



AI トランス フォーメーション

近未来を見据え、攻守一体の
AI トランスフォーメーションをともに推進

AI playbook

KPMG. Make the Difference.



はじめに ▶▶▶

KPMGでは、近未来における人とAIの共存・協働を見据え、AIトランスフォーメーションを推進しています。

収益向上を企図したAIトランスフォーメーションとともに推進

Chat GPTが一般公開されてから2年後の2024年冬、KPMGは26カ国の2,450人の上級管理職者を対象に、企業のデジタルトランスフォーメーションに関する調査*を行いました。

調査では、「AIをビジネスに実装して価値を生み出している」と回答した企業が全体で74%（日本66%）、「AIを本番環境にて大規模に展開して投資に見合う収益を上げている」が31%（日本23%）という結果が明らかになり、AI活用の進展に応じ収益性が増加するトレンドが確認できました。また、「自組織が競合他社に後れを取っていることに不満を感じている」が80%を占め、個々の企業におけるAI活用の進展がさらに加速することを予見させる結果となりました。同時に、「AIのブラックボックス化が従業員の不安を引き起こしている」が78%、「AIが既存のオペレーション構造に課題をもたらすと予想している」が77%となり、AI活用の進展に伴うリスクを継続的に見極めながら、常に先行アクションが必須となることが認識できます。

我々の日常生活は、ロボット掃除機やスマートフォンなど、AIが組み込まれた製品・商品に囲まれています。個々の企業におけるAI活用の進展に伴い、意識するか否かに関わらず、我々の日々の生活においてAIは当たり前の存在になっています。企業の視点では、AIサービス利用企業はAIを事業活動に実装して、顧客・サプライヤ・パートナー・従業員・株主など、すべてのステークホルダーへの価値創造を進め、AIサービス提供企業は汎用・特化型AI等のSLM・LLM開発を推進して、AIサービス利用企業の価値創造を牽引する未来が予想されます。また、国や地域の視点では、欧州、米国、中国等が各々のスタンスで投資と規制のバランスをとりながらAIを推進し、関税や安全保障などの国際関係に重要な影響を与えるアクションが、相互けん制しながら発動されと考えています。

我々の近未来とAIの共存・協働を見据え論点を整理してともに考える

本稿では、多元的で柔軟な仮想現実の構築により現実世界の不確実な経営環境に対するシミュレーションを高度化する世界モデル、さまざまな仕事を人間同等に実現する汎用AIや、物理的な身体と五感を持ち心理的な振る舞いを表現する身体AIが融合するロボティクス、これらのテクノロジーと我々が安心・安全に共存・協働するためのガバナンスとアラインメントにかかる論点を紐解きます。また、サステナブルなAI活用を進展するための電力問題、ならびに従前と異なる力学で変容しつつある地政学的なパラダイムを背景に、企業が今後考慮すべき論点をご紹介します。

※ 出典：「KPMGグローバルテクノロジーレポート2024」



KPMGコンサルティング
株式会社
中野 裕介



KPMGコンサルティング
株式会社
山邊 次郎

Contents ▶▶▶

- 01 ▶▶▶ 汎用AIの実現化に向けて
- 02 ▶▶▶ 生成AIの世界モデルと多元的リアリティ
- 03 ▶▶▶ AIガバナンスとアラインメントの関係
- 04 ▶▶▶ 生成AIの自己学習とエネルギー問題
- 05 ▶▶▶ 生成AI時代における地政学的パラダイムシフト

01 ...

汎用AIの実現化に向けて

はじめに

汎用AI (AGI: Artificial General Intelligence) は、「さまざまな仕事を人間と同等かそれ以上のレベルで遂行できる能力を持つAI」と定義され、現在のAIはこの実現に向けて進化しています。近年、機械学習モデル (LLM) の発達により、AIの汎用性は飛躍的に向上していますが、真の汎用AI実現のためには数多くの課題が存在します。

AIの発展段階



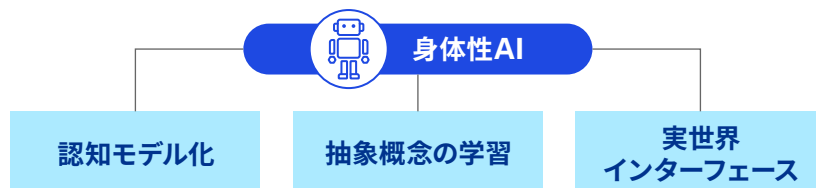
汎用AIに必要な身体性

「身体性AI」とは、物理的な身体を持ち、現実世界と相互作用しながら学習・行動するAIのことです。従来のAIは、デジタル空間内のデータ処理が中心でしたが、身体性AIは、ロボットなどの物理的な身体を通じて、環境を認識・操作・反応することで、実世界と直接的に連携することが可能となり、それによって理解がより深く体系的になると考えられています。

身体性AIは、以下の3つの利点から注目されています。

- 人間の認知や実体験をより正確にモデル化できる
- 抽象的概念を具体的な実世界を通じて学習していくことが可能となる
- 現実世界へのインターフェースの役割を果たせる

身体性AIの3つの利点



学習的アプローチの必要性

実現には、機械工学や情報科学に加えて、認知科学や心理学など多様な学問分野の視点が必要です。これらを取り入れることで、人間とロボットの相互作用 (Human-Robot Interaction: HRI) をより自然で持続的なものとすることができ、AIの行動や感情を理解し、それに適切に対応できるロボットの開発につながると期待されます。

課題

しかし、このような技術的な革新が急速に進み応用されるにつれ、技術的な課題だけでなく社会的な問題も浮き上がり上がっています。たとえば、AIの意思決定が倫理的に適切かどうか、AIが人間の仕事を奪う可能性があるといった懸念が指摘されています。

これらの問題を解決するためには、技術者だけではなく、社会学者、倫理学者、政策立案者、そして一般市民が参加する社会的な対話が必要です。AIの発展や応用が進むと同時に、それが社会にどのように組み込まれるべきかを議論する必要があります。

AIとロボティクス

AIとロボティクスの融合は、人間の生活をより豊かに、そして便利にする可能性を秘めた技術革新です。しかし、AIの能力拡張に向けて、いくつかの課題も残されています。たとえば、さまざまな環境に適応できる能力は学習データに大きく依存しており、すべての状況に対応できるデータを揃えることは困難です。

また、AIが判断をすべきことと、すべきでないことをはっきり定義するのは難しく、AIがその処理能力の限界を超えて無限に計算し続けることで、機能停止してしまう「フレーミング問題」が存在します。AIは現在のハードウェア技術の限界に能力を制約され、計算能力の進化だけでは克服が難しい状況があります。

AIの学習課題

- ✓ データ多様性
さまざまな環境への
適応
- ✓ 倫理的判断
適切な意思決定
- ✓ 説明可能性
AIの判断根拠を説明

社会的課題と未来展望

社会的な観点では、汎用AIが人間の労働を代替することで、労働市場への影響や、AIが人間に代わって倫理的な判断をする際の責任の所在が曖昧になるなどの課題が想定されます。これらの問題は技術的な問題に留まらないため、社会全体での慎重な議論が必要となります。

今後、社会が汎用AIをどう受け入れるかによって、その普及や実装のスピードが左右されるでしょう。

汎用AI社会実装の課題

- ✓ 労働代替
雇用への影響
- ✓ 倫理的判断
AIの意思決定
- ✓ 責任の所在
法的・社会的責任

ロボットAI技術企業の比較： Figure AI社とPhysical Intelligence社

Figure AI社は2022年設立の企業で、Brett Adcock氏が創業し、ヒューマノイドロボット「Figure 01」と「Figure 02」を開発しています。Figure 02はNVIDIA GPUを搭載、高い計算能力と重量物を持ち上げられる16自由度の手を備え、バッテリー性能も向上しています。BMW工場での実証が進み、将来は家庭支援や介護も視野に入れており、OpenAI社等から675百万ドルの資金調達に成功しました。

Physical Intelligence社は汎用ロボットAIシステム「π0」を開発し、多様なロボットを制御できる基盤モデルを構築しています。8種類のロボットデータと視覚言語学習を組み合わせ、テキスト指示で直接モーター制御を実現し、洗濯物の折りたたみや片付けなど、複雑なタスクを実証しています。

両社はAIロボット技術を開発していますが、Figure AI社は独自ハードウェアに、Physical Intelligence社は汎用ソフトウェア基盤に重点を置いている点が特徴的です。

結論

AI技術の進化は、今後さらに加速していくと予想されます。企業にとって、汎用AI、身体性を持つAIはさらに革新的な価値提供をもたらすでしょう。一方で、実装にあたっては、慎重に検討すべき多くの社会的責任やリスクが発生します。AI技術の進歩を私たちの社会に受け入れていくためには、社会全体が総合的・学際的な視点からAIに向き合っていくことが重要です。その先に、変革をもたらす真の汎用AIの実現が見えてくるでしょう。

02 ...

生成AIの世界モデルと 多元的リアリティ

世界モデルの基本概念

世界モデルとは、「外界から得られる観測情報に基づき、外界の環境を学習によって獲得するモデル」であり、人工知能研究において真の知能を実現するための最も重要な基盤です。実際の子どもの発達過程を見ても、子どもは外界と対話するなかで、周囲の因果的な連鎖を把握して学習していきます。私たちは直感的に「これを使いたい」、あるいは「これは間違いだ」という判断を日常的に行っていますが、これらの判断は「世界がこうなるはず」というという予測のもと行われており、その結果、将来予測や意思決定を行うことが可能になっています。

世界モデルの主要機能

予測機能

世界モデルの機能は、大きく「予測」と「推論」の2つに分類されます。予測とは、「現在わかっている状態から将来の状態がどうなるか」を特定することです。たとえば、雨が降りそうだと感じたり、道路の渋滞を見て到着時刻の遅れを見積もったりすることを指します。

推論機能

推論とは、現在の状態から過去の状態や原因を推定的に遡る能力です。たとえば、小包の重さから内容を推定したり、音楽の雰囲気からジャンルを推定したりといった、現状から過去の原因や特性を見出す能力が推論機能に該当します。この世界モデルによって、新商品をどの市場にいつ投入すべきか、顧客はどのような反応を示すか、競合はどの戦略を打ってくるかを、市場データという観測によって「推論」し、構造的かつ法則をはらんだ内部モデルを「予測」することで、従来の予測よりも複雑な要素を織り込んだ仮想シナリオを何度も試せるのです。かつては人間が経験と勘に頼って市場予測を行っていましたが、世界モデルの活用により、「微分可能なビジネス空間」すなわち包括的な環境での予測が数値化されるようになって考えられ、これによって新規事業開発、マーケティング戦略立案、サプライチェーン最適化などで、失敗コストや時間ロスを大幅に削減できる可能性が生まれます。

古典的人工知能との比較と世界モデルの意義

かつての古典的な人工知能は、言語を活用した学習を通じて高度な探索や推論を行う、人間にとって合理的な、いわば道具となることを目的として設計されていました。しかし、世界モデルがない状態では、計画データに基づいて機能する限定的な知能しか得られず、現実環境への適応力に限界があることがわかりました。

これに対し世界モデルは、個別に教えられずとも、環境と相互作用しながら直感的に包括的な構造を理解することができる、「子供の知能」のようなものです。

「大人の知能」をセミナーや講義における座学にたとえるならば、「子供の知能」は演習にあたるといえます。体験でしか得られない感覚や経験といった領域が、成熟した「大人の知能」の形成には不可欠であり、これができることで初めて、現実環境に適応した高度な数学問題の解決や複雑なタスク遂行などを可能とする人工知能が築かれると考えられています。

古典的AIと世界モデル型AIの比較



古典的AI

- ルールベースの学習
- 限定された特定領域に特化
- 環境変化への適応が困難
- あらかじめ設定された方法での処理



世界モデル型AI

- 環境からの学習と適応
- 広範な領域の知識統合
- 柔軟な推論と予測能力
- 創発的な問題解決能力

シミュレーション能力と実用的応用

これらの特性を備えた世界モデルは、内部に言語的・概念的 world を持ち、自由に未来を「シミュレーション」して、ユーザーや市場動向、ニュースなどのデータに基づいて将来シナリオを描くことが可能です。例えば、ロボットに制御を学習させる際、複雑な外部環境であればあるほど困難で、インプットされていない状況への対応はエラーを招くものでした。世界モデルによって内的な「微分可能な世界」が得られれば、複雑な現実世界を織り込んだシミュレーション空間内で繰り返し学習し、効率的に制御戦略を獲得できます。これは単にロボット制御に留まらず、ロジスティクス最適化や投資ポートフォリオの動的調整、店舗レイアウトやサプライチェーン管理などでの内部シミュレーションへと応用可能です。これは人間が頭の中で試行錯誤する「イメージトレーニング」に似たアプローチです。また、複雑な現実を理解しやすく、活用しやすい形に整理して推論できれば、高度な計画立案や問題解決も容易になります。

世界モデルの実用的応用



ロジスティクス 最適化

サプライチェーン予測



ポートフォリオ 最適化

市場シミュレーション



商品推薦 システム

顧客行動予測



アバター プレイヤー

仮想世界での代理行動



バーチャル トレーニング

安全なスキル習得環境



高度な 計画立案

多様な問題解決

現実認識への影響

世界モデルは、単一の未来予測ではなく、複数の可能性を同時に計算し、さまざまな状況をシミュレーションします。たとえば、販売戦略の検討において、従来は与えられたパラメーターの中から予測や選択肢を得ていましたが、世界モデルの活用により、織り込んでいなかった情報も「予測」により含まれ、想定していなかった選択肢まで「推論」によって提示される可能性があります。その結果、旧来よりも多く提示される選択肢からいずれを選ぶのか、意思決定がより高度になることが考えられます。

私たちには、最終的な結果を問うことだけでなく、どのような可能性があり、それぞれの可能性がどの程度確かで、異なる可能性がどのように関連しているのかを理解することが求められます。

課題と編纂展望

こうした世界モデルが創り出す多元現実（「唯一の正解」を追求するのではなく、複数の可能性を同時に理解し、それらの相互関係を把握しながら、より豊かな意思決定を可能にする新しい思考の枠組み）は、実世界との相互作用や適応に不可欠な要素となっています。課題としては、多くの実装上の問題が存在します。たとえば観測範囲が限られたなかで、いかに高性能な予測と推論を行うか、どのようなパラメーターを組み合わせで「何をしたらどうなるか」をモデル化するという点が挙げられます。現実世界を丸ごと再現することは不可能なため、範囲を限定しつつ効率的な学習を模索する必要があります。

総じて、世界モデルという考え方は、生成AIによる多元的で柔軟なリアリティを構築し、自社が直面する不確実な経営環境に対するシミュレーションの飛躍的な高度化につながります。従来見えていなかった市場やアプローチにまで光を当て、ビジネスを拡大できる機会を得ることができます。これにより、未来志向のビジネス判断が、実世界での大がかりな試行錯誤なしに行える道が拓かれます。

そしてこの大きな変化は巨視的には、多元的なリアリティと「真の知能」への扉を開く鍵となるのです。



現状の課題

- 計算リソースの制約
- 複雑な環境の完全モデル化の困難さ
- 多次元パラメーターの管理



将来の展望

- 多元的リアリティの実現
- ビジネス機会の拡大
- 真の人工知能への発展

Google DeepMindの「Genie 2」

2024年12月4日に公開されたGoogle DeepMindの「Genie 2」は、2D画像から多様な3D環境を生成することができる基盤世界モデルです。この技術は単なる画像生成にとどまらず、空間の物理的な構造と法則を理解し、その中で動作する環境を創造します。

Genie 2は2D画像の入力から、その画像が表す空間の3D的な広がりや物理的な性質を推測し、ユーザーが自由に探索できる3D環境として再構築することができます。これにより、例えばゲーム開発や仮想環境構築のプロセスが大幅に効率化される可能性があります。

自動運転分野での世界モデル：GAIA-1とGenAD

世界モデルは自動運転技術の進化においても重要な役割を果たしています。英国のスタートアップWayve社が開発した「GAIA-1」は、自動運転のための世界モデルとして注目されています。

GAIA-1は自動運転システムの一要素として、現在の交通状況に関する表現や未来の状況予測を行い、それを自動運転の意思決定に活用します。例えば、周囲の歩行者や車両の動きを予測することで、より安全で効率的な運転判断が可能になります。

同様に、OpenDriveLab社が開発した「GenAD」も、自動運転のための世界モデルです。これらのテクノロジーには、実世界の複雑な交通状況を理解し、予測することで、完全自動運転の実現に貢献することが期待されています。

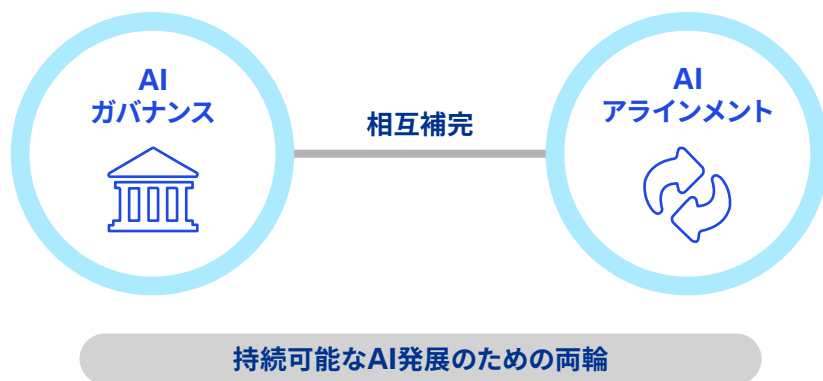
03 ...

AIガバナンスと アラインメントの関係

はじめに

現在、AI全般を運営するうえでガバナンスの重要性は広く認識されています。一方でAIシステムは人間の意図や価値観から逸脱する可能性があり、AIシステム構築において重要な方法で「倫理的にAIをコントロールする」という、「アラインメント」という考え方が注目されています。

アラインメントは、AIシステムの目標や価値を人間の意図や価値観と整合させることであり、この「価値を調整し一致させること」により、持続的に人間とAIが相互的に進化していく、「人間とAIの共進化」が可能となります。



AIガバナンスとアラインメント – 基本概念の整理

AIガバナンスとは

AIガバナンスは、AIシステムの開発・展開・利用に関する監視的・社会的な枠組みを指します。これは以下の要素から構成されます。

- AIの開発・利用に関する規制と指針
- 技術標準や倫理基準の設定
- 規制対策メカニズム

AIアラインメントとは

AIアラインメントとは、AIを人間の意図や価値観に整合させるために必要な方法です。主に以下の2つの側面があります：

- 技術的アラインメント（AIの目標や行動の調整）
- 価値観のアラインメント（人間の倫理観との整合）

両者の生成AI分野での関係においては、このアラインメントとAIガバナンスの両方を噛み合わせることで、汎用AI系の安全性を確保することが理想であるといえます。

アラインメントの必要性和その形態

アラインメントが確保されていないAIシステムは、意図せぬ価値観を助長したり、予期せぬ行動をとったりする可能性があります。そのためには、AIシステムの開発段階から運用段階に至るまでの包括的な管理と調整が必要になります。特に、AIやIoTなどのテクノロジーを用いて、人間の身体能力・知覚などを増強・拡張させる技術である人間拡張（Human augmentation）が進むに従って、AIと人間の境界が曖昧になり、結果、AIが生命や安全に危害を与える可能性や、その責任の所在が判断不能になるなど、これまで経験したことのない多くの新たなリスクが生じる可能性があります。企業におけるAI利用において、このようなリスクは看過できません。そこで、現在研究や実務が行われている、以下のアラインメントの3つの観点をふまえることが必須となります。

アラインメントの3つの観点

1 生成モデルの品質保証	出力の公平性／バイアス排除／説明可能性の確保
2 AIによるシステム行動の制御の目的	自律的なAIシステム行動の人間価値観との整合
3 社会システムの設計	AIを含む環境全体を考慮した包括的なガバナンス構築

これらの観点は、単なる技術的な制御を超えて、社会的な価値観の調和を目指すものです。実際、調査*では以下の結果が出ています。

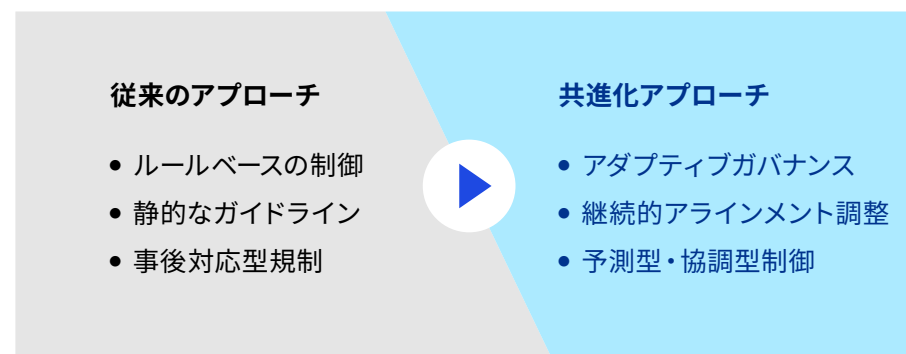
- 61%の企業がAIシステムを信頼することに警戒心を抱いている。
- 85%がAIはさまざまなメリットをもたらすと考えているが、メリットがリスクを上回ると考える人は約半数のみ。
- 現行の規制や保護措置がAIの安全利用に十分と考える人はわずか13%で最低水準。

このような社会的な不安に応えるためにも、アラインメントの確保は不可欠です。2019年3月に内閣府が発表した「人間中心のAI社会原則」や、2024年4月に経済産業省と総務省が公開した「AI事業者ガイドライン」も、このガバナンスとアラインメントの両方を考慮した包括的な枠組みの1つとして位置づけられます。

ガバナンスの再考ー共進化の視点から

このように、アラインメントの確保は企業の社会的責任として重要です。特に日本では、AIリスク管理体制の整備が十分とは言えず、アラインメントを考慮した新しいガバナンスの枠組みが求められています。AIの進展が加速度的に変化するなか、一般的なガバナンス（制度と規制）だけでは、人間と汎用AI（AGI：Artificial General Intelligence）の共生が実現される未来において、アラインメント不足や人間拡張の進展が、社会に対して何らかの危険を及ぼす可能性があります。

ガバナンスとアラインメントの共進化モデル



このような状況を踏まえると、企業には2つの重要な責任が伴います。

- 適切なアラインメントを確保するための枠組み・組織的能力を発展させる責任
- そのアラインメントをガバナンスの枠組みの中で実現すること

その実現には、従来型のPDCAサイクルではなく、変化掌握志向のボトム（現場）でのプランニングとモニタリング、そして経営の意思決定・対応スピードアップが求められます。

この新しいガバナンスの形は、予測困難な技術革新や社会変化にも柔軟に対応できる特徴を持っています。この現場主導の変化察知と迅速な意思決定を基盤としたアジャイル・ガバナンスの考え方は、この課題に対する1つの回答となりえます。

※ AIは信頼できるか ～AIへの社会的認識の変化に関するグローバル調査2023
<https://kpmg.com/jp/ja/home/insights/2024/01/trust-in-ai.html>

まとめ

AIの社会的影響力の大きさを考えると、企業にはアラインメントを確保しつつ、適切なガバナンスの下でAIを開発・運用する責任があります。そのためには、タイムリーかつ柔軟的に変化を取り入れ、将来を見据えた、技術革新に適合した継続的な枠組みの確保が必要です。KPMGの「Trusted AIフレームワーク」は、このような課題認識に基づき、アラインメントとガバナンスの両方の視点について、企業のAIガバナンス態勢構築を支援しています。

参考資料：

- 『人間と人工知能が協調するための新たなビジョン』KPMG（2023年）
- 『デジタルとAI革新時代における倫理課題』欧州AI倫理ガイドライン（2022年）
- 『AGIガバナンスとアラインメント - 共進化の視点から』日本AI学会（2023年）
- 『企業のためのAIガバナンスフレームワーク』経済産業省（2024年）

主要な国・地域、機関によるAIリスク対応策のまとめ

G7

行動規範履行状況の報告枠組みHAIP Reporting Framework(2025年2月稼働、初回19社公表)

- 「全てのAI関係者向けの広島プロセス国際指針」の策定
- 「高度なAIシステムを開発する組織向けの広島プロセス行動規範」の策定
- OECD ウェブサイトで、行動規範履行状況の報告枠組みが運用開始

国際連合

「グローバル・デジタル・コンパクト」の採択(2024年9月)

- AIのリスク・機会評価を通じた理解促進のための国際科学パネルの設置
- AIガバナンスに関するグローバル対話の開始

欧州評議会

人工知能と人権、民主主義及び法の支配に関する欧州評議会枠組条約

- EUに加え、米国、英国など15カ国が署名

OECD(経済協力開発機構)

AI原則の改定(2024年5月)

- 生成AIによる、偽・誤情報リスクへの対応を強化

GPAI(Global Partnership on Artificial Intelligence)

GPAIサミット2022 東京開催

- 人間中心の価値に基づくAI活用促進、違法・無責任な使用への反対、持続可能で平和な社会への貢献について合意、閣僚宣言として公表

AISI(AIセーフティ・インスティテュート)

- AIの安全性評価手法、基準の検討・推進を行う機関
- 国内外の安全性の知見ハブとして、各機関とのネットワーク構築、ガイダンス作成

EU

AI Act(欧州連合)

- リスクレベルに応じた4段階規制アプローチ
- 高リスクAIシステムには市場投入前の影響・適合性評価の義務
- 汎用AIモデル提供者への透明性義務(技術文書作成、学習データ開示等)
- 大規模モデルに対する追加義務

米国

1. 大手AI企業によるボランタリー・コミットメント(2023年7月～)

- リスク管理情報の共有、社会的リスク研究、サイバーセキュリティ投資等の実施

2. 大統領令(2023年10月→2025年1月20日撤回→1月23日発令)

- 安全保障上のリスク対応のため、国防生産法に基づき発令
- 大規模な基盤モデル開発企業に、政府への情報提出を義務化
- 25年1月23日に、政策や規制等の措置の見直しと、行動計画の策定指示

3. カリフォルニア州法(2024年9月)

- SB942: AI生成コンテンツの透明性向上
- AB2013: AI学習データの開示義務

日本

1. 「AIに関する暫定的な論点整理」「AI制度に関する考え方」(2023年5月～)

- AIに関する論点整理、AI制度の考え方を示す

2. 「統合イノベーション戦略2024」(2024年6月4日閣議決定)

- AI分野の競争力強化と安全・安心の確保を目的とした戦略を策定

3. 「AI事業者ガイドライン(第1.1版)」(2025年3月28日)

- 事業活動におけるAI開発、提供、利用指針の明記
- 関係府省庁が連携して検討、対応

4. 「人工知能関連技術の研究開発及び活用の推進に関する法律」

(AI新法法案: 2025年2月閣議決定、審議中)

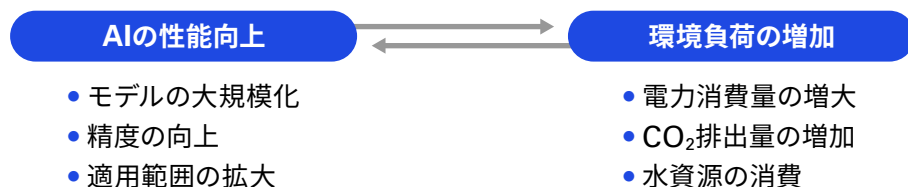
04...

生成AIの自己学習と エネルギー問題

はじめに

生成AIの発展的な進展はビジネスに革新をもたらす一方、環境面で深刻な課題を提起しています。AIの活用により、2022年から2026年の間に世界のデータセンターの電力需要は急増する見込みで、CO₂排出量も上昇すると予測されています。私たちは、より複雑な知能性能の向上には環境負荷を伴うというパラドックスに直面しています。

生成AIの発展と環境負荷のパラドックス



自己学習における技術的課題と環境影響

生成AIの性能向上はモデルの大規模化とデータ量増加に依存しており、いくつかのテクノロジー上の課題がエネルギー消費増大につながっています。

データ効率の問題点

AIは人間と比べて膨大なデータを処理する必要があり、多くが効率的でない処理を行っています。

例として：

膨大なパラメーターから一般化するプロセスに、大量の計算リソースを消費します。不必要なデータ処理のための計算資源が、電力消費の増加を招いています。

事実：生成AIへの質問は約0.1kWh、AIモデルの訓練には30kWhもの電力を消費します。

過学習の問題

特定のデータパターンへの過剰適応により、新しいデータへの対応力が低下する問題です。これを解決するためには、より多くのトレーニングデータと複雑な学習器が必要となり、エネルギー消費量が増加します。

AI処理の効率化プロセスがエネルギー消費削減につながりますが、同時に性能向上も必要となるため、現状では正確性と環境負荷の適切なバランスが実現できていません。

持続可能性への5つのアプローチ

生成AIのエネルギー消費を抑制しながら性能向上を実現するために、5つの主要なアプローチが存在します。



持続可能なAIへの5つのアプローチ

1 アルゴリズムの効率化	- 計算リソースの最適化 - 知識蒸留技術の活用
2 ハードウェア革新	- 半導体技術の発展 3Dアーキテクチャの採用
3 データ品質向上と倫理的活用	- 高品質データの効率的活用 - 社会的価値観との整合性確保
4 再生可能エネルギー活用	- データセンターの電力革新 - AIの計算効率と環境負荷の両立
5 次世代コンピューティング	- 量子コンピューティングの可能性 - ハードウェアとソフトウェアの緊密な連携

① アルゴリズムの効率化

計算効率を最大化するアルゴリズム開発は、AIの環境負荷低減において中核的役割を果たしています。知識蒸留技術により、大規模な教師モデルから得られた知識を小型モデルに効率的に転移させ、計算資源を削減します。また、転移学習を活用することで、特定の知識を新しいタスクに再利用して、トレーニングに必要な計算量を劇的に減少します。

② ハードウェア革新

半導体技術の発展が、技術的・環境的課題の解決に貢献しています。シリコンウェハー（チップが生産される20～30cmの薄平板）の改良や、3Dアーキテクチャの採用を通じて、チップの設計密度を高め、電力効率の向上に寄与しています。また、パッケージング技術も進歩しており、チップの小型化と性能向上を両立させながら、メモリアクセスの効率化と消費電力削減を実現しています。

③ データ品質向上と倫理的活用

高品質で多様なデータの効率的活用が、モデルの性能向上と必要なデータ量削減に直結します。そのため、インターネット上の偏りを持つデータだけでなく、目的特化型の高品質データセットにバランスの取れた価値を同めています。企業はデータ収集における透明性確保と公平性担保に注力し、社会的価値観との整合性確保を目指しています。

④ 再生可能エネルギー活用

大手テクノロジー企業を中心に、データセンター運用における再生可能エネルギーの活用が進んでいます。太陽光・風力発電の導入により、AIのトレーニングと運用における環境負荷を軽減するスマートグリッドや、特定の時間帯での計算処理の実施、電力需要の平準化などの取組みにより、AI関連の計算需要の増加が環境に与える影響を抑制する効果が期待されています。

注目データ：生成AIへの1質問あたり約0.1kWhの電力消費は、一般的なノートPCを1時間使用する電力量に相当します。

⑤ 次世代コンピューティング

現在のAI開発はGPUに大きく依存しており、特定企業への過度な依存や半導体供給の不安定さといった課題が顕在化しています。この課題に対してさまざまな新興技術が研究されています将来的には量子コンピューティングが革新的な可能性を秘めており、AIモデルのトレーニング効率向上と消費電力削減が期待されています。これらのアプローチは個別ではなく相互補完的な機能しており、ハードウェアとソフトウェアの緊密な連携を通じて大きな環境負荷削減効果を生み出しています。

今後は環境規制強化にともない、企業には環境負荷を考慮したAI開発が求められます。技術的可能性と倫理的側面をESG視点で検討し、長期的な環境への影響に配慮した持続可能な生成AIのエコシステム構築が必要となるでしょう。

持続可能なAI開発の未来



※本記事は生成AIの環境負荷と持続可能性に関する最新の研究と業界動向に基づいています。技術の進歩により、ここで述べた課題や解決策は変化する可能性があります。

シリコンカーバイド (SiC) 半導体は、従来のシリコン (Si) 半導体と比較して優れた省エネルギー特性を持っています。主な省エネルギー面での利点は以下の通りです。

- 1 圧倒的な電力損失の削減：従来のシリコン半導体と比較して、電力損失を約10分の1に抑制できる
- 2 高効率エネルギー変換：絶縁破壊電界強度がSiの10倍、バンドギャップが3倍、熱伝導率も3倍高いという特性により、電力変換効率が大幅に向上
- 3 冷却システムの簡素化：高温環境での安定動作により冷却に必要なエネルギーを削減

これらの特性は、特に生成AIの電力消費問題に大きな解決策を提供します。GPT-4のような大規模モデルの訓練には数百メガワットの電力が必要ですが、SiC半導体技術により、データセンターの消費電力を大幅に削減できる可能性があります。また、この技術はカーボンニュートラル目標にも貢献し、電力変換・利用の高効率化においてSiCパワー半導体の重要性が高まっています。

製造現場における生成AIと量子コンピュータの統合事例

2025年初頭に開催された岐阜県DX推進コンソーシアムワーキンググループ成果報告会では、製造現場で生成AIと量子コンピュータを連携させて活用する具体的な事例が発表されました。この発表は2025年2月12日に行われ、その後3月6日には中部経済新聞の取材も受け、3月7日にはハイパーネットワーク・ワークショップ2025でも「量子ビジネスの最前線とその未来」と題して紹介されています。

この事例で使用されている量子コンピュータは、公益財団法人ハイパーネットワーク社会研究所が管理・運用している量子コンピュータシミュレーターです。量子コンピュータは一般に最適化問題を解くのに適しており、この特性を生かして以下のような実務的な課題解決に応用されています。

- 1 従業員のスキルや経験を最大限に活かすための最適な人員配置
- 2 在庫管理の最適化による保管コストの削減
- 3 資材調達先の見直しによるコスト削減
- 4 営業訪問ルート最適化

05 ...

生成AI時代における 地政学的パラダイムシフト

パラダイムシフトの本質と生成AIの 地政学的影響

生成AIの急速な発展は、国際関係の力学を根本から変革しています。この革新は、国家安全保障、経済戦略、国際関係の本質的変容をもたらしています。この変容を理解するためには、私たちの認識を転換する必要があります。従来の地政学的分析が前提としてきた物理的な力の概念は、テクノロジーが進化する現在、その有効性を部分的に失いつつあります。代わりに、データ、アルゴリズム、計算能力といった非物理的要素が台頭しています。

1. パラダイムシフトの本質

デジタル主権の概念

「デジタル主権」という新しい概念が注目されています。データは物理的資源と異なり、同時に複数の場所に存在し、国境を超えた複雑な所有権と権力関係を創出します。各国はデータとアルゴリズムの分析と予測や意思決定に重要な価値をもたらします。

集合知としての国力

軍事や経済力は単なる政治力だけではなく、政府、企業、市民社会が一体となって初めて効果的に活用できる社会システム構築能力に依存しています。技術開発競争を超えた包括的な社会システムの競争が展開されています。

生成AI時代の地政学的パラダイムシフト

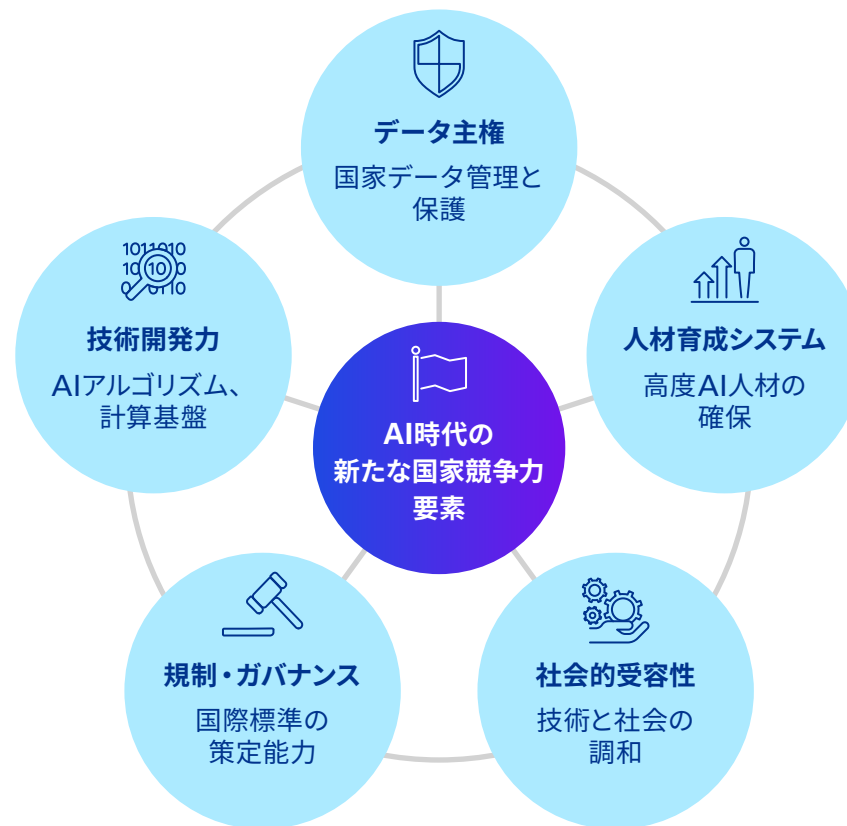
従来の国際秩序

- 経済力
- 軍事力中心
- 地理的境界の重要性

AI時代の国際秩序

- 技術力
- データ主権
- デジタル境界の台頭

2. 新しい国家競争力の定義



3. 主要国の対応と規制動向

世界的なAI開発競争において、米中両国は技術覇権をめぐる激しい対立を続けています。米国は2024年の対中投資規制や2025年のAI向け半導体輸出規制強化など、包括的な規制を実施。特に注目すべきは2024年10月発表の「AIの開発や利用に関する国家安全保障覚書」で、米国主導のAI開発、国家安全保障へのAI活用、国際的なガバナンス推進という明確な方針を示しました。

一方、中国はゲルマニウムやガリウムなどの重要鉱物管理で対抗し、「中国製造2025」計画下で技術的自立への投資を加速させています。

各国の規制アプローチ

米国

イノベーション重視型：比較的緩やかな規制と自主的な基準設定を重視し、開発をリスク分けで進行しています。民間主導の技術開発を促進しつつ、国家安全保障上の観点から選択的な規制を行っています。

中国

管理統制型：国家によるガバナンスを強化した厳格な規制枠組みを確立しています。技術開発と国家安全保障を一体的に捉え、中央政府主導の管理体制を構築しています。

EU

基準重視型：倫理的規範を重視した「第三の道」を構築し、包括的規制枠組みを推進しています。AIの利用に際して、人権や個人情報保護を重視した統合的アプローチを採用しています。

日本

ガイダンス型：米国の開発重視とEUの倫理観重視のバランスを取る法的拘束力のない指針を策定しています。産業振興と社会的受容性の両立を目指す柔軟なアプローチを採用しています。

4. 軍事・安全保障への影響

生成AI技術の発展は軍事・安全保障領域に大きな変革をもたらしています。各国はAI技術の軍事応用を積極的に進めており、将来の安全保障環境を大きく変える可能性があります。

主要国の軍事AI動向

米国防総省AI RCC (AI Rapid Capabilities Cell) を通じて、以下施策的展開を推進しています。国防分野でのAI活用を体系的に進め、次世代戦闘能力の強化を図っています。

中国は2035年までの国防の近代化完全を目指しAIを重視しています。軍民融合戦略の下、民間技術の軍事転用を積極的に進めています。

ウクライナ紛争ではAI音声型偽人媒体解証アルゴリズムなどAIの実践的活用が見られています。情報戦の新たな形態としてのAI活用が現実の紛争で検証されています。

日本も防衛省からAI活用推進基本方針とサイバー人材能力構築を策定しています。限られた防衛資源の効率的活用のためにAI技術の導入を検討しています。

5. 直面する課題

社会的課題

技術格差拡大、規制とイノベーションのバランス、国際協調と国家主権の調和などの課題が存在します。特にデジタルデバイドの拡大は社会的分断を深める恐れがあります。また、偽情報拡散、差別助長、サイバー犯罪者等化、著作権侵害などの問題も顕在化しています。これらの課題に対して、国際的な協調体制の構築が求められています。

倫理的課題

意思決定の説明責任と透明性、人間中心的な価値感への統合、アルゴリズムバイアス、システムの安全性確保、プライバシー保護など多くの課題があります。特に自律型AIシステムの意思決定プロセスの透明性確保は重要な課題となっています。また、AIシステムに人間の価値観をどのように組み込むかという問題も、技術的・哲学的な観点から検討が必要です。

6. 企業の対応戦略

生成AI時代の地政学的変化に企業はどのように対応すべきでしょうか。グローバルに事業を展開する企業にとって、各国の規制動向やAIガバナンスの違いを理解し、適切に対応することが不可欠です。

主な観点：

ガバナンス体制構築：クロスファンクショナルなAIガバナンス体制の確立が重要です。技術、法務、倫理、経営など、多様な視点からAI活用の方針を検討するチームを構成する必要があります。

データバランス：対象地域の偏りの是正やプライバシー配慮など、データ品質の確保が求められます。特に国際的に活動する企業は、各国・地域の規制に準拠したデータ管理体制を構築する必要があります。

パートナーシップ構築：技術移転・技術開発体制の維持を図るべきです。単独での技術開発が困難な場合、戦略的パートナーシップを通じて技術力を補完することが有効です。

人材・組織関係：全社的な教育と技術者・非技術者の協働促進が不可欠です。AI技術の理解は技術者だけでなく、経営層を含む全社的な取り組みとして推進する必要があります。

結論

生成AI時代の地政学的変容は、技術革新の成果だけでなく、人類の組織の在り方に対する根本的变化を示唆しています。企業には地政学的なインサイトに基づき、具体的な行動を通じて複雑な環境に体系的に対応し、持続可能な成長を実現することが求められます。今後、AI技術の発展とともに

国際秩序は一層複雑化していくと考えられます。各国の規制アプローチの違いを理解し、グローバルな視点と地域に適した戦略の両立が、AI時代を生き抜くための鍵となるでしょう。

地域	規制アプローチ	主要政策・取組み	軍事・安全保障面	特徴的な戦略
米国	イノベーション重視型	- 国家AI推進法(2020年) - AI権利章典のための青写真(2022年) - AIの開発や利用に関する国家安全保障覚書(2024年) - 対中投資規制(2024年) - AI向け半導体輸出規制強化(2025年)	- AI Rapid Capabilities Cell (AI RCC) の設立 - LLMの軍事応用推進 - 統合全領域指揮統制(JADC2) 戦闘構造の強化	- 緩やかな規制と自主的な基準設定 - AIのリスクより開発を優先 - グローバルサウス支援策 - 米国国立標準技術研究所(NIST) によるリスク管理フレームワーク
中国	管理統制型	- インターネット情報サービスアルゴリズム 推奨管理規定(2022年) - 中国標準2035 - 中国製造2025 - 国家人工知能産業総合標準化システム 構築ガイドライン	2035年までの国防近代化計画 - AI軍司令官の開発 - 軍民融合戦略 - 民間AI技術の軍事転用	- 国家によるガバナンス強化・厳格な規制枠組み - 重要鉱物(ゲルマニウム・ガリウム) 管理 - 技術的自立への投資加速 - 一帯一路構想
EU	倫理重視型	- AI法 - デジタルサービス法(DSA) - プライバシー関連規制	—	「第三の道」の模索・包括的な規制枠組み - 新興ソフトウェアへの規制・調査 - プライバシー保護重視
日本	ガイドライン型	AI事業者ガイドライン(2024年4月)	- 防衛省AI活用推進基本方針 - サイバー人材総合戦略	- 米国とEUによるアプローチのバランスをとる - 法的拘束力のないガイドライン - 産業界・利用者の調和を重視

融合する知能：AIと人間の新たな挑戦

AIはもう後戻りすることのない存在

事例で見てきたように、AIは日々急速に進化しています。この進化のスピードを生み出しているのは、これまでに類をみない、世界レベルでの多額の投資です。では、そもそも何故、AIが多額の投資に値するものと考えられているのでしょうか。

1つには、その中心的な技術であるディープラーニングが、人間の脳の神経ネットワークを模倣していることにあります。投資に関わる人は、自らの脳神経ネットワークを模したAIに大きな可能性を感じています。また、その進歩が一過性ではなく、決して後戻りしないことを直感的に理解しているといっても過言ではありません。2つ目は、既にAIがさまざまな領域で社会実装されており、自動化や高い効率性といった実証結果が確実に積み上がることで、多くの人が、AIを不可欠な技術として認識するようになってきていることです。さらに、ビジネスが近年、データ駆動型にシフトし、データと高い親和性を持つAIもビジネスの可能性を広げる存在と捉えられていることが挙げられます。

しかし、自らの脳神経ネットワークを模倣し、データと高い親和性を持つことで、自動化や高い効率性を生み出すAIの存在が、職業浸食への不安を掻き立てていることもまた事実です。AIは、もはや反復的で単純な作業に置き換わるだけでなく、より複雑なタスクをこなすようになっていて、適応領域は年々拡大しています。既に多くの職種で人間の労働力が不要になる可能性が囁かれています。一方、AIか人間かといった二者択一ではなく、AIと人間が融合することが必要であるとの考えも有力で、今後は、AIと協業できるスキルが労働者に求められます。しかし、多くの労働者はこれらの新しいスキルを持っておらず、再教育やスキルアップが必要です。しかし、そもそもどのようなスキルが必要か明確ではなく、教育整備が不十分であることも、労働市場における不安の高まりにつながっています。

AIと人の融合によりさらなる高みを目指す

2025年2月に国際AI交渉コンペティションであるANAC (Automated Negotiating Agents Competition) が開催されました。全世界から50を超えるチームが参加し、各チームのAI同士に交渉力を競わせる競技で、合理性を重視し、最新のAI技術や高度なデータアナリティクスに基づいた戦略を採る米国、カナダ、英国が上位を占めました。しかし、全く異なる戦略を取った東京大学チームが4位に食い込んだだけでなく、特別賞も受賞するという興味深い結果となりました。東京大学チームは、感情認識の精度や共感、協力を重視したAIを作り上げ、高い交渉成功率や満足度の向上といった、日本で日常的に行われている温かみあるアプローチにより、非常に高い評価を受けました。この東京大学のアプローチを発端に、AIの強みを活かしつつ、人間の交渉技術を補完することで、より効果的な交渉を実現するAI時代の新しい交渉理論の研究が既に始まっています。東京大学チームは、AIだけ、人間だけといった二者択一では決して到達できない、両者が融合することで初めて目指せる高みがあることを示してくれました。このことは労働者に対する再教育やスキルアップの内容を明確にするヒントになると考えています。

これからもKPMGでは、AIと人間の融合を推進し、AI技術の進化と人間の能力を最大限に活かした未来の企業経営の在り方を提案できるよう活動を続けていきます。共に新しい高みを目指していきましょう。

KPMGコンサルティング株式会社

kpmg.com/jp/kc

株式会社 KPMGアドバイザリーライトハウス

kpmg.com/jp/alh



ここに記載されている情報はあくまで一般的なものであり、特定の個人や組織が置かれている状況に対応するものではありません。私たちは、的確な情報をタイムリーに提供できるよう努めておりますが、情報を受け取られた時点及びそれ以降においての正確さは保証の限りではありません。何らかの行動を取られる場合は、ここにある情報のみを根拠とせず、プロフェッショナルが特定の状況を綿密に調査した上で提案する適切なアドバイスをもとにご判断ください。

© 2025 KPMG Consulting Co., Ltd., a company established under the Japan Companies Act and a member firm of the KPMG global organization of independent member firms affiliated with KPMG International Limited, a private English company limited by guarantee. All rights reserved. C25-1027

The KPMG name and logo are trademarks used under license by the independent member firms of the KPMG global organization.