多く、注意が必要だと指摘した。また、EU-AI法のドラフト版では「FLOP」に小文字の「s」を付けた表記も見られたが、混乱を避けるために最終版では削除された経緯があるという。

これを踏まえ、吉永氏は、改めて1点目の論点として、GPAIモデルの該当基準に関する説明を行った。今回のガイドライン案では、GPAIモデルの該当基準として 10^{22} FLOPのトレーニング量が設定されている。吉永氏は、2024年以降にリリースされたほとんどの主要なAIモデルがこの基準を超えていることから、この数値は現状の該否判断に大きな影響を与えるものではなく、妥当であると評価した。なお、 10^{22} の基準を若干下回る小規模モデルについては、パラメータ数が10億に達していないため、EU-AI法の前文98にある基準に照らしてもGPAIモデルには該当しないと述べた。

2点目の論点は、修正によるGPAIモデルのプロバイダー該当基準についてである。ファインチューニングなどの修正に必要となるトレーニング量が、 10^{22} FLOPの3分の1以上となると、新たなGPAIモデルやそのプロバイダーと見なされる。また、その修正によってシステミックリスクの基準（ 10^{24} FLOP）の3分の1以上となった場合も同様に、新たなシステミックリスクGPAIモデルやそのプロバイダーとなる。これにより、義務や責任の範囲が明確になり、小規模な修正でプロバイダーと見なされる懸念は軽減されたと述べた。一方で、例えばMetaのLlama3-70Bのように、修正前の段階でシステミックリスク基準の8割程度に達しているモデルを修正し基準を超えた場合には、修正・再トレーニングした者がモデル全体について責任を負う可能性がある。日本の産業界では既存の大規模モデルの再トレーニングが一般的であることから、こうした規定は影響が大きいと懸念を示した。加えて、こうした懸念を視覚的に示す図を紹介し¹⁰、基準を超えるか否かで修正者の義務の範囲が変わることを明示した。

続いて3点目の論点として、GPAIモデルの定義における「汎用目的（GP）」の基準について言及した。たとえトレーニング計算量が前述の基準を超えていたとしても、実行できるタスクが限定されている場合には、GPAIモデルとは見なされない可能性がある。実際、ガイドライン案では、特殊なデータセットやタスクに特化してトレーニングされたAIはGPAIモデルに該当しないと例とともに紹介されている。吉永氏は、日本のAI企業の多くが特定の業界や業種に特化したAIの開発を進めている現状を踏まえ、こうした特化型AIがGPAIモデルの定義から除外されるかどうか、今後の日本の産業界への影響を左右する重要な点になると述べた。

⁹ 講演資料:

https://www.tc.u-tokyo.ac.jp/blog/wp-content/uploads/2025/02/20250512_merged.pdf#page=16

¹⁰ 図:

https://www.tc.u-tokyo.ac.jp/blog/wp-content/uploads/2025/02/20250512_merged.pdf#page=20

最後に吉永氏は、3点をまとめて次のように整理した。第一に、GPAIモデルのトレーニング量の基準は大きな影響はなく、妥当である。第二に、システミックリスクの基準をやや下回るモデルを修正するには特に注意が必要となる。第三に、特化型AIがGPAIモデルの定義から除外されるか否かの判断基準は、特化型AIを多く開発している日本にとって重要な論点になると締めくくった。

次に、CoPの専門委員でもある日本電信電話株式会社の根本宗記氏¹¹から、今回のガイドライン案について話題提供があった。

まず、ガイドラインに示されたFLOPの基準について触れ、システミックリスク基準が 10^{25} FLOP、GPAIモデル基準が 10^{22} FLOP、そしてそのGPAIモデルの微修正時の基準が 3×10^{21} FLOPと明確に定められている点に関して高評価を示した。また、実務的な用語（例：市場投入）の解釈が具体的に示されている点も、実務上有益であると評価した。

特に注目すべきは、ガイドラインに記載された「モデルA～D」の例である。モデルAは、NVIDIA A100のGPUを用いた例で、FLOPの計算方法として「ハードウェアベースアプローチ」と「アーキテクチャベースアプローチ」の2つが提示されている。根本氏は、前者のハードウェアアプローチでは「学習時間（秒）× GPU性能 × GPU使用率」によりFLOPが算出され、非常に分かりやすいと評価した。

実際の計算例も示し、35,000時間の学習を行った場合、FLOPは 10^{22} の約2倍となり、これは先のGPAIモデルの基準を超える計算量になると解説した。仮に10枚のGPUで学習していた場合、約145日、つまり約5カ月弱の学習期間となる。これを基に「A100を10枚使用した場合、2.5カ月程度でGPAIモデルの基準に達する」と整理できるという。

一方で、根本氏はFLOP基準に対する懸念も表明した。具体的には、GPUの浮動小数点演算性能はAIモデルの実際の学習性能に直結しない点を指摘した。たとえば、NVIDIA H100は先ほどのA100よりも演算性能が約6倍あるが、実際の学習性能は2倍程度にとどまるという報告もある。そのため、今後より高性能なGPUが登場することで、FLOP基準が実質的に緩くなってしまう可能性があるとの懸念を示した。

そこで代替案として、モデルの「パラメータ数」と「学習トークン数」に基づいて学習ボリュームを評価する方法が提案された。こうした基準であれば、ハードウェアの進化による影響を受けにくく、より安定的な指標となるのではと述べた。

さらに、アルゴリズムの効率化によっても学習に要する計算量が大きく変わることを踏まえ、FLOP基準は定期的に見直すべきだと提言を行った。また、AIモデルの学習量と社会リスクとの関係性が十分に解明されていないことにも触れ、今後の検討課題として挙げた。

なお、 10^{25} FLOPというシステミックリスクの基準は既に法律で定められており、簡単には変更できない点にも言及した。そのうえで、「アーキテクチャベースアプローチ」の計算式（パラメータ数 × 学習トークン数 × 6 ÷ FLOP）を残すことで、より現実的な代替提案として提示できるのではないかとの見解を述べた。

¹¹ 講演資料:

https://www.tc.u-tokyo.ac.jp/blog/wp-content/uploads/2025/02/20250512_merged.pdf#page=24

続いて、株式会社 Preferred Networks の大野健太氏からは、三つの論点提示があった。

まず 1 点目として、特化型モデルが GPAI モデルと判断される基準として、タスクの実行可能性が一つの判断基準とされている点について言及した。そのため、たとえ医療用途に特化したモデルであっても、タスクの広がりによっては GPAI モデルと判断される可能性があり、その判断基準が特化型 AI の開発に大きく影響を及ぼすと述べた。

次に 2 点目として、合成データセットの生成にかかる計算量がモデルの評価に含まれる点について疑問を呈した。ワーキングドラフトや法律において、合成データの生成も FLOP の計算に加える旨が記されているものの、これは本当に妥当なのかという疑問を述べた。特に、非合成データと近いデータであっても、その生成に使ったモデルの計算量が含まれることとなれば、合成データの使用自体が大きく制約を受ける恐れがあるとの懸念を示した。

そして最後に 3 点目として、推論時の計算量について問題提起を行った。特に初期段階での In-Context Learning や、AI エージェントによる推論時間の増加が精度向上に寄与するケースが増えており、同じモデルでも推論にかかるコスト次第で汎用性が高まる可能性があるとして述べた。そのため、推論側の計算コストがモデル評価にどう影響するかについて、今回のワーキングドラフトでは明記がなかった点を懸念点として挙げ、今後の議論動向に注目していると述べた。

最後に、大野氏は改めて、特化型モデルの判断基準、合成データセット生成時の計算量、推論における計算量の 3 点が自身の懸念点であるとまとめ、発言を締めくくった。

2. パネルディスカッション

パネルディスカッションでは、前述の論点提供を行った 3 名に、公開アンケートの概要についての解説を行った工藤特任准教授と実積教授も加わり、江間有沙・東京大学准教授の司会で、議論が交わされた。

まず工藤特任准教授は、3 名の登壇者からの話題提供に対し感謝を述べた上で、自身が初めてガイドラインを読んだ際の戸惑いを振り返り、打ち合わせや今回のイベントでの登壇者による説明が理解を深める助けとなったと言及した。

続けて工藤特任准教授は、大野氏に対し質問を投げかけた。大野氏が挙げた、日本企業が注目すべき三つの論点 — ①特化型モデルと GPAI モデルの違い、②モデルの計算量における合成データセットの扱い、③推論に関する計算量 — に触れた上で、AI エージェント開発を想定した場合に、GPAI モデルの評価に推論時の計算量が含まれないほうが開発側にとって有利であるという理解でよいのかという点を確認した。また、どのような立場があり得るか、場合分けを求めた。

これに対し大野氏は、立場によって見解が分かれる点を認めた上で、計算量を多く見積もるほどシステムリスクを伴う GPAI モデルと見なされる可能性が高まり、対応すべき規制が増えることになるとして指摘した。そのため、開発側としては、推論時の計算量が規制判断に含まれない方が望ましい可能性があるとして述べた。また、推論時の FLOP 数を単純に足すべきかについては、議論の余地があるとし、使用する GPU によって FLOP 数が大きく変動する可能性があるため、単純加算による評価が妥当かは慎重に検討すべきであると指摘した。

これを受け工藤特任准教授は、大野氏の対応に感謝を示した上で、自社で LLM や SLM（小規模言語モデル）を開発する企業と、AI エージェントを通して LLM を使用している企業では利害が異なる点を

改めて確認した。加えて、今回の意見募集は誰でも提出可能であることから、多様な立場からの意見提出が望ましいと述べ、発言を締めくくった。

次に実積教授は、先の話題提供を通じて、登壇者たちが共通して、より詳細な基準設定の必要性を論じていたことに言及した。また、詳細化によって対象を絞る方向と、あえて最大公約数的に基準を曖昧に保ち、法律が対象としている GPAI モデルに自社サービスが該当しない可能性を維持しやすくするアプローチの両方が存在するのではとの見方を示した。そのうえで、どちらがビジネス上では有利なのかという点について登壇者に意見を求めた。

これに対し江間准教授が根本氏に意見を促すと、根本氏は、今回のガイドラインでは対象の明確化が非常に詳細に記載されており、実務的に大変参考になると評価した。その上で、技術の進歩によって基準がすぐに陳腐化してしまう点に懸念を示した。具体的には、GPUの種類が変わるだけで基準が大きく変化してしまうことから、より影響が少ない基準設定が望ましいとの見解を示した。

これに対し実積教授は、基準を現時点の GPU に基づいて明確に設定することが将来的な技術進展に対応しやすくなる一方で、それが日本企業にとって将来の制約となる可能性もあるのではないかと再度問いかけた。

それを受け根本氏は、現在のままでは技術進展により実質的に基準が緩和され、適用対象が拡大してしまうことを指摘した。そのため、基準を一定の水準以上に保つ、もしくは定期的に見直すことが、ビジネス展開の観点からも望ましいと述べた。さらに、GPU の進化により、たとえ基準の数値 (10^{25}) を変更しなくても、実質的には基準が下がり、適用対象が増えてしまう現状を説明した。

そして最後に、そのような状況を防ぐためには、基準を高めに設定し、将来的な技術変化にも対応できる設計とすることが、ビジネス上有利になるとの見解が両者で確認された。

続いて江間准教授は、企業ごとに戦略が異なる可能性があるとのことから、他の登壇者にも、実積教授の質問に対する回答を求めた。

まず吉永氏は、特化型 AI が GPAI モデルに該当する範囲についてのグレーゾーンが現状では広いため、より明確化されることで企業側としては対応しやすくなると回答した。

これに対し実積教授は、EU-AI 法の規制基準には計算量基準に加えて、EU-AI オフィスによる指定も含まれていると補足し、その AI オフィスに制約を求めるほうがビジネス上有利になるという理解で合っているか吉永氏に確認を求めた。

吉永氏は実積教授の理解に同意を示し、それに加え、特定のタスクに限定された AI がどこまで許容されるのかも重要なポイントになると続け、やはりグレーゾーンの明確化が望ましいとした。

続いて大野氏からも回答があった。まず大野氏は、基準に関して根本氏と同様の意見を持っていると述べた。そのうえで、計算量に関する基準 (10^{25}) が今後も適切であるかは疑問が残るとし、定期的な見直しの必要性を強調した。また、モデルが小型化しても精度が向上するケースが増えており、例えば整備されたデータセットによって 8B モデルが過去の 100B モデルと同等の精度を出すこともあるため、基準の適正性は技術の進展に応じて見直すべきだと述べた。

そのような大野氏からの回答を受け、実積教授は、今後のモデルアップグレード方針は企業によって

異なるため、例外規定（エスケープクローズ）に対応する意見を出すことも一つの方向性であるとまとめた。

これに対し大野氏は、企業の技術戦略が大きな影響を与えるとの実積教授のまとめに同意し、SLMを開発していく企業と、LLMを目指す企業とで対応が異なるだろうと指摘した。

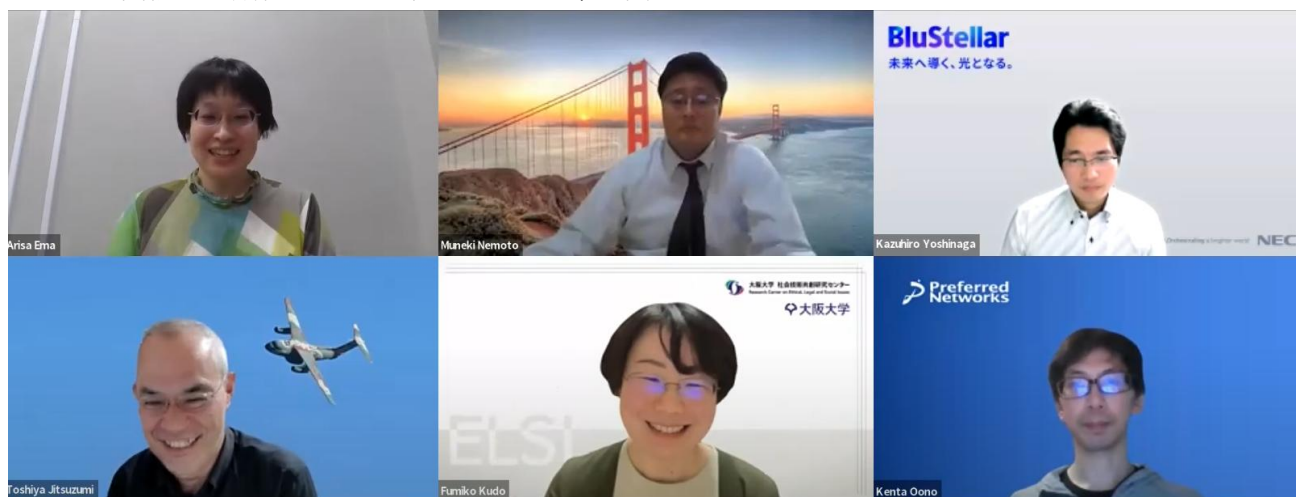
最後に実積教授は、意見集約を他社に依存するのではなく、各企業が個別にパブリックコメントを提出することの重要性を強調し、多様なアプローチが存在する現実を EU-AI オフィスに理解させるためにも、積極的な協力を呼びかけた。

その後、パネリストと司会者は、いくつかの寄せられた質問に答えた（ここでの議論の内容については、最後の [Q&A](#) を参照されたい）。

終わりに江間准教授は、限られた時間の中で話題提供者が要点を明確に提示してくれたことに改めて感謝を述べた上で、今回の公開アンケートは誰でも意見を提出できる機会であることを改めて強調した。また、細かな部分については各事業者の方針によって異なる可能性があるため、積極的な意見提出を促した。

さらに、来月6月には CoP の本体が公表される見通しであることを紹介し、今後も定期的に定義や基準の見直しが行われる予定であるため、引き続き注視する必要性を述べた。

そして、今回のように多様な視点を集め、共に考える場自体にも大きな意義があるとし、日本においても AI 法の制定が進む中で、今後もこのような意見交換の機会を継続していくことの重要性を訴えた。最後に、突然の開催決定にも関わらず参加してくださった話題提供者と視聴者への感謝を述べ、次回のイベント開催への期待とともに本イベントは幕を閉じた。



(上段左から) 江間准教授、根本氏、吉永氏
(下段左から) 実積教授、工藤特任准教授、大野氏

2. Q&A

以下では、視聴者から寄せられた質問とそれに対する登壇者の回答を、改めて整理して示す（なお、質問は一部を選び、適宜編集を行った）。

—— 計算量で汎用性を判断することは不適切であると考え。また、ガイドライン案によると、 10^{22} というのは画像・テキストの生成 AI の場合に限られている。したがって、顔認識 AI は、いくら計算しても GPAI モデルには該当しないということになる。そして、生成 AI 以外の場合についてガイドライン案には記載がないが、現状として、生成 AI 以外は原則として GPAI モデルではないということであるならば、その旨を記載すべきである。顔認識 AI を含め、どの AI も一定の汎用性を有しており、生成 AI の場合のみ記載されていても、不安の解消にはならない。

(補足)

3 条 63 号の GPAI の定義において、*is capable of competently performing* の部分は計算量で決められなくもないが、*displays significant generality* や *performing a wide range of distinct tasks* の部分（要件の主要な部分）については、あまり明示されていない。

工藤特任准教授：

全く同感である。先ほどのプレゼンにて、EU-AI 法は既に発効しており、条文は動かしがたいものであるため、今回の公開アンケートの対象とするのは無意味に近いと述べた。しかしながら、そもそも論でいうと、この GPAI モデルがターゲットとするシステムリスク、すなわち人間のコントロールが外れることや兵器開発の容易化、偽情報の拡散といったリスクに対し、計算量や汎用性という手段がどの程度合理的に関連しているかは非常に疑わしいと考えている。

報道されている通り、アメリカ・カリフォルニア州で類似の規制案である、フロンティア AI モデルに関する安全規制法案（「SB1047」）が提出された際にも同様の議論がなされており、目的とするリスクに対して計算量で閾値を設ける手段の合理性に疑問が呈され、憲法訴訟のリスクなども指摘された結果、ニューサム知事が署名せず廃案となった経緯がある。

繰り返しにはなるが、EU-AI 法に盛り込まれたものであっても、この点に関しては不適切であるという意見には強く共感する。

実積教授：

たしかに、今回の CoP を含めて生成 AI のみを対象としているかのように見えるが、法律自体には GPAI モデルが生成 AI に限られるとは明記されていない。単に「学習量が多く、複数かつ多様な目的に対応可能な AI」と定義されているため、生成 AI に限定すべきであることを明確にするには、今回の意見提出が有効になると言える。

根本氏：

汎用性の高さがリスクの高さに直結し、計算量で汎用性を判断するという論理は蓋然性が低いと考える。そのような状況下で、大きなコストを伴う義務を課している点が CoP に対する批判の要因につながっているのではないかと。

吉永氏：（補足に関して）

前文（99）において GPAI モデルの例として Large generative AI models のみが挙げられていることから、現時点では生成 AI モデルのみが GPAI モデルに該当する AI モデルとして認識されていると考えら

れる。ただし、将来的に生成 AI モデル以外にも汎用性を有する AI モデルが登場する可能性があるため、「GPAI モデル」という解釈に幅のある用語が用いられていると考えられる。

なお、EU-AI 法においては、AI モデルのリスクのみならず、AI モデルを組み込んだ AI システムのリスクおよび AI プラクティスのリスクに応じて要件や義務が課される。このため、顔認識 AI モデルのような生成 AI モデル（GPAI モデル）に該当しない AI モデルに対しても、製品・サービス及びその利活用に関して規制が及ぶこととなる。

EU AI法の概要

AIモデル(技術)、AIシステム(製品・サービス)、AIプラクティス(用途)におけるリスクレベルに応じて、要件や義務が定められている



—— 電力消費量等は規制の対象にならないのか？

実積教授：

CoP の Transparency の項目に環境性能の記載が求められており、これは規制というより情報公開の対象となっている。情報公開の対象である以上、今後の展開として各国の公共調達においてこの基準が参考にされる可能性が高いと考えられる。もしそうなれば、法律による直接的な規制ではなく、大きなユーザーが一定の基準を満たさない AI を採用しないことで、財務面からの規制がかかる方向に進むと予想される。

吉永氏：

実積教授からの回答にある通り、CoP にて情報開示が求められているため、政府調達等における基準として活用される可能性がある。また、EU-AI 法とは別に、エネルギー効率の高いコンピュータの開発促進を目的とした「Cloud and AI Development Act」という新たな法律が検討されている¹²。

この法律においては、minimum energy-efficiency criteria の設定も検討されている¹³。

¹² <https://digital-strategy.ec.europa.eu/en/faqs/ai-continent-action-plan-qa>

¹³ https://www.linkedin.com/posts/luca-bertuzzi-186729130_the-providers-of-cloud-services-for-ai-face-activity-7288525730085285888-vYJz/

—— EU-AI 法においては、テキストを対象とした GPAI モデルが主なターゲットになっており、Transformer のようなアルゴリズムが前提とされている。しかし、画像生成 AI は拡散モデルなど異なるアルゴリズムを用いているため、GPAI モデル向けのガイドラインは必ずしも適合しないと考えられる。これに関してどう考えるべきなのか？

画像生成 AI にシステミックリスクがないとみなすのか、あるいは GPAI モデルに該当しないと判断するのかについて知りたい。

大野氏：

私も同意見である。画像生成はフェイクニュースなどへの影響が大きいため、単に画像だから対象外とされるべきではないと考える。おそらくこの点はまだ未整備であり、例えば画像を FLOP 数の中で扱う際に 1 枚の画像を何トークンと見なすかといった粗い計算方法は存在するものの、詳細な整備や議論は十分に進んでいない状況であり、今後議論が展開される可能性がある。

吉永氏：

EU-AI 法の前文 (99) および (105) に示されているとおり、テキストに限らず、画像・音声・ビデオの生成 AI モデルも同法の対象とされている。また、今回のガイドライン案においても、拡散モデル (image diffusion model) における FLOP 計算の例が Annex A.1 に示されている。

大野氏：

吉永氏と同意見である。計算量を用いている点には（それが本当に適したものかという議論はある一方で）、アーキテクチャ非依存的に GPAI モデルを定義したいという意図が隠されているのだろうと推察している。

—— 生成 AI の犯罪的なリスト等は EU が作るのか、それとも国際的な AI 警察のような機関によって作られるのか？

根本氏：

生成 AI の犯罪的なリストの整備については、そのような動きは聞いたことはない。CoP においても海賊版サイトの学習禁止は記載があるため、海賊版サイトの情報運用は検討されると推測される。

—— 量子コンピュータ等が使えるようになった際、この規制は変わるのか？

大野氏：

当局側の考えは現時点で不明であるが、量子コンピュータに限らず、技術開発（利用する GPU の性能、アルゴリズム、データセットの質の改善など）進むことで、現在の規制の意味合いが変わってくる可能性があるため、 10^{25} FLOP などの基準は定期的に見直しが必要になるだろうと考えている。

—— 学習データということだけで考えると、中国やインドのような人口が多いところのデータに偏ることはないのか？

吉永氏：

データの量の観点で、そのような偏りが生じる可能性があると考えられるが、AI の学習においてはデータの質も重要となる。実際、情報通信研究機構（NICT）においても、高品質な日本語データを抽出する取り組みが行われている¹⁴。

大野氏：

人口は、モデルの学習を行う組織の数というよりは、どちらかというとそのモデルの利用者に直接的には影響すると思われる。ただ、GPAI モデルに該当するような LLM を開発できる体力のある企業（大手企業、テックスタートアップ）や研究機関が一定数ある、それを国などが支援できる体制がある、一定の利用者の見込みがある（特に特定の地域の言語や文化を学んだモデルなど）などの点を考慮すると、ある程度経済力の強い国や地域で学習が集中的に行われる可能性はあるのではないかと考えられる。

—— まだ研究開発段階ではあるが、ロボット用のビジョン言語アクションモデルも同様の計算量で考えるべきなのか？

「扱う道具によってリスクが劇的に変わる」と考えると、道具や環境、シーンで条件を設計してもらう道があるのか？

吉永氏：

ビジョン言語アクションモデルは、LLM に画像入力とアクション（ロボット制御）出力の能力を付与したものであるため、GPAI モデルの定義には該当すると考えられる。また、扱う道具により変わるリスクについては、AI システム（ロボット製品）および AI プラクティス（ロボット製品の用途や利用目的）に対して要件や義務を課すことで対応することになるのではないかと考えられる。

¹⁴ NICT における国産 LLM 開発の取り組み:

https://www.soumu.go.jp/main_content/000997288.pdf

EU AI法の概要

AIモデル(技術)、AIシステム(製品・サービス)、AIプラクティス(用途)におけるリスクレベルに応じて、要件や義務が定められている



—— 計算量を累積するのであれば多段の履歴を調べる必要があるように思うが、サプライチェーンを監視するような仕組みは考えられているのか？

大野氏：（質問の前半部分に関して）

今回の working draft では、「... the AI Office's preliminary approach is to allow providers to use reasonable estimates when precise information is impractical to obtain」とあり、履歴が追えない、もしくは、現実的でない場合には、代替として、推定量を利用することが認められている。

—— ガイドライン p.5 で対象の例示があるが、「スピーチ専用」は汎用に使えるモデルの機能の一部を利用するだけで、「汎用」と見なすという解釈がされている。そうすると、「医学用」や「学術研究用」も一般向け機能の一部の利用として例外にならない可能性を感じた。

吉永氏：

「汎用に使えるものを特定用途に利用」とみなされるか、「特定用途にしか使えないものを特定用途に利用」（p.5 の Examples of models out of scope）とみなされるかで、例外と扱われるかが変わると考える。この基準には不明確な部分があるため、本イベントの話題提供においては「明確化が必要」と提起した。

大野氏：

吉永氏と同意見である。「医療 LLM」がどのような条件で汎用的に使えるとみなされるのかは、議論が必要だと考えている。例えば、モデル自身の能力だけではなく、利用規約などで提供者側が汎用的に使うことを意図的に禁止しているなどは観点として考えられる。

—— 吉永氏の話提供にあった、システムックリスク GPAI モデルの基準をやや下回る Meta の Llama3-

70B に関して、これと同等の精度ではあるが、FLOP 数が非常に少ないモデルもあるのではないか？

吉永氏：

同意見である。例えば Gemma2 27B の FLOP やパラメータ数は Llama 3-70B の数分の一だが、性能はほぼ同じであると、昨年モデルリリース時に話題となっていた¹⁵。

そのような状況を踏まえると、FLOP をシステミックリスク GPAI モデルの指標とすることには些か疑問が残る。

¹⁵ 【Gemma 2】 Llama 3 を超えた Google のオープン LLM: <https://weel.co.jp/media/tech/gemma-2/>

参考リンク集

今回のイベント関連

[パブリックコンサルテーション提出フォーム（締切：2025 年 5 月 22 日正午 CET）](#)

[【講演資料】ブリュッセル効果への対応：日本企業は EU-AI 法にどう備えるべきか 4](#)

[【動画】ブリュッセル効果への対応：日本企業は EU-AI 法にどう備えるべきか 4](#)

過去のイベント関連

[【実施報告】ブリュッセル効果への対応：日本企業は EU-AI 法にどう備えるべきか 3](#)

[【講演資料】ブリュッセル効果への対応：日本企業は EU-AI 法にどう備えるべきか 3](#)

[【動画】ブリュッセル効果への対応：日本企業は EU-AI 法にどう備えるべきか 3](#)

[【実施報告】ブリュッセル効果への対応：日本企業は EU-AI 法にどう備えるべきか 2](#)

[【講演資料】ブリュッセル効果への対応：日本企業は EU-AI 法にどう備えるべきか 2](#)

[【動画】ブリュッセル効果への対応：日本企業は EU-AI 法にどう備えるべきか 2](#)

[【実施報告】ブリュッセル効果への対応：日本企業は EU-AI 法にどう備えるべきか](#)

[【講演資料】ブリュッセル効果への対応：日本企業は EU-AI 法にどう備えるべきか](#)

[【動画】ブリュッセル効果への対応：日本企業は EU-AI 法にどう備えるべきか](#)

開催報告作成協力：狩野愛歌