

EAJ 報告書 2024-02

一般財団法人 新技術振興渡辺記念会
令和 5 年度（2023 年度）下期科学技術調査研究助成

「5G/6G 時代の AI 利活用戦略」 / 「生成 AI をはじめとした
AI による社会変容とリスクマネジメントに関する調査研究」

成果報告書



令和 7 年（2025 年）2 月

公益社団法人日本工学アカデミー
「5G/6G 時代の AI 利活用戦略」プロジェクトならびに
「生成 AI をはじめとした AI による社会変容とリスクマネ
ジメントに関する調査研究」プロジェクト
(リーダー: 森健策)

2025 年 2 月 13 日

公益社団法人日本工学アカデミー

日本工学アカデミーは、工学・科学技術全般の発展に寄与する目的で設立された産学官の指導的技術者を会員とする団体です。会員の豊かな経験や知識、幅広いネットワークを活用したプロジェクトチームを中心に、広く会員外からの協力も得て、調査提言活動を進めています。その成果をまとめ、社会が目指すべき方向性に関して、官公庁、立法府、産業界、学会、研究機関等に先導的、創造的な施策を提言し、社会実装を目指します。

「5G/6G時代のAI利活用戦略」プロジェクトならびに「生成AIをはじめとしたAIによる社会変容とリスクマネジメントに関する調査研究」では、高速大容量情報ネットワーク時代におけるAI利活用の戦略を構築すべく、情報通信技術およびAI技術に関する国内外の代表的な研究開発施設・企業の専門家の方々へのヒアリングを通じて、最新の研究開発・社会実装動向ならびに現状の課題を整理し、次世代AI技術の適正開発・社会実装における必要事項を明らかにするための調査研究を実施しました。その結果、次世代AI技術の創出・普及に向け、本邦が世界標準から遅れを取らないように取り組むべき課題と提言を取りまとめました。今般、報告書原案につき、政策提言委員会で査読を受け、理事会での審査を経て、最終版を確定し、日本工学アカデミーとしての発出が理事会で承認されました。本提言が広く活用されることを期待します。

本報告は、公益社団法人日本工学アカデミーが進めてきた「5G/6G 時代の AI 利活用戦略」プロジェクト，ならびに本プロジェクトの一環として，一般財団法人新技術振興渡辺記念会の助成を受けて実施した「生成 AI をはじめとした AI による社会変容とリスクマネジメントに関する調査研究」の結果をとりまとめ公表するものである．表 1 に本プロジェクトの構成員を示す．

表 1 公益社団法人日本工学アカデミー「5G/6G 時代の AI 利活用戦略」プロジェクトならびに「生成 AI をはじめとした AI による社会変容とリスクマネジメントに関する調査研究」の構成員
(所属等は 2024 年 12 月 31 日時点)

役 割	氏 名	所属・役職
リーダー	森 健策	名古屋大学大学院情報学研究科 教授
副リーダー	浅間 一	東京大学・東京カレッジ 教授，日本工学アカデミー 理事
オーガナイザー	城石芳博	日本工学アカデミー 専務理事
オーガナイザー	中島義和	東京科学大学 教授，日本工学アカデミー 会員
オーガナイザー	中村道治	日本工学アカデミー 顧問
オーガナイザー	森本浩一	日本工学アカデミー 国際副委員長
幹事	杉野貴明	東京科学大学 助教
委員	漆谷重雄	国立情報学研究所 教授/副所長
委員	岡田 慧	東京大学 教授
委員	喜連川優	情報・システム研究機構 機構長 東京大学 特別教授
委員	金 太一	東京大学 特任教授
委員	高地雄太	東京科学大学 教授
委員	新保史生	慶應義塾大学 教授
委員	須藤 修	中央大学 教授
委員	谷口忠大	京都大学 教授
委員	寺田 努	神戸大学 教授
委員	萩田紀博	大阪芸術大学 教授
委員	原山優子	東北大学 名誉教授，日本工学アカデミー 国際委員長
委員	平野 晋	中央大学 教授
委員	間瀬健二	名古屋大学 名誉教授
委員	松尾 豊	東京大学 教授
委員	松田尚子	bestat 株式会社 代表取締役社長

要 旨

最近数 10 年に渡る AI の進化発展は目を見張るものがあり、AI は我々の生活において無視できない存在となり、欠かすことのできないインフラとなりつつある。今後も、高信頼化、汎用化、コモディティ化が進み、我々の生活のあらゆる場面で AI とヒトが協働し、社会を形成していることが予想される。一方、AI の社会導入が進むにつれて生じる懸念も明瞭化されつつある。AI とヒトとの共生による豊かな未来社会の創成、地域や人種、性別・ジェンダーなどを理由とした格差のない社会、戦争や侵略などの危機のない安全で平和な社会の実現をめざし、時期を遅れることなくこれらを検討し、AI の社会導入を進めるべく、国際貢献と国益貢献のバランスが取れた選択肢を積極的に提言し、政策化への橋渡しを行うことが、プロアクティブで適切な対策を行うために急務である。

そこで、本プロジェクトでは、我が国の技術立国としてのリカバリーと社会全体の安寧の実現に向けた、これからの AI 技術とその利活用、政策への期待と課題に関する調査・分析に資するため、国内外の通信技術ならびに AI 技術に関する研究開発施設などの社会実装の現場、AI による社会変容や社会学的影響に関する研究施設、あるいは AI 規制を検討する立法・行政組織に対し、ヒアリング・意見交換等を実施した。これにより、法規制の整備を待たずして急速に普及・発展する AI 技術に伴う社会変容がどのようなものであったか(これからどうなっていくか)、どのようなリスクマネジメントが取られてきたのか(取られようとしているのか)を正しく捉えた上で、AI の社会実装において内在する真の課題、我が国の強みは何か、これから何をするのが良いのかなど、生成 AI をはじめとする AI の技術開発・利活用戦略ならびに政策シナリオを作成し、国民個々人がそれらを理解できる形で報告書としてまとめ、政策提言として公表することを目的とする。

本報告書 兼 提言書では、以上の背景、問題意識、目的のもとでの調査研究の結果に基づき、まず、AI を取り巻く現状の概説と本書の作成に至った動機および立ち位置・目的について取り纏め、次いで現在までの AI の発展を支えている要素技術として、情報通信技術、人工知能の基盤・応用技術の研究開発動向について概説する。さらに、次世代人工知能技術に伴う技術的課題、倫理的課題、法的課題を含む課題とリスクについて整理し、これらの課題・リスクを踏まえた上で、次世代人工知能技術に求められる技術革新の方向性、倫理・法規制について概説し、最後に、「5G/6G 時代の AI 利活用戦略」/「生成 AI をはじめとした AI による社会変容とリスクマネジメントに関する調査研究」総括としての政策提言として、豊かな社会の実現に向けて AI の進化と深化、利活用を進めていくうえで、下記に示す 4 つの観点で AI の研究開発・社会を推進することが重要であることを明確にした。

1. AI の透明化, 高信頼化
2. AI の汎用化
3. ヒト-AI インタフェースのマルチモーダル化, AI インタフェースの進化, Zero UI 化
4. 法規制および産業教育振興など社会体制の整備, 倫理・価値観の多様化への対応

目次

1. はじめに.....	1
2. 現在までの情報通信技術および人工知能技術の研究開発動向	3
2.1. 情報通信技術の研究開発動向	3
2.1.1. 移動通信システムの歴史	3
2.1.2. 移動通信システムの現状と課題	6
2.1.3. 移動通信システムの今後の展望	8
2.2. 人工知能の基盤技術・応用技術の研究開発動向	10
2.2.1. 人工知能の定義・歴史	10
2.2.2. 人工知能の基盤技術の現状と課題	13
2.2.3. 人工知能の応用技術の現状と課題	15
3. 次世代人工知能技術の研究開発および社会実装の動向	22
3.1. 次世代人工知能技術に向けた技術革新の方向性	22
3.2. 次世代人工知能技術に向けた倫理・法規制	30
3.2.1. AI にかかる倫理・法的課題	31
3.2.2. AI にかかり懸念される社会的影響	34
3.2.3. AI にかかる過度な期待と報道	35
3.3. 次世代人工知能技術の社会実装に向けた施策・運用戦略	36
3.3.1. 諸外国の AI 政策と施策	36
3.3.2. AI 産業動向	50
3.3.3. AI にかかる学術教育，学術振興	52
4. 次世代人工知能技術の適切な社会実装・運用に向けた提言	54
参考文献	57
会議開催記録	59

1. はじめに

1970-80 年代に端を発した情報技術の発展は著しく、5G/6G 通信技術およびそれに伴う情報のネットワーク化、ならびにディープニューラルネットワークをはじめとした人工知能 (Artificial Intelligence, AI) によるデータ解析能の発展は、我々の社会や生活のあり方を一変した。日常業務において情報検索が占める割合は年々加速的に増加し、リモートでの情報交換や情報共有も進んだ。高速無線通信の社会実装は、情報をデータ容量制限から解放し、地理的物理的位置に依存しない自由な情報ネットワークを可能にした。インターネット上に広がるいわゆるビッグデータに対して AI を適用して、さまざまな情報分析や戦略決定に利用する動きも始まり、いまや多種多様な情報がネットワークを介して結合し、リモートで解析され、処理・制御されるようになっている。

また、ChatGPT (OpenAI 社、米国)をはじめとした AI 企業の台頭により、生成 AI は進化を遂げて社会実装が急速に進んでおり、AI の応用領域や社会的な影響度、ヒトの必要スキルが大きく変貌しつつある。生成 AI により AI 制御インタフェースが自然言語制御(プロンプト制御)になったことで、AI を使用するために必要とされた技術や知識に対する障壁が随分と取り除かれた。AI の使用人口は激増し続けており、その活用用途は多様化し続けている。一方、スマートフォンの普及や 5G 高速データ通信の確立によって各人が携帯コンピュータを常時身につけ、データを通信する時代になっている。今後、スマートフォンでの AI 利用が加速すると予想され、スマートフォンは現在のデータ端末(データを閲覧し送受信するための端末)としての利用から、AI 端末(AI をはじめとしたアルゴリズムによりデータを処理、加工して活用するための端末)への変容が期待される。

一方、社会体制は、2019 年末より COVID-19 感染症の急速拡大により、社会体制の各所で遠隔化やデジタル化が進み、今やオンラインコミュニケーションは日常となった。生成 AI の登場ならびに COVID-19 感染拡大による影響はどちらも最近 3-4 年の社会変容であり、また、6G 等高速通信の登場による社会変容も鑑みるに、これまでに検討されてきた AI 社会実装に関わる検討では内容として不十分と言わざるを得ず、情報を最新に更新した上で早急な再検討が必要である。

AI 技術の開発・利活用については、光と影の部分の両面での議論が活発である。米中では開発・利活用が先行して進んでいる一方、欧州では公平性・透明性を重視した人間中心の規制議論の中で慎重な開発・利活用が進められている。生成 AI をはじめとした AI 技術が高速大容量ネットワークとともにさらなる発展を遂げると、誰でも AI の恩恵を享受できるようになり、社会の根源的な進化につながる可能性もある。しかし、情報漏洩、改ざん、システム破壊、サービス妨害などへの配慮に加え、インクルーシブ性、個人を尊重しつつ、社会の利益に資する AI の提供価値が持続できるよう、セキュリティを重視した AI 技術の適正開発、維持管理システムの構築や運用レベルでの対策が今まで以上に求められることになる。また、AI の性能向上に比例して膨大に膨れ上がって

いる AI の学習に必要な電力消費量や AI の利活用の促進に伴うヒトの思考力低下などの問題点・懸念もある。さらには、高速通信技術、生成 AI 技術など革新的技術の登場に加え、COVID-19 感染症拡大、世界各地での戦争や紛争などによる地政学的な変容も考慮した検討を行わなければならない。

このように近年で急速に変容し、現在進行形で変容し続けている技術・社会背景を鑑み、次世代 AI 技術の適正開発・社会実装において、我が国が世界標準から後れをとらないよう、経済安全保障、地政学的リスク、規則、評価基準、安全・安心、公平性、尊厳・プライバシー、透明性等の倫理リスク、セキュリティ、ガバナンスなどの機微な問題に配慮しつつ、AI 技術に伴う効果・影響の多面性を正しく認識し、それに関する情報を広く社会化するとともに、社会全体の安寧の実現に向け、国際貢献と国益貢献のバランスが取れた選択肢を積極的に提言し、政策化への橋渡しを行うことがプロアクティブで適切な対策を行うために急務である。

そこで、本プロジェクトでは、我が国の技術立国としてのリカバリーと社会全体の安寧の実現に向けた、これからの AI 技術とその利活用、政策への期待と課題に関する調査・分析に資するため、国内外の通信技術ならびに AI 技術に関する研究開発施設などの社会実装の現場、AI による社会変容や社会学的影響に関する研究施設、あるいは AI 規制を検討する立法・行政組織に対し、ヒアリング・意見交換等を実施した。これにより、法規制の整備を待たずして急速に普及・発展する AI 技術に伴う社会変容がどのようなものであったか(これからどうなっていくか)、どのようなリスクマネジメントが取られてきたのか(取られようとしているのか)を正しく捉えた上で、AI の社会実装において内在する真の課題、我が国の強みは何か、これから何をするのが良いのかなど、生成 AI をはじめとする AI の技術開発・利活用戦略ならびに政策シナリオを作成し、国民個々人がそれらを理解できる形で報告書としてまとめ、政策提言として公表することを目的とする。

本報告書 兼 提言書は、全 4 章で構成され、まず本章にて序論として AI を取り巻く現状の概説と本書の作成に至った動機および立ち位置・目的について述べた。第 2 章では、現在までの AI の発展を支えている要素技術として、情報通信技術、人工知能の基盤・応用技術の研究開発動向について概説する。第 3 章では、次世代人工知能技術に伴う技術的課題、倫理的課題、法的課題を含む課題とリスクについて述べる。最後に、第 4 章では、第 3 章の課題・リスクを踏まえた上で、次世代人工知能技術に求められる技術革新の方向性、倫理・法規制について言及した後、次世代人工知能の社会実装に向けた施策や運用戦略について考察し提言する。

2. 現在までの情報通信技術および人工知能技術の研究開発動向

2.1. 情報通信技術の研究開発動向

現代の人工知能，特に現在のディープラーニングを基本とした AI は，データ駆動型であり，テキストデータ，各種センサデータ，画像・動画データなど，多種多様でかつ大規模なデータを効率よく伝送・集約するために，高速かつ大容量の通信環境が欠かせない．したがって，情報通信技術はこれまでの，そしてこれからの AI 技術の発展を支える重要な要素技術のひとつである．ここでは，情報通信技術，特に AI 技術と親和性の高い高速・大容量無線通信技術として現在第 5 世代 (5G) まで発展している移動通信システムを中心とした研究開発動向について述べる．

2.1.1. 移動通信システムの歴史

移動通信システムの発展の歴史を図 1 に示す．1979 年に第 1 世代となる移動通信システムの商用化が開始されて以降，移動通信システムは約 10 年周期で世代交代を繰り返し，第 5 世代に至る現在もなお，大容量化・高速化の進化・発展を続けている．

第 1 世代移動通信システム (1G) は，アナログ変調方式によって音声電波に載せることで音声の伝送を可能にし，主に音声通話サービスに利用されていた．自動車電話サービスから始まり，携帯電話サービスも開始された．1 つの基地局がカバーするエリア (セル) を小さめに設定した上で，基地局を多数設置してエリア全体をカバーするセルラー方式を採用するとともに，周波数資源の効率的活用と消費電力の省力化のために，離れたエリアに位置するセル同士で同じ周波数を繰り返し利用する周波数再利用を導入するなど，以降の携帯電話ならびに移動通信システムの根幹をなす基盤が確立された．

1990 年代に入ると，デジタル変調方式を取り入れた第 2 世代移動通信システム (2G) が登場し，パケット交換技術による通信が可能となり，音声通話サービスに加え，ショートメッセージサービスを中心としたデータ通信サービスも提供されるようになった．欧州の GMS (Global System for Mobile Communications) が先導役となって携帯電話の通信規格の標準化も進められ，米国では D-AMPS (Digital Advanced Mobile Phone System)，本邦では PDC (Personal Digital Cellular) と呼ばれる独自の通信方式が標準規格として採用されるなど，国内に限定されるものの，異なる事業者間でのローミングが容易になった．デジタル技術の導入によってデータの符号化・圧縮が可能となり音声品質・通信速度が向上したことで，携帯電話は当初，一部のビジネスユースにとどまっていたが，一般消費者の手に普及していく大衆化の局面を迎えた．

2000 年代には，国際電気通信連合 (International Telecommunication Union, ITU) によって標準化された「IMT-2000」(International Mobile Telecommunication 2000) 規格に準拠した，第 3 世代移動通信システム (3G) が導入され，異国間での国際ローミングも可能となっ

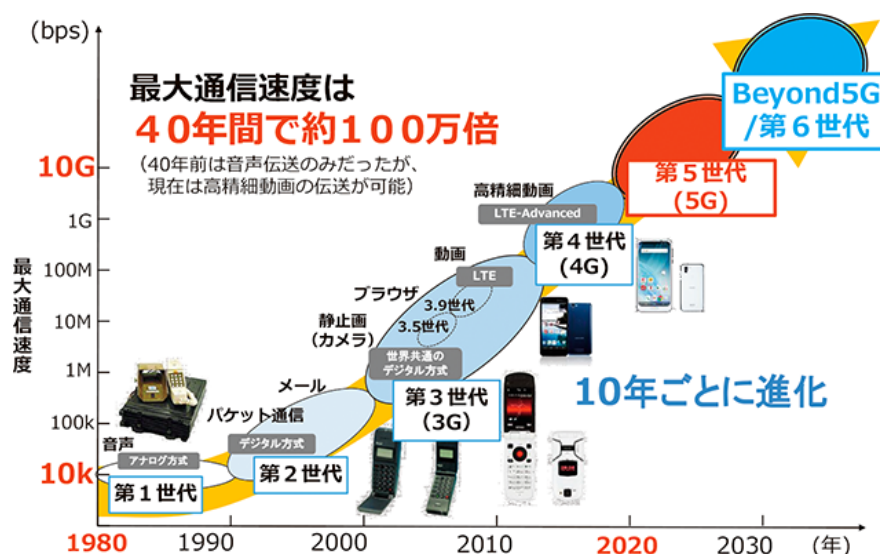


図1 移動通信システムの発展の歴史
(出典: 総務省 令和5年版情報通信白書^[1])

た. 3G では、ユーザを識別するための拡散符号と呼ばれるコードを割り当てた上でデータを広帯域化して送受信するスペクトラム拡散技術を利用した符号分割多元接続 (Code Division Multiple Access, CDMA) と呼ばれるアクセス方式を採用することで、多数のユーザが同一周波数を同時に利用できるようなした. これにより、2G よりも同時接続性とデータ伝送レートが飛躍的に向上し、数百 kbps から数 Mbps 程度までのデータ転送速度が確保されるようになったことで、それまで音声通話やショートメッセージのやり取りが中心であった携帯端末は、ウェブブラウジングやゲームなどの多様なコンテンツの利用を可能とする高機能端末へと進化を遂げた.

2007 年に Apple Inc. が世界に先駆けてスマートフォン「iPhone」を発表すると、フィーチャーフォンからスマートフォンへの本格的な移行が始まった. それまでは回線交換方式による音声通話とパケット交換方式によるデータ通信が併存していたが、2010 年代には、音声通話も含めてパケット交換方式に一本化したオール IP (Internet Protocol) ベースのアーキテクチャを採用した LTE (Long Term Evolution) が第4世代移動通信システム (4G) として登場した. 4G は ITU が定める「IMT-Advanced」規格に準拠している. オール IP ベースを採用することで、音声もデータもひとつのパケットコアで一元的に処理可能となり、ネットワーク運用管理がシンプルかつ柔軟になったことに加え、回線交換の制約を受けずに低遅延の高速通信を実現できるようになった. 4G では、周波数帯域を互いに直交する多数のサブキャリアに分割し、周波数利用効率を上げる OFDM (Orthogonal Frequency Division Multiplexing) というマルチキャリア変調方式と、複数のアンテナで並列にデータを送受信する MIMO (Multiple Input Multiple Output) を組み合わせることで、限られた周波数資源でも大容量かつ高速な通信を可能にした. 数十～数百

Mbps というブロードバンドレベルの通信が実現したことで、高解像度動画のストリーミングやオンラインゲームへのアクセス、クラウドサービスの活用が携帯端末上で当たり前に行われるようになった。さらに、複数の周波数帯を束ねて1つの通信回線としてデータの送受信を行うキャリアアグリゲーションの技術を導入することで、1Gbps 規模のさらなる高速通信を実現した LTE-Advanced も登場し、固定回線並みの通信品質・速度がモバイル環境で得られる時代が到来した。

携帯電話網の拡張とともに、無線 LAN の Wi-Fi や Bluetooth といった短距離無線技術も成熟期を迎えた。IEEE 802.11 シリーズで標準化された Wi-Fi は、11b から 11g, 11n, 11ac, そして 11ax (Wi-Fi 6) へと発展し、そのたびに通信速度と安定性を向上させることで、家庭やオフィス、公共空間での無線インターネット接続を当たり前の存在にした。Bluetooth はウェアラブルデバイスやヘッドセット、カーオーディオなどの機器間接続を支える近距離通信技術として数多くのデバイスと連携し、モバイル端末とのシームレスなやり取りを可能にする。こうした多様な無線技術の進歩は、携帯電話以外の端末がインターネットに接続する「IoT」(Internet of Things) への道を切り開くことにもなった。

さらに 2000 年代から現在にかけて、固定系通信においてもブロードバンド化が進み、電話回線を利用してインターネット通信を実現していた ADSL (Asymmetric Digital Subscriber Line) から光ファイバへの移行が一気に加速した。数百 Mbps から 1Gbps の常時接続を可能とする FTTH (Fiber To The Home) の普及に伴い、動画ストリーミングやオンライン会議、クラウドストレージといったサービスが社会基盤化し、場所を選ばない高速インターネットの利用がごく自然なものとして受け入れられるようになった。また、IP 化により、音声通話や映像通信もインターネットに完全に統合され、回線交換網からパケット交換網への移行が飛躍的に進展していった。その結果、通信事業者はネットワーク管理や新しいサービス創出を柔軟に行える体制を整え、スマートフォンや家庭向けブロードバンドサービスと組み合わせたパッケージを展開しやすくなった。

こうした流れをさらに先鋭化するように、2020 年前後からは第 5 世代移動通信システム(5G)が立ち上がった。5G は ITU が定める「IMT-2020」規格に準拠しており、最大通信速度(下り)20 Gbps の高速・大容量通信、1 msec 以下の低遅延、1 km²あたり 100 万台の同時接続する能力が目標とされている。これを支えるために、ミリ波帯の活用、4G の MIMO よりもはるかに多いアンテナを用いた Massive MIMO、高い指向性を持って一定方向に電波を送受信するビームフォーミング、ネットワークを仮想的に分割してさまざまな用途に合わせて制御するネットワークスライシング、分散処理のためのエッジコンピューティングなど、数多くの技術が組み合わせられている。こうした高度化は映像配信やモバイルインターネットにとどまらず、自動運転車両間のリアルタイム通信や遠隔医療手術、産業用ロボットの精密制御、没入型の AR/VR サービスなど、多様な領域に巨大なインパクトをもたらすことが期待されている。また、ローカル 5G という形態に

より、特定の企業や地域が独自の 5G ネットワークを構築し、高信頼かつ低遅延の通信環境を活かして課題解決やイノベーションを推進する事例も増え始めている。

このように、移動通信システムは、1G のアナログシステムから始まり、2G のデジタル化、3G の国際標準化とデータ通信の本格化、4G のスマートフォン普及に伴うブロードバンド化とオール IP 化、そして 5G の超高速・低遅延・多数同時接続化へと、約 40 年の間に大きな進化を遂げてきた。一人ひとりが常に持ち歩くようになった携帯端末は、人間同士のコミュニケーションを支える枠組みを超え、あらゆる情報の取得・送受信、さらには機械同士をつなぐネットワーク基盤の中心として機能しつつある。

2.1.2. 移動通信システムの現状と課題

現在、世界各国で 5G が導入され、大手通信事業者を中心に商用サービスが展開されている。5G では「高速・大容量」「低遅延」「多数同時接続」の 3 つが大きな特徴とされ、先述の通り、具体的には、「最大 20 Gbps のデータ通信速度」、「1 msec 以下の遅延」、「1 km² あたり最大 100 万台の端末の同時接続」が技術目標として掲げられている。ここでは、これら技術目標に対する現状の達成状況と、達成できていない部分に関する原因や課題について以下に述べる。

高速・大容量(最大 20 Gbps のデータ通信速度)については、限定的な条件下では既に実現されている。例えば、ミリ波帯(28 GHz 帯)を活用し、小セル内で基地局から近距離かつ障害物がほとんどない環境であれば、理論値に近い Gbps レベルの通信速度が得られることが実証されている。また、本邦における都市部の一部地域では、実際に 1～2 Gbps 程度の速度が計測された事例も報告されている。ただし、建物や樹木などの障害物に対して弱いミリ波帯の周波数帯を利用する場合、見通しの良い環境に限られ、かつエリアを細かく分割して多数の基地局を設置しなければならない。これらのインフラ整備には膨大なコストと時間を要するため、ローカルエリアやイベント会場など用途を絞った形での運用が中心となっているのが現状であり、まだ全国規模での高品質なミリ波帯の利用は実現できていない。一方、Sub6 帯(3.7 GHz 帯と 4.5 GHz 帯)と呼ばれる周波数帯の 5G も検討されているが、こちらはより広範囲のカバレッジが確保できるものの、ミリ波帯と比べると通信速度は抑えられ、高速化の伸びしろが制限される。そのため、こうした周波数帯の使い分けと設備投資の最適化が目下の課題となっている。

低遅延(1 msec 以下)については、現状の 5G 商用ネットワークではまだ十分に達成できていないのが実情である。この目標値を実現するには、無線区間での遅延削減だけでなく、端末からクラウドやサーバまでを含めたエンドツーエンドのネットワーク設計が重要となる。そこで導入が期待されているのが、基地局近傍にサーバを置くエッジコンピューティング(Multi-access Edge Computing, MEC)である。しかし、エッジサーバを大規模に配置・運用するには企業・地域ごとの導入費用や運用リソース確保が不可欠であ

り、まだ実証段階に留まっている事例が大半である。また、ネットワークスライシングによって低遅延が求められるサービスへ優先的にリソースを割り当てる技術も提案されているが、これらの仕組みは標準化や事業者間の運用ルール策定などが途上であり、社会実装は段階的に進められているのが現状である。

多数同時接続(100万台/km²同時接続)についても、現時点で実環境での本格的な導入事例は限られている。膨大なIoTデバイスが相互に通信し合うスマートシティや工場オートメーション、スマート農業などのユースケースが想定されているが、実際にはデバイスやセンサの導入コスト、データプラットフォームの整備、運用体制、セキュリティ対策など、技術以外の要因も絡むため、数十万～百万規模のデバイスを同時運用するユースケースは一部の先進的な実証プロジェクトに限られている。また、周波数帯や基地局の配置設計によっては、同時接続数が多いほど干渉制御やスループット確保が困難になり、技術面・コスト面での課題が顕在化している。

上述の目標の達成が不十分である要因をまとめると、大きく技術的課題と社会的課題に分けられる。技術的課題としては、周波数帯の特性に起因するカバレッジの問題や、エッジコンピューティングやネットワークスライシングなどの先進技術の標準化・実装が未成熟であることが挙げられる。基地局やアンテナの設置・運用に伴うハードウェア面のコスト負担、さらにソフトウェア側ではネットワーク機能を仮想化するNFV(Network Function Virtualization)やネットワーク制御をソフトウェアで行えるようにするSDN(Software-Defined Networking)の導入など運用管理が複雑化していることが開発・展開のボトルネックとなっている。

社会的課題としては、まずインフラ投資コストを回収するための明確なビジネスモデルの構築が追いついていない点が挙げられる。特にローカル5Gなどでは、大規模なネットワーク工事費用を誰がどのように負担するかが大きな論点となる。さらに、地方や過疎地への展開は都市部に比べて採算が取りにくいいため、地域格差が拡大する懸念も指摘されている。また、環境負荷やエネルギー効率に関する問題も無視できなくなっている。小セル化と基地局・エッジサーバの設置増加に伴う通信インフラ全体の電力消費量の増大は、持続可能な社会の実現に逆行しかねない。半導体の低消費電力化や再生可能エネルギー活用などの取り組みが進められているが、ネットワーク全体の省エネ化は喫緊のテーマであり、SDGsやカーボンニュートラルの要請とも融合しながら進化していくことが求められる。次に、セキュリティやプライバシーの確保も大きな課題として挙げられる。IoT端末の急速な普及に伴い、サイバー攻撃や個人情報漏洩などのリスクが高まる一方、対策コストや責任分界点などの整備は十分に進んでいない。さらに、社会的合意形成や規制の問題もあり、ネットワークスライシング等の新技術をどう位置付け、どのようなガイドラインやルールの下で運用するかが十分に定まっていない。5Gに関するセキュリティの枠組みとしては、総務省が策定した「5Gセキュリティガイドライン」^[2]などのガイドラインは存在するものの、今後の導入が検討されつつある超分散

型ネットワーク，仮想化技術，無数の IoT デバイス接続など，5G の先に想定される高度な運用形態をすべてカバーできているわけではないことから，継続的な見直しが欠かせない状況にある。

2.1.3. 移動通信システムの今後の展望

5G の導入により，「高速・大容量」「低遅延」「多数同時接続」という新たな可能性が切り開かれたものの，先述の通り，そのすべてを十分に達成するには技術的・社会的な課題がまだ残されている．こうした現状を踏まえた上で，5G の普及と並行し，2030 年代頃の導入を目指して Beyond 5G/6G を見据えた研究開発も世界規模で加速中である．Beyond 5G/6G で掲げられている要求条件を図 2 に示す．Beyond 5G/6G では，現在のミリ波帯よりもさらに高い周波数帯であるテラヘルツ帯の活用が検討されており，これにより最大通信速度 100 Gbps 超にまで高められる可能性がある一方，これまで以上に電波の直進性が強くなり小セル化が進むことから，カバレッジ改善のために基地局数のさらなる増加と高密度配置が求められる．その結果，インフラ投資負担はさらに大きくなることが予想されるが，こうした課題を解決する手段として，O-RAN (Open Radio Access Network) によるネットワーク機器の相互運用性の確保や，事業者・自治体・企業が共同でインフラ設備を整備する新たなビジネスモデルが模索されており，将来的には小規模・分散型のネットワーク設計と大規模ネットワークの融合も検討されている．

アンテナからの見通しが効かない遮蔽物が多い環境におけるカバレッジ改善の検討として，周波数帯を柔軟に組み合わせるマルチバンド運用や高度なビームフォーミング技術に加え，電波特性を制御可能なメタサーフェスレンズを取り付けることで窓ガラスなどの既存の構造物を電波の中継・制御に活用して電波環境を適応的・動的に制御する

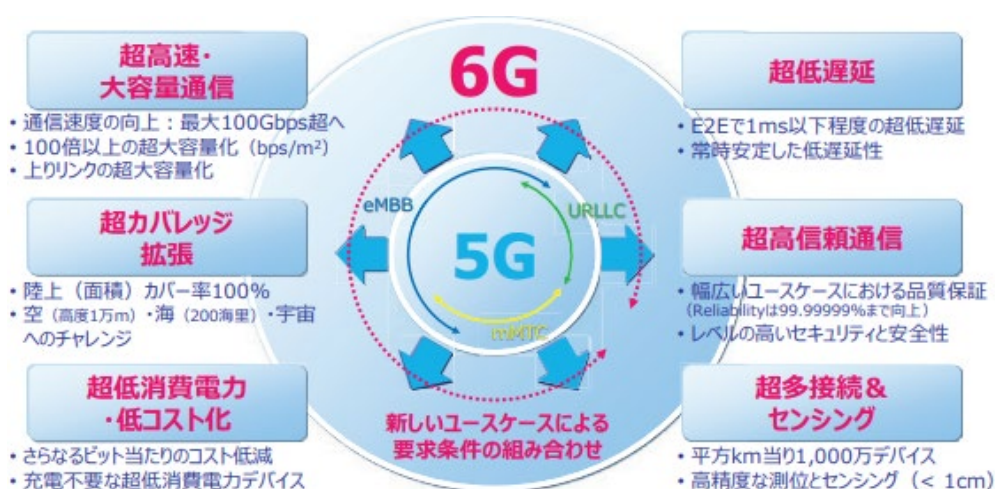


図2 Beyond5G/6G で掲げられている無線ネットワーク技術の要求条件
(出典: 株式会社 NTT ドコモ ホワイトペーパー 5G の高度化と 6G^[3])

IRE (Intelligent Radio Environment) 技術などの研究開発も進められている。また、地上だけにとどまらず、空・海・宇宙領域へカバレッジを拡張する取り組みも重要視され、静止衛星 (Geostationary Earth Orbit, GEO)、低軌道衛星 (Low Earth Orbit, LEO)、高高度疑似衛星 (High-Altitude Platform Station, HAPS) といったプラットフォームとの統合ネットワークを構築することで、船舶や航空機、離島・山間部などの従来カバーが難しかった地域を含めた広域での高品質通信を実現しようとする動きも活発化している。このような広域カバレッジ拡張は、防災・減災や物流など幅広い分野での利用が期待されている。

また、超低遅延と超高信頼化に関しては、エッジコンピューティングやネットワークスライシングといった 5G の基盤技術をさらに発展させ、端末からクラウド・サーバまでのエンドツーエンド区間を徹底的に最適化するアプローチが取られている。具体的には、基地局の近傍に高性能サーバを配置してデータ処理を分散する MEC のさらなる導入を通じ、IoT 機器を含む多様な端末に対して必要最小限の通信経路で応答できるようにすることで、1 msec を下回る超低遅延と、経路分散による障害の局所化やサービス継続性の向上といった高信頼化の実現が検討されている。こうした技術は自動運転、遠隔医療、ロボット制御など、特に安全性と可用性を要する領域での実用化が期待される反面、大規模なサーバ群やソフトウェア制御機構の整備には多大なコストと運用リソースが必要となることから、インフラ投資だけでなく事業者間の協調や標準化も不可欠である。

さらに、Beyond 5G/6G で掲げられている 1 km² あたり 1000 万台以上の端末の超多数同時接続の目標に対しては、IoT デバイスからの膨大なトラフィックを効率よく制御できるネットワーク運用とセキュリティが大きな鍵を握る。このために、AI を活用した自律的なトラフィック制御や異常検知をネットワーク管理に組み込むことが検討されており、必要に応じて仮想ネットワークを自在に生成・廃棄できるクラウドネイティブな基盤を構築することで、端末数の急激な増減にも柔軟に対応しようとしている。一方で、センサやデバイスが増えれば増えるほど、サイバー攻撃や情報漏洩のリスクが拡大するため、従来の集中管理ではなく、ブロックチェーンや分散型 ID などの仕組みを併用したセキュリティ強化が模索されている。これらの仕組みをクラウドネイティブに統合し、ネットワーク規模の拡張に応じて柔軟にスケールアウトする考え方が主流になりつつある。これらの実現には多様な企業・団体が協働し、標準化団体の協力や国際的な合意形成を通じてインフラ構築を推進する必要がある。技術開発と社会実装を同時に行う体制づくりも重要となる。

以上のような Beyond 5G/6G の研究開発は、周波数資源の限界を突破しつつネットワークの柔軟性と信頼性を高めることで、5G の導入とともに顕在化した課題を抜本的に解決するための方向性を示している。ただし一方で、サイバー攻撃や個人情報保護といった社会的リスクへの対策は一層困難化し、通信だけでなく AI・IoT など関連領域も含めた包括的な視点での法規制や標準化の整備が求められるだろう。さらに、地球規模の

環境問題やエネルギー問題とも無縁ではなく、超高速・多拠点通信をいかに省エネルギーかつ持続可能な形で運用できるかが今後の焦点となる。Beyond 5G/6G へ向けた今後の移動通信システムは、通信環境への依存度がますます高まる AI 社会においても不可欠なベースとなり、各国の産業や社会生活を支える原動力としての役割をより一層担っていくことが予想されるが、そのためには、技術の高度化はもとより、社会基盤としての通信インフラのあり方や価値観の再構築といった、より広範な視点からの検討が必要になっている。

2.2. 人工知能の基盤技術・応用技術の研究開発動向

高速・大容量の移動通信システムは、ビッグデータを含む膨大な情報のやり取りを可能にし、さらには超低遅延ならびに多数同時接続を実現することで、社会実装におけるデータ流通のあり方を根本的に変えつつある。現代社会において、AI は高度なデータ解析や知的タスクの自動化を担う基盤技術として急速に普及・進化してきたが、こうした情報通信環境の整備は、AI 技術のさらなる発展を支える不可欠な要素であり、現代のディープラーニングおよび生成 AI をはじめとするデータ駆動型 AI の可能性を大きく押し広げている。特に Beyond 5G/6G のような超高速通信が普及することで、大規模学習モデルのためのデータ集約や分散処理、学習済みモデルの迅速な配信・共有などが容易になるだけでなく、クラウドやエッジとの連携によるリアルタイム推論も大きく加速することが予想される。

ここでは、AI 技術を「基盤技術」と「応用技術」に大別し、基盤技術については AI そのものの歴史的経緯、現在の学術・産業界における技術開発の動向と課題を述べる。また、応用技術として、AI の社会実装の現状や具体的なユースケース、実際の課題などを整理した上で、次世代の AI 技術に関する将来展望について考察する。

2.2.1. 人工知能の定義・歴史

Artificial Intelligence(人工知能)とは、1956 年に開催されたダートマス会議にて、米国の計算機科学者 John McCarthy 教授によって提唱された言葉・概念である。同教授によれば、人工知能は「知的な機械、特に知的なコンピュータプログラムを作る科学と技術」^[4]として定義づけられている。現段階で人類はまだ「知能」の完全な解明と定義には至っていないことから、人工知能についても国際的な合意が得られた明確な定義は存在せず、専門家の間でもその定義は異なる^[5]のが現状であるが、本稿では上述の定義に従う。また、一般的に、人工知能は、特定のタスク・領域に特化した「目的特化型 AI (Specialized AI)」とあらゆるタスク・領域に対して汎用的に対応する「汎用型 AI (Artificial General

Intelligence, AGI)」の2つに大別されることがある。一部には、汎用化(Generalization)と、意識(Consciousness)を持つことを混同している意見も見られるが、汎用化／汎用性と意識はまったく別の概念である。汎用性とは、さまざまな広い範囲の目的に対して適用でき、特定の目的に特化していないことを指す。一方、意識は内的な経験や主観的な体験として内から生じるものであるとともに、心が知覚を有しているときの状態あるいはそのメカニズムを指し、はっきりと定義することが難しいとされている感情や情動などと関係を持つ。他者と区別した自分を意識することである「自我」とも関わりを持ち、AIが意識を持ったとき、我々の意図しない動作をAIが自ら選択する可能性が生じてくる。意識を持つことと、AIが環境や状況を識別して適切なタスクを行うプロセスを選択し実行することは別であり¹⁰⁾、意識の実装とは無関係に、適切な判断プロセス(アルゴリズム)の実装によってAIを汎用化できる。AIの汎用化、すなわちAIの適用範囲をさまざまな用途に広げるため、目的特化型AIを至適に選択して実行するフレームワークに関する研究開発について、OpenAI(米国)のSwarmなどいくつかの取り組みが始まっており、現時点でもっとも激しい競争を行っている研究開発領域のひとつでもある。一方で、汎用レベルを定量化することは難しく、AGIの実現に向けた指標を示すことを目的としてAGIのベンチマークや段階的なレベルを定義している文献^[6-8]も存在するが明確な定義には至っておらず、どこまでを目的特化型AIとし、どこからをAGIとするかについては、ケース毎に異なり細やかな解釈が必要である。AIはコモディティ化と多種多様相化に向かっており、さまざまな高性能な目的特化型AIが発表されている一方で、汎用化については一部を除いて現時点で十分とは言えず、今後の発展が期待される。汎用レベルについては議論の余地を残すが、今後、本稿ではAGI汎用レベルの議論は深く行わないこととし、ある特定の領域に限定された知的タスクを行うAIを目的特化型AI、あらゆる知的タスクに対して自律的にこなせる(適切なアルゴリズム・AIを選択して使いこなせる)ことを志向したAIをAGIと定義して議論を進める。

AIの研究領域は非常に多岐に渡り、そのアプローチや応用範囲も時代とともに大きく変遷してきた。人工知能技術の現在までの変遷の俯瞰図を図3に示す。1956年のダートマス会議に端を発する黎明期から現在に至るまで、AI研究は4度のブームを経ながら盛衰を繰り返し、大きく発展してきた。

1950年代後半～1960年代にかけての第1次AIブームは、推論・探索の時代と呼ばれ、コンピュータによる論理推論や探索アルゴリズムを通じて、記号的表現に基づくシンボリックAIの枠組みで知的行動を再現しようとする動きが活発化した。しかし、当時は計算機的能力やアルゴリズム上の限界から、「トイ・プロブレム」と呼ばれるような迷路やパズルなどの単純な問題しか扱うことができず、現実社会の複雑な問題には対応できないことが明らかになった。そのため、期待されていた実用に耐えるほどの成果が出せず、いわゆる冬の時代を迎えることとなった。

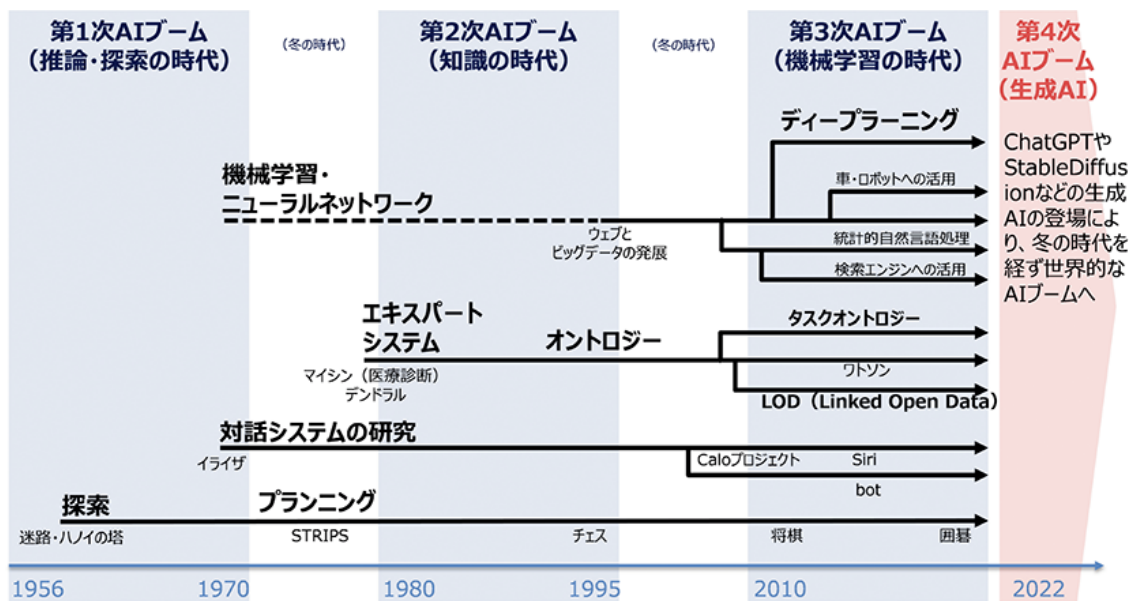


図3 人工知能・ビッグデータ技術の俯瞰図
(出典: 総務省 令和6年版情報通信白書^[9])

1980年代～1990年代に生じた第2次AIブームは、知識の時代と呼ばれ、エキスパートシステムなどの知識工学が中心に取り組みました。熟練者が持つ知識や推論ルールをコンピュータに組み込むことで、特定領域で専門家並みの判断を行うシステムの実現を目指したものであり、故障診断など一部の領域で成果を挙げた。しかし、知識獲得と維持に膨大なコストがかかり、さらに適用範囲が限定的であったことから、想定ほど実用が進まず、再び冬の時代を迎えることとなった。

ところが、2000年代に入るとインターネットの普及に伴うデータ流通量の爆発的な増加とコンピュータの高性能化によって、大規模データと十分な計算資源が得られるようになり、機械学習が飛躍的な進化を遂げ、機械学習の時代と呼ばれる第3次AIブームが到来した。機械学習手法の1つであるディープラーニングとして知られる、多層構造を持つニューラルネットワークに大量のデータを学習させることで、画像認識や音声認識、自然言語処理など幅広い分野で従来手法を大きく上回る成果が得られるようになった。特に画像解析の問題を主に扱う畳み込みニューラルネットワーク (Convolutional Neural Network, CNN) の一種である AlexNet^[10]が、2012年に ILSVRC と呼ばれる一般物体画像認識の精度を競うコンペティションにて圧倒的な性能で優勝したことで、ディープラーニングが大きな注目を集めるようになった。それまでの機械学習では、AIが識別するためのデータの特徴量をヒトが作成して与えていたが、ディープラーニングの登場により、AIが自ら識別に有用なデータの特徴量を獲得し予測できるようになり、識別モデルの飛躍的な性能向上をもたらした。また、2014年には、生成器と識別器を競わせながら学習するプロセスを通じて高品質な擬似データ生成を可能とする敵対的生成

ネットワーク (Generative Adversarial Network, GAN)^[11]が提案され、生成モデルも大きく注目されるようになった。以降、ディープラーニングをベースとした AI 技術は実用段階へと加速的に発展し、ビッグデータを高速・大容量でやり取りできる通信インフラの進化とも相まって、クラウドやモバイルサービスなど多様な領域での社会実装を急速に進める原動力となった。2017 年にはさらに、入力データの「どこに注目すべきか」を動的に学習して重要度の高い部分により強い重み付けを行う Attention 機構を導入した Transformer^[12]と呼ばれるモデルの登場により、機械翻訳や文章要約などの自然言語処理がさらなる質的な飛躍を遂げた。Transformer は Attention 機構という仕組みを軸に、離れた位置にあるデータの依存関係を効果的に学習できるという強力な能力を有することから、自然言語処理のみならず、音声・画像・マルチモーダル処理といった分野にも応用範囲が拡大し、現代の AI 研究・開発を支える中心的なアーキテクチャとして発展している。

そして 2022 年以降、Transformer をベースにした大規模言語モデル (Large Language Model, LLM) を利用した生成 AI が本格的に台頭することで、冬の時代を挟まずして、第 4 次 AI ブームが始まったとされる。Transformer ベースの LLM を大量のテキストデータで学習する GPT (Generative Pre-trained Transformer) シリーズをはじめとした、対話型の文章生成や高度な推論を可能とする生成 AI が急速に普及している。これにより、AI は単にデータを解析するだけでなく、ヒトのコミュニケーション領域にも深く入り込み、ユーザは自然言語によって AI を操作 (プロンプト制御) することが当たり前になってきている。また、拡散モデル (Diffusion Model)^[13]と呼ばれる新たな生成モデルが注目を集め、プロンプト制御で高品質な画像や動画も生成できるようになり、AI はコンテンツの制作面にも大きな影響を及ぼすようになった。こうした第 4 次 AI ブームでは、技術的ブレイクスルーと社会への急速な浸透が同時進行で進みつつある。

2.2.2. 人工知能の基盤技術の現状と課題

ディープラーニングや生成 AI を含む AI の基盤技術は、ビッグデータや高性能ハードウェア、そして 5G に代表される高速大容量通信環境の進展を背景に、これまでにないスピードで社会実装が進行している。代表的な生成 AI のひとつとして知られる GPT に関して言えば、2018 年 6 月に最初の GPT^[14]が公開されて以降、1 年経たずして 2019 年 2 月に GPT-2^[15]、2020 年 5 月に GPT-3^[16]、そして 2022 年 11 月には GPT-3.5 が ChatGPT (OpenAI Inc.) としてリリースされ、2023 年 3 月からはその後継モデルとなる GPT-4、さらに 2024 年 5 月からはテキスト以外の音声、画像、動画を含むマルチモーダル情報処理に対応した GPT-4o (GPT-4 Omni) が有料版として提供されている。GPT-1 から 4 までのシリーズを通してパラメータ数の飛躍的な増大とともに、Attention 機構を備えた Transformer アーキテクチャの高度化を重ねてきた。ChatGPT では現在、思考の連鎖

(Chain of Thought, CoT)を導入することで、より複雑で高度な推論能力を備え、これまでの GPT モデルのネーミングから脱却した新たな生成 AI モデルとして、o1 モデルがリリースされている。2025 年 1 月にはさらに、ARC-AGI と呼ばれる新たなベンチマークで o1 モデルを大きく上回り、特に数学や科学などの領域で卓越した性能を持つ o3 モデルのリリースが予定されるなど、生成 AI の進化は新たな局面を迎えている。また、生成モデルに限らず、認識モデルも含め、さらなる汎用性の獲得と応用範囲の拡張に向け、広範なタスクにおける基盤として利用可能な「基盤モデル(Foundation Model)」や、ロボットあるいは AI 自身を取り巻く環境の観測情報に基づき、環境(世界)の構造を学習によって獲得する「世界モデル(World Model)」の研究開発が加速化している。しかし、これらの基盤技術のさらなる大規模化・高度化とともに、多様かつ深刻な課題も浮上している。

AI 技術が進化している大きな要因として、インターネットやモバイル端末の普及・高度化により膨大なデータを取得できるようになった点が挙げられ、特に多様かつ大量なデータの事前学習により高い汎用性を獲得する基盤モデルの登場は、多くの領域でファインチューニングなどによる微調整を通じて効率的にタスクを実行できる可能性を示した^[17]。しかし、こうしたアプローチは、モデル規模や演算リソースを増やすほど性能が上がる、というスケーリング則に依存しており、学習に用いられるデータおよび計算量は膨大になる。結果的に電力消費やコスト、さらに環境負荷が増大し、持続的発展に対する懸念が高まっている。近年では、Green AI や再生可能エネルギーの活用などの取り組みが始まっているが、省電力化は依然として大きな課題として残っている。一方、AI の肝となる大規模学習を支えるには、強力なクラウドインフラと莫大な演算リソースが不可欠であり、資本力を持つメガテック企業が AI の開発・運用をリードする構造が固定化されつつある。学術研究やスタートアップ企業が同様の規模で競争するのは困難を伴うため、技術支配や市場独占の懸念も生じる。オープンソースコミュニティや公的機関の支援を通じて格差を是正しようとする試みはあるが、まだ十分とは言えない。

また、どのような根拠で特定の出力に至ったのかが不透明になるブラックボックスの問題は、深層学習ベースのモデルが持つ構造的な課題であり、基盤モデルや世界モデルのようにネットワーク構造が巨大かつ複雑化した場合には一層深刻化する可能性がある。医療や法務、公共分野など説明責任を要する領域では、モデルの決定プロセスがわからないことが社会実装の障壁となり得る。説明可能 AI(Explainable AI, XAI)の研究も進んではいるものの、高い精度と十分な解釈性を同時に満たす手法は依然として未成熟であり、さらにモデルが大規模化すればするほどこの問題の解決が難しくなるというジレンマを抱えている。

さらに、大量のデータを学習に用いる基盤モデルや世界モデルでは特に、データガバナンスやプライバシー保護の問題が浮上する。ウェブ上やセンサから収集した膨大な情報には、個人情報や著作物、企業の機密データなどが混在し得るため、適切なフィルタ

リングやライセンス管理を行わずに学習が進められると、情報漏洩や権利侵害を起こすリスクが高まる。さらに、精巧な生成 AI が普及したことで、フェイクニュースやディープフェイクが大量に生み出される危険性も増大しており、社会全体でのルール整備や技術的対策が追いついていないのが現状である。

このように現在の AI 基盤技術は、深層学習や生成 AI の革新に支えられながら急速に社会実装を進めている一方で、学習データと計算リソースの膨大化、開発主体の偏在、ブラックボックス化などの技術的課題と、データガバナンスやプライバシー保護、フェイクコンテンツや情報漏洩への対処、法規制の整備といった社会的課題が複合的に顕在化している。技術的側面では、演算不可の増大による環境負荷やコスト上昇を抑えつつ、説明可能性・透明性・安全性を担保した上で、AGI の実現に向けたさらなる技術革新・技術創発が求められる。一方、社会的側面では、データ利用ルールの整備や公正な研究開発体制の醸成、フェイクニュースやディープフェイクといったリスクへの対応が急務であり、ルール形成や協調的ガバナンスをめぐる議論が世界的に活発化している。Beyond 5G/6G との相乗効果によってさらに強力な AI 基盤技術が登場することが見込まれるが、持続可能かつ信頼性の高い AI 社会を構築するためには、これらの技術的・社会的課題を一体的に捉え、政策や倫理面を含めた多面的な取り組みを進めることが必要不可欠であるといえる。

2.2.3. 人工知能の応用技術の現状と課題

第 3 次ならびに第 4 次 AI ブームを牽引する現在の AI 技術は、多様な分野において、過去の AI ブームの際には到達できなかった実用レベルの推論能力を持つに至り、基盤技術としての枠組みを超えて、実社会の具体的な課題解決に活用されてきている。表 2 に社会領域ならびに AI 機能別の AI 応用技術の例を示す。ただし、実際には複数の AI 機能が横断的に組み合わせられていることも多く、ここで示す例はあくまで AI の応用範囲の動向を俯瞰するために位置付けた一例である点に留意されたい。各社会領域において、AI の利活用事例が増え、AI が従来の業務プロセスの効率化や最適化、あるいは新たな付加価値をもたらしている一方で、導入・運用においては多様なリスクや制約にも直面しているのが現状である。以下に各社会領域における AI 利活用の最新動向と、各社会領域が抱える AI 利活用に伴う課題について概説する。

医療・福祉・ヘルスケア

医療の現場では、ディープラーニングを活用した画像解析が病変検出をはじめとする診断支援で成果を挙げているほか、ゲノム情報を基にした個別化医療の試みが進展しており、患者一人ひとりの遺伝子変異や疾患リスクを AI が推定して治療法を提案する事例が増えてきている。また、福祉の現場では、介護施設や在宅ケアにおいて、見守りカ

表2 社会領域ならびに AI 機能別の AI 応用技術の例

社会領域	AI 機能				
	画像・映像解析	自然言語処理	ロボット・自動制御	データ生成	予測・最適化
医療・福祉・ヘルスケア	診断支援(病変検出・識別), 介護用見守りカメラ etc.	問診チャットボット, 相談窓口の自動応答 etc.	リハビリ支援ロボット, 生体情報モニタリング etc.	医療報告書生成, 治療計画作成 etc.	病床予測・管理, 薬剤・備品の在庫最適化 etc.
金融・保険	不正取引検知, 業務映像監視 etc.	チャットボット対応, 契約文書分析, RPA 文書処理 etc.	無人店舗の決済自動化(顔認証)・ATM 制御 etc.	顧客向け投資提案文書作成, 社内レポート作成 etc.	リスク解析, 顧客スコアリング, 保険料算定 etc.
教育	教材画像解析, オンライン授業映像の要点抽出 etc.	対話型学習支援, 自動作文添削ツール etc.	-	学習資料作成, 模擬試験・問題集の作成 etc.	学習到達度予測, カリキュラム最適化 etc.
マーケティング・広告・メディア	画像解析による視聴者分析, 広告効果検証 etc.	SNS 投稿分析, チャットボット(顧客対応) etc.	-	広告素材作成, 映像コンテンツ編集 etc.	需要予測・在庫最適化, 消費者行動分析 etc.
製造・産業	工場内品質検査, 外観検査の自動化 etc.	-	組立・搬送ロボット, 自動ライン制御 etc.	製造レポート作成, 手順書・マニュアル作成 etc.	生産スケジュール最適化, 在庫・輸送最適化 etc.
物流・交通	自動車の周囲環境認識, 倉庫内物品認識 etc.	配送問い合わせ応対チャットボット etc.	自動運転車両, 搬送ロボット, ピッキングロボット etc.	交通情報まとめ, 輸送効率化シミュレーション etc.	配送ルート最適化, 需要予測に基づく動的配車 etc.
農業・水産・畜産	ドローン映像・衛星画像解析, 家畜モニタリング etc.	-	自動運転トラクター, 養殖場の自動給餌制御 etc.	作物・畜産情報まとめ, 農業技術のテキスト化 etc.	作付・施肥・灌漑計画最適化, 収穫時期予測 etc.
防災・安全	災害時ドローン映像・衛星画像解析, 監視カメラ解析 etc.	SNS 投稿分析(避難誘導・災害報告), 通報自動応答 etc.	-	被災状況まとめ, 防災マニュアル作成 etc.	被害規模予測, 避難指示最適化, インフラ点検予測 etc.
行政・公共サービス	行政手続き用書類の画像読取・OCR, 公共施設監視 etc.	自動窓口対応チャットボット, 公文書の自動要約 etc.	-	報告書・白書の文章作成, 公報・広報資料作成 etc.	行政手続きの需給予測, マクロ経済推計・計画策定 etc.
エネルギー・環境	発電プラント映像監視, ドローンによる設備点検 etc.	-	エネルギー需給バランス制御, メンテナンスロボット etc.	エネルギー・環境関連のレポート作成 etc.	スマートグリッド最適化, CO ₂ 排出量予測 etc.
建設・不動産・都市開発	建設現場の映像解析(進捗管理, 安全確認) etc.	-	建設機械の自動運転・遠隔操作, ドローン測量 etc.	工事計画レポート作成, 都市設計イメージ作成 etc.	工事スケジュール・資材最適化, 都市インフラ計画 etc.
小売・EC・サービス	店舗内監視映像解析(行動分析, 棚卸し), 商品画像分類 etc.	チャット接客サポート, 顧客問い合わせ自動応答 etc.	自動レジ, 商品搬送ロボット etc.	レコメンド商品説明文作成, 広告バナー画像作成 etc.	在庫管理・需要予測, 価格最適化, 配送計画 etc.
セキュリティ・防衛	ドローン映像監視(境界警備), 侵入者検知 etc.	サイバー攻撃通信ログ解析 etc.	警備ロボット・兵器の制御, 無人車両制御 etc.	防衛報告書作成, シミュレーションによる脅威分析 etc.	リスク評価・戦略最適化 etc.

メラによる転倒検知や歩行リハビリロボットなど、高齢者や障害者の生活を支える技術が開発されている。こうした利点に加え、ウェアラブルデバイスが取得する心拍や血圧、活動量などの生体データをリアルタイムで解析し、利用者のヘルスケア管理や予防医療に活かす動きも加速している。

重要な判断を要する場面においては、あくまで AI は判断支援のために利用し、最終的な判断はヒトが下すことを前提としているが、一方で、ヒトの生命や福祉に直結する分野だけに、AI 導入時には安全性や説明責任を伴う厳格な検証が求められる。また、プライバシーが極めて重要な患者データの保護と活用をどのように両立させるかといった課題もある。さらに、ロボット支援やウェアラブルセンシングが高齢化社会で役立つ反面、技術的トラブルや倫理面でのコンセンサスが重要であり、現場の理解と丁寧な制度設計が求められる。AI の恩恵を広く享受するためには、研究開発側だけでなく、医療

施設や介護施設はもちろんのこと、行政など多くのステークホルダーが協力し合い、安全基準や法規制、利用者教育の充実を図ることも不可欠であるといえる。

金融・保険

これまで金融のアナリストや専門家が担ってきたリスク評価や投資判断などの業務に AI が導入され、不正取引の検知や信用スコアの算出にディープラーニングなどが活用されている。顧客対応においてもチャットボットが普及し、24 時間体制でローンや保険の商品説明を自動化するケースが増加している。一方、アルゴリズムトレーディングや高頻度取引が行き過ぎると、市場のボラティリティ(価格の変動性)を高める恐れが指摘されており、発注ミスや瞬時の暴落(フラッシュクラッシュ)などのリスク管理が大きな課題となりつつある。

また、信用度や保険リスクを AI で評価する場合、学習データのバイアスが不公平な結果を生む可能性があり、特定の属性を持つ顧客が不利な扱いを受けるリスクも否めない。さらに、顧客情報や金融取引データは機密性が高いため、サイバーセキュリティ対策と規制当局の監督の両面で慎重な運用が求められる。アルゴリズムがブラックボックス化し、金融機関や顧客が結果の根拠を把握できないまま信用判断やローン審査が行われると、誤判断の責任所在が不明確になり、大きな社会問題に発展する可能性もあることから、金融当局や国際機関との協力により、XAI の導入や公平性・透明性を担保する監査手法の普及が望まれる状況にある。

教育

学習者の進捗や理解度を AI が分析し、学習者に応じて最適化された学習プランを提示するアダプティブ学習システムが注目を集めている。学習管理システムと連携して、オンライン講義の動画要約や自動評価、作文添削、個別指導などが実用化されつつあり、大規模言語モデルによるチュータリング機能で学習者の疑問に対話的に答える事例も増えてきている。これにより、教育現場の負担の軽減と学習者の多様なニーズに応じた指導の実現が期待される。

一方で、AI による成績評価を無批判に受け入れることで、生徒の思考プロセスや創造性が軽視される恐れや、過度なプライバシー侵害(学習行動データの過度な解析)など懸念材料もある。また、教育はヒト同士のコミュニケーションを通じて成長を促す側面も大きいため、AI 導入と対面指導のバランスをどう取るかは現場ごとに異なる課題となっている。デジタルリテラシーの向上や学習データの扱いを明確化するガイドラインの整備が必須であり、教員研修やインフラ導入費の支援などが求められている。

マーケティング・広告・メディア

顧客の行動履歴や SNS での投稿内容を AI によって分析することで、ターゲット広告の精度を高め、消費者の潜在的ニーズを見極めるレコメンドエンジンを実装する事例が

増加している。さらに、生成 AI が広告コピー (Advertising slogan) やバナー画像を自動生成し、広告効果を検証するため複数の選択肢の中から最も良いものを選ぶ AB テストのプロセスを高速化することで、より成果の高い広告を効率的に生み出す試みも進展中である。メディア領域では、AI を用いてニュースやコンテンツの自動要約・編集を行い、報道やエンターテインメントの効率を高めるアプリケーションも登場している。

しかし、過剰なデータ収集がプライバシーを脅かす懸念や、AI が差別的もしくは不適切なコンテンツを生成・選択してしまうリスクも伴う。いわゆる「フィルターバブル」によってユーザが偏った情報のみを受け取ることになり、視野が狭まり価値観・考え方も偏る危険性も指摘されている。さらに、フェイクニュースやディープフェイクの技術を野放しにするとメディアの信頼性を大きく損ない得るため、コンテンツ認証や公開基準の厳格化など、技術面・制度面での取り組みが必須だと考えられている。

製造・産業

工場の生産ラインで外観検査を行う画像解析や、需要予測を組み込んだ生産スケジューリングなど、AI による自動化や最適化が急速に普及している。現実世界にあるものをデジタル上で再現するデジタルツインの概念を導入し、現実の生産プロセスを仮想空間でシミュレートすることで品質不良率を下げ、生産リードタイム(製品製造の開始から完成に至るまでの期間)を短縮するといった取り組みが注目されている。一方、工場ラインの大規模自動化に伴う雇用再配置の課題や、システム障害ならびにセキュリティ侵害が発生した場合の影響範囲が広いというリスクも大きい。

また、AI が複雑なロボット制御を行い、組立・搬送タスクを高度化する例もあるが、センサ情報に左右されやすく、イレギュラーが生じた場合に適切な判断ができないと事故やライン停止が起こる可能性がある。工事現場の技術者や管理者との連携およびフェイルセーフ機構の設計が不可欠であり、大規模な設備更新にはコストや導入期間もかかるため、導入計画の策定には長期的視点が求められる。

物流・交通

物流の領域では、倉庫内での自動搬送ロボットや在庫管理システム、路上での自動運転車両、交通渋滞予測や最適配車などのサービスが台頭している。EC(Electronic Commerce)市場の拡大に伴い、ラストワンマイル(顧客に商品やサービスが到達する最後の区間)配送の効率化やドローン・ロボット配送が実証されている一方、自動運転に関しては公道での安全性や責任分担の課題が大きく、法整備が先行する国とそうでない国とのギャップが顕著である。

交通システム全体を AI が制御する「スマートモビリティ」の概念も普及しつつあるが、車両から収集されるビッグデータをどう管理するか、通信インフラが混雑した際にリアルタイム性を維持できるかなど、技術的・社会的なハードルは多い。利用者データや位

置情報が無制限に蓄積されるとプライバシー侵害が懸念されるため、連携する公共機関やプラットフォームとの協調ガバナンスが鍵を握る。

農業・水産・畜産

スマート農業やアグリテックが注目され、ドローンや衛星画像、土壌センサなどからのデータを AI が解析して、生産管理や収穫時期の最適化、家畜飼育の健康モニタリング等を行う事例が広がっている。これにより生産効率や品質の向上、生産者にかかる労働負荷の低減が見込まれるが、高額な設備投資やインフラ整備、農家や漁業者の IT リテラシー確保など、実運用への障壁が少なくない。さらに気候変動や自然災害など、外部要因が予測しづらい環境では AI モデルの汎化性能に限界があり、推定が外れた際のリスク管理も課題となる。地域や国によって法制度や助成が異なるほか、データの標準化が進まずプラットフォームが乱立しがちな点も普及を阻む原因となっている。

防災・安全

災害発生の可能性を AI で予測したり、災害直後の被害状況をドローンや衛星画像から迅速に把握したりするシステムが開発されている。インフラ管理でも、橋梁や道路、下水道管などの劣化を画像解析で検知し、メンテナンス計画を自動作成する取り組みが進められている。これにより、災害やインフラ不全を未然に防ぎ、公共の安全性を高める効果が見込まれる。

しかし、災害時には通信・電力インフラそのものが機能しないシナリオを想定する必要がある。限られた状況下で AI をどこまで活用できるかが不透明である。被災地から得られるデータが断片的だった場合の推定精度や、自治体や民間企業間での情報共有ルールが未整備な点もボトルネックになり得る。システム障害や誤報が起きた際の責任所在や信頼確保の仕組みを事前に策定しておくことが不可欠である。

行政・公共サービス

文書の電子化や事務手続きの簡易化、住民サービスの最適化などを目的に AI が取り入れられ始めている。大量の書類画像を OCR (Optical Character Recognition) 処理し、自然言語処理で分類・要約する試みや、問い合わせ窓口をチャットボット化する動きが代表的なものとして挙げられる。税務処理や補助金申請などへの RPA (Robotic Process Automation) 連携も加わり、行政手続きの迅速化が期待されている。

ただし、行政は法や規則に基づいた厳格な手続きと説明責任が要求される世界であるため、AI の判断根拠が不透明だと、住民に十分な納得を与えられない恐れがある。また、公的データを活用する上でプライバシーや公平性の原則をどう担保するか、自治体ごとの IT インフラ格差や職員のデジタルリテラシー不足をどう補うかなど、広範な制度改革や教育が必要になる。

エネルギー・環境

再生可能エネルギーの効率的利用や発電・送電の需給バランスを AI が最適化する「スマートグリッド」が注目されている。太陽光や風力の発電量予測を高めることで、蓄電や送電を適切に制御し、電力ロスを減らす試みが進展している。環境監視では、温室効果ガス排出量の予測や、汚染源の特定に AI が活用される事例も増えている。一方、エネルギーや環境の問題は国際的な合意形成と密接にかかわるため、AI モデルの導入効果がどこまで国際的な枠組みに反映されるか、また高精度な予測モデルを作るためのデータが十分に公開・共有されるかが課題となる。電力インフラを AI 制御で最適化する場合、セキュリティ上のリスクも高まるため、サイバー攻撃への対策と耐障害性を組み込んだ設計が求められる。

建設・不動産・都市開発

ドローンや監視カメラの映像を AI が解析し、建設現場の進捗を管理したり安全性をチェックしたりする取り組みが進んでいる。測量や地形データの自動分析によって工事計画を効率化する事例も報告されている。また、都市開発では、人口流動や交通量データを AI が予測し、ゾーニング(地域や建物を利用目的・機能別に区画分けすること)やインフラ整備をサポートする例もある。

ただし、建設現場は多様なアクターが協働するため、AI を導入してもデータの標準化や共有フォーマットが整備されていなければ十分な効果を発揮しにくいことが指摘されている。安全管理面でも、大型機械の自動制御や作業者との協調が難しく、人的要因との相互作用を AI がどこまで正確に把握できるかが課題となる。さらに不動産評価などの自動化は、公平性や査定理由の説明責任が問われやすく、社会的受容を得るためにはモデルの透明性も重要になる。

小売・EC・サービス

小売業や EC では、店舗内監視映像解析による顧客行動の把握や無人レジ、レコメンドエンジンなどの AI 活用が進んでおり、パーソナライズドサービスを提供することで顧客満足度と売り上げ向上を狙う企業が増えている。また、対話型 AI の活用によって顧客からの問い合わせなどを自動対応するシステムの導入も加速化している。

ただし、チャットボットや音声アシスタントによって顧客とのコミュニケーションを自動化する場合でも、問い合わせ内容が複雑な場合に適切に対応できるか、誤回答やクレーム対応の責任が曖昧にならないか注意が必要である。大規模 EC プラットフォーマーは膨大なビッグデータを保持しており、市場支配力や競合他社との格差が拡大する恐れが指摘されるなど、競争政策や社会的ガバナンスの観点でも議論を要する。

セキュリティ・防衛

セキュリティ・防衛分野での AI 導入事例としては、サイバー攻撃通信ログの解析や、ドローン・カメラ映像を通じた境界警備、無人車両・無人兵器システムを AI が統括・制御する例などが挙げられる。脅威インテリジェンスの自動収集や不審行動検知といった機能により、安全保障上の即応力を高める意図がある。一方、このような用途での AI 利活用は、誤認識による友軍や民間への被害のリスクが極めて大きく、制御をどう行うかが重大な国際的課題になっている。

AI 判定によって自動攻撃が行われるような兵器システムは倫理面・法規面の深刻な論争対象であり、国際的な規制ルールや透明性の確保を求める声が強い。サイバーセキュリティにおいても、AI 同士の攻撃・防御が高度化する可能性があり、その対応に必要な人材や技術、標準・協定の策定が各国で模索されている。

上述の領域ごとの AI の利活用・社会実装が一段と加速している背景には、生成 AI や対話型 UI (User Interface) などインタフェース技術の革新が大きく寄与している。専門知識のない利用者でも、生成 AI を活用した GUI (Graphical User Interface) を介して、自然言語によって複雑な AI 機能を利用できるようになった結果、多様なビジネスシーンや公共サービスへ AI を導入しやすくなった。ただし、ユーザが AI の提示する結果を無条件に信頼してしまうリスクや、会話ログなどの秘匿情報ならびにセンシティブデータが学習データとして無意識に取り込まれる懸念も高まっている。

Beyond 5G/6G によって広帯域・低遅延のネットワーク環境が普及すれば、さらなる大規模データのリアルタイム分析・共有が可能となり、AI の応用の幅はさらに広がるとともに、AI が社会基盤として機能するスピードは一層高まることが見込まれる。その反面、プライバシー保護や説明責任、バイアス排除、フェイクコンテンツ対策など、課題の複雑化は避けられず、社会的・国際的なルール形成やガバナンスが不可欠になる。技術面・制度面・倫理面を一体的に検討しながら AI を適切に利活用していくことが、持続可能で包摂的な AI 社会を実現する鍵となるといえる。

3. 次世代人工知能技術の研究開発および社会実装の動向

3.1. 次世代人工知能技術に向けた技術革新の方向性

本章では、日々、急速に進展している AI に関し、そのコアとなる技術動向、産業界も含めた社会実装にかかる動向および研究開発段階のものも含めた未来へ向けた技術動向について述べる。

AI の将来を予想しようとするとき、いくつかのポイントをおさえれば効率良く状況を把握できる。まずは、AI の実体がどのような論理構造で動作しており、そこからどのような特徴が導かれ、どのような状況に置かれているのかという実態を知る必要がある。状況については、特に、その技術とそれ以外の技術あるいは我々の社会との関係性や相互作用について確認しておく必要があるが、これらを深く理解するためにも、実体、特徴および周囲と互いに及ぼし合う影響を正確に把握せねばならない。また、AI を取り巻く周辺技術それぞれについてもリテラシーを高めておく必要がある。ここまでで、そのような AI および周辺知識へのリテラシーを高めるための情報を述べさせていただいた。ここからはそれらを踏まえた上で、次世代 AI に期待される機能や社会実装などの未来像や、次世代 AI を取り巻く個々の課題、および社会実装に向けた取り組みについて述べる。

生成 AI、マルチモーダル AI

これまで、ヒトから AI への指示方法、コミュニケーション方法はプログラミングやデータによる指示が主流であり、これらは専門家のためのツールにとどまっていた。2019 年の GPT-2 (OpenAI, 米国) の登場により、プログラミングなど複雑な手順を経ることなく、自然文テキストで AI に指示が出せる「プロンプト」が実装された。さらに、2020 年の GPT-3 (OpenAI, 米国) への進化により、プロンプトの指示のみで多様なタスクをこなせるようになった。2021 年の Codex (OpenAI, 米国) ではプロンプトでプログラムコードの生成が可能になり、画像生成 AI である DALL-E (OpenAI, 米国) ではプロンプトで画像を生成できるようになった。2022 年には、GPT-3.5 をエンジンとした ChatGPT (OpenAI, 米国) が一般公開された。プロンプトという自然文テキスト対話型の直観的なインタフェースを AI が有するに至り、AI は人々との距離を縮め、これまで専門家のみが使いこなせるツールであった AI は一気に一般の人々が気軽に使いこなせる存在になった。文章やプログラミングコードといったテキストを提供する ChatGPT に加えて、さらに画像を生成する Stable Diffusion (オープンソース)、さらには動画を生成する Sora (OpenAI, 米国)、Midjourney (Midjourney, 米国)、Runway (Runway AI, 米国)、Pika Labs (Pika Labs, 中国) などの動画生成 AI も公開された。さらには、マルチモーダル化を推進する Bard/Gemini (Google Deepmind, 米国)、倫理的アプローチを重視した Constitution AI を基盤とした Claude (Anthropic, 米国) など、さまざまな特徴を持った生

成 AI が公開あるいは発表されており、今後、これらのコモディティ化が進むとともに、それらを統合したサービスや、それらを活用したより目的特化型のサービスなど、多種多様なサービスが提供されることが予想される。

Hallucination の解決

ハルシネーション(Hallucination)は、大規模言語モデル(Large Language Model, LLM)を活用したテキストや音声など文章を生成する AI において、AI が事実とは異なる内容の回答文を生成することである。古い内容や誤った内容の学習データを使用して AI が学習した場合、特例など AI の学習精度や推論アルゴリズムの限界を超えた場合に発生する。ハルシネーションの抑制が強い領域には、歴史や自然科学などの普遍的事実、(よく学習されているため)プログラミングコードなどのソフトウェア開発にかかる事項、公共データベースや辞書的な情報など、時間経過にともない内容が変化することのない事項や、社会的に総意が得られており、また広く同内容で発信されているものなどが多く含まれる。一方で、AI が学習を終えた時期以降の最新の情報、直接的記述が少なく複数の情報や背景にある知識を組み合わせで導き出さねば回答を得られない情報、公への発信が少ない企業、団体および個人の情報、社会の総意が得られておらず異なる複数の意見が発信されている情報などは、ハルシネーションを生じる可能性が高くなるため注意が必要である。

ハルシネーションの抑制には、

- 質問要求(クエリ)内容が当該 AI の適用限界を超えていないかを確認して対応する。
- 学習データにおいて、誤った内容を含まないように内容を確認し、また内容を最新に保ち続けて学習を更新する。
- 検索拡張生成(Retrieval-Augmented Generation, RAG)などの技術を活用して、データベースや検索エンジンなど他の情報ソースとの連携を図る。
- AI の学習や適応にかかるガイドラインを設定する。

などの対応が必要となる。ここで、RAG とは、データベースや検索エンジンなどを介して AI エージェントが外部データを動的に取得し、それらの情報を基に生成的なタスクの実行能力を強化する手法もしくはそれを実装したシステムを指す。実際に、企業などで生成 AI の活用において RAG を導入して社内データベースと連携させ、ハルシネーションを抑制することが検討されている。

学習の進化

AI の学習は、教師あり学習と教師なし学習の2つに分類される。教師あり学習では、人間があらかじめ入力データに対して正解のラベル(出力データ)を付け、それをまとめた学習データに基づいて学習を行う。ラベル付け作業のことをアノテーション(Annotation)とも呼ぶ。一方、教師なし学習では、入力データのペアとなるラベルは用

いず、入力データのみで AI は学習を行う。自己教師あり学習 (Self-Supervised Learning, SSL) は、2020 年に提案された AI の学習法であり、教師なし学習の一種である。ラベルを付与していないデータに対して AI 自身が擬似的なラベルを付与して、擬似的な問題 (Pretext Task) を解くことによって学習する。AI 自身により機械的にラベリングを行うことで、人間によるラベル付けの作業をなるべく減らしながらも教師あり学習の性能、あるいはそれ以上の性能を手に入れることを目的とする。画像を用いた認識や診断を行う AI など、膨大な画像データを学習に必要としてかつそれらのデータをラベル付けするために膨大な労力を要する場合に有効に働く。その中でも、入力データの画像に対してラベル付けを行ったのち、その入力画像のサイズ拡大縮小、並行移動、回転、反転などの画像を生成して入力データの数を増やし、それらに対してもオリジナルの入力データと同じラベルを自動で付けて学習を行うことを **Data Augmentation** と呼ぶ。また、AI 大規模モデルの学習にかかるエネルギー消費を低減する技術も研究開発されており、成果を期待されている。

インタフェースの進化

カメラ映像 AI (Vision AI) によるシーン認識やマイク音声 AI (Auditory AI) による言語認識などインタフェース AI が進化している。これにより、AI を利用者が増加し、AI の社会実装が進むことが予想される。

a. 自然言語インタフェース／自然言語処理 (Natural Language Processing, NLP) の進化

AI が人間の言語を適切に処理する能力が飛躍的に向上している。AI は人間の言語を理解しているかのように動作し、外部からのコミュニケーションからでは我々は AI が生成した文章と人間が生成した文章の区別がつかなくなった。ほぼ、チューリングテストに合格したと言って過言でない。これにより、ChatGPT-4o/o1/o1-Pro (OpenAI, 米国)、Siri (Apple, 米国) および Alexa (Amazon, 米国) をはじめとした、テキストベースのチャットボットや音声アシスタントおよび自然言語変換が、より自然で直観的なコミュニケーションを可能にしている。また、これらのスキームは、多言語対応が進められ、異なる言語間でのシームレスなコミュニケーションが可能になった。これにより、グローバルな利用が促進されている。

b. インタフェースのマルチモーダル化

テキスト、画像、音声および動画など複数のデータ形式を統合して処理する能力を持った、マルチモーダル AI が登場してきた。例えば、OpenAI の GPT-4 や DeepMind の Gemini は、テキスト、画像、音声、動画など複数のデータ形式を統合して処理する能力を持っている。これにより、ユーザは異なる形式のデータを使って AI とやり取りできるようになった。例えば、画像をアップロードして質問をしたり、音声で指示を出して画像や動画を生成したりするなど、複数の感覚を組み合わせたインタラクションが可能になっている。

c. 音声によるインタフェースの進化

Speech-to-Text(Google, 米国), Transcribe/Amazon-Web-Services(Amazon, 米国)や AmiVoice(アドバンスト・メディア, 日本)など音声認識技術や音声合成技術が進化し, 音声認識率や合成精度が向上し, 音声インタフェースがより自然で正確になっている. 音声からの文字起こしシステムや, リアルタイム翻訳が進化し, 異なる言語を話す人々の間でのコミュニケーションが容易になっている.

d. 画像/映像によるインタフェースの進化

カメラなど画像センサを活用して, 物体の種類やその配置などのシーン認識, 手の動きや体の動きを認識する技術が進化している. また, これによりジェスチャーの認識が進み, タッチレスで AI を操作することが可能になっている. さらに, 視線追跡技術を活用し, ユーザが見ている対象を AI が識別し, それに基づいて応答するインタフェースが開発されている.

e. 仮想現実(Virtual Reality, VR), 拡張現実(Augmented Reality, AR)/複合現実(Mixed Reality, XR)との統合

AI が VR や AR 環境でのインタラクションを支援することで, より没入感のある体験が可能になっている. 例えば, 仮想空間内で AI アシスタントがリアルタイムでガイドを提供するなどの応用が進んでいる. メタバース内で, AI キャラクターやアシスタントが, ユーザとのインタラクションを支援するなど, メタバースでの AI 活用も進んでいる.

f. インタフェースの個別化/パーソナライズ

ユーザの行動や嗜好を学習し, それに基づいてインタフェースをカスタマイズする AI が増えている. これにより, ユーザ個人の嗜好に合わせた体験が提供される. AI がコンテキストを認識して, ユーザの状況や環境に応じた応答や提案を行う能力が向上している.

g. 脳波を利用したインタフェース

脳波を直接読み取って AI とやり取りするブレイン・コンピュータ・インタフェース(Brain-Computer Interface, BCI/Brain-Machine Interface, BMI)と呼ばれる技術が進められている. これにより, 身体的な制約を持つ人々でも AI を操作できる可能性が広がっている. Neuralink(米国)などは, 脳とコンピュータを直接接続する技術を開発しており, 将来的には AI とのインタラクションがさらに直観的になると期待されている. また, さまざまな工夫により, 高齢者や障害者を含むすべての人々が利用しやすいインタフェースの設計が進んでいる.

h. ユーザインタフェース(User Interface, UI)の消失

前に述べた, 映像インタフェース AI や音声インタフェース AI を使い, ユーザが明示的に操作しなくても, AI が環境に溶け込んで動作し, 必要な情報やサービスを提供す

る「ゼロ UI (Zero User Interface, Zero UI)」の概念が注目されている。よりシームレスな人間と AI のインタラクションが可能になる。

意思決定プロセスの透明化

医療、金融や司法などの分野では判断が人々の生活に大きな影響を与えるため、AI 意思決定プロセスの透明化が求められ、研究開発が進められている。また、偏見や差別を助長しないようにするなど倫理的配慮を考慮した取り組みも進んでいる。各国や地域で AI に関する規制が進むなか、AI Act などにおいては法的要件として透明性が求められており、ユーザが AI の動作や意思決定プロセスを理解しやすいインタフェースが求められている。AI モデルがどのようにして結論を導き出したのかを説明する技術である説明可能 AI (Explainable AI, XAI) 技術、より透明性の高い AI モデルの選択、データの出所、品質およびバイアスの有無を明らかにするなど学習データの透明化、アルゴリズムの透明化、AI 意思決定プロセスの可視化など、さまざまな観点より、AI 意思決定プロセスの透明化が進められている。

多機能化、汎用化

AI 発展のひとつの方向性として、その機能の多様化および汎用化が示されている。一般に、帰納推論に基づく AI は汎用化を進めれば推論の精度が低下するため、推論の精度を維持するには汎用性を犠牲にしてでも目的を絞って AI を構成するのが現時点では一般的である。一方で、目的特化型 AI の目的の細分化が進みすぎて、まとまったひとつのタスクを遂行するためにいくつもの AI を選択し実行しなければならない状況も生じており、ひとつの AI からいくつものタスクを実行するいわゆる汎用 AI への期待も大きくなっている。目的特化型 AI の推論精度を維持しつつ汎用性を高める手段として、多数の目的特化型 AI を目的に応じて呼び出す AI エージェントを用意して汎用化を進めるフレームワークが提案されている。システム内に多数の目的特化型 AI を用意しておき、要求(クエリ)にしたがって適切な目的特化型 AI を選択して起動する。もっとも広く知られているフレームワークには、2024 年 10 月 12 日に OpenAI 社が公開したマルチエージェントオーケストレーションのためのフレームワークである Swarm がある。このフレームワークは複数の AI エージェントを協調的に実行することで、複雑なタスクを遂行させる。フレームワークの外部仕様は公開されているが、それを駆動している内部のアルゴリズムは非公開である。Swarm と同様の目的を持つ概念に、マイクロソフト社が 2024 年 1 月に発表した AutoGen、および同じく 2024 年 1 月に登場した sakana.ai 社（日本）のシステムがある。sakana.ai は、NVIDIA 社からの資金調達に加えて、2024 年 9 月には複数の日本のリーディングカンパニーから資金調達し、事業展開を加速させている。また、同様の概念に、2005 年に東京大学の研究チーム(現、東京科学大学の情報医工学分野)が提案した「Collaborative AI」および「AI (made) of AIs」の概念があり、その基盤となる技術である

再帰的等価変換 (Recursive Equivalent Transform, RET)は 2016 年に特許化され 2021 年に権利化されている。同技術は、ネットワーク上に配置された AI エージェントが、データベースやセンサ(リアルタイムデータ)を含むデータ群、シミュレータ、データ処理コンピュータおよび AI を含むアルゴリズム群の個々それぞれを担当し、それらが相互に自律的にコミュニケーションを取り、再帰的に情報接続し、データ処理・供給のためのプロセスツリーを自動生成するものである。AI エージェント間のやり取りには、セマンティック記述(セマンティックス)が用いられ、また必要な場合にはこのセマンティックスを展開して情報接続の再帰的展開を実現する。同技術は Autonomic Intelligence 社(日本)によって実装が進められ、2023 年 4 月に Intelligently Orchestrating Networking AI (iON AI) として発表された。また、日本工学アカデミー 科学技術イノベーション 2050 委員会の政策提言書「持続可能社会に向けた科学技術・イノベーションロードマップの提言」(2020)において、「マルチ AI ネットワーク」として複数の AI が協力して社会に貢献する将来像を提唱し、2021 年 5 月に国連 IATT (United Nations - Interagency Task Team on STI for the SDGs) の報告書に採択された。これは分散計算技術、いわゆるエッジコンピューティング (Edge Computing, EC) の一種であるが、従来のエッジコンピューティングがエージェント間であらかじめ決められた一種類の形式に従うデータセットをやり取りしていたのに対し、AI エージェントによる情報接続は、多種多様なデータセットを自律的に識別しながらやり取りし、データやアルゴリズムを適切に選択しながら情報接続する点が異なる。ここでは、タスクの種類をどのように識別し、AI を含むアルゴリズムをどのように適切に選択するかがシステムの性能に大きく関わる。また、自律分散ネットワーク型の AI エージェントによる協働においては、タスクの種類に見合ったアルゴリズムをネットワークから効率良く検索して情報接続しなければならない。前出の iON AI では、ネットワーク通信の仕組みを効率良く活用し、条件判断の並列処理を実現することで高速化に成功している。

人工意識 (Artificial Consciousness, AC)

現時点で存在する多くの AI は意識を持つことを意図して開発されておらず、これらの AI に意識は内在しないという意見が、研究者・科学者の意見の大半を占める。ひとつめの理由は、感情や情動を含めた意識がどのような現象であり、どのような機序で生じるのかが不明であり、少なくとも現存する AI が意識を持つとはいいきれないというものである。ふたつ目の理由は、現存する多くの AI は半導体チップ内に実装された NAND 論理演算回路で計算できる以上のものを持ち合わせておらず、確証はないが、ヒトの意識は NAND 論理演算回路で実装できる範囲を超えているであろうという意見に基づく (NAND 論理演算のみで完全系を構成する以外にも、OR 論理演算のみで完全系を構成する、あるいは AND, OR, NOT の各論理演算で構成するという考えもあるが、これらはいずれも NAND 論理演算の組み合わせに置き換えることができ、ここでは説

明の簡略化のため「NAND 論理演算回路で構成されている」との表現を用いた)。ふたつ目の理由について、意識の発生原理そのものでなくとも代替計算(エミュレーション, Emulation)によって実現可能であるという意見もあるが、いずれにせよ、意識を実装するための特別なアルゴリズムやフレームワークが必要であるとの意見に変わりはない。

将来において AI に意識の全体あるいはその一部を実装できるかについては現時点で不透明であるが、Artificial Consciousness を研究開発している研究者らが研究を進めていること、また意識の機序を解明するための脳科学についても、核磁気共鳴帰納画像法 (Functional Magnetic Resonance Imaging, fMRI) や光トポグラフィー (Optical Topography, OT) と呼ばれる近赤外線分光法 (Near-Infrared Spectroscopy, NIRS) など脳活性のリアルタイムイメージング手法の技術的進歩や、神経繊維の走行やつながりをみる拡散テンソルトラクトグラフィ (Diffusion Tensor Tractography, DTT) など解析手法の発展にともなって脳機能の解明が急速に進んでいることから、将来において人類が脳機能の全容を解明する可能性、および AI に意識を実装できる可能性を否定できない。感情や情動を含む意識の機序が解明されれば、人工意識の実装に向けた大きな助けになることは間違いなく、その成果が期待される。一方で、AI に意識を持たせることやそれにかかる研究開発を進めることは、それが現実味を帯びてきた時点で、AI に人格権を与えるべきか否かの判断も含めて慎重な検討が必要となる。

創造性

AI の機能について語るとき、意識とともによく議論されるのが創造性 (Creativity) である。人類の創造性をどのように定義し、それがどのような機序で生みだされ、さらにどのような基準で評価すべきかについては、見識が分かれる。一般に、現存する AI の創造性は、人間の創造性とは異なる。AI は大量のデータを分析し、創造前後のデータパターンの組み合わせや、オリジナルデータから新しいデータを創造するパターンを学習し、それをもとに新しいデータを生成する。その生成は「学習データの組み合わせ」や「確率的な推測」、あるいは「乱数的な組み合わせ」に基づいており、人間のように情動や、創造物に対する印象や評価など何かしらの意図や基準に従って創造するわけではない。創造性を「現存しない何かしらの新しいデータを生成するもの」とすれば AI はすでに創造性を獲得しているが、創造性を「創造物に対して抱かれる情動などの強さや印象の良し悪しに従って、新しい表現を創造するもの」とするならば AI は創造性を獲得しているとは言えない。一般的な AI は、現時点において創造物評価の意図的な実装を行っていない。現存する多くの AI は、学習データにないパターンの生成や、社会的あるいは哲学的な背景など複雑な関係を考慮した生成は行えず、現時点において「人間らしい創造性」は実装されていない。しかしながら、良し悪しの評価や取捨選択を人間が行うのであれば、あるいは過去に人間が生成した創造性に優れた結果を模倣するのであれば、AI は価値ある新しいデータを生成でき、我々が創造性を発揮するタスクの

一助となる処理を担当することができる。このような依頼は、とくに生成 AI に対するデータ生成において要求されることが多い。生成 AI は、記事、小説、詩などの文章生成、絵画やデザイン画などの画像作成、楽曲やメロディなどの音楽の作曲、プログラムなどのコードの生成など幅広い用途に使用される。また、創造性の先進的実装を行なった AI の研究開発も進められている。それらは、人間らしい創造性のパターンを学習し、あるいは人間の創造や思考における手順や処理を模倣して、新たなデータを生成できる可能性を有しており、人間らしい創造性の獲得に向けて今後の研究開発が期待される。一部において、生成と創造を区別せず混同している意見も見られるが、これらは別の概念である。また、汎用性と創造性を区別せず混同している意見も見られるが、汎用性と創造性は別の概念であり、汎用性を実現するために創造性は不要であるとともに、汎用性を実現しても創造性を実現したことにはならない。

高速無線通信など IT 分野における他の技術との連携、システム統合

これまで、大容量高速無線通信、Internet of Things (IoT) を含む通信の多様化、ネットワーク検索技術、ネットワークセキュリティ技術、および AI 技術は、個別に論じられてきた。これらの技術はいずれも IT 技術であって相性が良く、今後、これらの技術は個別でなく、複数あるいはすべての技術が連携およびシステム統合に向かうと予想される。現在主流となっている一箇所に学習データを集約して大型 AI を構成する「データ集約型 AI」に対して、現場で計測したリアルタイム計測データをエッジ AI 間の連携で統合して解析する「データ分散型 AI」の研究開発が進められている。特に、都市開発、医療など、リアルタイムあるいはユーザ個々の計測データが重要な領域においてデータ集約型 AI への期待が高まるとともに、知識を担うデータ集約型 AI と、現状との整合性を担うデータ分散型 AI の連携および統合が進められている。

量子コンピュータなど新しいハードウェア、プラットフォームの登場

オックスフォード大学の量子コンピュータ研究チームおよび QuantroLOx (<https://quantrolox.com/>, Oxford, UK) は、量子コンピュータ技術を研究開発するとともにそのアプリケーションとしてニューラルネットワークなど AI アルゴリズムを量子コンピュータに実装しようとしている。量子コンピュータの冷却にかかるエネルギー問題が解決されれば、量子コンピュータは高速かつ省電力での計算が期待でき、そのポテンシャルは大きい。また、量子コンピュータは、現在、主に使われているノイマン型コンピュータと動作原理が異なるため、比較演算を得意とするなど特徴も異なる。それらの差異を乗り越えて、量子コンピュータでニューラルネットワーク等の AI が実装できれば、我々の社会は大きな利益を得ることが期待される。

また、前章で述べたように、人工意識は現在の計算フレームワークでは実装が困難であるとする研究者が多いこともあり、さまざまな視点から新たなフレームワークの研

究開発が進められている。脳科学からのアプローチにより意識を解明し人工意識の研究開発に活かす試みもなされている一方で、人工意識の研究開発の成果を意識の解明、すなわち脳科学への成果に活かすというアプローチも試みられている。

ヒトとの協働，社会共創

現存する AI の多くは意識を持たないと前章に述べたが、一方で、AI は我々の言動をあたかも理解しているかのように、少なくとも AI とその外部とのやりとりの観察のみからは AI が我々を理解していると思えるくらい、また我々は AI が意識を持っているか否かを判断できないくらいに AI を強化し、AI を動作させることができるようになった。優れた人とのインタフェースを AI が有するに至り、AI と人々との距離を縮めて、これまで専門家のみが使いこなせるツールであった AI が一般の人々が気軽に使いこなせる存在になったといえる。インタフェース AI では、カメラ映像 AI (Vision AI) は我々も含めて映像に映る物体の種類を識別してその動作を識別し、マイク音声 AI (Auditory AI) はマイクで計測された音の種類や我々の会話を聞き分けてテキストデータに変換することができる。これらの情報に基づいて、あたかも AI が我々を理解し協働しているかのようなシステムを作り上げることが可能であり、多くの研究者がこれに取り組んでいる。社会導入時のリスクについてであるが、現存する AI に意識は内在しておらず、その反応や動作も事前に我々の都合に合わせてプログラミングあるいは AI を学習させたものであるため、現存する AI の動作は我々との協働に適するように制御可能であり、それを社会導入するリスクは現時点で極めて低い。ただし、その制御は学習データの質や内容に依るところが大きく、学習データやネットワークモデルについてしっかりとガイドラインを準備する必要がある。また、問題が起こった場合には、AI に責任が生じるのではなく、むしろ、当該 AI に対する知識の欠如により不適切な状態で AI を使用したり、不適切な目的で AI を使用したりするような、使用する人間の側に責任が生じることも意識しておく必要がある。

3.2. 次世代人工知能技術に向けた倫理・法規制

人工知能技術にかかる倫理・法規制について、価値観の多様性からさまざまな意見が出されている。立場や世代によっても異なるこれらの意見を取りまとめるのは、困難を極める作業であった。本章に限らず、本報告書全般において、著者らは意見を述べるに際して、できるだけ事実にならい真実に沿った意見になるよう、自身や所属組織あるいは所属する学術産業分野への不当に偏った利益誘導にならないよう、また極端に一部の立場に偏った意見にならないように、真正性および公平性に細心の注意を払った。読者の立場によっては不利益をもたらす意見と捉えられる可能性もあるが、本報告書は、我々が社会全般を鑑みて真実に沿うように検討を重ねた結果であり、誠意をもって検討

した現時点での成果である。著者らはさまざまな意見を真摯に受け止め、真実に沿うように誠実に検討したいと考えており、意見をお持ちの場合には著者らにご一報いただきたい。

3.2.1. AI にかかる倫理・法的課題

AI に関する倫理的および法的課題は、技術の進化とともにますます重要なテーマとなっている。これらに対処するためには、技術者、法律家、倫理学者、政策立案者が協力して包括的な枠組みを構築する必要がある。また、国際的な協調と調和が重要であり、特に生成 AI や自律型 AI の分野での規制が進むと予想される。これらの課題に対処することで、AI の社会的受容性を高め、持続可能な技術発展の実現が期待される。

倫理的課題

AI の利用が社会や個人に与える影響を考慮し、倫理的な問題を解決することが求められている。

a. 公平性とバイアス

AI モデルは、トレーニングデータに基づいて学習するため、データに含まれるバイアスがそのまま反映される可能性がある。そのため、特定の人種、性別、年齢、地域などに対する差別的な結果が生じるリスクがあり、採用プロセスやローン審査などで公平な判断に影響を与える可能性がある。

b. 透明性と説明可能性

AI システムの意思決定プロセスがブラックボックス化している場合、結果の根拠が不明瞭になる。特に、医療や司法などの分野では、AI の判断がどのように行われたのかを説明できるよう透明化を図る必要がある。

c. 個人情報保護、プライバシー

AI は大量のデータを処理する。学習データ構築のための個人情報の収集や利用がプライバシー侵害につながる可能性がある。現在、特に顔認識技術、監視システムの利用および画像生成 AI での問題が懸念されている。学習データにこれらを含むか、許可すべきかについては、慎重な判断が必要である。統計量として一定の平滑化にも似た処理がなされているものの、時として AI は学習データのオリジナルデータに極めて近いデータを提供することがあり、検索データすべてを自由に使用できないのと同様に、個人のプライバシーや権利を含む個人情報を学習データに使用することは、基本的に避けるべきである。ただし、個人情報が直接表面化しない形態での活用や、公共性が高い活用目的については、慎重を期した上で活用の検討を進めるべきである。

d. 責任の所在

AI が誤った判断を下した場合、誰が責任を負うべきかが不明確となっている。開発者、運用者、利用者の間で責任の分担が議論されている。AI の責任問題は、銃による殺傷の責任問題に似ている。銃を使った殺傷が起こった場合、直接的には銃の自動機構による殺傷であるが、一方で使用者はその行為が引き起こす結果を予測できていた可能性があり、また銃の使用に際してはその危険性を把握しておく必要がある。結果の責任は銃そのものにあるのか、銃を開発した人々にあるのか、銃を製造した人々にあるのか、銃を不適切に使用した人にあるのかを問う問題を考えたとき、AI の責任問題はこれに似ており、我々がこれまでに検討してきた判断や経験を活かして判断できる。使用者は AI プログラムの実行を命令しただけで使用者の直接的な行為でなく、直接的には AI により自動的に行われた行為であったとしても、使用者が誤った結果となることを意図し目的としていた場合、あるいはその結果を予想できていた未必の故意に該当する場合においては、使用者に責任が生じる可能性がある。使用者は AI の適用や導入前に使用目的と使用条件が適切であることを確認し、加えて、起こりうる不具合や不適切な結果に対してリスクを分析して危険防止策を施すなど、AI を適切に使用するための責任を持たなければならない。また、研究開発者も不適切な目的に向けた AI を開発してはならず、また、アルゴリズムの透明化や適用限界の明瞭化、ならびに AI の使用におけるハザード（Hazard, 危険要因）の検出やリスク分析などにかかる研究開発にも積極的に取り組むべきである。さらには、不適切な手段による学習を行わず、かつ不適切な結果を招く可能性を含まないように、学習においても細心の注意を払う必要がある。製造者や販売者もユーザの想定を外れた動作を AI が行わないよう、設計や製造および品質管理を行うとともに、ガイドラインやマニュアルを整えるなどして製造者や販売者としての説明責任を負うべきである。

e. 倫理的ジレンマ

さまざまな分野で AI の適用が検討されるに至り、自動運転車の「トロッコ問題」のように倫理的な判断を迫られる状況での判断が重要になってきた。どのような判断基準を適用すべきかを検討し、実装を進めるべきである。

f. 労働市場への影響

AI の普及により、特定の職業が自動化される一方で、新たな職業が生まれる可能性が生まれてきた。一方で、失業やスキルのミスマッチが社会的な不安を引き起こす可能性があり、懸念されている。

g. 人工意識研究開発

人工意識が研究開発され実装されるようになると、人工意識を持つ AI に対して人権を与えるべきかなど、新たな課題が生じる。

法的課題

AI の利用に関する法的枠組みには、現時点で以下のような課題が存在する。

a. 規制の枠組み

AI 技術の急速な進化により、既存の法律が適用できないケースが増えている。また、各国や地域で異なる規制が存在するため、国際的な調和が求められている。

b. 知的財産権

文章、画像、音楽など、AI が生成したコンテンツの著作権の帰属が不明確になることがある。AI モデル自体の特許やデータセットの利用権も議論の対象となる。

c. データ保護とセキュリティ

AI のトレーニングや運用に必要なデータの収集・利用が、EU 一般データ保護規則である GDPR (General Data Protection Regulation) などのデータ保護法に違反する可能性がある。また、サイバーセキュリティの観点から、AI システムがハッキングや悪用されるリスクも懸念されている。

d. 責任と賠償

AI が引き起こした損害に対する責任の所在が現時点で不明確である。製造物責任法 (Product Liability) や契約法の適用範囲が議論されている。

e. 生成 AI の規制

生成 AI の利用により、偽情報やディープフェイクが拡散するリスクがあり、これに対する法的規制や対策が求められている。

f. 国際的な競争と規制の調和

各国が AI 技術の開発競争を進める中で、規制の不一致が国際的な課題となっている。国連や OECD (経済協力開発機構) などの国際機関が調整役を果たすことが期待されている。

各国の具体的な取り組み

国連では、UNESCO の AI 倫理勧告や国連 AI 諮問委員会を通じて、国際的な枠組みを構築している。EU (欧州連合) では、AI Act (AI 規制法) や AI 倫理ガイドラインを策定し、リスクベースのアプローチを採用している。米国では、AI Bill of Rights や国家 AI 戦略を通じて、倫理的・法的課題に対応している。日本では、AI 戦略 2023 で、倫理的課題や法的課題への対応を強化している。

3.2.2. AIにかかり懸念される社会的影響

新しい技術が台頭してきたとき、しばしば我々は混乱を生じてきた。蒸気動力織機、車製造のオートメーション化、駅自動改札の導入など失職への懸念や、極端な場合は機械による我々社会の支配を懸念する意見も聞かれた。近年の急速なAIの発展と普及に伴い、これらと同様にさまざまな社会的影響や懸念が指摘されている。AIが普及するにつれて、その不適切な使用による危険や不安にさらされたり、特定の職業が自動化されて失業など社会的不安を引き起こしたりする可能性がある一方で、新たな可能性や職業も生まれている。AIの社会的な負の影響を最小限に抑えつつ、その利点を最大限に活用するためには、技術者、政策立案者、一般市民が協力して取り組むことが重要である。

a. 雇用への影響

AIの自動化技術が進むことで、多くの職業が機械やアルゴリズムに置き換えられる可能性が懸念されている。特に単純作業やルーチンワークが影響を受けやすいとされている。これにより、失業率の上昇や特定のスキルを持つ労働者の需要減少が懸念されている。一方で、新しい職業やスキルの需要も生まれる可能性がある。

b. 倫理的問題

AIが意思決定に関与する場合、その判断が倫理的に適切であるかどうかが問題となる。例えば、AIが医療や司法、軍事などの分野で使われる場合、誤った判断が重大な結果を招く可能性がある。AIの透明性や説明責任が求められる一方で、AIの判断基準がブラックボックス化している場合、信頼性が損なわれる可能性があり、重要あるいは複数要因を考慮すべき複雑な局面における最終判断は、人間の責任において行うべきである。

c. 偏見と差別の拡大

AIは学習データに基づいて動作するため、データに含まれる偏見や差別がそのまま反映される可能性がある。これにより、特定の人種、性別、地域などに対する不公平な扱いが生じることがあり、社会的不平等が拡大する可能性があるため、AIの利用およびAIの判断の採用について慎重な判断が必要とされる。

d. プライバシーの侵害

AIが個人データを収集・分析することで、プライバシーが侵害されるリスクがある。特に、顔認識技術や監視システムの利用が広がることで、監視社会など個人の自由が制限される可能性がある。また、個人情報の不正利用について対策の強化が必要となる。

e. 情報操作とフェイクニュース

AIを利用して生成されたフェイクニュースやディープフェイク(偽の映像や音声)が、社会に混乱をもたらす可能性がある。誤情報の拡散により、個人や企業の社会的な信頼が不当に損なわれたり、政治的・経済的な混乱が引き起こされたりする可能性がある。

f. 安全性とセキュリティの問題

AI システムがサイバー攻撃によりハッキングされたり悪用されたりするリスクがある。また、AI に限らず機械システムをはじめとした技術全般に言えることであるが、AI が人間の意図に反した予期しない行動を取らないよう、AI システムの設計(AI モデル選択)ならびに構築(学習)は慎重に行う必要がある。

g. AI の軍事利用

AI が兵器や軍事システムに利用されることで、戦争の自動化や倫理的な問題が生じる可能性がある。人間の関与が減少することで、戦争開始／維持のハードルが下がり、被害が拡大するリスクがある。一方で、科学技術の軍事利用については世界的には倫理観点の枠組みの外で扱われており、AI についても同様であるという専門家の意見がある。

h. 人間性の喪失

AI が人間の感情や意思決定を模倣することで人間と AI の区別が曖昧になり、我々が人間性を軽視して扱ってしまう可能性がある。また、人間同士のコミュニケーションや共感が減少し、孤立感が増す可能性もある。

i. AI の独占と格差の拡大

AI 技術が一部の企業や国家に独占されることで、経済的・技術的な格差が拡大する懸念が生じている。グローバルな不平等が深刻化し、社会的な緊張が高まる可能性がある。

3.2.3. AI にかかる過度な期待と報道

新しい科学技術が発明発見されたとき、エコノミスト、投資家あるいはメディアジャーナリストは、過度な期待を抱き、行き過ぎた解説や報道を行うことがある。コンセプトや内容を理解するためにデフォルメされた比喻やメタファーが用いられることも多く、これが正しい理解へと早く導く助けになっているのも事実であるが、意図せずとも誤った解釈や過度な期待へと導く危険性も無視できない。誤った解釈を生じる可能性が高い発信についてはメディアが責任を負うべきであるが、解説を受け取る側にも課題がある。解説を受け取りその真偽や是非を判断するのは受け取る側の役目であるが、すべてにおいて演繹的な理解に基づき真偽を判断するのは負担である。そのため我々は時として思考停止に陥り、ウェブを含めたメディアで目にする頻度が高いという理由や権威者の意見という理由だけで真実と認識するような帰納的判断を行うことがある。類似した他のケースで正しかった例にならっているという意味において、条件が適合すれば、その多くの場合で帰納的判断は正しく動作する。しかしながら、演繹的判断と異なり、帰納的判断は論理展開に基づいておらず、判断の根拠が脆弱であることには留意しなければならない。

現存する多くの AI の仕組みは帰納的フレームワークに基づいている。情報を 1 次元あるいは 2 次元以上の多次元のデータ値の増減パターンとみたとき、帰納的フレームワークはあるデータのパターンから確率（厳密にはデータに基づいて算出されるので頻度確率）が高い別のデータのパターンへと対応付けして変換する。言葉を選ばなければ「もっとも出現確率の高い類似した前例になろう」変換であり、AI は、学習データ（ここでのいう前例）に含まれていないことは正しく推論できない。しかしながら、AI の機能（やれること）や性能（達成点）は、我々の予想を上回って進化した。AI はこの 20 年で大きく進化した。この確率に基づく変換を並列化と深層化により、膨大なパターン計算を極めて効率良く高速で行っており、20 年程度前の専門家でさえ想像しなかった領域や内容のタスクまでを高い精度でこなす。深層化による推論の精度向上と効率化に加えて、さらに、従来であればそれぞれの問題に合わせて慎重に計算モデルを構築して検証を行わなければならなかった種類の推論問題に対しても、適切な学習のみで適用が可能になった点は特筆に値する。これは、モデルの深層化やモデル構造の複雑化にともない計算収束が困難になっていった学習において、深層学習（Deep Learning, DL）として説明される AI の学習にかかる一連の技術が提案され、実用化されたことによる貢献が大きい。深層学習技術は AI モデルの深層部にまで至る学習を可能にし、深層 AI を実用化に大きく引き寄せて AI を新しいステージに押し上げた。また、現在ではさまざまな AI モデルが提案、研究開発されており、今後、AI は、より高精度、高信頼、高効率および複雑なタスクをこなす科学技術に進化していくことが予想される。たとえば、帰納推論（Inductive Inference）の限界を越えるためにそれに加えて演繹推論（Deductive Inference）および仮説推論（Abductive Inference）（仮説推論は大きくは帰納推論に含まれることも多い）を行う AI が研究開発されたり、前出の東京科学大学の研究チームによりそれらを統合して適切に選択しながら使用する統合推論（Integrated/Federated Inference）が提案されたりしている。あくまで過度な期待は禁物であり、現時点での可能性と将来も含めたポテンシャルとしての可能性は区別して議論すべき、すなわち現時点で適用可能な領域と適用が難しい領域を正しく区別して社会導入を進めていく必要があるが、現時点の推論レベルでみても AI は役に立たないどころか多くの可能性を秘めており、さまざまな領域やタスクでの活躍が期待される。

3.3. 次世代人工知能技術の社会実装に向けた施策・運用戦略

3.3.1. 諸外国の AI 政策と施策

諸外国における現時点での AI 政策と施策について紹介する。諸外国の AI 政策と施策を整理し、これらについて知っておくことは、各国が AI 技術の開発競争を進める中で表面化してきた規制の不一致という国際的課題を解決する目的においても、世界と協調しながらより豊かな社会を目指すという目的においても、我が国の未来に向けた AI 政

策と施策を検討し決定する上で重要である。ここで紹介する組織や国家以外にも AI 政策と施策を積極的に行っている諸国はたくさんあるが、ここでは、AI 政策と施策において世界のスタンダードに大きく影響すると思われる国際連合、EU、英国、米国および中国の AI 政策と施策について簡単にまとめて報告する。

国際連合 (United Nations, UN) の AI 政策と施策

国連は AI の国際的な規制と倫理的な利用を推進するための枠組みの構築に取り組んでいる。

a. UNESCO の AI 倫理勧告

UNESCO (国際連合教育科学文化機関) は、2021 年 11 月に「AI 倫理に関する勧告 (Recommendation on the Ethics of Artificial Intelligence)」を採択した。この勧告は、AI 技術の開発と利用において倫理的な枠組みを提供し、AI が人類や地球にとって有益で持続可能な形で活用されることを目指している。ここでは、

- AI の透明性と説明可能性の確保
- プライバシーとデータ保護の強化
- AI による不平等の是正
- 環境への影響の軽減

などについて検討された。

勧告の背景と目的

AI 技術は、社会や経済、環境に大きな影響を与える可能性を持っている。一方、その急速な発展に伴い、倫理的な課題やリスクも懸念されている。UNESCO の勧告は、AI の開発・利用における倫理的な原則を明確にし、各国がこれを基に政策や規制を策定するための指針を提供することを目的とする。

勧告の主要な原則

UNESCO の AI 倫理勧告では、以下のような主要な原則が掲げられた。

(1) 人権の尊重

- AI 技術は、すべての人々の人権、基本的自由、尊厳を尊重しなければならない。
- 差別や偏見を助長することなく、公平性を確保することが求められる。

(2) 包摂性と公平性

- AI は、すべての人々に利益をもたらすべきであり、特に社会的に弱い立場にある人々や地域に配慮する必要がある。
- デジタル格差を縮小し、AI の恩恵を広く共有することが重要である。

(3) 持続可能性

- AI 技術は、環境への影響を最小限に抑え、持続可能な開発目標 (SDGs) の達成に貢献すべきである。
- エネルギー効率や資源の持続可能な利用が求められる。

(4) プライバシーとデータ保護

- AI の開発・利用において、個人のプライバシーとデータ保護が確保されるべきである。
- データの収集、保存、利用において透明性と責任が求められる。

(5) 透明性と説明責任

- AI システムの設計、アルゴリズム、意思決定プロセスは透明であるべきである。
- 開発者や運用者は、AI の影響に対して説明責任を負う必要がある。

(6) 安全性とセキュリティ

- AI システムは、安全で信頼性が高く、悪用されないように設計されるべきである。

実施のための具体的な提言

UNESCO の勧告は、各国政府や企業、研究機関に対して、以下のような具体的な行動を求めている。

- 政策と規制の整備：AI 倫理に基づいた法律や規制を策定し、実施する。
- 教育と能力開発：AI に関する倫理的な教育を推進し、すべての人々が AI を理解し活用できる能力を育成する。
- 国際協力：AI 倫理に関する国際的な協力を強化し、グローバルな課題に対応する。
- 研究とイノベーション：AI 倫理に基づいた研究とイノベーションを奨励する。

勧告の意義

UNESCO の AI 倫理勧告は、初めての国際的な AI 倫理の枠組みであり、各国が AI 技術を倫理的に活用するための基盤を提供している。この勧告は法的拘束力を持たないものの、各国の政策立案や国際的な議論において重要な指針となっている。

今後の課題

各国がこの勧告をどのように具体的な政策や規制に反映させるかが課題である。また、技術の進化に伴い、新たな倫理的課題が生じる可能性があるため、継続的な議論と更新が必要である。

UNESCO の AI 倫理勧告は、AI 技術が人類全体にとって有益であるための重要なステップといえる。今後、各国や企業がこの勧告をどのように実践していくかが注目される。

b. AI と持続可能な開発目標 (SDGs)

国連は AI を持続可能な開発目標 (SDGs) の達成に活用することを推進している。特に、教育、医療、気候変動対策、貧困削減などの分野で AI の活用が期待されている。

c. 国連 AI 諮問委員会

2023 年、国連は AI の国際的な規制とガバナンスを議論するための「AI 諮問委員会」を設立した^[18]。この委員会は、AI のリスクと機会を評価し、国際的な枠組みを提案する役割を担っている。国連 AI 諮問委員会 (United Nations Advisory Body on Artificial Intelligence) は、国連が AI に関する倫理的、法的、社会的な課題に取り組むために設立を計画している組織である。この委員会は、AI 技術の急速な進展がもたらす影響を評価し、国際的な政策や規範を策定するための助言を行うことを目的としている。

背景

AI 技術の進化に伴い、国際社会では AI の利用がもたらす倫理的課題やリスク (プライバシー侵害、不平等の拡大、偏見の助長、軍事利用など) に対する懸念が高まっている。これに対応するため、国連は AI に関する国際的な枠組みを構築し、各国が協力して AI の開発と利用を管理する必要性を認識している。

主な目的

1. 倫理的ガイドラインの策定

AI の開発と利用における倫理的原則を明確化し、各国や企業がそれに従うよう促す。

2. 国際協力の促進

各国間で AI に関する知識や技術を共有し、国際的な規範を整備する。

3. リスク管理

AI の悪用や予期せぬ影響を防ぐためのリスク評価と管理方法を提案する。

4. 持続可能な開発目標 (SDGs) との連携

AI を活用して、貧困削減、教育の普及、気候変動対策など、SDGs の達成を支援する。

活動

2023 年 10 月 26 日、アントニオ・グテーレス国連事務総長は AI に関する諮問機関を立ち上げた。彼らは、AI の国際的な規制と倫理的枠組みの必要性を強調し、国連が AI に関する議論を主導する意向を示すとともに、AI にかかるテクノロジーのリスク、機会、そして国際的なガバナンスについて検討を行っている。ひとつの成果として、2024 年 3 月 21 日、国連総会は、すべての人々の持続可能な開発にも恩恵をもたらす、「安全、安心で、信頼できる」AI システムの推進に関する画期的な決議を採択した。また、国連教育科学文化機関 (UNESCO) や国際電気通信連合 (ITU) などの国連機関も、AI に関

する倫理的ガイドラインや政策提言を行っている。これらの活動が、AI 諮問委員会の設立に向けた基盤となる可能性がある。

d. AI と平和維持活動

国連は AI を平和維持活動に活用する可能性を模索している。例えば、AI を用いた紛争予測や人道支援の効率化が検討されている。

EU の AI 政策と施策

a. AI 法 (AI Act)

2021 年 4 月に提案された「AI 法 (Artificial Intelligence Act)」は世界初の包括的な AI 規制法であり、AI システムをリスクベースで分類し、高リスクな AI の使用を規制することを目的とする。AI の開発、利用、普及に関する規制を目的として提案した法律であり、AI 技術の安全性、透明性、倫理性を確保し、同時にイノベーションを促進することを目指している。2024 年 7 月 12 日に EU において公布され、2026 年 8 月 2 日より適用が開始される。AI Act は、AI システムをリスクのレベルに基づいて分類し、それに応じた規制を適用する「リスクベースのアプローチ」を採用している。AI Act は、EU 内だけでなく、世界的な AI 規制の基準となる可能性がある。特に、GDPR (一般データ保護規則) が世界中のプライバシー規制に影響を与えたように、AI Act も他国の AI 政策に影響を与えると考えられる。

リスク分類

分類されたリスクを以下に示す。

1. 禁止される AI システム (Unacceptable Risk)

公共の安全や基本的人権を侵害する可能性がある AI システムは全面的に禁止される。例としては、「社会信用システム」など社会的スコアリングや公共空間での顔認識技術などリアルタイムの遠隔生体認証が挙げられる。

2. 高リスク AI システム (High Risk)

医療、教育、雇用、法執行、重要インフラなど、人々の生活や権利に重大な影響を与える可能性がある AI システムは「高リスク」とされる。これらのシステムには、厳格な要件 (透明性、説明責任、データ管理など) が課される。

3. 限定的リスク AI システム (Limited Risk)

チャットボットや AI アシスタントなど、比較的风险が低いシステムには、透明性の要件が課される。例えば、ユーザが AI と対話していることを明示する必要がある。

4. 最小リスク AI システム (Minimal Risk)

ゲームやスパムフィルタなど、リスクがほとんどない AI システムは規制の対象外と

される。

高リスク AI システムへの要件

高リスク AI システムには、以下の要件が課される。

1. データの品質管理

AI システムのトレーニングデータは、バイアスを最小限に抑え、公平性を確保する必要がある。

2. 透明性と説明責任

ユーザが AI システムの動作を理解できるようにするための情報提供が求められる。

3. 監査と認証

高リスク AI システムは、EU の基準に基づく監査や認証を受ける必要がある。

4. 人間の監視

AI システムの使用において、人間が最終的な意思決定を行う仕組みが求められる。

期待される効果

期待される効果として、データの共有と再利用が容易になることで、新しい製品やサービスの開発が加速し、イノベーションが促進される。また、データ経済の活性化により、EU 全体の競争力が向上し、経済成長につながる。EU 内でのデータの流通を強化し、国際的なデータ競争における独立性を確保し、データ主権を確立する。

b. データガバナンス法 (Data Governance Act)

EU の Data Governance Act (データガバナンス法, DGA) は、EU がデータの共有と利用を促進し、データ経済を強化するために策定した法規制である。この法律は、データの透明性、公平性、信頼性を確保しつつ、データの流通を促進することを目的とする。DGA は、2020 年 11 月に欧州委員会によって提案され、2022 年 6 月に正式に採択された。DGA は、データの共有と利用に関する新しい枠組みを提供するものであり、AI やデジタル経済の発展において重要な役割を果たすと考えられる。

目的

DGA の主な目的は、EU 内でのデータ共有を促進し、データ経済を活性化させることである。これにより、企業、研究機関、公共部門がデータをより効率的に活用できるようにし、イノベーションや経済成長を支援する。

具体的な目標を以下に示す。

- データの透明性と信頼性を向上させる。
- 公共部門が保有するデータの再利用を促進する。

- データ仲介サービス (Data Intermediaries) の役割を明確化し、信頼できるデータ共有の枠組みを提供する。
- 個人や企業が自分のデータをよりコントロールできるようにする。

規定

DGA が規定する主な内容を以下に示す。

1. 公共部門データの再利用

公共部門が保有する特定のデータ (例：商業的に利用可能なデータや機密性の高いデータ) を、透明性のある条件で再利用できるようにする。ただし、個人データや機密情報の保護を確保するため、適切なセキュリティ対策が求められている。

2. データ仲介サービス (Data Intermediaries)

データ仲介サービスプロバイダー (例：データ共有プラットフォーム) を規制し、信頼性の高いデータ共有を促進する。これにより、データの提供者と利用者の間で公平かつ透明な取引が行われるようになる。

3. データのアルゴリズム的利用と個人の権利

個人や企業が自分のデータを第三者に提供する際の権利を強化する。データの利用に関する同意や透明性を確保するための仕組みを導入する。

4. データ共有のための欧州データスペース

特定の分野 (例：医療、エネルギー、農業) におけるデータ共有を促進するため、欧州データスペースを構築する。これにより、分野ごとにデータの相互運用性を確保し、イノベーションを加速させる。

期待される効果

1. イノベーションの促進

データの共有と再利用が容易になることで、新しい製品やサービスの開発が加速する。

2. 経済成長

データ経済の活性化により、EU 全体の競争力が向上する。

3. データ主権の確立

EU 内でのデータの流通を強化し、国際的なデータ競争における独立性を確保する。

4. AI 発展への貢献

DGA は、AI の発展にも大きな影響を与えと考えられる。AI モデルの開発には大量のデータが必要であり、DGA によってデータのアクセスと共有が容易になることで、AI の研究開発が加速する可能性がある。また、DGA はデータの透明性や信頼性を重視しているため、AI システムの倫理的な利用や説明可能性の向上にも寄与すると期待さ

れる。

c. AI 倫理ガイドライン

EU の「信頼できる AI のための倫理ガイドライン」(Ethics Guidelines for Trustworthy AI) は、欧州委員会 (European Commission) が 2019 年 4 月に発表したもので、AI システムの開発と利用において倫理的で信頼できる AI (Trustworthy AI) を実現するための枠組みを提供している。このガイドラインは、AI の倫理的な側面を重視し、技術の進歩と社会的価値の調和を目指している。

ガイドラインの 3 つの主要な要素

EU の AI 倫理ガイドラインは、信頼できる AI を実現するために以下の 3 つの主要な要素を定義している。

1. 合法性 (Lawful AI)

AI システムは、適用されるすべての法律や規制を遵守する必要がある。

2. 倫理性 (Ethical AI)

AI システムは、倫理的原則や価値観に基づいて設計され、運用されるべきである。

3. 堅牢性 (Robust AI)

AI システムは、技術的および社会的に堅牢であり、リスクや不確実性に対処できるものでなければならない。

信頼できる AI の 7 つの要件

ガイドラインでは、信頼できる AI を実現するために、以下の 7 つの要件を提案している。

1. 人間の介入と監視 (Human Agency and Oversight)

AI システムは、人間の意思決定を補完し、支援するものであり、人間の介入や監視が可能であるべきである。

2. 技術的堅牢性と安全性 (Technical Robustness and Safety)

AI システムは、信頼性が高く、安全であり、悪意のある攻撃や誤動作に対して耐性を持つべきである。

3. プライバシーとデータガバナンス (Privacy and Data Governance)

AI システムは、個人のプライバシーを保護し、データの管理と処理が適切に行われるべきである。

4. 透明性 (Transparency)

AI システムの動作や意思決定プロセスは、説明可能で透明性があるべきである。

5. 多様性, 不差別, 公平性 (Diversity, Non-discrimination, and Fairness)

AI システムは、偏見を排除し、すべての人々に公平に利益をもたらすべきである。

6. 社会的・環境的福祉 (Societal and Environmental Well-being)

AI システムは、社会全体の福祉や持続可能性に貢献すべきである。

7. 説明責任 (Accountability)

AI システムの開発者や運用者は、そのシステムの影響に対して説明責任を負うべきである。

実践ツールと評価リスト

ガイドラインには、AI システムが信頼できるかどうかを評価するための「自己評価リスト (Assessment List for Trustworthy AI)」も含まれている。このリストは、AI 開発者や運用者が自らのシステムを評価し、信頼性を確保するための実践的なツールとして提供されている。

ガイドラインの意義と影響

EU の AI 倫理ガイドラインは、AI 技術の開発と利用における倫理的な基準を確立する重要なステップとされている。このガイドラインは、EU 内外の政策立案者や企業、研究者に影響を与え、AI の倫理的な利用を促進するための基盤となっている。また、EU の AI Act (AI 規制法) やその他の政策にも影響を与えている。

d. AI 研究と投資

EU は、AI の研究と投資において世界的に重要な役割を果たしており、AI 技術の開発と倫理的な利用を推進するためのさまざまな取り組みを行っている。EU は、Horizon Europe において、2021 年から 2027 年までの研究開発プログラムで、AI 関連の研究に多額の資金を投入する。また、Digital Europe Programme は、AI、サイバーセキュリティ、デジタルスキルの向上を目的とした支援プログラムである。

1. EU の AI 戦略

EU は、AI を経済成長、社会的利益、競争力の向上に活用するため、包括的な AI 戦略を策定している。2018 年に「AI for Europe」という戦略を発表し、2021 年には「AI に関する調整計画 (Coordinated Plan on AI)」を改訂した。この計画では、AI 研究と投資を強化し、倫理的で信頼できる AI の開発を目指している。

2. Horizon Europe プログラム

EU の主要な研究資金プログラムである「Horizon Europe」(2021-2027) は、AI 研究に多額の資金を提供している。このプログラムでは、AI を含むデジタル技術の研究開発に約 150 億ユーロ (約 2 兆円) が割り当てられている。特に以下の分野に重点を置いている。

- AI の基礎研究 (アルゴリズム、モデル、データ処理技術など)

- AI の応用研究(医療, 農業, 製造業, エネルギーなどの分野)
- AI 倫理と信頼性の確保

3. Digital Europe プログラム

「Digital Europe プログラム」(2021-2027)は, AI の実装と普及を支援するために設計された EU の別の主要プログラムである。このプログラムでは, AI, サイバーセキュリティ, データインフラストラクチャなどの分野に約 75 億ユーロ(約 1 兆円)を投資している。特に以下の取り組みが含まれる。

- AI テストベッドの設立
- 中小企業(SMEs)への AI 技術の導入支援
- AI スキルの向上を目的とした教育とトレーニング

4. AI 研究センターとネットワーク

EU は, AI 研究を促進するために, 複数の研究センターやネットワークを設立している。例えば,

- AI4EU プラットフォーム: AI 研究者, 企業, 政策立案者を結びつけるためのオンラインプラットフォーム
- 欧州 AI 研究ネットワーク(CLAIRE): AI 研究の協力を促進するためのネットワーク
- 欧州イノベーションハブ(Digital Innovation Hubs): AI 技術の普及を支援する地域拠点がある。

5. AI 投資の強化

EU は, AI 分野への投資を増やすために, 官民パートナーシップを推進している。例えば,

- AI Public-Private Partnership(PPP): AI 研究とイノベーションを促進するための官民連携
- InvestEU プログラム: AI スタートアップや中小企業への資金提供を支援

がある。EU は, 2021 年から 2027 年の間に, AI とロボティクス分野に少なくとも 200 億ユーロ(約 3 兆円)の投資を目指している。

6. 国際協力

EU は, AI 研究と投資において国際的な協力を重視している。例えば, 米国や日本, カナダなどの国々と AI 分野での協力を進めている。また, OECD や G7, G20 などの国際フォーラムを通じて, AI の倫理的利用に関する議論をリードしている。

英国の AI 政策と施策

英国は AI 分野でのリーダーシップを目指しており, 政府や関連機関が積極的に政策を策定し, 施策を実施している。

a. AI 戦略

英国政府は 2021 年 9 月に「国家 AI 戦略(National AI Strategy)」を発表した。この戦略は、AI 分野での英国の競争力を強化し、経済成長や社会的利益を促進することを目的とする。主な目標を以下に示す。

- AI 研究とイノベーションの推進：AI 分野の研究開発(R&D)を支援し、世界的なリーダーシップを維持。
- AI のスキルと人材育成：AI 分野での専門人材を育成し、教育プログラムを拡充。
- AI の倫理と規制：AI の安全性、透明性、倫理的利用を確保するための規制枠組みを整備。

b. AI に関する具体的な施策

英国が実施している具体的な施策の一部を以下に示す。

1. AI 研究とイノベーションの支援

- Alan Turing Institute：英国の国家 AI 研究機関であり、AI とデータサイエンスの研究を推進。
- UKRI(UK Research and Innovation)：AI 関連の研究プロジェクトに資金を提供。
- 産業戦略チャレンジファンド(Industrial Strategy Challenge Fund)：AI を含む先端技術の商業化を支援。

2. AI スキルと教育

- AI スキル育成プログラム：AI 分野の大学院プログラムやトレーニングコースを拡充。
- AI フェローシップ：優秀な研究者を支援するための奨学金制度。
- STEAM 教育の強化：AI 分野での将来の人材を育成するため、科学、技術、工学、芸術、数学(STEAM)教育を推進。

3. AI の倫理と規制

- AI 倫理委員会(Centre for Data Ethics and Innovation, CDEI)：AI の倫理的利用を促進し、政策提言を行う。
- AI 規制の枠組み：AI の透明性、公平性、安全性を確保するための規制を策定中。

4. AI の産業応用

- 産業界との連携：医療、金融、製造業などの分野で AI の応用を促進。
- スマートシティと AI：都市計画や交通管理に AI を活用。

c. 国際的な連携

英国は、AI 分野での国際的な協力を重視している。

- OECD の AI 原則：英国は OECD の AI 原則を支持し、国際的な AI 規制の調和を目指す。
- G7/G20 での AI 議論：英国は AI の倫理的利用や規制に関する議論をリードする。

- EU との協力：Brexit 後も，AI 分野での研究協力を維持する．

d. AI に関する予算と投資

英国政府は，AI 分野への投資を増加させている．2021 年には，AI 研究とイノベーションに数億ポンドの予算を割り当てた．民間セクターとの連携を強化し，AI スタートアップや中小企業への支援を拡大している．

e. 今後の課題

- 規制のバランス：イノベーションを阻害せずに，AI の安全性と倫理性を確保する必要がある．
- 人材不足：AI 分野での専門人材の需要が高まる中，教育とトレーニングの拡充が必要である．
- 国際競争：米国や中国などの AI 大国との競争にどう対応するか，検討と対策が必要である．

英国は，AI 分野でのリーダーシップを目指し，研究，教育，規制，産業応用の各分野で積極的に施策を展開している．特に，AI の倫理的利用や規制において，国際的な議論をリードする役割を果たしている．今後も AI 分野での投資と政策の進展が期待される．

米国の AI 政策と施策

AI に関する米国の政策と施策は，政府，議会，民間セクターが協力して，AI 研究開発への投資，倫理的利用の推進，人材育成，国際協力など，多岐にわたる分野で進展している．特に，生成 AI や軍事 AI の分野では，他国に先駆けてリーダーシップを発揮している．一方で，AI 規制の枠組みについては，今後の議論が重要となる．

a. 国家 AI 戦略

米国は AI 分野でのリーダーシップを維持するために，国家 AI 戦略を策定している．

1. 国家 AI イニシアティブ法 (National AI Initiative Act of 2020)

2020 年に成立したこの法律は，AI 研究，開発，教育，倫理的利用を推進するための包括的な枠組みを提供している．この法律に基づき，以下のような取り組みが進められている：

- 国家 AI イニシアティブオフィス (National AI Initiative Office) の設立：AI 政策の調整と推進を担当．
- AI 研究機関の設立：大学や研究機関での AI 研究を支援．
- 連邦政府の AI 活用：政府機関での AI 導入を促進．

2. AI に関する大統領令 (Executive Orders)

トランプ政権(1 期目)下では「American AI Initiative」が発表され、バイデン政権もこれを引き継ぎ、AI 研究開発への投資を拡大した。トランプ政権(2 期目)では、さらなる加速が予想される。

b. 研究開発への投資

米国政府は、AI 研究開発に多額の予算を投じている。

1. 国家科学財団(NSF)

AI 研究のための新しいセンターを設立し、基礎研究を支援している。特に、倫理的 AI, 説明可能 AI, ロボティクスとの連携、自然言語処理などの分野に注力している。

2. 国防総省(DoD)

米国国防総省は、AI を軍事用途に活用するための「Joint Artificial Intelligence Center (JAIC)」を設立し、AI 技術の開発と導入を進めている。

3. DARPA(国防高等研究計画局)

AI の最先端技術を開発するための「AI Next」プログラムを推進している。

c. AI 倫理と規制

1. AI 倫理ガイドライン

米国政府は、AI の透明性、公平性、説明可能性を確保するためのガイドラインを策定している。これには、プライバシー保護やバイアスの排除が含まれる。

2. NIST(国立標準技術研究所)

NIST は、AI システムの信頼性と安全性を評価するための基準を開発している。

3. AI 規制の議論

米国議会では、AI の規制に関する議論を進めている。特に、生成 AI の影響や、AI によるディープフェイクの規制が注目されている。

d. 国際協力

米国は、AI に関する国際協力も積極的に行っている。

1. OECD の AI 原則

米国は、OECD が策定した AI 原則(透明性、公平性、説明可能性など)を支持している。

2. G7 や Quad での AI 議論

米国は、G7 や Quad(米国、日本、インド、オーストラリア)などを通じて、AI の国際的な枠組みを議論している。

e. 教育と人材育成

AI 分野の人材育成も重要な政策課題である。

1. STEM 教育の強化

AI に関連する科学，技術，工学，数学(STEM)教育を強化し，次世代の AI 人材を育成する。

2. AI リテラシーの普及

一般市民や労働者に対して，AI リテラシーを向上させるためのプログラムを推進する。

f. 民間セクターとの連携

米国政府は，民間セクターとの連携を重視している。

1. Google, Apple, Meta, Amazon, Microsoft, OpenAI などの企業との協力

民間企業が開発する AI 技術を政府機関で活用する取り組みが進行している。

2. AI 倫理に関する共同研究

民間企業と連携して，AI の倫理的利用に関する研究を進めている。

中国の AI 政策と施策

中国における AI 政策と施策は，政府主導で積極的に推進されている。

a. 国家 AI 戦略

中国政府は 2017 年に「新一代人工知能発展計画（Next Generation Artificial Intelligence Development Plan）」を発表した。この計画では，2030 年までに中国を AI 分野で世界のリーダーにすることを目標としている。具体的には以下の段階的な目標が設定されている。

2020 年まで：AI 技術とアプリケーションで世界の先進国と肩を並べる。

2025 年まで：AI 産業の主要分野で世界のリーダーとなる。

2030 年まで：AI 分野で世界の中心地となり、主要な技術革新をリードする。

b. 政府の投資と支援

AI 研究開発に多額の投資を行っている。とくに以下の分野に注力している。

基礎研究：AI アルゴリズム，機械学習，自然言語処理，コンピュータビジョンなど

応用分野：医療，交通，教育，農業，軍事，スマートシティなど

インフラ整備：AI 研究のためのスーパーコンピュータセンターやデータセンターの構築

また，地方政府も AI 産業の育成に積極的であり，北京，上海，深圳，杭州などの都市が AI ハブとして発展している。

c. AI 倫理と規制

AI の倫理や規制に関して、2021 年に「新世代 AI 倫理規範」を発表した。これには AI の開発と利用における倫理的なガイドラインが示され、以下の原則が含まれる：

人間中心：AI は人間の福祉を向上させるために使用されるべきである。

透明性：AI システムの意思決定プロセスは透明であるべきである。

安全性：AI の利用は安全で信頼できるものであるべきである。

また、アルゴリズムの透明性や公平性を確保するための規制を強化するため、2022 年に「アルゴリズム規制」に関するガイドラインを発表した。

d. 軍事利用

各種ドローンを含む無人機の制御（ドローン本体は AI でなく機械工学の範疇）、サイバーセキュリティ、監視技術など、AI を活用した軍事技術の開発が進められている。

e. 監視技術と社会的影響

AI を監視技術に活用しており、顔認識技術やビッグデータ分析を用いた社会信用システム（Social Credit System）を展開している。これにより、公共の安全や秩序の維持を目指す一方で、プライバシーや人権の観点から国際社会からの批判もある。

f. 国際協力と競争

AI 分野で国際的な協力を進める一方で、米国や欧州との競争も激化している。とくに、AI チップや GPU など半導体技術の分野では、米国の輸出規制が中国の AI 開発に影響を与えている。

g. 教育と人材育成

AI 人材の育成にも力を入れている。大学や研究機関で AI 関連のプログラムを拡充し、次世代の AI 研究者やエンジニアを育成している。また、企業と連携して実践的なスキルを持つ人材を育てる取り組みも行われている。

中国の AI 政策と施策は政府主導で進められており、その迅速かつ強力な対応から国際的に影響を強めている。一方で、国際的な規制との協調や倫理的な課題も存在する。

3.3.2. AI 産業動向

現代社会において科学技術への期待は大きい。人々は、科学技術は何か新しい発見や技術発明をし、我々の社会を豊かにしてくれるという期待を抱いている。過去において科学技術は産業として社会実装され、実際、経済社会を発展させる原動力となり、社会

を牽引してきた。近年の 20-30 年間においては、AI をはじめとした IT 技術が大きな発展を遂げ、産業および経済を牽引している。

IT と AI の融合

これまで、大容量高速無線通信、Internet of Things (IoT) を含む通信の多様化、ネットワーク検索技術、ネットワークセキュリティ技術、および AI 技術は、個別に論じられてきた。これらの技術はいずれも IT 技術であって相性が良く、今後、これらの技術は個別でなく、複数あるいはすべての技術が連携およびシステム統合に向かうと予想される。

6G(第 6 世代移動通信システム)は、現在の 5G の次世代となる通信技術で、2030 年頃の商用化を目指して研究・開発が進められている。6G は、通信速度、接続性、低遅延性、エネルギー効率など、5G をさらに進化させた特性を持つと期待されている。6G は、通信速度が最大で 1Tbps(テラビット毎秒)に達すると予測されており、これは 5G の最大速度(約 10Gbps)を大幅に上回る。遅延時間は 1 ミリ秒未満からマイクロ秒レベルにまで短縮されることを目指しており、これにより、自動運転や遠隔手術などリアルタイム性が求められるアプリケーションへの応用が期待される。接続の規模は、IoT デバイスの爆発的な増加に対応し、6G は 1 平方キロメートルあたり数百万台のデバイスを接続できる能力を持つとされている。さらには、エネルギー消費を最小限に抑える設が進められ、位置情報についてもセンチメートル単位の高精度な位置情報を提供することが可能になるとされている。さらには、ゼロトラストネットワークに対応するため、暗号通信技術、ブロックチェーンなどのネットワーク上におけるデータ保護技術などとの融合が進められている。

AI との技術統合については「データ分散型 AI」への活用として進められている。現在主流となっている一箇所に学習データを集約して大型 AI を構成する「データ集約型 AI」に対して、現場で計測したリアルタイム計測データをエッジ AI 間の連携で統合して解析する「データ分散型 AI」の研究開発が進められている。特に、都市開発、医療など、リアルタイムあるいはユーザ個々の計測データが重要な領域においてデータ集約型 AI への期待が高まるとともに、データ集約型 AI とデータ分散型 AI の連携および統合が進められている。

日本における AI 産業と周辺分野

AI をはじめとした計算機科学(Computer Science)および情報科学(Information Science, Informatics)は概念や思考(数学やアルゴリズム)に基づいており、計算やオペレーションのためのハードウェアを除き、物理的な実体をともしない。すなわち、その本質部分を目で見て、手で触ることができず、実体験としての感覚的経験を得にくい。これが情報技術(Information Technology, IT)や AI に対する理解、IT リテラシーや AI リテラシー

が社会で進みづらい要因のひとつとなっている。一部にこれを理由として社会における IT や AI の軽視論や不要論に結びつける意見も見られるが、少々乱暴な論理である。（英訳に備えて以下を注釈する：自然科学が含まれる「形而下学」に対して、物事の本質や真理および実在といった抽象的な概念を扱って物理的な現象を超えた領域を考察する「形而上学」など概念にかかる哲学や科学の重要性が広く認識されている欧米をはじめとした国々では、このような意見は多く見られない。）また、衣・食・住など生命に関わる産業でないことも軽視や不要論につながる要因のひとつであるが、こちらも自動車産業や電気通信業の発展をみれば、軽視すべきでないことは明らかである。ある産業分野強化の影響はそれ以外の周辺分野にも波及し、産業界全体の強化、ひいては国力の強化につながる。勢いのある産業を強化、発展させることは基本である。加えて、勢いのある産業を強化、発展させることで、その他の産業の活性化と発展につなげることが重要であり、全体のバランスをとった産業振興が重要となる。また、現在の産業のポテンシャルを考慮することも重要であり、将来性と現在の強みの両方を活かしたバランスの産業振興が必要とされる。

3.3.3. AI にかかる学術教育、学術振興

AI にかかる学術教育および学術振興について述べる。

AI にかかる基礎教育と応用教育のバランス

AI に関する基礎教育と応用教育のバランスは、AI 技術の進化や社会的なニーズに応じて重要性が変わるため、技術の進化や社会のニーズに応じて柔軟に調整されるべきある。基礎と応用の両方を重視することで、AI 技術を安全かつ効果的に活用できる人材を育成することが重要である。

1. 基礎教育の重要性

基礎教育は、AI の理論や基本的な仕組みを理解するための土台を築くものである。これには、数学、統計学、データサイエンス、ソフトウェア工学などの計算機科学などの基礎理論が含まれる。AI の基礎教育には、以下のメリットが考えられる。

- AI 技術の限界やリスクを正しく認識できる。
- 新しいアルゴリズムやモデルの開発に貢献できる。
- 応用技術のブラックボックス化を防ぎ、透明性を高めるとともに、適切に使用されるようシステムを設計および設定できる。

2. 応用教育の重要性

応用教育は、AI 技術を実際の課題解決に活用するためのスキルを養うものである。これには、AI ツールやプラットフォームの使用法、特定の産業分野での AI の応用、

プロジェクト管理などが含まれる。AI の応用教育には、以下のメリットが考えられる。

- AI 技術を迅速に実務に導入できる。
- 産業や社会のニーズに応じた AI ソリューションを提供できる。
- 非技術者でも AI の恩恵を享受できるシステムを設定でき、実務に導入できる。

AI 教育の課題と教育強化の効果

AI にかかる教育の課題として、大きく 3 つの課題が考えられる。ひとつは、教育格差の解消であり、地域や経済状況による AI 教育へのアクセスの格差を解消する必要がある。ふたつ目は、倫理教育の推進であり、AI の応用が進む中で倫理的な問題に対処する能力を育成することの重要性が増している。三つ目は、AI 教育の標準化であり、各国や地域で AI 教育の内容や水準を標準化する取り組みが求められている。また、これらを IT や AI 技術で解決しようとする試みも始められている。

AI 教育を効率化することで、以下のメリットが期待できる。

- リテラシーの普及：初等・中等教育で AI の基本概念を教えることで、社会全体の AI リテラシーを向上させる。
- 専門教育の深化：高等教育や専門教育で、基礎と応用の両方をバランスよく学べるカリキュラムを提供する。
- 産業界との連携：実務に直結した応用教育を強化し、産業界のニーズに応える人材を育成する。

AI 教育と産業

学術教育と産業は、人材リソースの育成供給とその活躍の場・社会接点という意味で、深い関係にある。中長期的視野に立ったとき、各分野の学術教育規模のバランスがそのまま産業界への人材供給のバランスとなり、ひいてはその国における産業分野のバランスや振興に直結する。学術振興を考えたとき、すなわち学科や専攻のバランスを考えたとき、その全体のバランスが産業界への人材リソース供給のバランスとなることを考慮すべきであり、長期的には社会ニーズに合わせる必要がある。一方で、急激なバランス変更は社会への影響が大きく、検討を重ねて、最適な状態とそこへの変革の速度を慎重に決定すべきである。

4. 次世代人工知能技術の適切な社会実装・運用に向けた提言

ここまでで、次世代人工知能技術について、状況や諸外国が発信した提言等についてまとめた。本章では、我々が次世代人工知能技術の目指すべき姿について提案する。

我々が目指す AI とヒトの共生社会は以下に示す 3 つの目標を掲げる(図 4)。

1. ヒトと AI の共生による豊かな未来社会の創成

AI はヒトとのインタフェースを洗練し続け、我々が AI インタフェースの操作を学んだり意識したりしなくとも、我々と共生し協働できる社会が実現する。ヒトと AI は得意とするタスクは異なるため、互いの能力を補完して活かすことで、より豊かな未来社会の創成が期待できる。

2. 地域や人種、性別・ジェンダーなどを理由とした不当な格差のない社会

現在ではネットワークを介したコミュニケーションが可能になり国家や地域の境界が排除されつつある。さらには、自然言語処理 AI の発展にともない、言語の違いによる境界までも排除しつつある。我々は、地域や人種、性別・ジェンダーといったボーダーを超えたコミュニケーションが可能になったが、一方で文化や価値観の差異など今まで明瞭化されてこなかった問題が浮き彫りになってきた。これらの課題は IT 化や AI 化を原因とするものではないが、IT 化社会および AI 化社会において議論を避けることができない。逆に、これまで自由に交流できてこなかったこれらの差異あるコミュニティが、IT 技術や AI 技術によって交流を促進し、不当な格差の解決に向けた道筋をつけるチャンスであると捉えることができ、また IT と AI はそれを解決するポテンシャルを秘めた技術であると考ええる。

3. 戦争や侵略などの危機のない安全で平和な社会

戦争や侵略などの紛争が世界各地で起きている。AI が兵器や軍事システムに利用されることで、戦争の自動化や倫理的な問題が生じる可能性がある。兵器の自動化など人間の関与が減少することで、戦争開始／維持のハードルが下がり、被害が拡大するリスクがある。科学技術は、本来、人類の幸福と繁栄を志向した平和的利用に適用されるべきであり、AI が文化や価値観の差異を超えたコミュニケーションを促進し、世界平和と人類の幸福および繁栄に帰することを旨とする。

これら 3 つの目標を達成するため、下記に示す 4 つの観点で AI の進化と深化を進めていくことが重要である。

1. AI の透明化、高信頼化

各国や地域で AI に関する規制が進むなか、AI Act などにおいては法的要件として透明性が求められており、ユーザが AI の動作や意思決定プロセスを理解しやすいインタ

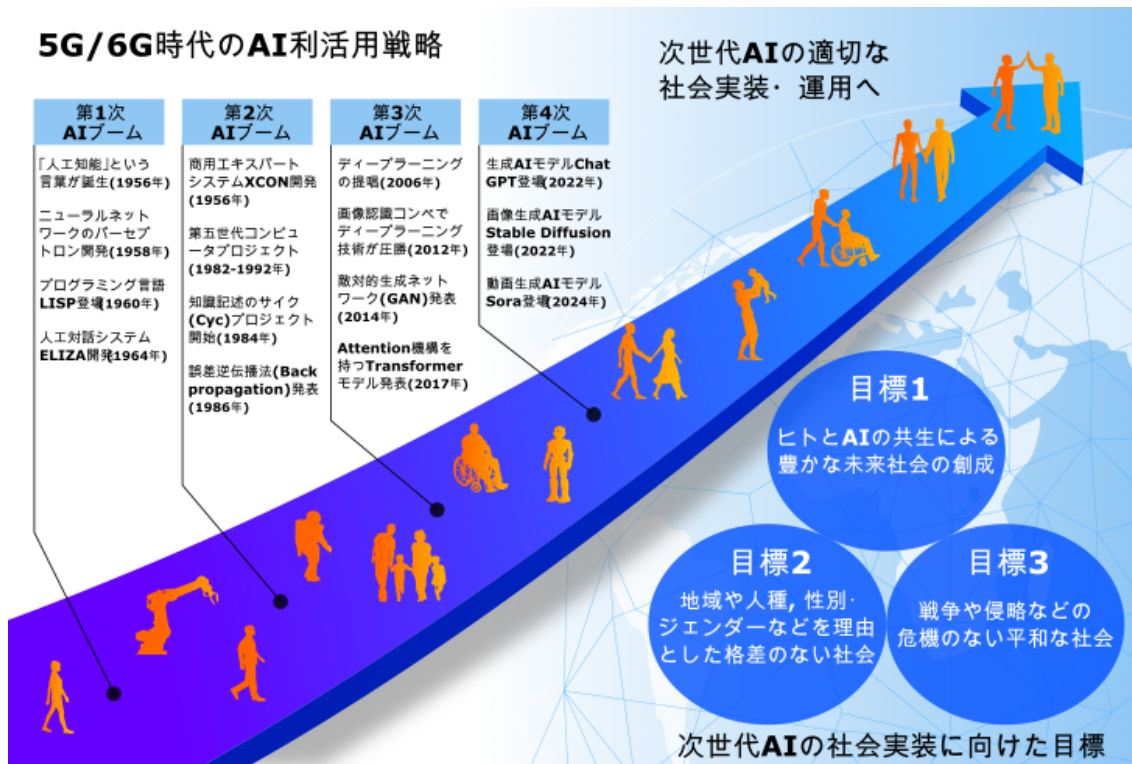


図4 AIの適切な社会実装・運用に向けた3つの目標

フェースが求められている。AIモデルがどのようにして結論を導き出したのかを説明する技術である説明可能AI(Explainable AI, XAI)技術、より透明性の高いAIモデルの選択、データの出所、品質およびバイアスの有無を明らかにするなど学習データの透明化、アルゴリズムの透明化、AI意思決定プロセスの可視化など、さまざまな観点より、AI意思決定プロセスの透明化が進められている。透明化、高信頼化の先には、制御不能の回避にとどまらず、AIにかかる格差や差別など倫理問題の解決などが期待される。

2. AIの汎用化

現存するAIの多くは目的特化型であることが多くAIの社会実装を考えたとき、汎用化、すなわちAIの自動検索および選択は、その種類や仕様の選択におけるコストを低減するとともに、潜在リスクの発見やハザードへの気づきなど安全性と信頼性の向上に寄与する。また、ヒトがコンピュータを意識せずとも協働できる環境の構築により、技術スキル格差をなくしてAIのユーザ数を指数的に増大させて社会導入が進み、誰もがAIによる恩恵を傍受できるようになる。AIに加えて、5G/6Gなど大容量高速無線通信、暗号化技術、ブロックチェーンなどのデータ保護技術、量子コンピューティング技術など、あらゆるIT技術がそれぞれで進化発展するとともに、互いの特徴や優位性を活かしながら統合された形態でも社会実装され、人類の豊かさに貢献するようになる。相性の良いこれらの技術が独立のままでいるとは思われず、高い確率をもって統合利用に向かうと予想される。さらには、データ集約型(知識集約型)AIで学習し獲得した知

識と、データ分散型（知識分散型）AI で計測したリアルタイム／準リアルタイムなデータを情報接続 AI エージェントによりフレキシブルに連携させ、限界のないフレキシビリティと汎用性を実現する。

3. ヒト-AI インタフェースのマルチモーダル化、AI インタフェースの進化、Zero UI 化

インタフェースを進化させることにより、より（人間にとって）自然かつシームレスな人間と AI のインタラクションが可能となる。インタフェース AI では、カメラ映像 AI (Vision AI) は我々も含めて映像に映る物体の種類を識別してその動作を識別し、マイク音声 AI (Auditory AI) はマイクで計測された音の種類や我々の会話を聞き分けてテキストデータに変換することができる。これらの情報に基づいて、あたかも AI が我々を理解し協働しているかのようなシステム、コンピュータであることを意識せずともヒトとコンピュータが協働できる社会を作り上げることが可能となる。

また、自然言語変換 AI を搭載したインタフェース AI の登場により、地域や人種、母国語の違いを超えて、さまざまなコミュニティの人々がコミュニケーションを取れるようになり、互いの持つ文化や価値観の相互理解を促進するとともに、差別問題の解決、ひいては世界平和につながるものと考ええる。

4. 法規制および産業教育振興など社会体制の整備、倫理・価値観の多様化への対応

これまでで述べてきたように、AI にかかる倫理問題の多くは AI に特化し限定されたものでなく、「個人の正当な権利を守り、不当な不利益を生じさせない」原理原則に則って判断されるべきものである。その際、原理原則を犯すほとんどのケースにおいてその原因は人間にあり、不適切な目的や条件での使用である。個人情報保護、社会倫理、倫理・価値観の多様化への対応など、科学技術が人類にとってパワフルであればあるほど不適切使用におけるデメリットも大きく、社会一般の AI リテラシーを向上させるとともにしっかりとした法規制などの社会体制の整備が必要である。また、産業および教育の振興においても、他の分野や産業を牽引し、社会全体としてトータルに人類の幸福と豊かさに直結するよう政策を決定し推進すべきである。

最近数 10 年に渡る AI の進化発展は目を見張るものがあり、AI は我々の生活において無視できない存在となり、欠かすことのできないインフラとなりつつある。今後も、高信頼化、汎用化、コモディティ化が進み、我々の生活のあらゆる場面で AI とヒトが協働し、社会を形成していくことが予想される。一方、AI の社会導入が進むにつれて生じる懸念も明瞭化されつつある。適切な時期に遅れることなく課題への対策や AI 推進施策を検討し、適切な速度で AI の社会導入を進めるべきである。

参考文献

- [1] 総務省. “令和 5 年版 情報通信白書”. 総務省 情報通信統計データベース. 2023 年 7 月.
<https://www.soumu.go.jp/johotsusintokei/whitepaper/ja/r05/pdf/00zentai.pdf> (参照 2024 年 12 月 20 日)
- [2] 総務省. “5G セキュリティガイドライン第 1 版”. 総務省 サイバーセキュリティタスクフォース (第 37 回). 2022 年 4 月 22 日.
https://www.soumu.go.jp/main_content/000812253.pdf (参照 2024 年 12 月 20 日)
- [3] 株式会社 NTT ドコモ. “ホワイトペーパー 5G の高度化と 6G (5.0 版)”. ドコモ 6G ホワイトペーパー. 2022 年 11 月.
https://www.docomo.ne.jp/binary/pdf/corporate/technology/whitepaper_6g/DOCOMO_6G_White_PaperJP_20221116.pdf (参照 2024 年 12 月 20 日)
- [4] John McCarthy. “WHAT IS ARTIFICIAL INTELLIGENCE?” November 12, 2007.
<https://www-formal.stanford.edu/jmc/whatisai/node1.html> (参照 2024 年 12 月 20 日)
- [5] 松尾豊編著. 中島秀之ほか. “人工知能とは”. 近代科学社, 2016.
- [6] Morris MR, Sohl-Dickstein J, Fiedel N, et al. “Position: Levels of AGI for operationalizing progress on the path to AGI.” 41th Int Conf Mach Learn, 2024.
- [7] Butlin P, Long R, Elmoznino E, et al. “Consciousness in artificial intelligence: insights from the science of consciousness.” arXiv preprint arXiv:2308.08708, 2023.
- [8] Yue X, Ni Y, Zhang K, et al. “MMMU: A massive multi-discipline multimodal understanding and reasoning benchmark for expert AGI.” IEEE Comput Soc Conf Comput Vis Pattern Recognit, 2024.
- [9] 総務省. “令和 6 年版 情報通信白書”. 総務省 情報通信統計データベース. 2024 年 7 月.
<https://www.soumu.go.jp/johotsusintokei/whitepaper/ja/r06/pdf/00zentai.pdf> (参照 2024 年 12 月 20 日)
- [10] Krizhevsky A, Sutskever I, Hinton GE. “ImageNet classification with deep convolutional neural networks.” Adv Neural Inf Process Syst, 25, 2012.
- [11] Goodfellow I, Pouget-Abadie J, Mirza M, et al. “Generative adversarial nets.” Adv Neural Inf Process Syst, 27, 2014.
- [12] Vaswani A, Shazeer N, Parmar N, et al. “Attention is all you need.” Adv Neural Inf Process Syst, 30, 2017.
- [13] Ho J, Jain A, Abbeel P. “Denoising diffusion probabilistic models.” Adv Neural Inf Process Syst, 33, 2020.
- [14] Redford A, Narasimhan K, Salimans T, et al. “Improving language understanding by generative pre-training.” 2018.
https://cdn.openai.com/research-covers/language-unsupervised/language_understanding_paper.pdf (参照 2024 年 12 月 20 日)

- [15] Redford A, Wu J, Child R, et al. “Language models are unsupervised multitask learners.” 2019.
https://cdn.openai.com/better-language-models/language_models_are_unsupervised_multitask_learners.pdf
(参照 2024 年 12 月 20 日)
- [16] Brown TB, Mann B, Ryder N, et al. “Language models are few-shot learners.” Adv Neural Inf Process Syst, 33, 2020.
- [17] Kirillov A, Mintun E, Ravi N, et al. “Segment Anything.” IEEE Int Conf Comput Vis, 4015–4026, 2023.
- [18] United Nations, “New UN Advisory Body aims to harness AI for the common good.” 2023.
<https://news.un.org/en/story/2023/10/1142867>
(参照 2024 年 12 月 20 日)

会議開催記録

「5G/6G時代のAI利活用戦略」プロジェクト、ならびに「生成AIをはじめとしたAIによる社会変容とリスクマネジメントに関する調査研究」において実施した会議および現地施設訪問の記録を以下の表にまとめる。

表2 「5G/6G時代のAI利活用戦略」プロジェクトにおける会議開催記録
(全てオンライン会議にて実施)

会 議	日 時	参加者(敬称略)	内 容
運営会議	2022/03/03 19:30-20:30	森健策/浅間一/中島義和/杉野貴明(4名)	プロジェクトメンバー構成, 検討事項等のすりあわせに関するミーティング
第1回全体会議	2022/05/26 19:00-20:30	森健策/浅間一/城石芳博/中島義和/中村道治/杉野貴明/漆谷重雄/岡田慧/喜連川優/金太一/高地雄太/新保史生/須藤修/谷口忠大/寺田努/萩田紀博/原山優子/間瀬健二/松田尚子(19名)	キックオフミーティング(メンバー紹介, プロジェクトの目的および検討事項・スケジュール確認等)
第2回全体会議	2022/08/03 19:00-20:30	森健策/浅間一/城石芳博/中島義和/中村道治/杉野貴明/漆谷重雄/中村武宏/喜連川優/金太一/高地雄太/新保史生/谷口忠大/寺田努/萩田紀博/間瀬健二(16名)	委員および外部有識者による話題提供 ①「5Gの更なる高度化と6G」 中村 武宏 (NTT ドコモ株式会社) ②「5G/6G時代のAI基盤技術に向けて～大規模言語モデル、世界モデル、記号創発システム～」 谷口 忠大 (立命館大学)
第3回全体会議	2022/09/13 19:00-21:00	森健策/浅間一/城石芳博/中村道治/杉野貴明/漆谷重雄/岡田慧/喜連川優/金太一/寺田努/萩田紀博/原山優子/間瀬健二/松尾豊/松田尚子(15名)	委員による話題提供 ①「5G/6G時代におけるメディカルイメージングAI」 森 健策 (名古屋大学) ②「5G/6G時代のAI応用技術～知能ロボットの情報基盤システム～」 岡田 慧 (東京大学)
第4回全体会議	2022/10/12 19:00-21:30	森健策/浅間一/城石芳博/中島義和/中村道治/杉野貴明/漆谷重雄/岡田慧/喜連川優/金太一/高地雄太/新保史生/谷口忠大/寺田努/萩田紀博/間瀬健二/松尾豊/松田尚子(18名)	委員による話題提供 ①「深層学習の進展と世界モデル」 松尾 豊 (東京大学) ②「ウェアラブルセンシング・情報提示とAI」 寺田 努 (神戸大学) ③「知的インタフェースから人間とコンピュータの共生へ」 間瀬 健二 (名古屋大学)
第5回全体会議	2022/12/12 19:00-21:00	森健策/浅間一/城石芳博/中島義和/中村道治/杉野貴明/漆谷重雄/岡田慧/喜連川優/金太一/高地雄太/新保史生/谷口忠大/寺田努/萩田紀博/間瀬健二/松田尚子(17名)	委員による話題提供 ①「自律自動ネットワーク連携型医療データ・AIシステム」 中島 義和 (東京医科歯科大学(現 東京科学大学)) ②「AI規制の動向 ―EUにおけるAI規制政策の構造と今後の展開を中心に―」 新保 史生 (慶應義塾大学)
第6回全体会議	2023/05/10 17:00-19:00	森健策/浅間一/城石芳博/中島義和/中村道治/杉野貴明/岡田慧/喜連川優/金太一/谷口忠大/寺田努/原山優子/間瀬健二(13名)	委員による話題提供 ①「5G/6G通信技術, 次世代AI技術の可能性と社会実装上の課題」 浅間 一 (東京大学) ②「5G/6G通信技術, 次世代AI技術の社会実装(医療)に関する取り組み」 金 太一 (東京大学)
第7回全体会議	2023/10/31 18:00-20:00	森健策/浅間一/城石芳博/中島義和/杉野貴明/岡田慧/喜連川優/金太一/高地雄太/寺田努/萩田紀博/原山優子/間瀬健二/松田尚子(14名)	委員による話題提供 ①「AI技術が切り開くゲノム精密医療」 高地 雄太 (東京医科歯科大学(現 東京科学大学)) ②「bestat株式会社における取り組み」 松田 尚子 (bestat株式会社)

「生成 AI をはじめとした AI による社会変容とリスクマネジメントに関する調査研究」における調査記録（国外調査は現地施設訪問，国内調査はオンライン会議にて実施）

国外調査				
訪問都市	訪問施設	日 時	参加者(敬称略)	内 容
London (the United Kingdom)	University College London (UCL)	2024/07/03 14:00-15:50	森健策/浅間一/Sarah Spurgeon/ Polina Bayvel/Cyril Renaud/James Seddon (6名)	UCL の Optical Networks Group と Photonics Group におけるミーティング・意見交換と研究施設見学を通じた通信技術開発動向の調査
	Royal Academy of Engineering (Prince Philip House)	2024/07/04 14:00-17:30	森健策/浅間一/中島義和/ Mike Short / Rahim Tafazolli / Dritan Kaleshi / Muhammad Miah / Simon Rowell / Andrew Smith / Alexandra Smyth / Shane McHugh / Alaka Bhatt / Taylor Huson (13名)	AI 規制政策, 高速通信技術や AI 技術の有識者である英国の王立工学アカデミー会員の方々とのミーティング・意見交換を通じた, 英国における AI・通信技術の社会実装のための取り組みに関する動向調査.
	I-X Imperial, Imperial College London	2024/07/05 14:00-17:00	森健策/浅間一/中島義和/ Sophia Yaliraki (4名)	Imperial College London の AI 研究施設 I-X におけるミーティング・意見交換と研究施設見学を通じた AI 技術開発動向の調査
San Diego (the United States)	University of California San Diego (UC San Diego)	2024/07/18 9:30-17:00	森健策/浅間一/中島義和/ Albert P. Pisano / Bill Lin / Xinyu Zhang / Ian Galton / Sheng Xu / Mohan Paturi / Chris Longhurst / Miwako Waga (11名)	UC San Diego の 5G/6G 通信技術, ウェアラブルセンサ, AI の専門家の先生方とのミーティング・意見交換を通じた AI・通信技術開発動向の調査.
	Qualcomm Research Center	2024/07/19 10:00-11:00	森健策/浅間一/中島義和/ John E. Smee (4名)	アメリカを拠点とする通信技術の大手企業である Qualcomm Inc.の研究施設見学およびワイヤレス研究部門の方々とのミーティング・意見交換を通じた 5G/6G 通信技術の社会実装動向の調査.
	Sanford Burnham Prebys Medical Discovery Institute (SBPDiscovery)	2024/07/19 13:00-14:00	森健策/浅間一/中島義和/ David A. Brenner (4名)	バイオメディカル研究機関である SBPDiscovery 研究所の所長である David A. Brenner 氏とのミーティング・意見交換を通じた医療分野における AI 技術応用に関する動向調査.
国内調査				
対象機関		日 時	参加者(敬称略)	内 容
理化学研究所 革新知能統合研究センター		2024/10/22 10:00-12:00	森健策/中島義和/杉野貴明/中川裕志 (4名)	社会における AI 利活用と法制定チームのリーダーを務める中川裕志氏とのミーティング・意見交換を通じた AI 利活用と法規制を中心とした動向調査.
東京大学 次世代知能科学研究センター		2024/11/13 17:00-18:30	森健策/中島義和/杉野貴明/國吉康夫/原田達也 (5名)	次世代知能科学研究センターのセンター長を務める國吉康夫氏と機械知能部門長を務める原田達也氏からの話題提供・意見交換を通じた AI 研究開発・人材育成に関する動向調査.
ソフトバンク株式会社 (AI 戦略室 産学連携事業推進統括部 AI 事業研究推進部)		2024/12/06 17:00-18:00	森健策/中島義和/杉野貴明/國枝良/飯沼直祥/井上陸太/對馬清元 (7名)	ソフトバンク株式会社の AI 戦略室の方々からの話題提供・意見交換を通じた AI 研究開発・ビジネスモデル・社会実装などに関する動向調査.

本資料の転載を希望される場合は、（公社）日本工学アカデミー事務局までご相談ください。

編集発行

公益社団法人日本工学アカデミー

〒101-0064 東京都千代田区神田猿樂町二丁目7番3号HK パークビルIII 2F

Tel: 03-6811-0586 Fax: 03-6811-0587

E-mail: desk@ej.or.jp

URL: <https://www.ej.or.jp/>