



海外標準化動向調査(9月)

令和6年度エネルギー需給構造高度化基準認証推進事業費(我が国の国際標準化戦略を強化するための体制構築)

2024年9月1日

一般財団法人日本規格協会

ピックアップ：人工知能(AI) (関連ニュース番号14)



トピック	商務省、バイデン大統領のAIに関する大統領令を実施するための新たな措置を発表
推進組織	米国商務省、米国国立標準技術研究所（NIST）
内容	ポイント - 国立標準技術研究所（NIST）は、人工知能（AI）システムの安全性、セキュリティ、信頼性の向上に役立つことを目的とした4つの草案出版物をリリースした。NISTはまた、人間が作成したコンテンツとAIが作成したコンテンツを区別する方法の開発を支援するチャレンジシリーズも開始した。 背景 - 米国商務省は、バイデン大統領による安全でセキュリティが高く信頼できるAI開発に関する大統領令（EO）から180日が経過したのを受けて、EOに関連するいくつかの新たな発表を行った。 概要 - NIST の出版物は、AI テクノロジーのさまざまな側面をカバーしている。 - 生成AIのリスクの軽減する：NIST AI 600-1 - AIシステムの訓練に使用されるデータへの脅威を減らす：NIST SP 800-218 - 合成コンテンツリスクの低減する：NIST AI 100-4 - AI標準に関するグローバルな取り組み：NIST AI 100-5 - NISTの出版物に加えて、商務省の米国特許商標庁（USPTO）は、米国法の下で発明が特許を受けることができるかどうかを判断するために行われる技術における通常の技能レベルの評価にAIがどのように影響するかについてのフィードバックを求めるパブリックコメント要請（RFC）を公表しており、今年初めにはAI支援発明の特許性に関するガイダンスを公表している。 4つの文書に加え、NISTはジェネレーティブAI技術を評価・測定する新しいプログラム「NIST GenAI」も発表している。このプログラムは、大統領令に対するNISTの対応の一環であり、その取り組みは、NISTの米国AI安全研究所の作業に情報を提供するのに役立つ。生成AI技術の能力と限界を評価・測定するために設計された一連のチャレンジ問題を発行する。これらの評価は、情報の完全性を促進し、デジタルコンテンツの安全で責任ある利用を導くための戦略を特定するために使用される。このプログラムの目標のひとつは、与えられたテキスト、画像、ビデオ、オーディオ録音が、人間によって作られたのか、AIによって作られたのかを人々が判断できるようにすることである。パイロット評価への参加登録は5月に開始され、人間が制作したコンテンツと合成コンテンツとの違いを理解することを目指す。

出所: [米国商務省](#)、[NIST](#)等の情報に基づきJSAグループ作成

ピックアップ：人工知能(AI) (関連ニュース番号25)



トピック	欧州委員会、安全で信頼できる人工知能におけるEUのリーダーシップを強化するため、AIオフィスを設置
推進組織	欧州委員会
内容	ポイント - 欧州AIオフィスは、AIのリスクから保護しながら、信頼できるAIの開発と使用をサポートする。AIオフィスは、AIの専門知識の中心として欧州委員会内に設立され、単一の欧州AIガバナンス システムの基盤を形成する。 背景 2024年1月、欧州委員会は、信頼できる人工知能（AI）の開発において欧州の新興企業と中小企業を支援するための一連の措置を開始。これらの措置の一環として、AIオフィスを設立するという欧州委員会の決定が採択された。 概要 - AIオフィス、リスクを軽減しつつ、社会的・経済的利益とイノベーションを促進する形で、AIの将来の開発、展開、利用を可能にすることを目的としている。同オフィスは、特に汎用AIモデルに関して、AI法の実施において重要な役割を果たす。また、信頼できるAIの研究と技術革新を促進し、EUを国際的な議論のリーダーとして位置づけることにも取り組む。 - AIオフィスは、AIに対するEUのアプローチを支援する独自の機能を備えている。加盟国のガバナンス機関を支援することで、AI法の実施において重要な役割を果たしている。また、汎用AIモデルに関する規則を施行する。これは、AI法が欧州委員会に与えた権限（汎用AIモデルの評価、モデル提供者への情報提供や措置の要求、制裁措置の適用など）に支えられている。AI室はまた、社会的・経済的利益を享受するため、信頼できるAIの革新的エコシステムを促進する。AIオフィスは、AIに関する欧州の戦略的、首尾一貫した効果的なアプローチを国際レベルで確保し、世界的な参照点となる。 - 組織レベルでは、AIオフィスは、加盟国の代表者によって構成される欧州人工知能委員会および欧州委員会の欧州アルゴリズム透明性センター(ECAT) と緊密に連携している。 AIオフィスの構造

出所: [欧州委員会](#)等の情報に基づきJSAグループ作成

【人工知能(AI)】関連記事詳細（1/35）

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
1	国際	ISO/IEC JTC 1/SC 42(Artificial intelligence)における規格開発状況	2024/7/26	ISO/IEC JTC 1/SC 42(人工知能)は、活動範囲を「人工知能に関するJTC 1の標準化プログラム」の中心となり、推進役を務めること及び人工知能アプリケーションを開発するJTC 1、IEC、ISO委員会にガイダンスを提供すること」として、人工知能分野の標準化を行っている。 SCの幹事国はアメリカ(ANSI)。国内審議団体は(一社)情報処理学会内の情報規格調査会。2024年7月26日現在、発行済みの有効な規格は31規格。 2024年1月以降に発行された規格は、次の10規格である。 • ISO/IEC TR 5469 機能安全と AI システム（1/8発行） ➢ 機能性を実現するために安全関連機能内でのAI使用、AI制御機器の安全性を確保するためにAI以外の安全関連機能の使用、安全関連機能を設計・開発するためのAIシステム使用に関連するリスク要因やプロセスについて規定。 • ISO/IEC 5339 AI応用のガイダンス（1/15発行） ➢ AIアプリケーションに関するガイダンスを提供し、利害関係者の関与とAIアプリケーションのライフサイクルに重点く。AI システムの作成、使用、影響の観点を含むフレームワークを提供することで、複数の利害関係者間のコミュニケーションと受け入れを強化することを目的としている。 • ISO/IEC TS 25058 システムおよびソフトウェアの品質要件と評価 (SQuaRE) - AIシステムの品質評価のガイダンス（1/24発行） ➢ AIシステム品質モデルを使用したAIシステムの評価に関するガイダンスを提供。AIの開発と使用に携わるあらゆる種類の組織に適用される。 • ISO/IEC 5392 知識工学のリファレンスアーキテクチャ（3/15発行） ➢ AIにおける知識工学 (KE) のリファレンス アーキテクチャを定義。リファレンス アーキテクチャでは、体系的なユーザーおよび機能の観点から、KEの役割、アクティビティ、構成レイヤー、コンポーネント、およびそれらの相互関係と他のシステムとの関係について説明する。	International Organization for Standardization (ISO) https://www.iso.org/committees/6794475/x/catalogue/p/1/uo/w/0/d/0

【人工知能(AI)】関連記事詳細（2/35）

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
1 (続き)	国際	ISO/IEC JTC 1/SC 42(Artificial intelligence)における規格開発状況	2024/7/26	2024年1月以降に発行された規格。(前スライドの続き) • ISO/IEC TR 24030 ユースケース (4/4発行) ➢ 様々な分野にわたるAIの使用事例をまとめた包括的な文書。幅広いアプリケーションを網羅し、さまざまな分野におけるAIの適用性と可能性を示す。 • ISO/IEC 8200 自動化された人工知能システムの制御可能性 (4/10発行) ➢ 自動化された人工知能 (AI) システムの制御可能性の実現と強化のための原則、特性、アプローチを含む基本的なフレームワークを指定。 • ISO/IEC TR 17903 機械学習コンピューティングデバイスの概要 (5/6発行) ➢ MLコンピューティングデバイスの用語と特性や、MLコンピューティングデバイスのパフォーマンスを最適化するための特性の設定と使用に関する既存のアプローチについて説明。 • ISO/IEC 5259-1 分析と機械学習 (ML) のためのデータ品質 - パート 1: 概要、用語、例 (7/2発行) ➢ 分析とMLのデータ品質に焦点を当てたISO/IEC 5259 シリーズの基礎部分。信頼性の高い分析とMLの結果に不可欠な、データ ライフ サイクルのさまざまなフェーズにわたってデータ品質を評価および強化するためのフレームワークを確立。 • ISO/IEC 5259-3 分析と機械学習 (ML) のためのデータ品質 - パート3: データ品質管理の要件とガイドライン (7/2発行) ➢ 分析とMLの分野で使用されるデータの品質を確立、実装、維持、継続的に改善するための要件を指定し、ガイダンスを提供。 • ISO/IEC 5259-4 分析と機械学習 (ML) のためのデータ品質 - パート4: データ品質プロセスフレームワーク (7/15発行) ➢ 分析とMLのトレーニングと評価のためのデータ品質を確保するために、適用する組織の種類、規模、性質に関係なく、一般的な組織的アプローチを確立。	International Organization for Standardization (ISO) https://www.iso.org/committees/6794475/x/catalogue/p/1/uo/w/0/d/0

【人工知能(AI)】関連記事詳細（3/35）

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 （機関・団体名／URL）	
2	アメリカ	NIST、AIシステムの動作を操作するサイバー攻撃の種類を特定	2024/1/4	敵対者は人工知能（AI）システムを意図的に混乱させたり、「毒」を盛ったりして誤動作させることができる。米国立標準技術研究所（NIST）のコンピューター科学者とその共同研究者たちは、AIと機械学習（ML）のこうした脆弱性とその他の脆弱性を、新しい論文の中で明らかにしている。 そのタイトルは「Adversarial Machine Learning（敵対的機械学習）」である： A Taxonomy and Terminology of Attacks and Mitigations (NIST.AI.100-2)と題されたこの出版物は、信頼できるAIの開発を支援するためのNISTの幅広い取り組みの一環であり、NISTのAIリスクマネジメントフレームワークの実践に役立つものである。政府、学术界、産業界の協力により作成されたこの出版物は、AIの開発者やユーザーが、銀の弾丸は存在しないことを理解した上で、想定される攻撃の種類とそれを軽減するためのアプローチを把握できるようにすることを目的としている。 このレポートでは、回避攻撃、ポイズニング攻撃、プライバシー攻撃、不正使用攻撃という 4 つの主要な攻撃タイプについて検討している。また、攻撃者の目的や目標、能力、知識など、複数の基準に従って攻撃を分類している。	NIST	https://www.nist.gov/news-events/news/2024/01/nist-identifies-types-cyberattacks-manipulate-behavior-ai-systems
3	中国	「データ要素×」3か年行動計画（2024年～2026年）の発行に関する国家データ管理局およびその他の部門からの通知	2024/1/5	近年、中国のデジタル経済の急速な発展は、デジタルインフラ規模の容量が大幅にジャンプしている、デジタル技術と産業システムは、より良いデータ要素の役割を果たすために、より成熟している強固な基盤を築いた。同時に、データ供給の質の低さ、流通メカニズムの貧弱さ、応用可能性のリリース不足などの問題もある。 データ要素に様々なテーマを掛け合わせる行動の実施は、中国の超大型市場、膨大なデータ資源、豊富な応用シナリオなどの複数の利点をフルに発揮し、データ要素と労働や資本などの要素とのシナジーを促進することであり、データの流れが技術、資本、人材、材料の流れをリードし、資源要素に対する従来の制約を打破し、全要素生産性を向上させることができるようにすることである。 多様化するデータの統合を加速し、データ規模の拡大とデータの種類の充実により、生産ツールの革新とアップグレードを促進し、新産業と新モードを生み出し、経済発展の新たな運動エネルギーを開拓する。	中央サイバーセキュリティ情報化委員会事務局	https://www.cac.gov.cn/2024-01/05/c_1706119078060945.htm

【人工知能(AI)】関連記事詳細（4/35）

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
4	アメリカ	SAE J3312 (人工知能 - 地上車両アプリケーションにおけるユースケース)の開発	2024/1/9	現在、SAEではJ3312を開発しており、地上車両および交通インフラへのAI技術応用のユースケースを特定している。SAE地上車両人工知能（GVAI）委員会の他のタスクフォースからのインプットと組み合わせられた後、ベストプラクティスまたは標準となり得る情報報告書が発行されることが期待される。 該当する場合はいつでも、機能的な定義や指摘された問題や懸念事項が、現在の業界のモビリティプラクティスや公開された査読付き文献と整合する形で提供される。このタスクフォースは、AI 委員会を通じて、必要に応じて他の SAE または標準（ISO、IEEE など）グループと連携する。	SAE https://www.sae.org/standards/content/j3312/
5	国際	AIは世界経済を変革する。人類に利益をもたらすようにしよう	2024/1/14	新しい分析では、IMFスタッフがAIが世界の労働市場に与える潜在的な影響を検証している。多くの研究が、AIによって仕事が代替される可能性を予測している。しかし、多くの場合、AIが人間の仕事を補完する可能性が高いことも分かっている。IMFの分析は、この両方の力を捉えている。 世界の雇用のほぼ40%がAIにさらされている。歴史的に、自動化と情報テクノロジーは定型的な仕事に影響を与える傾向があったが、AIが他と異なる点のひとつは、高技能職に影響を与える能力である。その結果、先進国は新興市場や発展途上国に比べ、AIによるリスクが大きい、その恩恵を活用する機会も多い。 先進国では、約60%の雇用がAIの影響を受ける可能性がある。影響を受ける仕事の約半分は、AIの統合による生産性向上の恩恵を受ける可能性がある。残りの半分については、AIアプリケーションが現在人間が行っている主要なタスクを実行することで、労働需要が低下し、賃金の低下や雇用の減少につながる可能性がある。最も極端なケースでは、これらの仕事の一部が消滅する可能性もある。 一方、新興市場と低所得国では、AIの普及率はそれぞれ40%と26%になると予想される。これらの調査結果は、新興市場や発展途上国の経済がAIによる直接的な混乱に直面する可能性が低いことを示唆している。同時に、これらの国々の多くは、AIの恩恵を享受するためのインフラや熟練した労働力を有していないため、この技術が長期的に国家間の不平等を悪化させるリスクがある。	IMF https://www.imf.org/en/Blogs/Articles/2024/01/14/ai-will-transform-the-global-economy-lets-make-sure-it-benefits-humanity

【人工知能(AI)】関連記事詳細（5/35）

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
6	オーストラリア	オーストラリア政府の安全で責任あるAIに関する協議に対する暫定的な対応	2024/1/17	オーストラリア政府は、2023年に開催された安全で責任あるAIに関する協議に対する中間回答を発表した。 この中間回答は、安全で責任あるAIについて、オーストラリア国民、学界、企業から寄せられた意見をまとめたものである。また、政府が現在、そして長期的にどのような行動をとるかについても詳述している。 協議では、AIシステムとアプリケーションが、ウェルビーイングと生活の質を向上させ、経済を成長させるのに役立っていることが明らかになった。しかし、現在の規制の枠組みでは、AIのリスクに十分に対処できていない。 提出された意見では、AIの合法的だがリスクの高い用途や、強力な「フロンティア」モデルによる予期せぬリスクについて、さらなるガードレールを設けることが求められている。	オーストラリア産業科学資源省 https://www.industry.gov.au/news/australian-governments-interim-response-safe-and-responsible-ai-consultation

【人工知能(AI)】関連記事詳細（6/35）

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
7	欧州	欧州委員会、人工知能の新興企業と中小企業を支援するAIイノベーションパッケージを開始	2024/1/24	欧州委員会は、EUの価値観と規則を尊重した信頼できるAIの開発において、欧州の新興企業や中小企業を支援するための施策パッケージを発表した。これは、2023年12月にEU AI法について政治的合意に達したことを受けたもので、EUにおける信頼できるAIの開発、導入、普及を支援するものである。フォン・デル・ライエン欧州委員会委員長は、2023年の一般教書演説で、欧州の革新的なAI新興企業が信頼できるAIモデルを訓練するために、欧州のスーパーコンピューターを利用できるようにする新たな構想を発表した。このパッケージは、AIスタートアップやより広範なイノベーション・コミュニティにスーパーコンピューターへの特権的なアクセスを提供する提案を含む、AIスタートアップやイノベーションを支援するための広範な措置を通じて、このコミットメントを実践するものである。その内容は以下の通りである： - EU のスーパーコンピューター共同事業活動の新たな柱となる AI ファクトリーを設立するための EuroHPC 規則の改正 - 欧州委員会内にAI オフィスを設立する決定 - 追加の主要な活動を概説したEU AI スタートアップおよびイノベーションコミュニケーション 欧州委員会はまた、加盟国とともに、2つの欧州デジタル基盤コンソーシアム(EDIC)を設立する： - 言語技術のための同盟」(ALT-EDIC) は、AIソリューションの訓練に必要な欧州言語のデータ不足に対処するとともに、欧州の言語多様性と文化的豊かさを維持するため、言語技術における欧州共通のインフラストラクチャーを開発することを目的としている。これにより、欧州の大規模言語モデルの開発が支援される。 - CitiVERSE」EDICは、最先端のAIツールを応用して、スマートコミュニティのためのローカル・デジタル・ツインを開発・強化し、交通管理から廃棄物管理まで、都市のプロセスのシミュレーションと最適化を支援する。 欧州委員会は、人工知能の利用に関する欧州委員会独自の戦略的アプローチを概説するコミュニケーション（AI@EC Communication）も採択した。この戦略的ビジョンにより、欧州委員会は、EU AI法の施行に向けた準備を社内を進めている。その中には、信頼に足る、安全で倫理的な人工知能の開発と利用を確保するために、欧州委員会がどのように制度上および運用上の能力を構築していくかについての具体的な行動も含まれている。	欧州委員会 https://ec.europa.eu/commission/presscorner/detail/en/ip_24_383

【人工知能(AI)】関連記事詳細（7/35）

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 （機関・団体名／URL）
8	イタリア	ChatGPT：個人情報保護団体、OpenAIに個人情報保護規定違反の異議申し立て行為を通告	2024/1/29	個人情報保護団体は、人工知能プラットフォーム「ChatGPT」を運営するOpenAI社に対し、個人情報保護規則違反による異議告知を通知した。 昨年3月30日にGaranteが同社に対して行った暫定的な処理制限命令を受け、実施された予備調査の結果、当局は、取得した要素がEU規則の規定に関して1つ以上の違法行為に該当する可能性がある判断した。 OpenAI社は、30日以内に違反の疑いに関する答弁書を提出しなければならない。 手続きを定めるにあたり、団体は、EUデータ保護当局（Edpb）を集めた理事会が設置した特別タスクフォースの進行中の作業を考慮する。	イタリアデータ保護機関 https://www.garanteprivacy.it/home/docweb/-/docweb-display/docweb/9978020
9	欧州/アメリカ	EUと米国、貿易と技術協力を評価	2024/1/30	欧州連合と米国はワシントンDCでEU・米国貿易技術協議会（TTC）の第5回会合を開催した。この会合では、閣僚らがTTCの活動の進捗状況を評価し、春にベルギーで開催される次回のTTC閣僚会合の主要優先事項について政治的な指針を示した。TTCは大西洋横断貿易と技術問題に関する緊密な協力のための主要フォーラムである。 参加者は、二国間の貿易と投資を引き続き拡大し、経済安全保障と新興技術で協力し、デジタル環境における共通の利益を推進するという強い共通の願望を示した。このTTC会議の傍らで、双方は、適合性評価に関する協力を強化することを含め、グリーン移行に不可欠な物品と技術の貿易を促進する方法を引き続き模索することで合意した。 前回のTTC閣僚会議での約束に続き、EUと米国はG7で採択された人工知能（AI）に関する国際指導原則とAI開発者向け自主行動規範を歓迎し、国際的なAIガバナンスで協力し続けることに合意した。両者はまた、この重要な技術を開発するための指導原則と次のステップを示す6Gの業界ロードマップを歓迎した。双方は、次回のTTC閣僚会合を、理事会議長国ベルギーの主催により、春にベルギーで開催することに合意した。	欧州委員会 https://ec.europa.eu/commission/presscorner/detail/en/IP_24_575

【人工知能(AI)】関連記事詳細 (8/35)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
10	欧州/ カナダ	EUとカナダ、新たな課題に対処するため戦略的デジタルパートナーシップを強化	2024/2/1	ティエリー・ブルトン欧州委員会域内市場担当委員とカナダのフランソワ・フィリップ・シャンパーニュ革新・科学産業大臣は2月1日に会談し、11月23日と24日にカナダで開催された首脳会談で締結されたEU・カナダデジタルパートナーシップの実施に向けた作業を開始した。 新しいデジタルパートナーシップは、研究、産業、社会、そして経済全体に影響を及ぼすデジタル変革における新たな課題にEUとカナダが取り組むことを支援する。 人工知能（AI）、量子科学、半導体、オンラインプラットフォームに関する公共政策、安全な国際接続、サイバーセキュリティに関する協力の強化に重点を置くことを目指している。これらの優先事項は、2月にデジタル対話を通じて当局者レベルで議論される予定。また、EUとカナダは、AIガバナンスや国際基準などに関するワークショップを通じて、定期的なコミュニケーションチャネルを構築し、情報交換を行う予定。 両者は、本日の議論に沿って進捗状況を評価し、次のステップを決定するため、春に閣僚級のデジタルパートナーシップ協議会を開催することで合意した。	欧州委員会 https://ec.europa.eu/commission/presscorner/detail/en/IP_24_614
11	アメリカ	NISTのAIリスクマネジメントフレームワークを活用し、人工知能のリスクを認識する文化を形成する	2024/2/2	米国国立標準技術研究所（NIST）が作成した人工知能リスクマネジメントフレームワーク（AI RMF）は、組織がAI技術が個人、地域社会、社会、環境にもたらすリスクを管理するための自主的なリソースである。 米国商務省の一機関であるNISTは、2021年の議会指令に基づき、米国および世界の民間企業、市民社会、学界、政府機関と協力し、オープンで透明性が高く、コンセンサスに基づくプロセスでAI RMFを策定した。AI RMFは、チェックリストを提供するのではなく、組織文化の変革を促すことを意図しており、AI技術のリスク管理を目指す機関の幅広いエコシステム全体で実施することができる。	NIST https://www.nist.gov/publications/informing-artificial-intelligence-risk-aware-culture-nist-ai-risk-management-framework

【人工知能(AI)】関連記事詳細 (9/35)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
12	アメリカ	バイデン・ハリス政権、AIの安全性に特化した初のコンソーシアムを発表	2024/2/8	ジーナ・ライモンド米商務長官は本日、米国AI安全研究所コンソーシアム（AISIC）の設立を発表した。同コンソーシアムは、安全で信頼できる人工知能（AI）の開発と普及を支援するため、AIの開発者と利用者、学者、政府・産業界の研究者、市民団体を結びつける。 このコンソーシアムは、米国AI安全研究所（USAISI）の下に置かれ、レッドチーム、能力評価、リスク管理、安全・セキュリティ、合成コンテンツの電子透かしに関するガイドラインの策定など、バイデン大統領の画期的な大統領令で示された優先行動に貢献する。このコンソーシアムには、最先端のAIシステムやハードウェアを開発・使用する最前線にいる200以上の企業や組織、国内最大手企業や最も革新的な新興企業、AIがどのように我々の社会を変革しうるかについての基礎的な理解を構築している市民団体や学術チーム、そして今日のAIの使用に深く関与している専門職の代表者が参加している。 このコンソーシアムは、これまでに設立されたテスト・評価チームの最大の集合体であり、AIの安全性における新たな測定科学の基盤を確立することに重点を置く。コンソーシアムには、州政府や地方自治体、非営利団体も参加し、世界中の安全に関する相互運用可能で効果的なツールの開発において重要な役割を担う、志を同じくする国々の組織とも協力する。	NIST https://www.nist.gov/news-events/news/2024/02/biden-harris-administration-announces-first-ever-consortium-dedicated-ai
13	中国	TC260-003「生成型人工知能サービスの基本セキュリティ要件」がリリース	2024/3/1	リリースされた「生成型人工知能サービスの基本セキュリティ要件」では、適用範囲として、「本文書は、コーパスの安全性、モデルの安全性、安全対策など、生成AIサービスの安全性に関する基本的な要件を規定し、安全性評価に関する要件を与える。本書は、サービス提供者が安全性評価を実施し、安全レベルを向上させるために適用可能であり、また、関係当局が生成AIサービスの安全レベルを判断する際の参考資料となる。」と記載され、総則と5つの観点からの要求事項が規定されている。	国家サイバーセキュリティ標準化専門委員会事務局 https://www.tc260.org.cn/front/postDetail.html?id=20240301164054

【人工知能(AI)】関連記事詳細（10/35）

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
14	欧州	AI法：議員が画期的な法律を採択	2024/3/13	欧州議会は、安全性と基本的権利の遵守を確保し、イノベーションを促進する人工知能法を承認した。2023年12月に加盟国との交渉で合意されたこの規則は、賛成523票、反対46票、棄権49票で欧州議会により承認された。 同規則は、基本的権利、民主主義、法の支配、環境の持続可能性をリスクの高いAIから保護する一方、技術革新を促進し、欧州をこの分野のリーダーとして確立することを目的としている。同規則は、AIの潜在的なリスクと影響の度合いに基づいて、AIに対する義務を定めている。 新規則は、市民の権利を脅かす特定のAIアプリケーションを禁止している。これには、敏感な特徴に基づくバイOMETRICS分類システムや、顔認識データベースを作成するためのインターネットやCCTV映像からの顔画像の非標的なスクレイピングなどが含まれる。職場や学校での感情認識、ソーシャルスコアリング、予測的取り締まり（人物のプロファイリングや特性の評価のみに基づく場合）、人間の行動を操作したり、人々の脆弱性を悪用したりするAIも禁止される。 この規則は弁護士と言語学者による最終チェックを受けており、立法府が終了する前に最終的に採択される見込みである（いわゆる正誤表手続きによる）。また、同法は理事会の正式承認も必要となる。発効は官報公布の20日後となり、発効から24ヵ月後に完全に適用される。ただし、禁止行為の禁止（発効から6ヵ月後に適用）、実践規範（発効から9ヵ月後）、ガバナンスを含むAI汎用規則（発効から12ヵ月後）、ハイスクスシステムに関する義務（36ヵ月後）は例外となる。	欧州議会 https://www.europarl.europa.eu/news/en/press-room/20240308IPR19015/artificial-intelligence-act-meps-adopt-landmark-law
15	アメリカ	国土安全保障省が人工知能ロードマップを発表、技術のメリットを最大化し国土安全保障ミッションを推進するためのパイロットプロジェクトを発表	2024/3/18	アレハンドロ・N・マヨルカス国土安全保障長官とエリック・ハイセン最高情報責任者兼最高人工知能責任者は、国土安全保障省（DHS）初の「人工知能ロードマップ」を発表した。このロードマップには、個人のプライバシー、市民権、市民的自由を確実に保護しつつ、米国民に有意義な利益をもたらす、国土安全保障を前進させる技術の利用をテストすることを含む、DHSの2024年計画が詳述されている。 ロードマップの一環として、DHSは特定の任務分野にAIを導入する3つの革新的なパイロットプロジェクトを発表した。国土安全保障省捜査局（HSI）は、フェンタニルの検出と児童の性的搾取対策に関連する捜査の効率化に焦点を当てた捜査プロセスを強化するため、AIを試験的に導入する。連邦緊急事態管理庁（FEMA）は、地域社会がレジリエンスを構築し、リスクを最小化するための災害軽減計画を立案するのを支援するためにAIを導入する。また、米国民権・移民局（USCIS）は、AIを使って移民局員の訓練を改善する。	米国国土安全保障省 https://www.dhs.gov/news/2024/03/18/department-homeland-security-unveils-artificial-intelligence-roadmap-announces

【人工知能(AI)】関連記事詳細（11/35）

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
16	国際	国連総会、人工知能に関する画期的な決議を採択	2024/3/21	国連総会は、すべての人々の持続可能な開発にも役立つ「安全で安心かつ信頼できる」人工知能（AI）システムの推進に関する画期的な決議を採択した。 米国主導の決議案を無投票で採択した総会は、AIの設計、開発、配備、使用における人権の尊重、保護、促進にも焦点を当てた。この文章は、他の120以上の加盟国によって「共同提案」または支持された。 総会はまた、17の持続可能な開発目標の達成に向けた進展を加速させ、可能にするAIシステムの可能性を認めた。 総会がこの新興分野の規制に関する決議を採択したのはこれが初めてである。米国の国家安全保障顧問は今月初め、この採択はAIの安全な利用にとって「歴史的な前進」を意味すると述べたと報じられている。	国際連合 https://news.un.org/en/story/2024/03/1147831
17	欧州/韓国	EUと韓国は包括的かつ強靱なデジタル変革に向けたパートナーシップを再確認	2024/3/26	EUと韓国はブリュッセルで第2回デジタル・パートナーシップ協議会を開催した。閣僚会合では、EUと韓国は、市民と経済の利益のために主要なデジタル技術で協力するという約束を再確認した。EUと韓国は、第1回デジタル・パートナーシップ協議会以降の進捗状況を確認し、さらなる協力のための主要分野のリストについて合意した。 EUと韓国は、半導体、5Gおよびそれ以降、量子技術、プラットフォーム、AI、サイバーセキュリティに関する協力を継続することに合意し、ネットワーク接続などの他の協力分野についても定めた。これらの共同プロジェクトの実施を促進するため、欧州委員会と韓国は昨日、EUの研究・技術革新のための資金援助プログラムであるホライゾンヨーロッパに韓国を加盟させるための交渉を終えた。EUと韓国は、サイバーセキュリティの動向、中小企業のデジタル化に関するベストプラクティスに関する情報を共有し、共通のビジョンを推進するために国際的な場でICT標準化について協力するための措置を講じた。	欧州委員会 https://ec.europa.eu/commission/presscorner/detail/en/ip_24_1708

【人工知能(AI)】関連記事詳細（12/35）

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
18	アメリカ	ファクトシート: ハリス副大統領が連邦政府機関の人工知能の利用におけるガバナンス、イノベーション、リスク管理を推進するためのOMB 政策を発表	2024/3/28	カマラ・ハリス副大統領は、ホワイトハウスの行政管理予算局（OMB）が、人工知能（AI）のリスクを軽減し、その利点を活用するためのOMB初の政府全体の方針を発表することを発表した。 この大統領令は、AIの安全性とセキュリティの強化、米国人のプライバシーの保護、公平性と公民権の向上、消費者と労働者の支援、イノベーションと競争の促進、世界における米国のリーダーシップの推進など、多岐にわたる行動を指示している。連邦政府機関は、大統領令によって課せられた150日間の行動をすべて完了したと報告している。 主な内容として、下記が挙げられている。 • AI使用によるリスク対処 • AI利用の透明性拡大 • 責任あるAIイノベーション推進 • AI人材育成 • AIガバナンス強化	THE WHITE HOUSE https://www.whitehouse.gov/briefing-room/statements-releases/2024/03/28/fact-sheet-vice-president-harris-announces-omb-policy-to-advance-governance-innovation-and-risk-management-in-federal-agencies-use-of-artificial-intelligence/

【人工知能(AI)】関連記事詳細（13/35）

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
19	アメリカ	NTIA が AI 説明責任政策レポートを公開	2024/3/29	米国商務省電気通信情報局（NTIA）は、人工知能（AI）の説明責任に関する政策報告書を発表し、安全、安心、信頼できるAIのイノベーションを支援するための政策提言を行った。 この報告書は、NTIAのAI説明責任政策意見募集に回答した関係者からの1,400件を超える意見を考慮して作成された。同報告書では、政府がAIシステムに対する指針、支援、規制を打ち出すよう求めており、AIシステムに対する透明性の向上、これらのシステムに関する主張を検証するための独立した評価、許容できないリスクを課したり根拠のない主張をしたりした場合の結果などを求めている。 報告書は基準の役割について掘り下げており、国際的な技術基準は極めて重要であり、ある種の監査の手法を定義するために必要かもしれないと指摘している。未整備の基準は、「コンプライアンスを求める企業にとって不確実性を意味し、監査の有用性を低下させ、顧客、政府、国民に対する保証を低下させる」。AI用語を使用する分野は多岐にわたり、その多くが独自の用途、リスク、用語を持っているため、AI標準の開発には困難が伴う。 報告書は、国際標準化作業を加速させ、技術標準や標準設定プロセスへの参加を拡大する必要性を指摘している。同報告書は、政府が説明責任を果たすために標準の有用性を促進するためには、次のような方法があると提言している： 1. 市民社会、非産業界参加者、AIシステムによって不本意に影響を受ける人々を含む多様な利害関係者の参加を奨励し、育成する 2. 従来は十分に代表されていなかった関係者による標準出版物へのアクセスの改善と拡大を支援する 3. 業界標準を社会的価値と整合させる方法を支援する 4. 適切な状況においては、標準策定に寄与するガイドラインやその他のリソースを開発する	ANSI https://www.ansi.org/standards-news/all-news/2024/03/3-29-24-ntia-publishes-ai-accountability-policy-report

【人工知能(AI)】関連記事詳細（14/35）

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
20	アメリカ	ユタ州、米国初のAIに焦点を当てた消費者保護法案を制定	2024/4/1	ユタ州知事コックス氏は、2024年3月13日に、民間部門のAI利用を規制する主要な人工知能法（SB 149）に署名した。こうした法制定はアメリカで初めての州となる。 主な規定内容は以下の通り。 ・人工知能（AI）の使用が適切に開示されていない場合、消費者保護法に違反した場合の責任を定める ・人工知能政策室（Office of Artificial Intelligence Policy）および規制AI分析プログラムの創設 ・人工知能（AI）試験運用中の規制への影響を一時的に緩和することを可能にする ・技術、リスク、政策を評価する人工知能学習ラボラトリー・プログラムを設立する。 ・個人が規制対象の職業でAIと相互作用する場合、情報開示を義務付ける。 ・AIプログラムおよび規制適用除外に関する規則制定権限を事務局に付与する。	ユタ州議会 https://le.utah.gov/~2024/bills/static/SB0149.html
21	イギリス / アメリカ	英国と米国、AIの安全性に関する科学で提携を発表	2024/4/2	英国と米国は2024年4月1日、昨年11月に開催されたAI安全サミットでのコミットメントを受け、最先端の人工知能（AI）モデルのテスト開発に向けて協力する覚書に調印した。 ミシェル・ドネラン技術長官とジーナ・ライモンド米商務長官が署名したこのパートナーシップは、両国の科学的アプローチを一致させ、AIモデル、システム、エージェントのための堅牢な評価スイートを加速し、迅速に反復するために緊密に協力する。 英国と米国のAI安全性研究所は、AIの安全性テストに共通のアプローチを構築し、これらのリスクに効果的に取り組めるよう能力を共有する計画を打ち出した。両機関は、一般にアクセス可能なモデルで少なくとも1回の合同テストを実施する意向だ。また、両研究所間の人材交流を模索することで、専門知識の集合体を活用する意向だ。 このパートナーシップは直ちに発効し、両機関がシームレスに連携できるようにすることを目的としている。AIは急速な発展を続けており、両政府は、この技術の新たなリスクに対応できるAIの安全性に対する共通のアプローチを確保するために、今行動する必要性を認識している。両国はAIの安全性に関するパートナーシップを強化するとともに、世界全体でAIの安全性を促進するため、他の国とも同様のパートナーシップを構築することを約束した。 ※本件に関する米国商務省からのリリースはこちら	GOV.UK https://www.gov.uk/government/news/uk-united-states-announce-partnership-on-science-of-ai-safety

【人工知能(AI)】関連記事詳細（15/35）

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
22	欧州/ アメリカ	EU・米国貿易 技術理事会の 共同声明	2024/4/5	貿易技術理事会（TTC）の第6回閣僚会合は、2024年4月4日と5日にベルギーのルーヴェンで開催された。人工知能に関する声明は以下の通り。 欧州連合（EU）と米国は、人工知能（AI）に対するリスクベースのアプローチと、安全、安心、信頼できるAI技術の推進に対するコミットメントを再確認する。また、AIの責任ある管理を進めるため、G7、OECD、G20、欧州評議会、国連プロセスなどの多国間イニシアティブへの貢献を引き続き調整していく。我々は、米国と欧州の先進的AI開発者に対し、それぞれのガバナンスと規制システムを補完する「先進的AIシステムを開発する組織のための広島プロセス国際行動規範」の適用を促進するよう奨励する。 AIに関する継続的かつ影響力のある協力を確保するため、欧州AI事務局と米国AI安全研究所の指導者は、それぞれのアプローチと任務について互いに説明した。両機関は本日、特にベンチマーク、潜在的风险、将来の技術動向などのトピックについて、それぞれの科学機関や関連機関の間で科学的な情報交換を促進するため、協力関係を深めるためのダイアログを設置することを約束した。 この協力は、信頼できるAIとリスク管理のための評価・測定ツールに関する共同ロードマップの実施を進展させることに資するものであり、これは、それぞれの新たなAIガバナンスと規制制度における乖離を適切に最小化し、相互運用可能な国際標準について協力するために不可欠なものである。ステークホルダーとの協議の結果、我々は、相互に受け入れ可能な共同定義を持つ主要AI用語のリストをさらに作成し、更新版を公表した。 我々は、英国、カナダ、ドイツのパートナーとともに、AI for Developmentドナー・パートナーシップにおいて、教育者、起業家、一般市民がAIの可能性を活用できるよう支援するため、アフリカにおける我々の対外支援を加速させ、連携させる機会を引き続き探っていく。	欧州委員会 https://ec.europa.eu/commission/presscorner/detail/en/state_ment_24_1828

【人工知能(AI)】関連記事詳細（16/35）

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
23	アメリカ	国家安全保障局、AIシステムのセキュリティ強化に関するガイドラインを公開	2024/4/15	米国家安全保障局（NSA）は本日、サイバーセキュリティ情報シート（CSI）「AIシステムを安全に展開する」を発表する：安全で弾力性のある AI システムを展開するためのベストプラクティス "を発表した。このCSIは、外部団体によって設計・開発されたAIシステムを導入・運用する国家安全保障システム所有者や防衛産業基盤企業を支援することを目的としている。 CSIは、NSAの人工知能セキュリティ・センター（AISC）が、サイバーセキュリティ・インフラストラクチャ・セキュリティ局（CISA）、連邦捜査局（FBI）、オーストラリア信号総局のオーストラリア・サイバーセキュリティ・センター（ACSC）、カナダ・サイバーセキュリティ・センター、ニュージーランド国家サイバーセキュリティ・センター（NCSC-NZ）、英国国家サイバーセキュリティ・センター（NCSC-UK）と連携して初めて発表したもの。 このガイダンスは、国家安全保障を目的としているが、AI機能を管理環境に導入するすべての人、特に脅威が高く、価値の高い環境にある人に適用できるものである。このガイダンスは、以前に発表された「安全なAIシステム開発と人工知能の活用のためのガイドライン」をベースにしている。 これは人工知能セキュリティ・センター（AISC）が主導する最初のガイダンスであり、同センターの中心目標の1つであるAIシステムの機密性、完全性、可用性の向上をサポートする態勢を整えている。	国家安全保障局 https://www.nsa.gov/Press-Room/Press-Releases-Statements/Press-Release-View/Article/3741371/nsa-publishes-guidance-for-strengthening-ai-system-security/
24	韓国/イギリス	国表院-BSI量子会議	2024/4/17	韓国は4月16日、英国のBSI標準部門代表であるScott Steedman博士をソウルで迎え、量子技術、AIなど両国間の先端技術標準協力の高度化について話し合った。 5月28日～30日に開かれるJTC3量子技術創立総会及び5月初旬に英国と共同で開催する「AI for All : Standardization and Industrialisation」イベントに関する準備状況について議論した。また、今年英国で開催されるIEC総会と連携して、英国はスマートシティ/量子技術/AI/ネットゼロ関連イベントを開催する予定であり、両国は専門家の参加、産業界広報などについて協力していくことにした。両国は、標準化は単独で進めることができないため、友好国と一緒に行うことが重要であるとの意見で一致し、今後、先端技術に対する標準化協力をより緊密に展開することにした。	韓国国家技術標準院 https://www.kats.go.kr/content.do?cmsid=283&cid=24450&mode=view

【人工知能(AI)】関連記事詳細（17/35）

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
25	アメリカ	カリフォルニア州、州政府機関による生成AIツール調達ガイドラインを発表	2024/4/26	カリフォルニア州は、2023年11月に発表された「Generative Artificial Intelligence Report（生成AIの利点とリスク）」を基に、州政府機関による生成AIの調達、利用、トレーニングに関するガイドラインを発表した。 カリフォルニア州技術局、総合サービス局、データ・イノベーション局、人的資源局によって作成されたこのガイドラインは、生成AIの倫理的、透明性、信頼性のある利用を実施・監督する機関を支援することを目的としたベストプラクティスを示している。このガイドラインは、現時点では学区やその他の地方機関には適用されないが、このガイドラインに含まれる概念と制限は、将来の生成AIガイドライン、命令、法律の先駆けとなる可能性がある。	カリフォルニア州学校管理者協会 https://content.acsa.org/california-issues-guidelines-for-state-agencies-purchasing-generative-ai-tools/
26	アメリカ	商務省、バイデン大統領のAIに関する大統領令を実施するための新たな措置を発表	2024/4/29	米国商務省は、バイデン大統領による安全でセキュリティが高く信頼できるAI開発に関する大統領令（EO）から180日が経過したのを受けて、EOに関連する新たな発表を行った。国立標準技術研究所（NIST）は、人工知能（AI）システムの安全性、セキュリティ、信頼性の向上に役立つことを目的とした4つの草案出版物をリリースした。NISTはまた、人間が作成したコンテンツとAIが作成したコンテンツを区別する方法の開発を支援するチャレンジシリーズも開始した。 NISTの出版物に加えて、商務省の米国特許商標庁（USPTO）は、米国法の下で発明が特許を受けることができるかどうかを判断するために行われる技術における通常の技能レベルの評価にAIがどのように影響するかについてのフィードバックを求めるパブリックコメント要請（RFC）を公表しており、今年初めにはAI支援発明の特許性に関するガイダンスを公表している。 NIST の4つの出版物は、以下の通り。 - AI RMF 生成AI プロファイル：NIST AI 600-1 - 生成AIおよびデュアルユース基盤モデル向けのセキュアなソフトウェア開発プラクティス：NIST特別出版物（SP）800-218A - 合成コンテンツによってもたらされるリスクの軽減：NIST AI 100-4 - AI標準に関するグローバルな取り組み計画：NIST AI 100-5	NIST https://www.nist.gov/news-events/news/2024/04/department-commerce-announces-new-actions-implement-president-bidens

【人工知能(AI)】関連記事詳細（18/35）

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
27	アメリカ	AIリスクマネジメントフレームワーク	2024/4/29	NISTは、民間および公的セクターと協力して、人工知能（AI）に関連する個人、組織、社会に対するリスクをより適切に管理するためのフレームワークを開発した。NIST AIリスクマネジメントフレームワーク（AI RMF）は、AI製品、サービス、システムの設計、開発、使用、評価に信頼性への配慮を組み込む能力を向上させ、自主的に使用することを目的としている。 2023年1月26日に発表されたこのフレームワークは、情報提供の要請、パブリックコメントのための草案、複数のワークショップ、その他の意見提供の機会を含む、コンセンサス主導の、オープンで透明性のある、協力的なプロセスを通じて策定された。このフレームワークは、他の機関によるAIリスクマネジメントの取り組みを基礎とし、それと整合し、支援することを意図している。 NISTは、AI RMFロードマップ、AI RMFクロスワーク、様々な展望とともに、NIST AI RMF Playbookも発行している。さらに、NISTはAI RMFに関する説明ビデオを公開している。2023年3月30日、NISTはTrustworthy and Responsible AI Resource Centerを立ち上げ、AI RMFの実施と国際的な整合性を促進する。	NIST https://www.nist.gov/itl/ai-risk-management-framework
28	国際	岸田総理大臣の生成AIに関するOECD閣僚理事会サイドイベント出席	2024/5/2	OECD閣僚理事会（MCM）に出席した岸田文雄内閣総理大臣は、生成AIに関するサイドイベント「安全、安心で信頼できるAIに向けて：包摂的なグローバルAIガバナンスの促進」において以下のとおり発言した。 - 安全、安心で信頼できるAIの実現のための国際ガバナンスの形成が急務であり、その観点から広島AIプロセスを立ち上げ、国際指針や行動規範を策定し、生成AIを巡る具体的なリスクの低減に取り組んだ。OECDにおいてもAI原則の改定という具体的な成果が生み出されることを歓迎する。 - また、49か国・地域の参加を得て、広島AIプロセスの精神に賛同する国々の自発的な枠組みである「広島AIプロセス フレンズグループ」を立ち上げ、国際指針等の実践に取り組み、世界中の人々が安全、安心で信頼できるAIを利用できるよう協力を進めていく。 - さらに、AIによるリスク低減のための技術的措置も重要であり、日本はGPAI東京センターを新設して専門家による技術実証等のプロジェクトを支援するとともに、生成AIがもたらす偽情報等のリスクに対応するため、発信者情報を確認するための技術の社会実装に向けた取組も支援していく。	外務省 https://www.mofa.go.jp/mofaj/ecm/ec/pageit_000001_00601.html

【人工知能(AI)】関連記事詳細（19/35）

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
29	国際	OECD、急速な技術発展に対応するためAI原則を更新	2024/5/3	2024年OECD閣僚理事会は、AIに関する画期的なOECD原則の改訂版を採択した。AI技術の最近の発展、特に汎用的で生成的なAIの出現に対応して、改訂された原則は、プライバシー、知的財産権、安全性、情報の完全性に関わるAI関連の課題をより直接的に取り上げている。 OECDの改訂の主要素は、この原則が引き続き適切で、堅固で、目的に適合したものであることを保証するものであり、以下を含む： - AIシステムが不当な危害をもたらす危険性や望ましくない挙動を示した場合、それを安全に無効化、修復、および／または廃止するための強固なメカニズムとセーフガードが存在するように、安全性への懸念に対処する。 - 誤った情報や偽情報に対処し、生成的AIの文脈で情報の完全性を保護することの重要性が高まっていることを反映する。 - AIの知識やAIのリソースの供給者、AIシステムの利用者、その他の利害関係者との協力を含め、AIシステムのライフサイクル全体を通じて責任ある事業活動を行うことを強調する。 - 透明性と責任ある情報開示を構成するAIシステムに関する情報を明確にする。 - 過去5年間で重要性が著しく高まった環境の持続可能性について明確に言及する。 - 世界中でAI政策の取り組みが急増する中、各国・地域が連携してAIに関する相互運用可能なガバナンスと政策環境を推進する必要性を強調する。	OECD https://www.oecd.org/en/about/news/press-releases/2024/05/oecd-updates-ai-principles-to-stay-abreast-of-rapid-technological-developments.html

【人工知能(AI)】関連記事詳細 (20/35)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
30	中国/フランス	人工知能とグローバルガバナンスに関する中華人民共和国とフランス共和国間の共同声明	2024/5/7	2023年4月7日の「中仏共同声明」での合意に基づき、中仏両首脳は特に人工知能の分野において合意した。主な合意内容は以下の通り。 - 中国とフランスは、開発とイノベーションにおける人工知能の重要な役割を認識し、人工知能の開発と使用がもたらす可能性のある一連の課題を考慮し、人工知能の開発と安全性を促進し、適切な人工知能を推進することに合意する。 - 中国とフランスは、人工知能技術の急速な発展による重大な影響と、この技術に関連する潜在的かつ実際のリスクを十分に認識し、これらのリスクに対処するための効果的な措置を講じ、世界的なガバナンスを強化することにコミットする。 - 人工知能がもたらす機会を最大限に活用するため、中国とフランスは、人工知能の国際統治モデルに関する議論を深めることに尽力する。このガバナンスでは、テクノロジーの継続的かつ迅速な開発に必要な柔軟性を考慮するだけでなく、個人データ、人工知能ユーザーの権利、人工知能によって作品が使用されるユーザーの権利に必要な保護も提供する必要がある。	中華人民共和国中央人民政府 https://www.gov.cn/yaowen/libiao/202405/content_6949586.htm
31	国際	OECD閣僚理事会の結果	2024/5/8	5月2日～3日にパリにおいてOECD閣僚理事会が開催された。 生成AIに関するサイドイベント「安全、安心で信頼できるAIに向けて：包摂的なグローバルAIガバナンスの促進」にて岸田総理大臣がスピーチを行い、広島AIプロセスの精神に賛同する国々の自発的な枠組みである「広島AIプロセス フレンズグループ」の設立を発表した。 また、議題6「新興課題に対する解決志向型アプローチ」前半AIパートについて松本総務大臣が議長を務め、2016年のG7香川高松情報通信大臣会合における日本からの提言を契機にOECDが検討を開始し、2019年に公表した「人工知能（AI）に関する勧告」（OECD AI原則）の見直しにあたり、近年急速に発展した生成AIのガバナンスの在り方について、日本が昨年のG7議長国として主導した「広島AIプロセス」の成果を反映する方向で議論がなされた。 閉会式では、本会合の成果として、OECD AI原則改定版及び閣僚声明等が採択された。	総務省 https://www.soumu.go.jp/menu_news/s-news/01tsushin06_02000289.html

【人工知能(AI)】関連記事詳細 (21/35)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
32	欧州	欧州評議会、人工知能に関する初の国際条約を採択	2024/5/17	欧州評議会は、人工知能（AI）システムの使用における人権、法の支配、民主主義の法的基準の尊重を確保することを目的とした、初の国際的な法的拘束力のある条約を採択した。この条約は欧州以外の国にも適用され、AIシステムのライフサイクル全体をカバーし、責任ある技術革新を促進しながら、それらがもたらす可能性のあるリスクに対処する法的枠組みを定めている。 この条約は、AIシステムの設計、開発、使用、廃棄にリスクベースのアプローチを採用しており、AIシステムの使用による潜在的な悪影響を慎重に検討することを求めている。 この条約は、政府間組織である人工知能委員会（CAI）による2年間の作業の成果であり、条約の草案作成には、欧州評議会の46の加盟国、欧州連合、および11の非加盟国（アルゼンチン、オーストラリア、カナダ、コスタリカ、ローマ教皇庁、イスラエル、日本、メキシコ、ペルー、アメリカ合衆国、ウルグアイ）、さらにオブザーバーとして参加した民間部門、市民社会、学界の代表者が集まった。	欧州評議会 https://www.coe.int/en/web/portal/-/council-of-europe-adopts-first-international-treaty-on-artificial-intelligence
33	イギリス/アメリカ	政府の先駆的なAI安全研究所がサンフランシスコに開設	2024/5/20	英国政府の先駆的なAI安全研究所は、今夏サンフランシスコに初の海外事務所を開設し、国際的な視野を広げると、ミシェル・ドネラン技術長官が発表した。この拡張は、英国がベイエリアの豊富な技術人材を活用し、ロンドンとサンフランシスコの両方に本社を置く世界最大のAI研究所と連携し、公益のためにAI安全性を推進する米国との関係を強化することを可能にする重要な一歩となる。このオフィスは今年の夏に開設される予定で、リサーチ・ディレクターが率いる技術スタッフの最初のチームを募集する。 米国での足場を拡大することで、研究所は米国との緊密な協力関係を確立し、同国の戦略的パートナーシップとAIの安全性に対するアプローチをさらに強化するとともに、研究を共有し、世界中のAI安全性政策に情報を提供できるAIモデルの共同評価を実施する。 この協定の一環として、両国は専門知識を共有し、既存の試験・評価作業を強化することを目指す。このパートナーシップはまた、両国間の出向ルートを可能にし、研究協力のための分野を共同で特定する作業も行う。	GOV.UK https://www.gov.uk/government/news/government-news/government-news-trailblazing-institute-for-ai-safety-to-open-doors-in-san-francisco

【人工知能(AI)】関連記事詳細 (22/35)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
34	アメリカ	ジーナ・ライモンド 米国商務長官 がAI安全に関する戦略ビジョンを 発表、AI安全 機関間の世界的協力計画を 発表	2024/5/21	AIソウル・サミットが開幕する本日、ジーナ・ライモンド米商務長官は、AI安全性に対する同省の取り組みについて、米国人工知能安全研究所（AISI）の戦略的ビジョンを発表した。バイデン大統領の指示により、商務省内の国立標準技術研究所（NIST）は、長年にわたるAIに関する研究を基にAISIを発足させた。また、AI安全性研究所やその他の政府系科学機関との有意義な関わりを通じて、AIの安全性のための世界的な科学ネットワークと協力し、AISIが最近設立したサンフランシスコ地域で今年後半に研究所を招集するという同省の計画も発表された。 同時に、同省とAISIが、AIの安全性に焦点を当て、国際協力に取り組むAI安全性研究所やその他の政府系科学機関との有意義な関わりを通じて、AIの安全性に関する世界的な科学ネットワークの立ち上げを支援すると発表した。 このネットワークは、AIソウル・サミットで韓国と他のパートナーが「AI安全科学に関する国際協力に向けたソウル声明」を通じて達成した基礎的理解に基づき、AISIが以前に発表した英国、日本、カナダ、シンガポールのAI安全研究所、欧州AI事務局およびその科学的構成機関・関連機関との協力関係を強化・拡大し、AI安全科学とガバナンスに関する国際協調の新たな段階を促進するものである。このネットワークは、戦略的研究と公的成果物に関する緊密な協力を可能にすることで、世界中の人々にとって安全、安心、信頼できる人工知能システムを促進する。	米国商務省 https://www.commerce.gov/news/press-releases/2024/05/us-secretary-commerce-gina-raimondo-releases-strategic-vision-ai-safety

【人工知能(AI)】関連記事詳細 (23/35)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
35	国際	世界のリーダーたちは、AIの協力を強化するために、AI安全研究所の初の国際ネットワークを立ち上げることに合意	2024/5/21	本日開催されたAIソウル・サミットにおいて、10カ国と欧州連合が新たな合意に達し、AI安全科学の発展を加速させるための国際ネットワークの立ち上げに向けて各国が協力することが約束された。 AI安全科学の国際協力に向けたソウル声明」は、英国が第1回AI安全サミットで世界初のAI安全研究所を設立して以来、米国、日本、シンガポールを含む、英国のAI安全研究所と同様の公的支援機関を結集する。 このネットワークは、AIの安全・安心で信頼できる開発を促進するため、それぞれの技術的作業とAIの安全性へのアプローチの間に「補完性と相互運用性」を構築する。これには、モデル、その限界、能力、リスクに関する情報の共有、具体的な「AIの危害と安全インシデント」が発生した場合の監視、AIの安全性に関する科学への世界的理解を深めるためのリソースの共有が含まれる。 これは、AIソウル・サミットのリーダーズセッションで合意されたもので、世界のリーダーや主要なAI企業が一堂に会し、AIの安全性、革新性、包括性について議論した。 協議の一環として、各国首脳はより広範なソウル宣言に署名し、「人間中心で、信頼でき、責任ある」AIを開発するための国際協力強化の重要性を確認した。 これらは、AIテクノロジー企業16社が発表したばかりの「フロンティアAI安全性コミットメント」に続くもので、主要なAI開発企業が、リスクを管理できないと判断する閾値を設定する際に、政府やAI安全性機関からの意見を取り入れることを定めている。世界初となるこのコミットメントには、米国、中国、中東、ヨーロッパを含む世界中のAI企業が署名している。 ※岸田総理大臣の発言はこちら	GOV.UK https://www.gov.uk/government/news/global-leaders-agree-to-launch-first-international-network-of-ai-safety-institutes-to-boost-understanding-of-ai

【人工知能(AI)】関連記事詳細 (24/35)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
36	欧州	人工知能（AI）法：評議会がAIに関する世界初の規則に最終承認	2024/5/21	本日、欧州理事会は人工知能に関する規則の調和を目指す画期的な法律、いわゆる人工知能法を承認した。この主要法案は「リスクベース」のアプローチに従っており、社会に危害を及ぼすリスクが高ければ高いほど、規則も厳しくなる。この種の法律は世界初であり、AI規制のグローバル・スタンダードとなる可能性がある。適切な執行を確保するため、いくつかの管理機関が設置されている： ・ 欧州委員会内にAI事務局を設置し、EU全域で共通の規則を施行する。 ・ 執行活動を支援する独立専門家からなる科学委員会 ・ AI法の一貫した効果的な適用について欧州委員会と加盟国に助言し、支援する ・ 加盟国代表によるAI委員会AI委員会と欧州委員会に技術的な専門知識を提供する利害関係者のための諮問フォーラム 欧州議会と欧州理事会の両議長が署名した後、同法は近日中にEU官報に掲載され、掲載から20日後に発効する。新規則は発効から2年後に適用されるが、特定の条項については例外がある。	欧州理事会 https://www.council.europa.eu/en/press-releases/2024/05/21/artificial-intelligence-act-council-gives-final-green-light-to-the-first-worldwide-rules-on-ai/
37	カナダ/イギリス	カナダと英国による人工知能の安全性向上に関する共同大臣声明	2024/5/21	カナダ政府と英国政府は本日、人工知能（AI）の安全性に関して、それぞれのAI安全性研究所（AISI）を含む協力を進める意向を発表した。この協定の一環として、両国は、既存の試験・評価作業を強化するための専門知識共有の道筋を作り、研究協力のための他の優先分野を共同で特定することを約束する。 両国は、専門家の交流や出向を促進し、それぞれの計算資源を活用・共有するための措置を講じる。さらに、英国AISIは、共同研究に関して、英国AI研究リソースへの優先アクセスの割り当てをカナダAISIと共有する。カナダAISIについては、今後数ヶ月の間に詳細が発表される予定であり、カナダと英国は、両国間のAI安全性に関する協力に関する覚書の締結に向けた作業を含め、両国の協力の詳細をさらに具体化するために関与を続けていく。 両国のAISIはまた、「システムのAI安全性」（AIが導入される社会システムの保護）の分野を促進するための研究プログラムにおいて、米国AISIと協力する意向である。 両政府は、AIの潜在的な利益を解放とうとする一方で、AIモデルの開発を安全に保つことを目指すこの作業の価値を認識している。この新しいパートナーシップは、国境を越えた開かれた対話を促進し、基礎研究を推進し、AIの安全性に関する連携を促進するのに役立つ。	カナダ政府 https://ised-isde.canada.ca/site/ised/en/joint-ministerial-statement-between-canada-and-uk-advancing-artificial-intelligence-safety

【人工知能(AI)】関連記事詳細 (25/35)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
38	中国	国家標準「サイバーセキュリティ技術における生成型人工知能サービスのセキュリティに関する基本要件」草案に対するコメント募集のお知らせ	2024/5/23	国家規格「ネットワークセキュリティ技術の生成的人工知能サービスのセキュリティのための基本要件」の草案が形成され、意見を求めている。標準の品質を確保するため、サイバーセキュリティ標準委員会事務局は、広く社会に開かれた協議の場を設けている。 国家標準化管理委員会によれば、本規格は SAC/TC260（国家サイバーセキュリティ標準化技術委員会）の管轄下であり、管轄当局は国家標準委員会とされている。	国家サイバーセキュリティ標準化専門委員会事務局 https://www.tc260.org.cn/front/bzzqyjDetail.html?id=20240523143149&norm_id=20240430101922&recorde_id=55010
39	韓国	韓国、ETRIにAI安全研究所を設立へ	2024/5/23	韓国は人工知能（AI）技術に関する安全性の問題を研究することに特化した研究所を韓国電子通信研究院(ETRI)に設立する。この計画は、水曜日にソウルの韓国科学技術研究院（KIST）で開催されたAIソウルサミットの閣僚級セッション後の記者会見で発表された。 科学技術情報通信部のLee Jong-ho氏は「AI安全機関間の協力を強化し、AI生成コンテンツに透かしを入れるなどAI生成コンテンツを識別する対策を講じ、国際標準の策定に向けた協力を強化していく」と述べた。 大臣声明における安全性の中核は、AI のリスクを測定するためのフレームワークを構築することであり、特にAI が安全レベルを超える重要な閾値を特定することとし、さらに、行政、福祉、教育、医療などの分野でもAIを積極的に導入し、イノベーションを推進していくとしている。	韓国電子通信研究院 (ETRI) https://www.etri.re.kr/eng/bbs/view_etri?board_id=ENG02&b_idx=19282

【人工知能(AI)】関連記事詳細（26/35）

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
40	アメリカ	NIST、AIの社会技術的テストと評価を推進する新プログラムARIAを開始	2024/5/28	米国国立標準技術研究所（NIST）は、人工知能の能力と影響についての理解を深めることを目的とした新しい試験・評価・検証・検証（TEVV）プログラムを開始する。 AIのリスクと影響の評価（ARIA）は、あるAI技術が導入された後、有効で、信頼性が高く、安全で、プライバシーが守られ、公正であるかどうかを、組織や個人が判断できるようにすることを目的としている。このプログラムは、信頼できるAIに関する大統領令が施行されてから180日が経過し、米国AI安全研究所が戦略的ビジョンと国際的な安全ネットワークを発表した直後に、NISTが発表したものである。 ARIAは、NISTが2023年1月に発表したAIリスク管理フレームワークを発展させたもので、フレームワークのリスク測定機能の運用を支援するもので、AIのリスクと影響を分析・監視するために定量的・定性的手法を用いることを推奨している。ARIAは、社会的文脈の中でシステムがどの程度安全な機能を維持できるかを定量化するための新しい一連の方法論と測定基準を開発することで、こうしたリスクと影響の評価を支援する。 ARIAの成果は、米国AI安全研究所を含め、安全、安心、信頼できるAIシステムの基盤を構築するためのNISTの総合的な取り組みを支援し、情報を提供する。	NIST https://www.nist.gov/news-events/news/2024/05/nist-launches-aria-new-program-advance-sociotechnical-testing-and
41	欧州	欧州委員会、安全で信頼できる人工知能におけるEUのリーダーシップを強化するため、AIオフィスを設置	2024/5/29	欧州委員会は、欧州委員会内に設置されたAI室を発表した。AI室は、リスクを軽減しつつ、社会的・経済的利益とイノベーションを促進する形で、AIの将来の開発、展開、利用を可能にすることを目的としている。同室は、特に汎用AIモデルに関して、AI法の施行に重要な役割を果たす。また、信頼できるAIの研究と革新を促進し、EUを国際的な議論のリーダーとして位置づけることにも取り組む。 AI室はAI室長によって率いられ、モデルや革新的なアプローチの評価において科学的な卓越性を確保するための主任科学アドバイザーと、信頼に足るAIについて国際的なパートナーと緊密に協力するという我々のコミットメントをフォローアップするための国際問題担当アドバイザーの指導のもとに活動する。上記の組織変更は、6月16日に発効する。AI理事会の初会合は6月末までに開催される予定である。AI事務局はAIシステムの定義と禁止事項に関するガイドラインを準備している。また、発効から9ヵ月後に予定されている、汎用AIモデルの義務に関する実施規範の作成に向けた調整も進めている。	欧州委員会 https://ec.europa.eu/commission/presscorner/detail/en/ip_24_2982

【人工知能(AI)】関連記事詳細（27/35）

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
42	中国	技術基準向上のための行動計画	2024/5/30	中国は、人工知能チップ、生成AI、量子情報など、さまざまな最先端技術の標準を強化するための3カ年行動計画を発表した。中央サイバースペース委員会事務局のニュースリリースによると、行動計画は、情報化のための標準、つまりシステムを通しての情報の流れが国家標準システムの重要な一部であり、質の高い発展を推進するための重要なサポートとして機能することを強調している。 2027年までには、質の高い情報化標準が一挙に発表され、標準化分野における「専門化、専門化、国際化」された人材チームが育成されることが期待されている。さらに、その3年後には、技術革新を導き、経済・社会の発展を推進する標準の役割が完全に実現され、国際標準における中国の貢献と影響力が大幅に増大するはずである。 行動計画では、情報化標準化の作業メカニズムを革新し、8つの重要分野における標準の策定を促進するなど、4つの分野における主要な任務の概要を示している。例えば、重要情報技術分野では、集積回路の主要分野に焦点を当て、先端コンピューティング・チップと新型ストレージ・チップの標準策定への取り組みを強化し、AI等の応用標準策定を推進することを求めている。さらに、生成AIをはじめとする主要最先端技術も、標準化努力の優先対象としている。この計画では、ネットワークインフラの標準を改善し、コンピューティングパワーのインフラ標準の開発を進めること、データリソースの基本標準の構築を強化すること、公共データリソースの利用標準を革新することを求めている。さらに、同計画によると、今後3年間で、中国は国際標準化交流と協力を深化させ、国際標準化団体の活動に積極的に関与していく。	中華人民共和国中央人民政府 https://english.www.gov.cn/news/202405/30/content_WS6657d8dcc6d0868f4e8e7a18.html
43	国際	国連AI会議はすべての人に発言権を与える - ITUのAI for Goodグローバルサミットで包括的なAI標準と能力構築に関するコミットメントが発表	2024/5/31	スイスのジュネーブで開催されたAI for Goodグローバル・サミットで、政府、産業界、世界の人工知能コミュニティのリーダーたちが、AIをより包括的なものにするための大胆な一歩を踏み出した。 一連の行動、コミットメント、新たなイニシアチブは、多様な視点を強化し、AIをすべての人のために機能させるというITU（国連デジタル技術機関）のビジョンを反映したものだった。 ITUは、40の国連パートナーとともに、グローバルなデジタルの未来に関する議論に、あらゆる視点、特に発展途上国の視点を含める必要性を強調した。	ITU https://www.itu.int/en/mediacentre/Pages/PR-2024-05-31-AI-for-Good-Global-Summit.aspx

【人工知能(AI)】関連記事詳細 (28/35)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
44	アメリカ	ANSIはNISTに対し、「AI規格に関するグローバルな取り組みのための計画」について協調的な回答を提出	2024/5/31	米国規格協会 (ANSI) は、米国国立標準技術研究所 (NIST) が発表した草案「NIST AI 100-5、AI 標準に関する世界的な取り組み計画」に対して、標準化コミュニティを代表して協調的な回答を提出した。この草案は、AI 関連のコンセンサス標準、協力と調整、情報共有の世界的な開発と実装を推進することを目的としている。この回答は、ANSIのメンバーや構成員から提出されたものをまとめたもので、オープンで透明性が高く、コンセンサスによって推進されるプロセスで開発されるAI標準が、市場や政府のニーズに最も合致するというNISTの認識を称賛した。 ANSIは、連邦政府に対し、幅広い分野横断的な応用が可能なAIの技術標準やツールの開発への関与を引き続き優先するよう勧告しており、国際標準化機構/国際電気標準会議 (ISO/IEC) 合同技術委員会JTC1、情報技術、小委員会SC42、人工知能の中で進行中の作業について助言している。その他の分野としては、国内のキャパシティ・ビルディングとグローバルなキャパシティ・ビルディングが挙げられ、これには低・中所得国が重要技術や新興技術の標準開発に参加できるようにすることも含まれる。	ANSI https://www.ansi.org/standards-news/all-news/2024/05/5-31-24-ansi-submits-coordinated-response-to-nist-on-draft-publication-a-plan-for-global-engagement
45	中国	国家人工知能産業総合標準化システム構築のためのガイドラインの発布に関する通達 (2024年版)	2024/6/5	国家標準化発展綱要と人工知能ガバナンスに関するグローバル・イニシアティブを実施し、AI標準化業務の体系的な計画をさらに強化し、AI産業の高品質な発展と「AI+」のハイレベルなエンパワーメントのニーズを満たす標準体系の構築を加速する。AI標準化業務の体系的な計画をさらに強化し、AI産業の高品質な発展及び「AI+」のハイレベルな権限付与のニーズを満たす標準システムの構築を加速し、技術進歩を促進し、企業の発展を促進し、産業の高度化を先導し、産業の安全を守る標準の支持的な役割を強化し、AIを活用した新産業化をよりよく促進するため、ガイドラインを策定する。 このガイドラインにおける一般的な要件として、「AI産業の発展の最初のチャンスをつかむことを目指し、AI標準作業を完成させる。AI標準業務のトップレベル設計を行い、全産業チェーンの標準業務の相乗効果を強化し、標準の研究、策定、実施、国際化を協調して推進し、中国のAI産業の高品質な発展を促進するための確かな技術サポートを提供する」としている。 また、標準化の方向性として、基本的共通基準としての「用語の標準化」、「リファレンスアーキテクチャの標準化」、「テストの評価基準」、「管理基準」、「持続可能性基準」を挙げており、個別分野においても標準化の方向性が記載されている。	中華人民共和國中央人民政府 https://www.gov.cn/zhengce/zhengceku/202407/content_6960720.htm

【人工知能(AI)】関連記事詳細 (29/35)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
46	アメリカ	G7プーリア首脳 声明	2024/6/14	この声明では、AIに関する規制・標準の文脈について下記のように述べられている。 「我々は、アプローチや政策手段がG7各国間で異なる可能性があることを認識しつつ、確実性、透明性、説明責任を高めるため、AIガバナンスのアプローチ間の相互運用性を高めるための努力を強化する。我々は、イノベーションと強固で包摂的かつ持続可能な成長を促進するため、これらの取り組みにおいてリスクに基づくアプローチを採用する。この目標を達成するため、我々はベストプラクティスの共有を含め、ガバナンスと規制の枠組みの進化に関する調整を強化する。我々は定期的な協議を強化する。我々はまた、リスク管理の共通理解に向けて取り組み、AIの開発と展開に関する国際基準を前進させるために、AIに焦点を当てたそれぞれの研究所と事務所間の調整を深めることにコミットしている。我々は、先進的なAIシステムを開発する組織のための国際行動規範を監視するための報告枠組みの開発を含む、昨年発表された広島AIプロセスの成果を前進させるための産業・技術・デジタル大臣の努力を歓迎する。我々は、10月の産業・技術・デジタル大臣会合を見据え、OECDと協力して開発された報告枠組みのパイロット版を期待している。我々は、この規範の今後の報告枠組みに自発的に参加し、それを実施する組織を識別するために使用できるブランドの開発に取り組んでいく。」	THE WHITE HOUSE https://www.whitehouse.gov/briefing-room/statements-releases/2024/06/14/g7-leaders-statement-8/
47	国際	インパクトをもたらす AI: 社会イノベーションにおける責任あるAIのためのPRISMフレームワーク	2024/6/25	人工知能（AI）は、組織が社会的・環境的課題を克服する上で大きな可能性を秘めている。本シリーズの最初のレポート「AI for Impact」は、社会的イノベーションにおけるAIの役割についてまとめたものである：『ソーシャル・イノベーションにおけるAIの役割』では、社会的インパクトの領域、地域、ジェンダー別にAI導入の状況をマッピングし、インパクトのためにテクノロジーを導入して成功したソーシャル・イノベーターのケーススタディを紹介している。 このホワイトペーパーは、最初の報告書に基づき、世界経済フォーラムのAIガバナンス・アライアンス（AIGA）のプレシディオ・フレームワークを活用している。PRISMフレームワークを紹介し、ソーシャル・イノベーター、インパクト・エンタープライズ、仲介者が、そのインパクト・ミッション、能力、リスクとテクノロジーの利用を照らし合わせることができるよう、様々な採用経路と実際のケーススタディを提供している。本レポート・シリーズは、社会起業家グローバル・アライアンス（Global Alliance for Social Entrepreneurship）のAI for Social Innovation ワークストリームの成果である。	世界経済 フォーラム https://www.weforum.org/publications/ai-for-impact-the-prism-framework-for-responsible-ai-in-social-innovation/

【人工知能(AI)】関連記事詳細 (30/35)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
48	中国	中国を拠点とする発明家が生成AIの世界的な特許をリード：国連報告書	2024/7/3	国連世界知的所有権機関（WIPO）の新しい報告書によると、中国を拠点とする発明家らが生成AIの特許を最も多く申請している。同庁の特許状況報告書によると、2014年から2023年の間に中国から3万8000件以上の生成AI特許が発行され、2位の米国の発明者による出願件数の6倍に達した。 報告書によると、生成AI特許は現在、世界全体のAI特許の6%を占めている。出願者のトップ10には、テンセント（2,074件）、平安保険（1,564件）、バaidu（1,234件）、中国科学院（607件）、IBM（601件）、アリババ・グループ（571件）、サムスン電子（468件）、アルファベット（443件）、バイトダンス（418件）、マイクロソフト（377件）が含まれている。場所別では、中国（38,210件）がトップで、米国（6,276件）、韓国（4,155件）、日本（3,409件）、インド（1,350件）を大きく引き離れた。画像・映像データが生成AI特許の大半を占め（17,996件）、テキスト（13,494件）、音声・音楽（13,480件）と続く。分子、遺伝子、タンパク質ベースのデータを使用した生成AI特許は急速に増加し（2014年以降1,494件）、過去5年間の年平均成長率は78%であった。	国際連合 https://news.un.org/en/story/2024/07/1151761
49	メキシコ	UNESCOがメキシコの人工知能準備状況評価を発表	2024/7/4	UNESCOは、Centro-i for the Society of the Future (Centro-i)の支援を受け、National Alliance for Artificial Intelligence (ANIA)と共同で、連邦政府、各州政府、自治体、市民団体、学界、民間セクターの250人以上の代表者を交えたオープンで複数のプロセスを通じて作成された「メキシコの人工知能準備アセスメント」を発表した。 メキシコはまた、プライバシー、透明性、情報へのアクセス権に関する規制の枠組みもしっかりしている。AIシステムやデジタル・コンテキストによって引き起こされた損害に対処し修復するための具体的な法律や政策がないとしても、現行の基準やメカニズムによっていくつかのケースに取り組むことができる。 しかし、例えば、68の国内固有言語のAIのライフサイクル段階やデジタル表現全体を通してのインクルージョン、デジタルプラットフォームの大幅な利用促進、メディアと情報のリテラシーを教育システムに組み込むことなどを通じて、十分に代表されないグループへの対応を強化しなければならない。 この報告書は、2021年にユネスコ加盟国が採択した「人工知能の倫理に関する勧告」の一手段である「準備状況評価法（Readiness Assessment Methodology：RAM）」を実施したもので、AIの倫理的方向付けにおける進捗段階を特定するのに役立つ。	UNESCO https://www.unesco.org/en/articles/unesco-presents-artificial-intelligence-readiness-assessment-mexico

【人工知能(AI)】関連記事詳細 (31/35)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
50	中国	人工知能のグローバルガバナンスに関する上海宣言	2024/7/4	7月4日、2024年世界人工知能会議および人工知能のグローバルガバナンスに関するハイレベル会議は、「人工知能のグローバルガバナンスに関する上海宣言」を発表した。 その冒頭では、「人工知能技術の開発と応用を共同で推進しながら、その開発プロセスの安全性、信頼性、制御性、公平性を確保し、人工知能技術を促進して人間社会の発展を促進する必要性を強調」したうえで、 1. 人工知能の開発を促進 2. 人工知能のセキュリティの維持 3. 人工知能ガバナンスシステムを構築 4. 社会参加の強化と国民のリテラシーの向上 5. 生活の質と社会福祉の向上 について触れ、「世界中の政府、科学技術界、産業界、その他の関係者が積極的に対応し、人工知能の健全な発展を共同で促進し、人工知能の安全性を共同で維持し、人類の共通の未来に力を与えることを期待する」と結んでいる。 ※この会議に出席した李強首相のスピーチ概要はこちら	中華人民共和国中央人民政府 https://www.gov.cn/yaowen/liobiao/202407/content_6961358.htm
51	国際	NATOが改訂されたAI戦略を発表	2024/7/10	NATOは人工知能（AI）戦略の改訂版を発表した。同戦略は2021年に発表されたものを基礎とし、ジェネレーティブAIやAI対応情報ツールなどAI技術の最近の進歩を考慮に入れている。 この戦略では、NATOの責任ある使用原則の実施を推進すること、同盟国全体のAIシステム間の相互運用性を高めること、AIと他の新たな破壊的技術との組み合わせ、同盟国の産業界や学界、NATOの防衛イノベーション・アクセラレータDIANA、NATOイノベーション・ファンド、志を同じくするパートナーとの緊密な協力を通じてNATOのAIEコシステムを拡大することなど、いくつかの優先事項が挙げられている。 この戦略では初めて、AIを活用した偽情報、情報操作、ジェンダーに基づく暴力も同盟、われわれの社会、民主主義にとっての懸念事項として挙げている。 新たなAI戦略の下で、NATOは戦略的な先見性と分析を強化することを含め、敵対的なAIの利用から保護するために取り組む。	NATO https://www.nato.int/cps/en/natohq/news_27234.htm?selectedLocale=en

【人工知能(AI)】関連記事詳細 (32/35)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
52	アメリカ	アメリカ上院議員が、透明性を高め、AIディープフェイクと戦い、ジャーナリスト、アーティスト、ソングライターがコンテンツを管理できるようにする法案を提出	2024/7/11	上院商務委員会委員長のマリア・キャントウェル（ワシントン州選出）、商務委員会委員のマーシャ・ブラックバーン（テネシー州選出）、上院AI作業部会委員のマーティン・ハインリッヒ（ノースカロライナ州選出）は、有害なディープフェイクの台頭と闘うため、「Content Origin Protection and Integrity from Edited and Deepfaked Media Act（COPIED ACT）」を提出した。 この法案は、AIによって生成されたコンテンツの表示、認証、検出に関する連邦政府の新たな透明性ガイドラインを設定し、ジャーナリスト、俳優、アーティストをAIによる窃盗から保護し、違反者に悪用の責任を負わせるものである。	米国上院商務科学運輸委員会 https://www.commerce.senate.gov/2024/7/cantwell-blackburn-heinrich-introduce-legislation-to-combat-ai-deepfakes-put-journalists-artists-songwriters-back-in-control-of-their-content
53	欧州	欧州AI規則の官報公布	2024/7/12	7月12日、AI規則はEU官報で公示された。この第113条には発効と適用について、以下のように記載されている。 本規則は、欧州連合官報に掲載された翌日から20日目に発効する。 2026年8月2日より適用される。ただし、 (a)第1章および第2章は2025年2月2日から適用される； (b)第3章第4節、第5章、第7章、第12章および第78条は、第101条を除き、2025年8月2日から適用される； (c)第6条1項および本規則の対応する義務は、2027年8月2日より適用される。 本規則は、その全体を拘束し、すべての加盟国に直接適用されるものとする。	EUR-Lex https://eur-lex.europa.eu/legal-content/EN/TXT/?uri=OJ:L_202401689

【人工知能(AI)】関連記事詳細 (33/35)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
54	韓国	データ分析とデータ品質標準の開発 AI信頼性の強化と産業発展への寄与を期待	2024/7/16	韓国の経済メディアである毎日経済によれば、韓国電子通信研究院(ETRI)は人工知能(AI)分野の国際標準化を推進し開発してきた「データ分析および機械学習のためのデータ品質標準」が国際標準として制定されたと16日明らかにした。 研究陣は計6つの部分で構成された「データ分析および機械学習のためのデータ品質」シリーズの中でAIに使われるデータ品質に対する概要と共通概念を確立した。「データ品質測定」「データ品質管理要求事項」「品質管理手続き」などシリーズ標準に対する理解と適用を助ける「第1部:概要および用語、例題」開発を主導した。制定された標準は、データ基盤の意思決定時代において、組織がAI開発過程でデータ品質を評価、管理、改善できるツールと方法を提供する。データが目的に合うように使われるように助け、一貫して信頼できる共通の用語と実践方案を提示している。 今回の国際標準制定によってデータ品質に対する測定と評価ができるようになったことで、品質可否を簡単に知ることができるようになった。顧客は品質に基づいてデータの品質に対する信頼性に応じて購入可否を判断できるだけでなく、AI開発中にも持続的に使用中のデータ品質を改善できる基準を持つことになる。	韓国の経済メディア（毎日経済） https://www.mk.co.kr/jp/it/11068249
55	欧州	AI法の枠組みにおけるデータ保護当局の役割に関する声明	2024/7/16	7月12日、人工知能に関する調和のとれた規則（AI法）を定め、特定の連邦立法法を改正する規則（EU）2024/1689が官報に掲載された。各国のデータ保護当局（DPA）は、こうした技術開発に関して積極的な姿勢を見せており、AI法に関する立法過程を注意深く見守ってきたEDPBは、すでにEUデータ保護法との（多面的な）相互関係の検討を開始している。 この声明により、欧州データ保護委員会（EDPB）は、個人データ保護に密接に関連する分野において、加盟国が管轄当局を指定することによって生じうる監督・調整の問題を強調する。国家レベルでは、AI法の施行後1年以内に、1つまたは複数の「国家管轄当局」を中心にAI法施行の枠組みを構築する必要があることを考慮すべきである。 EDPBは、AI法第74条に記載されている高リスクのAIシステムについては、加盟国がDPAを市場監視当局（MSA）として指定すべきであると勧告する。さらに、EDPBは、加盟国に対し、国内DPAの意見を考慮した上で、特に、これらの高リスクAIシステムが、個人データの処理に関して自然人の権利と自由に影響を与える可能性が高い部門にある場合、当該部門がAI法で義務付けられた選任の対象になっていない限り（金融部門など）、附属書IIIに記載された残りの高リスクAIシステムのMSAとしてDPAを選任することを検討するよう勧告する。	欧州データ保護委員会 https://www.edpb.europa.eu/our-work-tools/our-documents/statements/statement-32024-data-protection-authorities-role-artificial-en

【人工知能(AI)】関連記事詳細 (34/35)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
56	アメリカ	アスペンセキュリティフォーラムがCoSAIの立ち上げを主催	2024/7/18	Coalition for Secure AI (CoSAI) は、コロラド州アスペンで開催されたアスペン・セキュリティ・フォーラム (ASF) で、このプロジェクトの創設プレミア・スポンサーであるグーグル、IBM、インテル、マイクロソフト、エヌビディア、ペイパル、追加創設スポンサーであるアマゾン、アンソロピック、チェインガード、シスコ、コヒア、ジェンラボ、OpenAI、ウィズを含む業界リーダーからの支援を受けて設立された。 OASIS Open Projectの傘下で発足したこの共同イニシアチブは、共同イノベーションと標準化を通じて、AIの開発と展開における信頼とセキュリティを強化することを目的としている。CoSAIは、AIのセキュリティとプライバシーの技術、専門知識、ソリューションを民主化し、あらゆる規模の組織におけるAIシステムの開発、実装、利用を改善することに焦点を当てている。	Coalition for Secure AI (CoSAI) https://www.coalitionforsecureai.org/latest-news-about-your-project/
57	アメリカ	重要技術と新興技術に関する国家標準戦略の実施	2024/7/26	バイデン-ハリス政権は、2023 年 5 月に米国政府が策定した「重要・新興技術 (CET) に関する国家標準化戦略 (USG NSSCET)」の実施ロードマップを発表した。このロードマップは、民間部門が主導し、公的機関とのパートナーシップによって強化された標準化開発に対する米国政府のコミットメントを維持・強化するものであり、米国の国家安全保障と経済安全保障を守るため、CETの標準化に積極的に関与することを求めている。 AIに関しては、以下の通り記載されている。 国立標準技術研究所 (NIST) は、標準および適合性評価に関連する国家技術移転促進法 (National Technology Transfer and Advancement Act) の規定を連邦機関が実施するための調整を行っている。さらに NIST は、重要技術や新興技術を含め、科学技術事業全体にわたる標準開発努力に技術的専門知識を提供している。例えば 一般的な新技術、特にAI研究に対する信頼に貢献してきたNISTの経験は、バイデン大統領の安全、安心、信頼できる人工知能に関する大統領令 (EO) (14110) の履行を支援する連邦機関に選ばれた理由である。成果物には、人工知能の開発と利用の形成における技術標準の重要性を認識した「AI標準に関するグローバルな関与のための計画」が含まれる。	THE WHITE HOUSE https://www.whitehouse.gov/briefing-room/statements-releases/2024/07/26/fact-sheet-implementing-the-national-standards-strategy-for-critical-and-emerging-technology/

【人工知能(AI)】関連記事詳細 (35/35)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
58	アメリカ	ファクトシート：バイデン・ハリス政権が新たなAI対策を発表し、AIに関する追加の主要な自主的な取り組みを承認	2024/7/26	ハリス副大統領が「AIの安全性に関するグローバル・サミット」での主要政策演説の一環として行った大統領令と一連の行動要請を受け、政府各機関は果敢に行動を起こしている。 AIの安全性とセキュリティのリスクを軽減し、アメリカ人のプライバシーを保護し、平等と公民権を促進し、消費者と労働者のために立ち上がり、イノベーションと競争を促進し、世界におけるアメリカのリーダーシップを促進するなどの措置を講じた。本日、各機関が完了を報告したアクションが本ファクトシートにまとめられている。	THE WHITE HOUSE https://www.whitehouse.gov/briefing-room/statements-releases/2024/07/26/fact-sheet-biden-harris-administration-announces-new-ai-actions-and-receives-additional-major-voluntary-commitment-on-ai/
59	イギリス	AI専門家が英国が人工知能の恩恵を確実に享受できるよう行動計画を主導	2024/7/26	英国のピーター・カイル新科学長官は、変革、持続的な経済成長、公共サービスの向上を実現する政府のアジェンダの中心にAIを据えた。 同長官は、技術系起業家でAdvanced Research And Invention Agency（ARIA）会長のマット・クリフォード氏を任命し、この活動を開始した。彼は、サービスを改善し新製品を開発することによって、人々の生活を向上させるためにAIの利用を加速させる方法を特定するための新しいAI機会行動計画を提供する。 同計画は、グローバルな舞台で規模を拡大し競争できる英国のAIセクターを構築する方法を模索するだけでなく、経済のあらゆる部分でこの技術の利用を促進する方法を定め、公共部門と民間部門による採用を促進するために必要なインフラ、人材、データアクセスを検討する。 アクションプランでは、2030年までに英国が必要とするコンピュータ・インフラや、より広範なインフラ、スタートアップやスケールアップ企業がこのインフラを利用できるようにする方法、官民で優秀なAI人材を育成・誘致する方法など、重要なイネーブラーについても検討する。	GOV.UK https://www.gov.uk/government/news/ai-expert-to-lead-action-plan-to-ensure-uk-reaps-the-benefits-of-artificial-intelligence



海外標準化動向調査(2月)

令和6年度エネルギー需給構造高度化基準認証推進事業費(我が国の国際標準化戦略を強化するための体制構築)
2025年2月1日

一般財団法人日本規格協会

ピックアップ：人工知能(AI) (関連ニュース番号31)



トピック	米国とアラブ首長国連邦の人工知能に関する協力の政府間覚書（MoU）が締結	
推進組織	アメリカ政府・アラブ首長国連邦（UAE）政府	
内容	ポイント	・「開発途上国におけるAIの支援」が覚書の事項として掲載されており、UAE政府がイニシアチブを持ち中東・北アフリカ地域での能力構築を担う事を目指す。
	背景	・ 以下概要の項目について協力する事を明記。 ・ ニュース番号No21に記載の、イギリスとカタールが共同研究プロジェクトを発表した事から、中東と西側各国の協働が今後進むことが推測される。
	概要	
	安全で信頼できるAIの推進	説明可能で公平性があり、人権と基本的自由を保護し、国際的な規範、ベストプラクティス、AIガバナンスにおける相互運用性を促進するAI技術の責任ある信頼性の高い開発と利用を確保するために、国際的なAIの枠組み、原則、基準の受容を促進する。
	イノベーションのエコシステムを強化するための規制枠組みの整合	AIおよび関連技術に関する規制枠組みと規則の整合性をさらに高め、国家安全保障上の利益を保護し、信頼できる投資と起業を可能にし、国境を越えたイノベーションを促進すると同時に、先進技術の保護と国際的な原則およびベストプラクティスの尊重を促進する
	倫理的なAI研究開発の推進	AIアルゴリズムにおけるバイアスや差別に対処し、公平性を確保するための研究を優先することで、倫理的なAI研究開発を行う
	AIの保護とサイバーセキュリティにおける協力の拡大と深化	オープンで相互運用可能、安全かつ信頼性の高いサイバーセキュリティ環境と、関連する新技術のリスクを管理しつつ、重要インフラの効率性と回復力を促進するサイバーインシデント対応戦略を推進する。
	信頼できる貿易と投資の機会を促進	強固で安全なAIインフラを開発する機会を捉えるため、二国間投資と効率的なライセンスを支援し促進する。
	人材開発と交流	AI研究者、エンジニア、政策立案者向けの共同トレーニングプログラムやワークショップを通じて、各国間の知識の交換と発展を促進するための人材開発を促進する。
	AIの未来のためのグリーンエネルギーを推進する	気候変動対策への共通の取り組みに沿って、グリーンエネルギー源でAIシステムのエネルギー需要を満たすため、米国とUAEのグリーンエネルギー加速パートナーシップ（PACE）を通じて現在進行中の二国間協力をさらに発展させる。
	開発途上国における持続可能な開発のためのAIの支援	世界最大の課題に対処し、デジタル格差を世界全体で、特に中東および北アフリカ地域で解消するために、AIおよびAIインフラに関する包括的で責任ある持続可能な能力構築を促進する。

【人工知能(AI)】関連記事詳細 (1/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
1	国際	ISO/IEC JTC 1/SC 42(Artificial intelligence)における規格開発状況	2024/12/18	ISO/IEC JTC 1/SC 42(人工知能)は、活動範囲を「人工知能に関するJTC 1の標準化プログラムの中心となり、推進役を務めること及び人工知能アプリケーションを開発するJTC 1、IEC、ISO委員会にガイダンスを提供すること」として、人工知能分野の標準化を行っている。 SCの幹事国はアメリカ(ANSI)。国内審議団体は(一社)情報処理学会内の情報規格調査会。2024年12月17日現在、発行済みの有効な規格は33規格。 2024年7月以降に発行された規格は、次の2規格である。 • ISO/IEC TR 5259-2 分析と機械学習 (ML) のためのデータ品質 — パート 2: データ品質の測定 ➢ 分析と機械学習 (ML) のコンテキストにおけるデータ品質モデル、データ品質の測定、およびデータ品質のレポートに関するガイダンスを指定。 • ISO/IEC 12791 情報技術 — 人工知能 — 分類および回帰機械学習タスクにおける不要なバイアスの処理 ➢ 機械学習を使用して分類および回帰タスクを実行する AI システムにおける不要なバイアスに対処する方法について説明。このドキュメントでは、不要なバイアスに対処するために AI システムのライフサイクル全体に適用できる緩和手法を紹介する。 なお、SC42の年次総会において、新プロジェクトとしてCASCOとのJWG6が設置。このJWG6の適用範囲は、以下のとおり。 • JTC1の人工知能標準化プログラムの焦点および推進者としての役割を果たす人工知能アプリケーションを開発すること • JTC、IEC、およびISO委員会にガイダンスを提供すること	International Organization for Standardization (ISO) https://www.iso.org/committee/6794475/x/catalogue/p/1/u/0/w/0/d/0

【人工知能(AI)】関連記事詳細 (2/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
2	イギリス	AI安全研究所がAI安全性評価の新プラットフォームをリリース ※前回の収集に含まれていなかったため、今回追加。	2024/5/10	英国AI安全研究所の評価プラットフォームが本日（5月10日金曜日）世界中のAIコミュニティに公開され、グローバルなAIの安全性評価が強化されることになり、AIモデルの安全なイノベーションへの道が開かれる。 世界初の国が後援するAI安全研究所を設立した英国は、AI安全評価に関するよりグローバルな協力体制の構築に向けた取り組みを継続しており、AI安全研究所が開発した評価プラットフォーム「Inspect」を公開した。Inspectを世界中のコミュニティに公開することで、同研究所は世界中で実施されているAI安全評価の取り組みを加速させ、より安全なテストとより安全なモデルの開発につながるよう支援している。これにより、世界中で一貫したAI安全評価のアプローチが可能になる。 Inspectは、スタートアップ企業、学術機関、AI開発者から各国政府まで、さまざまな立場のテスターが個々のモデルの特定の能力を評価し、その結果に基づいてスコアを算出することを可能にするソフトウェアライブラリである。Inspectは、コアとなる知識、推論能力、自律能力など、さまざまな分野におけるモデルの評価に使用できる。オープンソースライセンスでリリースされたことで、InspectはAIコミュニティが自由に利用できるようになった。 このプラットフォームは本日より利用可能となる。国家支援機関が主導するAI安全性テストプラットフォームが広く利用可能となるのは今回が初めてである。 英国を代表するAIの専門家たちによって推進された今回のリリースは、AI開発にとって重要な時期に行われる。2024年にかけてより強力なモデルが市場に投入されることが予想される中、安全で責任あるAI開発の推進はこれまで以上に急務となっている。	GOV.UK https://www.gov.uk/government/news/ai-safety-institute-releases-new-ai-safety-evaluation-platform

【人工知能(AI)】関連記事詳細 (3/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
3	日本	防衛省 A I 活用 推進基本方針と防 衛省サイバー人材 総合戦略の策定に ついて ※前回の収集に含 まれていなかったた め、今回追加。	2024/7/6	防衛省は、AIの活用促進や、サイバー人材の確保・育成の羅針盤となる文書として、「防衛省AI活用推進基本方針」と「防衛省サイバー人材総合戦略」という2つの文書を、それぞれ策定した。 ・ 防衛省AI活用推進基本方針 基本方針は、第1部「基本方針策定の背景と目的」、第2部「AIの活用分野と方向性」、第3部「AI活用推進に向けた取組」から構成される。国内・国外で行われている最新の議論を参考に、AIの活用分野、AIの利用に伴うリスクへの対応、データ基盤の構築、AI・データ人材の確保・育成、研究開発など、幅広い事柄について防衛省の考え方を示すもの。 ・ 防衛省サイバー人材総合戦略 人材施策を、①特定、②確保、③育成、④維持・管理、これら全てを統合するための⑤総合的な強化といった5つの柱に沿って検討。新たな取組を実施するとともに、既存の取組を抜本的に強化する。社会全体のサイバー人材の循環の中に防衛省・自衛隊を位置づけ、好循環を生み出すとの視点を持ち、施策を検討。	防衛省（日本） https://www.mod.go.jp/j/press/news/2024/07/02a.html

【人工知能(AI)】関連記事詳細 (4/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
4	アメリカ	商務省、バイデン大統領のAIに関する大統領令から270日後に新たな指針とツールを発表	2024/7/26	米国商務省は本日、バイデン大統領による人工知能（AI）の安全、セキュリティ、信頼性の確保に関する大統領令（EO）から270日目にあたるこの日、人工知能（AI）システムの安全性、セキュリティ、信頼性の向上に役立つ新たな指針とソフトウェアのリリースを発表した。 同省の国立標準技術研究所（NIST）は、4月にパブリックコメント用に最初に公開された3つの最終ガイダンス文書に加え、リスク軽減を目的とした米国AI安全研究所の草案ガイダンス文書を公開した。NISTはまた、敵対的攻撃がAIシステムのパフォーマンスを低下させる度合いを測定するソフトウェアパッケージも公開している。さらに、商務省の米国特許商標庁（USPTO）は、AIを含む重要な新興技術におけるイノベーションに対応するため、特許対象適格性に関するガイダンスの更新を発表した。また、全米通信情報庁（NTIA）は、広く利用可能な重み付けを持つ大規模なAIモデルのリスクとメリットを検証した報告書をホワイトハウスに提出した。 NISTの文書公開は、AI技術のさまざまな側面をカバーしている。そのうちの2つが本日初めて公開された。1つは、米国AI安全研究所によるガイダンス文書の最初の公開草案であり、生成型AIおよびデュアルユース基礎モデル（有益な目的にも有害な目的にも使用できるAIシステム）に起因するリスクをAI開発者が評価し、軽減するのを支援することを目的としている。もう一つは、AIシステムユーザーと開発者が特定の種類の攻撃がAIシステムのパフォーマンスにどのような悪影響を及ぼすかを測定するのを支援するために設計されたテストプラットフォームである。残りの3つの文書のうち、2つは生成型AIのリスク管理を支援するガイダンス文書であり、多くのチャットボットやテキストベースの画像・動画作成ツールを可能にする技術である生成型AIのリスク管理を支援し、NISTのAIRisk管理フレームワーク（AI RMF）およびセキュアソフトウェア開発フレームワーク（SSDF）の補完的なリソースとなることを目的としている。3つ目は、米国の利害関係者が世界中の人々と協力してAIの標準化に取り組むための計画を提案している。 NTIAがまもなく発表する報告書では、広く利用可能なモデルウェイト（「オープンウェイトモデル」）を持つデュアルユース基礎モデルのリスクとメリットを検証し、リスクを軽減しながらそのメリットを最大限に引き出す政策提言を策定する。オープンウェイトモデルにより、開発者は過去の成果を基に改良を加えることができ、小規模企業、研究者、非営利団体、個人などへのAIツールの普及が拡大する。	U.S. Department of Commerce https://www.commerce.gov/news/press-releases/2024/07/departments-commerce-announce-new-guidance-tools-270-days-following

【人工知能(AI)】関連記事詳細 (5/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
5	国際	AIの導入の勢いはインドと中国で最も強い	2024/7/30	BSIが新たに発表した「グローバルAI成熟度モデル」によると、中国とインドの企業はAIを業務にスムーズに統合し、従業員の業務上の変化への対応を準備し、AIを良い方向に活用する準備が整っている。一方、英国、日本、オランダは、投資、トレーニング、サプライヤーの関与などの分野で、より大きな進歩を遂げる必要がある。 BSIのモデルは、世界中の企業におけるAI導入に対する組織の自信と準備態勢を含む一連の指標を評価し、重み付けを行い、単一のスコアを算出する。インドが最もAI成熟度の高い市場であると特定し、中国（4.25）を上回る4.58のスコアを付けた。9カ国7セクターの932人のビジネスリーダーからの洞察に基づき、投資、トレーニング、社内外のコミュニケーション、安全性などに関する姿勢や行動を指標として採用した。BSIの「AIへの信頼」レポートの一部として発表されたこの分析では、英国と日本が他国と比較して成熟度が低いことが明らかになった。その要因として、政策の方向性や、機会よりもリスクに焦点を当てたメディアの論調などが影響している可能性がある。すべての指標において、中国とインドが他国をリードしており、米国が3位、オーストラリアがそれに続いている。 全体的に見ると、AI への企業の関与は高いが、大きな差異も見られる。中国とインドでは、自社のビジネスが AI の利用を奨励していると答えた割合がそれぞれ 96% と 94% に達しているのに対し、日本では 40%、英国では 65% にとどまっている。日本では、自社のビジネスが AI のメリットを活用できる能力に対する自信についても、同様の傾向が見られる（中国 96% に対して日本 50%）。大企業の方が中小企業よりも、自社の事業がAIの利用を奨励していると答える割合が高く（84%対67%）、AIの利用能力についてもより自信を持っている（89%対76%）。 セクター間には明確な違いがある。医療セクターでは、雇用主が現在AIに投資していないと答えた人が40%に上ったが、テクノロジー関連の職務ではわずか4%だった。これは、医療セクターではAIが効率性と生産性を向上させるという楽観的な見方が強い（62%）のに対し、運輸（51%）、小売（53%）、農業（46%）のリーダーはより慎重な見方をしていることとは対照的である。	英国規格協会 (BSI) https://www.bsigroup.com/en-GB/insights-and-media/media-centre/press-releases/2024/july/momentum-of-ai-adoption-strongest-in-india-and-china-finds-bsi/

【人工知能(AI)】関連記事詳細 (6/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
6	欧州	欧州AI法が発効	2024/8/1	本日、人工知能に関する世界初の包括的な規制である欧州人工知能法（AI法）が施行される。AI法は、EUで開発・利用されるAIが信頼に足るものであることを保証することを目的としており、人々の基本的権利を保護するための安全対策が盛り込まれている。この規制は、EUにおけるAIの調和のとれた域内市場の確立を目指しており、この技術の普及を促進し、イノベーションと投資を支援する環境を創出することを目的としている。 AI法では、EUにおける製品安全とリスクベースのアプローチに基づき、AIの将来を見据えた定義が導入されている。 AI対応のレコメンデーションシステムやスパムフィルターなど、ほとんどのAIシステムはこれに該当する。これらのシステムは、市民の権利や安全に対するリスクが最小限であるため、AI法に基づく義務は課されない。企業は自主的に追加の行動規範を採用できる。 加盟国は2025年8月2日までに、AIシステムの規則の適用を監督し、市場監視活動を実施する国内の管轄当局を指定しなければならない。欧州委員会のAI局は、EULEVELでのAI法の実施の主要機関となるほか、汎用AIモデルの規則の施行機関ともなる。3つの諮問機関が規則の実施を支援する。欧州人工知能委員会は、EU加盟国全体でAI法が統一的に適用されることを確保し、欧州委員会と加盟国間の協力の主要機関としての役割を果たす。独立した専門家で構成される科学パネルは、執行に関する技術的な助言や意見を提供する。特に、このパネルは汎用AIモデルに関連するリスクについてAIオフィスに警告を発することができる。AIオフィスは、多様な利害関係者で構成される諮問フォーラムからの指導を受けることもできる。 規則に準拠しない企業には罰金が科せられる。禁止されたAIアプリケーションの違反に対しては最大で年間世界売上高の7%、その他の義務の違反に対しては最大で3%、不正確な情報の提供に対しては最大で1.5%の罰金が科せられる可能性がある。 AI法の規則の大半は2026年8月2日に適用が開始される。ただし、容認できないリスクがあるとみなされたAIシステムの禁止は6か月後から、いわゆる汎用AIモデルに関する規則は12か月後から適用される。完全実施までの移行期間を埋めるため、欧州委員会はAI協定を開始した。このイニシアティブは、AI開発者に対して、法的な期限に先立ってAI法の主要な義務を自主的に採用するよう呼びかけている。	欧州委員会 https://ec.europa.eu/commission/presscorner/detail/en/ip_24_4123

【人工知能(AI)】関連記事詳細 (7/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
7	イギリス	BSI、エジンバラ大学とAIタレントアカデミーを立ち上げる	2024/8/21	エジンバラ大学の学生を対象に、人工知能（AI）研究を支援する新たな奨学金が授与された。この奨学金は、BSIが同大学と提携し、未来のAI専門家を育成する一環として学生に提供される。 AI タレント・アカデミー・スキームの一環として、エジンバラ大学に通う学生5名に奨学金が授与された。奨学金は授業料の全額と生活費をカバーするもので、2024年9月に始まるエジンバラ・フューチャーズ・インスティテュートのAIとデータ倫理の修士課程、およびコンピュータサイエンスと認知科学の学部課程AI関連コースの学生に授与される。 AI奨学金のうち3つは、Cowrie Scholarship Foundation（CSF）との提携の一環として、社会経済的に厳しい環境にあるブラックアフリカおよびカリブ海地域の遺伝的背景を持つ学部生に提供される。ビジネス改善と標準化企業BSIの規制サービス部門が主導するアカデミーでは、5人の学生全員が卒業後にBSIに入社できる機会も提供される。 AI タレントアカデミーは、エジンバラ大学で開催されたイベントで発足し、AI の専門家、カウリ奨学財団の代表者、両者が一堂に会した。	英国規格協会（BSI） https://www.bsigroup.com/en-GB/insights-and-media/media-centre/press-releases/2024/august/bsi-launches-ai-talent-academy-with-the-university-of-edinburgh/
8	中国/ロシア	中国とロシア、二国間経済・貿易協力の拡大へ	2024/8/22	中国とロシアは、二国間の経済・貿易協力の拡大に向けて協力していくと、水曜日に発表された第29回中国・ロシア政府首脳会談の共同声明で述べられた。 中国の李克強首相とロシアのミハイル・ミシュスティン首相が共同議長を務めた。文書によると、双方は貿易構造の最適化、両国の経済および二国間貿易量における新たな成長点の創出、電子商取引の発展促進に向けた共同の取り組みを行うことで合意した。共同声明によると、双方は2025年にロシアで開催される第9回中国・ロシア博覧会と、同博覧会の枠組み内で行われる第5回中国・ロシア地域間協力フォーラムの開催を支持する。 BRICS諸国が人工知能分野における能力構築、共同研究、グローバル・ガバナンスで協力することを支援することでも合意した。	中国人民政府 https://english.www.gov.cn/news/2024/08/22/content_W566c6338cc6d0868f4e8ea23c.html

【人工知能(AI)】関連記事詳細 (8/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
9	欧州	欧州評議会、AIに関する史上初の国際条約に署名開始	2024/9/5	欧州評議会枠組み条約（人工知能と人権、民主主義、法の支配）（CETS No. 225）は、欧州評議会法務大臣会議（ヴィリニウス）で署名のために開放された。これは、AIシステムの利用が人権、民主主義、法の支配と完全に一致することを保証することを目的とした、史上初の国際的な法的拘束力のある条約である。 この枠組み条約には、アンドラ、ジョージア、アイスランド、ノルウェー、モルドバ共和国、サンマリノ、英国、そしてイスラエル、米国、欧州連合（EU）が署名した。 この条約は、AIシステムのライフサイクル全体をカバーする法的枠組みを提供する。人権、民主主義、法の支配に対するAIのリスクを管理しながら、AIの進歩とイノベーションを促進する。時代に耐えうるものとするため、この条約はテクノロジー中立である。 この枠組み条約は、2024年5月17日に欧州評議会の閣僚理事会によって採択された。欧州評議会の46加盟国、欧州連合、および11の非加盟国（アルゼンチン、オーストラリア、カナダ、コスタリカ、バチカン市国、イスラエル、日本、メキシコ、ペルー、米国、ウルグアイ）が条約の交渉を行った。民間部門、市民社会、学術界の代表者もオブザーバーとして参加した。 条約は、欧州評議会の加盟国3カ国以上を含む5カ国が批准してから3カ月が経過した翌月の初日に発効する。条約への加入と条項の順守を誓うことができるのは、世界各国である。	council of europe https://www.coe.int/en/web/portal/-/council-of-europe-opens-first-ever-global-treaty-on-ai-for-signature

【人工知能(AI)】関連記事詳細 (9/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)	
10	アメリカ	AIの安全性を支援する米国AI安全研究所、AnthropicとOpenAIと協定を締結し、共同研究を可能に	2024/9/9	米国商務省国立標準技術研究所（NIST）の米国人工知能（AI）安全研究所は最近、安全で信頼性の高いAIの推進を目的として、Anthropic社およびOpenAI社と協定を締結したことを発表した。これにより、米国AI安全研究所は、各社が公開する主要な新モデルについて、公開前および公開後にアクセスできる枠組みが確立された。 NISTは、各社の覚書により、能力と安全リスクの評価方法、およびそれらのリスクを軽減する方法についての共同研究が可能になることを報告している。 2023年11月に発足したAI Safety Instituteは、AIの安全、セキュリティ、信頼性の高い開発と利用に関するバイデン政権の行政命令に対するNISTの対応の一部である。2024年5月に発表された戦略ビジョンを読む。 OpenAI（ChatGPTのメーカー）とAnthropicは、AI Safety Institute Consortiumのメンバーである。AI Safety Institute ConsortiumはAI Safety Instituteの一部であり、AIシステムが安全で信頼できることを保証するために、政府機関、企業、影響を受けるコミュニティ間の緊密な連携を可能にすることを目的としている。 また、NISTは、米国AI安全研究所が英国AI安全研究所のパートナーと緊密に連携しながら、AnthropicとOpenAIのモデルの潜在的な安全性向上について両社にフィードバックを提供することを計画していると報告している。	ANSI	https://www.ansi.org/standards-news/all-news/2024/09/9-9-24-us-ai-safety-institute-signs-agreements-with-anthropic-and-openai
11	中国	中国、AIに関するセキュリティガバナンスの枠組みを公表	2024/9/10	中国南部の広東省の省都、広州で開催中の今年の中国サイバーセキュリティ・ウィークのメインフォーラムで、人工知能（AI）のセキュリティガバナンスに関する枠組みが月曜日に発表された。 中国国家標準化管理委員会が設立した「サイバーセキュリティに関する国家技術委員会260」が発表したこの枠組みは、安全性を確保するための融和性と慎重さの維持、迅速なガバナンスのためのリスク管理、共同ガバナンスのための協力関係の開拓など、AIのセキュリティ管理における原則を特徴としている。 AIの技術的特性を踏まえ、同フレームワークではAIに関するリスクの発生源と形態を分析し、AIの応用におけるセキュリティリスクを対象とした技術的対応と包括的な予防・抑制措置を提示するとともに、AIの安全な開発と応用に関する指針を定めている。 技術委員会の関係者は、「この枠組みは、AIのセキュリティガバナンスへの社会参加と進展を促進し、AIの開発と応用における安全で信頼性が高く、公平かつ透明性の高い環境の構築に役立つだろう」と述べた。	中国人民政府	https://english.www.gov.cn/news/202409/10/content_W566df9f30c6d0868f4e8eac91.html

【人工知能(AI)】関連記事詳細 (10/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)	
12	国際	LLMで論理的に考えることを学ぶ	2024/9/12	当社は、回答を生成する前により多くの時間をかけて思考するように設計されたAIモデルの新シリーズを開発しました。新モデルは、科学、コード生成、数学の分野で、前身モデルよりも複雑なタスクを通して推論し、難解な問題を解決できます。 ChatGPTと当社のAPIでは、このシリーズの最初のモデルが現在公開されています。これはプレビューで、定期的なアップデートと改良を提供予定です。この新モデルの公開と共に、現在開発中の次のアップデート向けの評価も含めています。 当社は、人間が何かを問われた時にそうするように、問題に対する回答を生成する前に問題についてより長い時間をかけて考えるよう、新モデルを訓練しました。訓練を通してモデルは、思考プロセスを改善し、さまざまな戦略を試し、自らの誤りを認識することを学習しています。 自社試験において、次のモデルのアップデートは、物理、科学、生物の分野でベンチマークタスクに挑戦する博士課程の学生と類似のパフォーマンスを発揮しています。さらに、数学とコード生成においても優れた能力を発揮していることが認められました。国際数学オリンピック（IMO）の予選試験において、GPT-4oが正解したのは問題のわずか13%でしたが、推論モデルは83%のスコアを達成しました。この推論モデルのコード生成能力は、各種コンテストで評価され、Codeforcesのコンテストでは、89パーセンタイルを達成しました。詳細については、当社の技術研究に関する投稿をご覧ください。 初期モデル同様、情報を入手することを目的としたウェブブラウジングやファイルと画像のアップロードなど、ChatGPTを有用なものにする機能はまだそこまでたくさん搭載されていません。多くの一般的なケースにおいて、GPT-4oは近い将来、今以上に有能なものとなる予定です。 ただし複雑な推論タスクに関してこれは大きな進歩であり、AIの新たなレベルの能力を示すものです。以上を踏まえ、当社はカウンターを1にリセットし、このシリーズを「OpenAI o1」と名付けました。	Open AI	https://openai.com/ja-introducing-openai-o1-preview/

【人工知能(AI)】関連記事詳細 (11/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
13	国際	オラクル、業界初となるゼタスケールのクラウド・コンピューティング・クラスタを発表	2024/9/13	オラクルは本日、NVIDIA Blackwellプラットフォームによって高速化された、業界初のゼタスケールのクラウド・コンピューティング・クラスタを発表しました。「Oracle Cloud Infrastructure (OCI)」は現在、最大131,072基のNVIDIA Blackwell GPUを搭載可能なクラウド最大クラスのAIスーパーコンピューターを受注開始しています。 OCIは現在、クラウドで最大クラスのAIスーパーコンピューターを受注開始しています。このスーパーコンピューターは、最大131,072基のNVIDIA Blackwell GPUを搭載可能で、これまでにない2.4ゼタFLOPSのピーク性能を実現します。「OCI Supercluster」は、最大でFrontierスーパーコンピューターの3倍以上のGPU数、他のハイパースケーラーの6倍以上のGPU数を提供します。 「OCI Supercluster」は、NVIDIA H100/H200 Tensor Core GPU、またはNVIDIA Blackwell GPUを搭載した「OCI Compute」とともに注文可能です。H100 GPU搭載の「OCI Supercluster」は、最大16,384 GPUまで拡張可能で、最大65エクサFLOPSの性能と13Pb/sの集約ネットワーク・スループットを実現します。H200 GPU搭載の「OCI Supercluster」は、最大65,536 GPUまで拡張可能で、最大260エクサFLOPSの性能と52Pb/sの集約ネットワーク・スループットを実現し、今年後半に提供開始予定です。NVIDIA GB200 NVL72を搭載した液冷ベアメタル・インスタンスの「OCI Superclusters」は、NVLinkおよびNVLink Switchを使用して、単一のNVLinkドメイン内で最大72基のBlackwell GPUが互いに通信できるようにし、129.6 TB/sの合計帯域幅を実現します。2025年前半に提供開始予定のNVIDIA Blackwell GPUは、第5世代のNVLink、NVLink Switch、およびクラスタ・ネットワーキングにより、単一のクラスタ内でシームレスなGPU-GPU通信を可能にします。 AIファーストのコラボレーション・プラットフォームのリーダー企業であるZoomは、追加料金なしで利用できる同社のAIパーソナル・アシスタントであるZoom AI Companionの推論処理を実行するためにOCIを活用しています。Zoom AI Companionは、電子メールやチャット・メッセージの下書きの作成、会議やチャット・スレッドの要約の作成、ブレインストーミング中のアイデアの生成などを支援します。OCIのデータおよびAI主権機能は、同社が顧客データを地域内に保持し、OCIのソリューションが最初に展開されるサウジアラビアにおけるAI主権要件をサポートするのに役立ちます。	ORACLE https://www.oracle.com/jp/news/announcement/ocw24-oracle-offers-first-zettascale-cloud-computing-cluster-2024-09-11/

【人工知能(AI)】関連記事詳細 (12/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
14	中国	中国の科学者、地理科学における世界初のマルチモーダル大規模言語モデルを発表	2024/9/20	地理科学のマルチモーダル大言語モデル（LLM）が木曜日に北京で発表された。これは世界初の試みであり、地理学と人工知能の統合を支援し、地理学の発見を加速させる可能性がある。 シグマ地理学と名付けられたこのモデルは、中国科学院の地理科学・自然資源研究機構（IGSNRR）、チベット高原研究機構、オートメーション研究所の研究者チーム、およびその他の組織によって開発された。シグマ地理は、地理に関する専門的な質問に答え、地理に関する記事を分析し、地理データの照会や詳細な分析を行い、テーマ別マップを作成することができると、IGSNRRの副所長である蘇鳳珍氏は述べた。一般的なLLMと比較すると、シグマ地理は地理分野における言語パターン、専門用語、専門知識をより深く理解しており、専門的な問題にもより適切に対応できると蘇氏は述べた。 地理に関する質問に答えるだけでなく、シグマ地理は生成されたテキスト回答を地理的な風景写真、テーマ別地図、または模式図とマッチングさせることもでき、ユーザーがテキスト回答をより視覚的かつ想像的に理解する手助けとなる、と彼は付け加えた。 シグマ地理を基に研究チームが開発した研究アシスタント機能は、ユーザーの指示に従って概念の理解、データ取得、情報分析、マッピングなどのプロセスを完了し、最終的にユーザーが必要とする専門的な地理図表を生成することができる。 このモデルは、地理科学に対する一般の人々の理解を深めるのに役立つ可能性があり、地理科学における重要な発見を明らかにするための学術研究や調査を支援することができる。 これまでに、Sigma Geographyは、Nature、The Innovation、Earth’s Futureなどの学術誌に掲載された10以上の論文で引用されている。 チームはSigma Geographyの改良を継続し、地図の理解を可能にし、科学者や研究チームがデータ、モデル、研究アイデアを共有することで共同作業を円滑に行うためのプラットフォームの構築を目指している。	中国科学院 https://english.cas.cn/newsroom/cas-media/20240920_689827.shtml

【人工知能(AI)】関連記事詳細 (13/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)	
15	国際	国連、グローバルなAIガバナンスの枠組み案を発表、標準の重要な役割を強調	2024/9/23	人工知能（AI）を人類のために統治するにはどうすればよいのか？ 国連事務総長直属のAIに関するハイレベル諮問委員会は、今月発表された報告書の中で、調整力を強化し、グローバルなAIリスクに対処するための主要な提言として、AI標準の交換を提案している。 「Governing AI for Humanity（人類のためにAIを統治する）」は、世界中の地域から2,000人以上の参加者の洞察に基づいており、AIの変革の可能性を強調し、新たな国際的なAIガバナンス機関の形成を導く原則を概説している。 特に、この報告書では、人権の保護と発展を促すために、すべての政府および関係者がAIの管理において協力することなどを含め、AI関連のリスクに対処するための青写真を概説している。グローバルなAIガバナンスに関する主な提言の中で、諮問委員会は、報告書の55ページに記載されているように、AI標準の交換を推奨している。 枠組みに関する提言では、現在のAIガバナンス体制におけるギャップに対処するためのさらなる調整が提案されている。報告書（2023年12月に発表された中間報告書に基づく）を作成するために、諮問機関は「AIリスクグローバルパルスチェック」（AIリスクに関する最も包括的なグローバルなホライズンズキャンニングの取り組み）と「AIオポチュニティスキャン」（AIの新たなトレンドに関する専門家の評価をクラウドソーシングする）を委託したと、国連は報告している。	ANSI	https://www.ansi.org/standards-news/all-news/2024/09/23-24-un-releases-proposed-framework-for-global-ai-governance
16	中国	2024年世界インターネット会議・烏鎮サミットはAIに焦点を当てる	2024/10/1	11月に開催予定の2024年世界インターネット会議（WIC）烏鎮サミットでは、人工知能（AI）に焦点が当てられる予定であると、主催者は発表した。 このサミットでは、グローバル開発イニシアティブ、デジタルとグリーン開発の調整、デジタル経済、AI技術革新とガバナンスなど、さまざまなトピックについて20以上のサブフォーラムが開催される予定である。 主催者によると、WICはサミット中にAIに関する特別委員会を設置する予定である。また、今年のイベントでは、世界のインターネット分野および関連分野に多大な貢献をした個人や企業を表彰する「WIC 特別貢献賞」など、新たな活動も予定されている。サミットは11月19日から22日まで、中国東部の浙江省の烏鎮で開催される予定である。	中国人民政府	https://english.www.gov.cn/news/2024/10/01/content_WS66fbf90ec6d0868f4e8eb77c.html

【人工知能(AI)】関連記事詳細 (14/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
17	中国	中国、AI搭載ロボットの開発を促進するために力を結集	2024/10/11	中国は、新しい質の高い生産力の台頭を推進する取り組みの一環として、身体性人工知能（AI）ロボットの開発を促進するために、中央政府と地方政府間の協力を強化している。 木曜日、北京身体化人工知能ロボットイノベーションセンター（HUMANOID）が、中国全国・地方共同身体化人工知能ロボットイノベーションセンターに格上げされた。これは、人型ロボットに代表される身体化AIロボットの研究開発やその他の分野において、北京と中央政府の両方から支援を受けることを意味する。HUMANOIDは2023年11月、首都の南東に位置する経済技術開発区である北京Eタウンに設立された。これは、身体化AIロボットのコア技術、製品開発、アプリケーションエコシステムに焦点を当てた中国初のイノベーションセンターである。 同センターは、すでに2つのヒューマノイドロボット製品を展開しており、今後は業界標準の確立を推進し、企業の研究開発コストをさらに削減し、応用シナリオを拡大し、関連製品の発売を早めるなど、さまざまな課題に取り組んでいくと述べた。 同センターの発足に伴い、基礎技術や重要共通技術に関するイノベーションの加速、科学技術成果の転換と応用を促進し、重点産業のインテリジェント化と高度化を推進するための共同開発が行われる予定である。工業情報化部のジン・ジュアンロン部長は、同センターの発表セミナーでこのように述べた。 工業情報化部の指針によると、中国は2025年までに人型ロボットの初期のイノベーションシステムを確立することを目指している。2027年までに、中国は安全で信頼性の高い産業およびサプライチェーンシステムを構築し、関連製品は実体経済に深く統合されることになる。 今年に入ってから、浙江省、山東省、安徽省、四川省などの製造業が盛んな地域では、産学研究応用一体化のイノベーションセンターを設立することで、国家の発展に追いついた。5月には、国家と地方が共同で設立したヒューマノイドロボットのイノベーションセンターが上海に設立された。	中国人民政府 https://english.www.gov.cn/news/202410/11/content_WS67092284c6d0868f4e8ebb88.html

【人工知能(AI)】関連記事詳細 (15/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
18	国際	AI標準化に向けた取り組み：ISO、IEC、ITUが2025年国際AI標準サミットを発表	2024/10/15	先ごろ開催された世界標準の日（World Standards Day）において、国際標準化機構（ISO）、国ANSI 際電気標準会議（IEC）、国際電気通信連合（ITU）は、人工知能（AI）に関する課題に取り組む ことを目的とした共同イニシアティブ「2025年国際AI標準サミット（International AI Standards Summit 2025）」を発表した。このイニシアティブは、AIの安全性、透明性、包括性を促進する標準の 開発と維持に重点的に取り組むことを目的としている。 2025年国際AI標準サミットは、韓国技術標準院（KATS）の主催により、2025年12月2日～3日に ソウルで開催される。 このイニシアティブは、9月に世界のリーダーたちによってグローバル・デジタル・コンパクトが採択されたことを受 けたものであり、国際標準を通じてガバナンスを推進するよう国連が呼びかけたことへの直接的な回答である。 国連は最近、ハイレベル諮問機関の報告書「人類のためのAIガバナンス」を発表し、標準化の連携の重要 性を強調し、その提言の中でAI標準の交換を求めている。 このイニシアティブの発表の席上、ISO事務局長のセルジオ・ムヒカ氏は、国際標準規格を通じて効果的な AIガバナンスを実現するには、協調的なアプローチが不可欠であると強調した。「国際標準規格を協調的に 採用することは、AIの責任ある利用の未来を確保する上で重要な役割を果たす」とムヒカ氏は述べた。「標 準規格は、グローバルなガバナンスが不可欠な政策目標を支援し、有益なシステムや慣行の普及を促進し、 先進的なAI技術の効率的な開発を促進することができる。」	https://www.ansi.org/standards-news/all-news/2024/10/10-15-24-advancing-ai-standards-collaboration-iso-iec-and-itu-announce-ai-standards-summit

【人工知能(AI)】関連記事詳細 (16/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
19	欧州	KPMGオーストラリアは、BSIのAI認証を最初に取得した企業である。	2024/10/16	KPMGオーストラリアが、人工知能（AI）の安全な管理を強化し、社会全体におけるAIの安全かつ責任ある利用を可能にする信頼を構築することを目的とした、世界的に適用可能な規格の認証をBSIから取得した世界初の組織として発表された。 国際規格である「情報技術 - 人工知能 - マネジメントシステム（BS ISO/IEC 42001）」は、AIを責任を持って使用する組織を支援することを目的としており、透明性のない自動意思決定、システム設計における人間がコード化したロジックではなく機械学習の活用、継続的な学習といった課題に対処する。 英国規格協会（BSI）が1年未満前に英国で発表したこのガイダンスは、AI管理システムの確立、実施、維持、継続的な改善の方法を、保護策に重点を置いて定めている。これは影響に基づく枠組みであり、AI製品およびサービス（社内および社外）のリスク処理と管理の詳細とともに、状況に基づくAIリスク評価を促進するための要件を規定している。このフレームワークは、組織が品質重視の文化を導入し、AIを搭載した製品やサービスの設計、開発、提供において責任ある役割を果たし、組織や社会全体に利益をもたらすことを目的としている。 ISO 42001規格の発表を受け、KPMGオーストラリアは2024年2月、BSI認証プロセスに向けた事業推進チームを立ち上げた。これには、ガバナンスからプロジェクト管理、KPMGオーストラリアの社員の業務への信頼できるAIの実践の組み込みに至るまで、ビジネスがAIにどのようにアプローチしているかについての詳細な監査が含まれていた。ISO規格は最低限の要件を満たすだけでなく、継続的な改善の文化を根付かせることでもある。	英国規格協会（BSI） https://www.bsigroup.com/en-GB/insights-and-media/media-centre/press-releases/2024/october/kpmg-australia-is-the-first-organization-globally-to-achieve-bsi-ai-management-system-certification/

【人工知能(AI)】関連記事詳細 (17/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名/URL)
20	欧州	消費者はAIには親しみを感じているが、量子技術についてはあまり知識がない。	2024/10/21	世界1万人の成人を対象とした調査で、2023年以降、自身の仕事のいくつかの部分をAIが代替できると考える人の数が13%増加したことが判明した。量子コンピューティングがもたらす機会がリスクを上回ると感じているのはわずか38%にとどまった。 英国エディンバラで開催された約1,200人の技術専門家による国際会議に先立ち、BSIがデータを公表。BSIによる新たな調査では、AIの能力に対する一般の人々の信頼が高まっていることが明らかになった。回答者の半数以上（51%）が、AIが自分の仕事のいくつかの側面を遂行できると回答しており、2023年の38%から増加している。同様に、53%がAIが自分の仕事の単純な部分を遂行できると感じており、16%から増加している。AIに対する信頼感が高まっているにもかかわらず、世界全体では5人中3人以上（62%）が、AIツールの問題や不正確な点を指摘するための標準化されたシステムが必要だと考えており、57%がAI駆動型テクノロジーとのやり取りに際してプライバシー上の懸念を抱いている。 今回の調査では、導入に関する懸念もいくつか指摘されているが、AIの能力については楽観的な見方が示された。対象となった4か国[1]の労働者のほぼ半数（49%）が、2050年までにAIと「同僚」として協働することを期待しており、その期待はインド（62%）で最も高く、英国（34%）で最も低かった。これは、2055年までに自動運転車が道路を走るようになることを期待する人々（42%）よりも高い割合であった。このデータは、このような社会がもたらすであろうものに対する人々の認識を理解することを目的とした報告書『Innovating for our future: AI, quantum and an all-electric and connected society [2]』の一部として発表された。これは、国際社会がAIやテクノロジーなどの分野における標準規格について議論する今年の国際電気標準会議（IEC）総会に先立って発表された。1985年以来初めて英国のBSIが主催し、10月21日から25日までエディンバラで開催される。学术界、産業界、政府から約1,200人の専門家が集まり、最先端技術を管理する技術標準について議論する。 この調査では、AIに対する信頼感が高まっている兆候も見られたが、世界中の消費者は、量子コンピューティングなどの他の先進技術にはあまり詳しくないことが分かった。スーパーコンピューティングの機会とリスクについて、政府や専門家が積極的に十分なコミュニケーションを行っていると感じていると答えたのはわずか40%で、英国では5分の1（21%）に減少した。また、スーパーコンピューターや量子コンピューターの可能性がリスクを上回ると答えたのは38%にとどまり、英国では24%、ドイツでは30%に減少した。世界全体では34%が、スーパーコンピューターの普及拡大によりコンピューター技術への依存度が高まることを懸念しており、29%が、世界全体の二酸化炭素排出量が劇的に増加し、その結果、気候変動の影響が大きくなると懸念していることが分かった。	英国規格協会（BSI） https://www.bsigroup.com/en-GB/insights-and-media/media-centre/press-releases/2024/october/consumers-feel-increasingly-familiar-with-ai-but-less-informed-on-quantum-technology/

【人工知能(AI)】関連記事詳細 (18/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
21	アメリカ	バイデン＝ハリス政権、米国の国家安全保障におけるAIの力を活用するための協調的アプローチを概説	2024/10/24	本日、バイデン大統領は人工知能（AI）に関する初の国家安全保障覚書（NSM）を発令する。NSMの基本的な前提は、AIの最先端における進歩が近い将来、国家安全保障および外交政策に重大な影響を及ぼすというものである。NSMは、AIの安全、セキュリティ、信頼性の高い開発を推進するために大統領および副大統領が講じた主要な措置を基にしており、その中には、AIの可能性を最大限に引き出し、リスクを管理する上で米国が主導的な役割を果たすことを確実にするためのバイデン大統領による画期的な大統領令も含まれている。 特に、NSMは、以下の重要な行動を指示している。 ・米国が安全でセキュアかつ信頼性の高いAIの開発を世界的に主導することを確実にする ・米国政府が、人権と民主的価値を守りつつ、国家安全保障の目標を達成するために、最先端のAIを活用できるようにする： ・AIに関する国際的な合意とガバナンスの推進： 本日発表されたNSMは、バイデン＝ハリス政権の責任あるイノベーションに関する包括的な戦略の一環であり、バイデン大統領とハリス副大統領がこれまでに実施してきた施策を基盤としている。	THE WHITE HOUSE https://www.whitehouse.gov/briefing-room/statements-releases/2024/10/24/fact-sheet-biden-harris-administration-outlines-coordinated-approach-to-harness-power-of-ai-for-u-s-national-security/

【人工知能(AI)】関連記事詳細 (19/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
22	アメリカ	懸念国における特定の国家安全保障技術および製品に対する米国の投資への対応	2024/10/28	国境を越えた投資の流れと米国の開放的な投資政策は、米国の経済活力に貢献している。しかし、懸念国は、米国からの特定の海外投資を、米国の国家安全保障上の利益を損なうような機微技術や製品の急速な発展を加速させるような方法で利用している。バイデン＝ハリス政権は、懸念国（すなわち中華人民共和国）が軍事近代化に不可欠な主要技術の進歩を遂げることを阻止することで、米国の安全を確保することに尽力している。 本日、米国財務省は、2023年8月9日付のバイデン大統領の行政命令第14105号「懸念国における特定の国家安全保障技術および製品への米国投資への対応」を実施するための最終規則を発行した。最終規則は、その意図と適用に関する運用規則と詳細な説明を提供している。 大統領令の指示に従い、最終規則は、米国に対して特に深刻な国家安全保障上の脅威をもたらす特定の技術および製品に関する一定の取引への米国人（U.S. person）の関与を禁止している。また、最終規則は、米国の国家安全保障に対する脅威につながる可能性がある特定の技術および製品に関する一定の取引について、米国人（U.S. person）が財務省に通知することを義務付けている。 対象となる技術は、半導体およびマイクロエレクトロニクス、量子情報技術、人工知能の3つのカテゴリーに分類される。この限定的な技術セットは、次世代の軍事、サイバーセキュリティ、監視、および情報アプリケーションの中核となるものである。 米国はすでに、最終規則の対象となる技術や製品の多くについて、懸念国への輸出を禁止または制限している。このプログラムは、米国の投資が懸念国における機微技術や製品の進展を促すことを防ぐことで、米国の既存の輸出規制および輸入審査ツールを補完するものである。 本日の発表は、数百の利害関係者、超党派の連邦議会議員、業界関係者、および外国の同盟国やパートナーとの広範かつ徹底的な協議、および一般市民からの2回にわたる正式な意見提出を経たものである。この発表は、2023年8月にバイデン大統領が大統領令に署名した際に発表したプロセスの最終段階である。	THE WHITE HOUSE https://www.whitehouse.gov/briefing-room/statements-releases/2024/10/28/fact-sheet-addressing-u-s-investments-in-certain-national-security-technologies-and-products-in-countries-of-concern/

【人工知能(AI)】関連記事詳細 (20/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
23	アメリカ	オープンソースイニシアティブ、業界初のオープンソースAI定義のリリースを発表	2024/10/28	1年間にわたるグローバルなコミュニティ設計プロセスを経て、オープンソース定義（OSAID） v.1.0が一般公開された。 バージョン1.0のリリースは、本日、世界中のオープンソースコミュニティが関心を寄せる共通の課題に焦点を当てた業界カンファレンス「All Things Open 2024」で発表された。OSAIは、AIシステムがオープンソースAIと見なされるかどうかを検証するために、コミュニティ主導のオープンかつ公開された評価を行うための基準を提供する。OSAIの最初の安定版は、オープンソースを定義する権威として個人、企業、公共機関から世界的に認知されているOpen Source Initiative（OSI）が主導した、複数年にわたる研究と協力、国際的なワークショップのロードショー、1年間にわたる共同設計プロセスの成果である。	the Open Source Definition https://opensource.org/blog/the-open-source-initiative-announce-s-the-release-of-the-industrys-first-open-source-ai-definition

【人工知能(AI)】関連記事詳細 (21/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
24	イタリア	リスク評価と認定認証の間の人工知能	2024/11/7	紛れもなく話題性のあるトピックがあるとすれば、それは間違いなく人工知能である。 一方ではテクノロジーの新たなフロンティアを象徴し、他方では研究とイノベーションが私たちの社会構造にこれまで以上に厳しい形で課す倫理的な意味合いを象徴するものとして、人工知能は今や誰にとっても議論の対象となる分野である。また、標準化、認証、認定の世界にとっても、決して小さくない問題である。 10月16日にローマのサピエンツァ大学で開催された会議「リスクアセスメントと認定認証の間の人工知能」は、アクレディアが 国立大学間情報学コンソーシアム (CINI) の人工知能・知能システム国立研究所 (AIIS) と共同で主催したもので、多くの示唆を与えてくれた。 議論の中心となったのは、「人工知能システム開発のための技術標準と認定適合性評価」に関する天文台の最新研究であり、AIがイノベーションと生活の質の向上にもたらす機会を探求すると同時に、セキュリティ、プライバシー、公平性、倫理に関する最も重要な問題を浮き彫りにした。 Accredia-CINIの調査は、最近のEU規則2024/1689（いわゆるAI法、人工知能に関する世界初の具体的かつ完全な法的介入）から始まり、AIに関する規制の変遷を探り、3つの具体的なPoC（概念実証、すなわち標準適用のケーススタディ）を詳しく説明している。INAILが関与した後者のケースでは、UNI CEI ISO/IEC 42001:2024（「情報技術-人工知能-マネジメントシステム」）の役割が研究された。この規格は、組織におけるAIマネジメントシステムの構築、実施、維持、継続的改善のための要求事項を規定し、ガイダンスを提供している。この規格に準拠することで、自動化された意思決定プロセスが実施される前に厳格な管理下に置かれることが保証される。 また、UNI標準化研究センターのメンバーであるダニエレ・ジェルンディーノ氏も会議に出席し、「標準化と適合性評価を戦略的に推進することが非常に重要である」と述べた。技術進化の複雑さとスピードを考えると、これは非常に難しい課題であることは承知している」が、G7諸国は「異なる管轄区域で利用できる枠組みを作るために、国際標準化に非常に有益な貢献をすることができる」と述べた。 標準化と適合性評価の役割の促進、科学的協力の促進、人間とその価値観を中心とした人工知能を創造するための開発者の意識向上、個人データを保護する基本的権利の保証は、社会が直面する最も緊急な課題である。	UNI https://www.uni.com/intelligenza-artificiale-tra-valutazione-del-rischio-e-certificazione-accreditata/

【人工知能(AI)】関連記事詳細 (22/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名/URL)
25	欧州	独立専門家が執筆した汎用人工知能（AI）行動規範の草案が発表される	2024/11/14	独立した専門家が汎用AI行動規範の最初の草案を発表し、来週、約1000人の利害関係者と議論する予定である。AI事務局は、利害関係者向けに専用の質問と回答を用意し、関連するAI法の理解を促進している。 汎用AI行動規範の反復的な草案作成は、9月30日のキックオフ後、2025年4月までの4回の草案作成ラウンドの最初のラウンドを終え、重要なマイルストーンに到達した。行動規範の最初の草案は、4つのテーマ別作業部会の議長および副議長に任命された独立した専門家によって作成された。欧州AIオフィスは、行動規範の策定を促進する立場として、本日、草案を公表する。専門家らは、汎用AIモデルの提供者を行動規範の対象者として、その貢献を基にこの初期バージョンを策定した。また、草案の策定にあたっては、国際的なアプローチも考慮した。 最終文書は、信頼性が高く安全な汎用AIモデルの今後の開発と展開を導く上で重要な役割を果たす。最終文書では、汎用AIモデルの提供者に関する透明性と著作権に関する規則を詳細に規定すべきである。システムリスクをもたらす可能性のある最先端の汎用AIモデルを提供する少数のプロバイダーについては、システムリスクの分類、リスク評価の尺度、技術的およびガバナンス上の緩和措置についても、行動規範で詳細に規定すべきである。 来週、行動規範全体会議の一環として、議長らは1000名近い利害関係者、EU加盟国の代表者、欧州および国際的なオブザーバーとともに、草案について専門のワーキンググループ会議で議論する。4つのワーキンググループのうち1つが毎日会議を行い、それぞれの議長が最近の草案作成の進捗状況について最新情報を提供する。利害関係者の中からバランスの取れたステークホルダーが招待され、口頭での意見を述べる。並行して、本会議の参加者は、専用プラットフォーム（Futurium）を通じて草案を受け取り、11月28日（木）12:00 CET（中央ヨーロッパ時間）までに書面によるフィードバックを提出する。このフィードバックに基づき、議長は、行動規範にさらに詳細を追加しながら、第1草案の対策を調整する可能性がある。起草の原則では、措置、副措置、KPIはリスクに見合ったものでなければならず、汎用AIモデルプロバイダーの規模を考慮し、中小企業や新興企業には簡素化されたコンプライアンスオプションを認めるべきであると強調している。AI法に続いて、本規範ではオープンソースモデルのプロバイダーに対する重要な免除事項も反映される。	欧州委員会 https://digital-strategy.ec.europa.eu/en/library/first-draft-general-purpose-ai-code-practice-published-written-independent-experts

【人工知能(AI)】関連記事詳細 (23/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
26	アメリカ	米国AI安全研究所、国家安全保障能力とリスクを管理するためのAIモデルの研究とテストで協力する米国政府のタスクフォースを新設	2024/11/20	本日、米国商務省国立標準技術研究所（NIST）の米国人工知能安全研究所は、急速に進化するAI技術が国家安全保障および公共の安全に及ぼす影響を特定、測定、管理するために、米国政府のパートナーを結集した「国家安全保障のためのAIのリスクに関する試験（TRAINS）タスクフォース」の結成を発表した。この発表は、米国がサンフランシスコで史上初の「AI安全研究所国際ネットワーク」の会合を開催する予定であることを受けて行われた。 このタスクフォースは、放射線・核セキュリティ、化学・生物セキュリティ、サイバーセキュリティ、重要インフラ、通常戦力など、国家安全保障および公共安全の重要な領域全体にわたって、高度なAIモデルの協調的な研究と試験を可能にする。 これらの取り組みは、米国政府がAI開発における米国のリーダーシップを維持し、敵対者が米国のイノベーションを悪用して国家安全保障を脅かすことを防ぐという重要な使命を推進するものである。 TRAINSタスクフォースは米国AI安全研究所が議長を務め、以下の連邦政府機関が初期段階の代表として参加している。 国防総省（最高デジタル・人工知能責任者（CDAO）および国家安全保障局を含む） エネルギー省およびその国立研究所10ヶ所 国土安全保障省（サイバーセキュリティ・インフラ保護庁（CISA）を含む） 保健社会福祉省の国立衛生研究所（NIH） 各メンバーは、独自の専門知識、技術インフラ、リソースをタスクフォースに提供し、新たなAI評価方法およびベンチマークの開発、さらには共同での国家安全保障リスク評価やレッドチーム演習の実施において協力する。 この取り組みは、最近の「AIに関する国家安全保障覚書（National Security Memorandum on AI）」で指示されたとおり、AIの安全性に対する政府一体となったアプローチを具体化するものであり、TRAINSタスクフォースは、その活動が継続するにつれて、連邦政府全体にわたってメンバーシップを拡大していくことが期待されている。	NIST https://www.nist.gov/news-events/news/2024/11/us-ai-safety-institute-establishes-new-us-government-taskforce-collaborate

【人工知能(AI)】関連記事詳細 (24/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
27	カナダ	AI安全協会国際ネットワーク：ミッションステートメント	2024/11/21	オーストラリア、カナダ、欧州委員会、フランス、日本、ケニア、大韓民国、シンガポール、英国、米国のAI安全性と評価を促進するAI安全性研究所および政府機関は、2024年11月21日にサンフランシスコに集まり、AI安全性研究所の国際ネットワークを立ち上げた。2024年5月21日に開催されたAIソウルサミットで発表された「AI安全科学に関する国際協力に向けたソウル声明」を承認し、これを基盤として、このネットワークはAIの安全性に関する新たな国際協力の段階を促進することを目的としている。 私たちの研究所および事務局は、AIの安全性を向上させ、政府や社会が高度なAIシステムがもたらすリスクを理解し、そのリスクに対処するための解決策を提案することで、被害を最小限に抑えることを目的とした技術組織である。リスクの緩和にとどまらず、これらの研究所や事務局は、AIシステムの責任ある開発と展開を導く上で重要な役割を果たす。AIは、社会のニーズに応え、人間の幸福、平和、繁栄を向上させる能力という大きな可能性をもたらす一方で、世界的なリスクをもたらす可能性もある。AIの安全性、セキュリティ、包括性、信頼性を促進するための国際協力は、これらのリスクに対処し、責任あるイノベーションを推進し、世界中でAIの恩恵へのアクセスを拡大するために不可欠である。 AI安全研究所国際ネットワークは、世界中の技術的専門知識を集結させるフォーラムとなることを目的としている。文化や言語の多様性の重要性を認識し、私たちは、AIの安全性に関するリスクと緩和策について、私たちの研究所やより広範な科学コミュニティの取り組みを基に、技術的な共通理解を促進することを目指している。これにより、相互運用可能な原則やベストプラクティスの国際的な開発と採用が促進されるだろう。また、AIの安全性に関する一般的な理解とアプローチを世界的に促進し、AIのイノベーションの恩恵が開発のあらゆる段階において各国間で共有されることを目指す。そのため、各AI安全性研究所および事務局のリーダーシップが協力し、AIの安全性に関する研究、試験、指導に関する技術的な整合性を推進するネットワークメンバーとして結集した。	カナダ政府 https://ised-isde.canada.ca/site/ised/en/international-network-ai-safety-institutes-mission-statement

【人工知能(AI)】関連記事詳細 (25/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
28	アメリカ	米国、グローバルパートナーと国際的なAI安全ネットワークを立ち上げ	2024/11/25	米国商務省（DOC）と米国国務省は最近、安全な人工知能イノベーションに関する国際的な協調を推進する取り組みとして、国際AI安全機関の初会合を共同開催した。このイニシアティブは、サンフランシスコで開催された2日間の会議で発足し、合成コンテンツのリスク管理、基礎モデルのテスト、高度なAIシステムに対するリスク評価の実施という3つの重要な分野に焦点を当てることを目的としている。 米国は初代議長およびグローバルネットワークのメンバーを務める。その他の初期メンバーには、オーストラリア、カナダ、欧州連合、フランス、日本、ケニア、韓国、シンガポール、英国が含まれる。 2025年2月にパリで開催されるAIアクションサミットに先立ち、11月20日～21日に開催された会議では、各メンバーのAI安全研究所（または政府支援の科学機関）から技術専門家が集まり、優先すべき作業分野について意見を一致させた。 この広範な取り組みを支援するため、複数の国や組織から多額の寄付金が寄せられ、ネットワークは1100万ドルを超える世界的な研究資金調達を確保した。米国は、USAIDを通じて、本年度、海外のUSAIDパートナー諸国における「安全で責任あるAIの能力構築、研究、展開の強化」に380万ドルを割り当て、合成コンテンツのリスク緩和に関する研究支援も行っている。 また、今月発表された米国商務省標準技術研究所（NIST）の米国人工知能安全研究所は、「国家安全保障のためのAIのリスク試験（TRAINS）タスクフォース」の結成を発表した。このタスクフォースは、急速に進化するAI技術が国家安全保障や公共の安全に及ぼす影響を特定、測定、管理するために、米国政府のパートナーを結集する。TRAINSは、放射線・核セキュリティ、化学・生物セキュリティ、サイバーセキュリティ、重要インフラ、通常戦力など、国家安全保障および公共安全の重要な領域全体にわたって、高度なAIモデルの協調的な研究とテストを可能にする。NISTの報告によると、	ANSI https://www.ansi.org/standards-news/all-news/2024/11/11-25-24-us-launches-international-ai-safety-network-with-global-partners

【人工知能(AI)】関連記事詳細 (26/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
29	中国	研究不正の抑制を目的とした新しい学術AIガイドライン	2024/11/27	中国科学アカデミー（CAS）は火曜日、科学研究におけるデータのねつ造や盗用などの不正行為を減らし、科学の信頼性を高めるための取り組みの一環として、科学研究における人工知能（AI）の使用に関する新たなガイドラインを発表した。 中国科学院の管理機関が発行したこの文書は、技術者、研究者、学生が透明性、規範性、責任感を持って科学研究を行うよう導くことを目的としていると、中国科学院のウェブサイトは説明している。この動きは、AI技術の急速な発展と広範な応用に対応するものである。AI技術はイノベーションに新たな機会をもたらしたが、科学研究の誠実性に対する重大な課題も提起した。アナリストらは、このガイドラインが科学界におけるAI利用の倫理的環境形成において重要な役割を果たすことが期待されていると指摘した。 ガイドラインでは、学術評価、科学賞の推薦、学術成果の発表に関わる個人が、倫理的にAIを業務で使用し、学術的誠実性を達成することが求められる、と整合性基準が提案された。このガイドラインでは、研究動向の追跡、参考文献の収集、資料の整理にAIを使用することは可能だが、その技術を使用した内容や結果は、そのように明確に表示しなければならない。 倫理の監督を強化し、倫理審査に関する関連法規や規則を改善することは、イノベーションや創造性を制限するものではなく、特定の科学技術活動における倫理的风险を明確にするものである、と趙氏は述べた。人工知能を専攻する中国科学院の大学院生、李開氏は水曜日、グローバル・タイムズ紙に対し、生成型AIツールは学生や研究者の研究生産性を高め、効率性を向上させるのに役立つと語った。しかし、AIは学術的不正行為や盗作を可能にするために利用されるべきではない。このガイドラインは、AIの倫理的な利用文化を促進し、技術的進歩の信頼性を維持するために不可欠であり、関与する研究者の正当な利益を保護しながら、と李氏は付け加えた。	中国科学院 https://english.cas.cn/newsroom/cas-media/202409/t20240913_689462.shtml

【人工知能(AI)】関連記事詳細 (27/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
30	国際	Amazon Nova ファンデーションモデルのご紹介：最先端のインテリジェンスと業界をリードする価格性能	2024/12/3	Amazon Web Servicesは、Amazon Bedrock限定で提供される、最先端のインテリジェンスと業界をリードする価格性能比を実現する次世代の最先端のファンデーションモデルであるAmazon Novaを発表した。 Amazon Nova を使用すると、ほぼあらゆる生成型 AI タスクのコストとレイテンシを削減でき、基盤として構築することで、複雑な文書や動画の分析、チャートや図表の理解、魅力的な動画コンテンツの生成、企業向けワークロードに最適化された幅広いインテリジェンスクラスから高度な AI エージェントの構築が可能になる。このうち、理解モデルは、テキスト、画像、または動画の入力を受け取り、テキスト出力を生成。また、クリエイティブコンテンツ生成モデルは、テキストと画像の入力を受け取り、画像または動画出力を生成する。Amazon Novaモデルには、異なるニーズに対応する3つの理解モデル（まもなく4つ目が追加される予定）が含まれている。 • Amazon Nova Micro • Amazon Nova Lite • Amazon Nova Pro • Amazon Nova Premier（2025年提供開始予定）	Amazon Web Services https://aws.amazon.com/jp/blogs/aws/introducing-amazon-nova-frontier-intelligence-and-industry-leading-price-performance/
31	イギリス カタール	英国とカタール、人工知能の共同開発を促進するプロジェクトを開始	2024/12/5	英国とカタールは、両国に利益をもたらすAIに関する英国・カタール共同研究のロードマップ策定を目指し、人工知能（AI）の共同研究委員会を発足させた。 この共同研究は、英国のクイーン・メアリー大学が主導し、カタールのハマド・ビン・ハリファ大学（HBKU）と提携して実施される。研究は、クイーン・メアリー大学デジタル環境研究所の倫理・テクノロジー・社会教授であり、アラン・チューリング研究所の倫理・責任あるイノベーション研究ディレクターであるデビッド・レスリー教授が主導する このプロジェクトは、両国がAIの分野で成し遂げてきた素晴らしい進歩を基盤とし、AIおよびテクノロジー戦略に沿って、この分野における両国の協力関係を強化するための現実的かつ野心的な方法を特定し、その範囲を明らかにする。生態系の開発、政策および規制、セキュリティ、国際的な関与など、幅広い分野が検討される カタール外務省、カタール通信・情報技術省（MCIT）のAI委員会、カタール研究・開発・イノベーション評議会（QRDI）、在ドーハ英国大使館の協力により企画・開発された この共同事業の発表は、カタール首長の英国公式訪問と時を同じくして行われた。 このプロジェクトは、英国政府の湾岸戦略基金（Gulf Strategy Fund）から資金提供を受けている。同基金は、外務・英連邦・開発省（FCDO）の国際プログラムの一部である。	GOV.UK https://www.gov.uk/government/news/uk-and-qatar-launch-project-to-boost-artificial-intelligence-collaboration

【人工知能(AI)】関連記事詳細 (28/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名/URL)
32	欧州	委員会、AIによる科学の進歩に関するハイレベル円卓会議を開催	2024/12/12	欧州委員会は本日、さまざまな科学分野における人工知能（AI）の応用と、AIの基礎研究における第一人者による円卓会議を開催した。技術主権、セキュリティ、民主主義担当のヘンナ・ヴィルクネン副委員長と起業、研究、イノベーション担当のエカテリーナ・ザハリエヴァ委員が主催したこの円卓会議では、欧州における科学技術の革新と卓越性を加速させるためにAIを活用するためのギャップの特定と道筋の模索に焦点が当てられた。欧州AI研究評議会と欧州の科学者によるAIの活用拡大戦略に関する今後のイニシアティブを踏まえて、専門家の意見は時宜を得たものとなった。 円卓会議の参加者は、次の3つの主要なトピックについて議論した。 研究におけるAIの統合の促進：科学分野におけるAIの普及を可能にするための重要な課題と機会の特定。 AI研究の推進：AI研究における根本的な問題に対処し、欧州がこの分野における世界のリーダーとなることを目指す。 リソースの共有：欧州のリソースを統合し、研究能力の向上と協力関係の促進を図るための枠組みを構築する。	欧州委員会 https://research-and-innovation.ec.europa.eu/news/commission-hosts-high-level-roundtable-advancing-science-ai-2024-12-12_en

【人工知能(AI)】関連記事詳細 (29/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名/URL)
33	イギリス	新しい記録は、AIが政府による迅速かつ正確な意思決定を支援し、貿易の促進、対応の迅速化など、さまざまな成果をもたらしていることを詳細に説明している。	2024/12/17	新しい記録は、市民を支援する公共サービス、迅速かつ正確な意思決定、そして政府の重要なサービスを改善し、未処理案件を削減するためのテクノロジー利用を推進するために使用されているアルゴリズムツールの内部構造を示している。本日（12月17日）発表された「アルゴリズムの透明性記録基準（ATRS）」は、例えば以下のような方法で、AIを利用していると述べた。 外務・英連邦・開発省が英国国民が海外で助けを必要とする際に、より迅速に情報を提供するためにAIを使用していること 法務省がアルゴリズムを使用して、研究者が司法制度と人々がどのように関わっているかをよりよく理解できるようにしている 他の省庁がAIを使用して求人広告を改善している これは、科学技術大臣の省が、公共部門全体でテクノロジーの採用を促進し、公共サービスを改善するために、政府の新しい「デジタルセンター」の形成を継続していることによるものであり、政府の5つのミッションすべてを支援し、「変革計画（Plan for Change）」の下で公共サービス改革を推進している 本日公開された記録の中で、ビジネス・貿易省は、どの企業が他国に商品を輸出しているかを予測し、経済成長を促進し、政府の「変革計画（Plan for Change）」を支援するために、アルゴリズムツールを使用していることを明らかにした。 これにより、同省の担当者は、どの企業に手を差し伸べ、支援を提供すべきかについて、よりのめを絞った決定を下すことができるようになり、輸出の可能性が高い企業がより多くの国際的な顧客を迅速に獲得できるようになる。このツールが導入される前は、担当者は、支援の対象となる企業を絞り込むために、Companies Houseに登録されている500万社以上の企業に関するデータをより手作業で精査する必要があったため、政府が提供できる支援が遅れ、成長性の高い企業を支援する機会を逃していた。 本日新たに、政府はアルゴリズムの透明性に関する記録の対象となるツールの明確な条件も定めた。これにより、政府がAIをどのように活用しているかが国民に分かるようになる。これは、国家安全保障など限られた例外が適用されない限り、中央政府機関が国民と直接やりとりするアルゴリズムツール、または国民に関する意思決定に重大な影響を与えるアルゴリズムツールの記録を公開することを確認するものである。また、ツールが公開試験段階にある場合、または稼働中および実行中の場合にも記録が公開されることも確認するものである。	GOV.UK https://www.gov.uk/government/news/new-records-detail-how-ai-helps-government-make-quick-accurate-decisions-to-boost-trade-speed-up-responses-and-more

【人工知能(AI)】関連記事詳細 (30/30)

番号	地域・国	情報記事・タイトル	発行日	要旨	情報源 (機関・団体名／URL)
34	欧州	AIモデルに関するEDPBの意見：GDPRの原則は責任あるAIをサポートする	2024/12/18	欧州データ保護委員会（EDPB）は、AIモデルの開発と展開における個人データの利用に関する意見を採択した。この意見では、1）AIモデルが匿名とみなされる場合と方法、2）AIモデルの開発または利用の法的根拠として正当な利益が使用できるかどうかと、その方法、3）違法に処理された個人データを使用してAIモデルが開発された場合の対応について検討している。また、ファーストパーティデータおよびサードパーティデータの利用についても検討している。 この意見は、欧州全域での規制の調和を求めるという観点から、アイルランドデータ保護委員会（DPA）が要請したものである。社会に重要な影響を与える急速に進化するテクノロジーを取り扱うこの意見の参考とするため、EDPBはステークホルダーのイベントを企画し、EU AI Officeと意見交換を行った。 匿名性に関しては、AIモデルが匿名であるかどうかは、DPAsがケースバイケースで評価すべきであるという意見が述べられている。モデルが匿名であるためには、（1）モデルの作成に使用されたデータを持つ個人を直接的または間接的に特定する可能性が極めて低く、（2）問い合わせを通じてモデルからそのような個人データを抽出することができないものでなければならない。意見書では、匿名性を証明するための方法として、規定的でも網羅的でもないリストが提示されている。 正当な利益に関しては、意見書は、AIモデルの開発および展開における個人データの処理について、正当な利益が適切な法的根拠であるかどうかを評価する際に、データ保護当局が考慮すべき一般的な事項を提示している。 3段階テストは、正当な利益を法的根拠として使用するかどうかを評価する際に役立つ。EDPBは、ユーザーを支援する会話型エージェントや、サイバーセキュリティの向上を目的としたAIの利用を例示している。これらのサービスは個人にとって有益であり、法的根拠として正当な利益に依拠できるが、処理が厳密に必要なものであり、権利のバランスが尊重されている場合に限られる。 最後に、違法に処理された個人データを用いてAIモデルが開発された場合、そのモデルが適切に匿名化されていない限り、その展開の適法性に影響を及ぼす可能性がある。	European data protection board https://www.edpb.europa.eu/news/news/2024/edpb-opinion-ai-models-gdpr-principles-support-responsible-ai_en