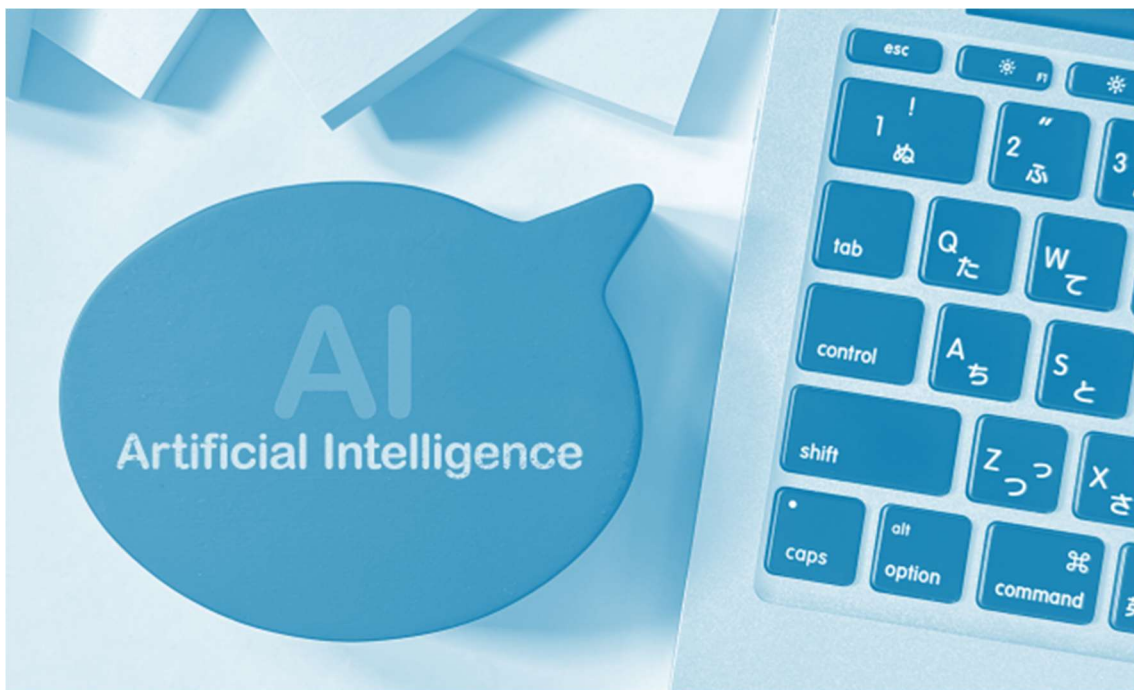


生成 AI 旋風に企業経営者はどう向き合うか

第二次人工知能(LLM：生成 AI)委員会



2025 年 8 月

一般社団法人日本経済調査協議会

Japan Economic Research Institute

「第二次人工知能(LLM：生成 AI)委員会」 委員名簿

(2025 年 3 月現在 五十音順・敬称略)

委 員 長	岩本 敏男	株式会社 N T T データグループ シニアアドバイザー
主 査	須藤 修	中央大学国際情報学部 教授
副 主 査	伏田 享平	株式会社 N T T データグループ グローバルガバナンス本部 A I ガバナンス室長
委 員	安宅 和人	慶應義塾大学環境情報学部 教授／ L I N E ヤフー株式会社 シニアストラテジスト
	江間 有沙	東京大学国際高等研究所東京カレッジ 准教授
	栗原 聡	慶應義塾大学理工学部 教授／ 一般社団法人 人工知能学会 会長
	土田 剛	株式会社 I H I 代表取締役 副社長執行役員
	徳永 俊昭	株式会社日立製作所 代表執行役 執行役副社長 デジタルシステム & サービス統括本部長
	西山 圭太	東京大学未来ビジョン研究センター 客員教授
	村野 剛太	東京海上日動火災保険株式会社 理事 I T 企画部 部長
	山口 真一	国際大学グローバル・コミュニケーション・センター 准教授
顧 問	杉浦 哲郎	当会 調査委員長

外部講師 名簿

(敬称略、所属・役職は講演当時)

後藤 信一 東海大学医学部 講師／
ブリガムアンドウィメンズホスピタル循環器内科 講師

事務局 小田 寛一 当会 専務理事
石井 浩之 当会 主任研究員
北島 基子 当会 リサーチアシスタント

第二次人工知能(LLM：生成 AI)委員会 報告書目次

はじめに	1
第 1 章 AI 開発・研究の歴史	3
1.1 人工知能の起源と初期の研究	3
1.2 第一次 AI ブーム：探索と推論（1950 年代～1960 年代）	5
1.3 第二次 AI ブーム：エキスパートシステム（1980 年代）	6
1.4 第三次 AI ブーム：機械学習（2000 年代～2010 年代）	6
1.5 第四次 AI ブーム：生成 AI（2020 年代）	7
第 2 章 AI 技術の基礎と最新動向	9
2.1 機械学習（MACHINE LEARNING）	9
2.2 深層学習（DEEP LEARNING）	10
2.3 大規模言語モデル（LLM）	11
第 3 章 AI のリスクと国際的な対応状況	15
3.1 OECD AI 原則	15
3.2 生成 AI のリスク	16
3.3 AI 規制のアプローチ：ソフトローとハードロー	17
3.4 AI リスク管理の歴史と各国の動向	18
3.5 広島 AI プロセスに基づく取り組み	18
3.6 日本の「AI 事業者ガイドライン」	19
第 4 章 生成 AI の社会実装	21
4.1 町田市 AI ナビゲーター	24
4.2 七十七銀行 生成 AI によるデータ分析自動化	25
4.3 日立製作所 製造設備保全業務への利用	26
4.4 セブンイレブン 生成 AI 商品企画システム	27
第 5 章 生成 AI 活用の戦略：日本への提言	29
5.1 企業経営者に向けた提言	29
5.2 AI・LLM 活用に向けた政府関係者への提言	32

はじめに

2024 年のノーベル物理学賞は、プリンストン大学（米国）のジョン・ホップフィールド教授と、トロント大学（カナダ）のジェフリー・ヒントン教授の 2 人に与えられた。彼らの功績は、現在の AI 技術の中核を担う「機械学習」と、その後の「深層学習」の基礎となる人工ニューラルネットワークの基礎的発見と発明である。

さらに、その翌日発表になったノーベル化学賞は、ワシントン大学（米国）のデイビッド・ベイカー教授と、グーグルのグループ会社で、ロンドンに本社のあるディープマインド社のデミス・ハサビス CEO、と研究チームのジョン・ジャンパー氏の 3 人に与えられた。ベイカー教授はアミノ酸を使って、まったく新しいたんぱく質を設計することに成功した功績であり、ハサビス氏とジャンパー氏は、たんぱく質の複雑な構造を予測するという 50 年来の問題を解決する AI モデルを開発した功績である。ハサビス氏は、2016 年、韓国囲碁界のプロ棋士イ・セドル氏を破った AI ソフト「アルファ碁」を開発したことでも良く知られている。二つのノーベル賞が AI に関係した受賞であることは、今日の AI ブームをよく表している。

AI（Artificial Intelligence）という用語が初めて使われたのは、1956 年ダートマス大学（米国）で開催された会議であり、この会議で AI の学術研究分野が確立したと言われる。以来、60 年以上にわたりコンピュータやアルゴリズムの技術革新とともに発展を続けてきた。この間、探索と推論、エキスパートシステム（知識表現）、機械学習という三つの技術的アプローチを経て AI は着実に進化を遂げてきた。

2022 年 11 月末に ChatGPT が公開されて以来、世界中が生成 AI に沸き立った。わずか 2 か月でアクティブユーザーが 1 億人を突破し、第二のインターネットの到来と考える人も多い。AI が単に学術的な分野だけでなく、政治・経済・社会のすべての領域で大きな影響力を発揮し始めた。AI ブームは次の世代に入ったと言える。

AI 技術の進展は産業界に大きな機会とリスクをもたらしている。企業の経営層は従来型の業務効率化の範疇を超えたビジネスモデルの革新や産業構造の変革といった重要な経営判断を迫られている。特に、グローバルな技術開発競争の加速は日本の産業競争力に大きな影響を及ぼす可能性がある。

一方で、AI 技術の社会実装には様々な課題が存在する。品質や信頼性の確保、プライバシーの保護、法規制への対応、人権侵害問題、戦争や詐欺等犯罪への利用、環境負荷の軽減など、技術的および倫理的な課題は依然として山積している。さらに各企

業への導入に際し、企業経営の視点からは少子高齢化による労働力不足の解消や、生産性の向上が主要な目的であり、その可能性は格段に高いと言われている。しかし、期待通りの実効性が確実に得られるとは断言できない。

そしてヒントン教授が指摘するように、上記のリスク以外に人類が未だ経験したこともないような AI ならではのリスク、人類の学習スピードをはるかに上回るパワーによって人類が及びもつかないような高度な知能を獲得することや、獲得した知能を AI 同士の凄まじいスピードでの共有することなど、新たなリスクの存在も未知数である。そして、こうしたリスクに備えるべく各企業における適切なガバナンス体制の構築も急務となっている。

そこで本委員会では各界の有識者を招いて委員会を組成し、設立趣意書に基づき、合計 11 回にわたる議論を重ねた。本委員会では、生成 AI という新技術との向き合い方について企業経営者目線で議論し、技術的な論点は最小限に留めつつ、国内は勿論であるが、米国を始めとする諸外国において、どのような活用方法やリスクが認識され、ガバナンスのあり方が議論されているかを共有した。

本報告書はこれらの調査活動を通じて得られた知見と考察を取りまとめたものである。特に、企業が直面する機会とリスク、実践的な活用方法、ガバナンスのあり方について、具体的な事例を交えながら論じている。既に生成 AI を活用し一定の成果を得ている企業の先行事例を紹介し、導入に当たっての課題や工夫された点なども共有した。それによって、企業経営者に対して生成 AI にどのように対処すべきかの示唆を提示することができた。

さらに、個人レベルでの利用が進むと、いわゆる AI エージェント機能が社会生活の中で重要な存在となり、計り知れない効能をもたらすと同時に、現在の SNS が抱える課題と同様かそれ以上のリスクの発生が想定され、これに備える必要が生じてくる。こうしたリスクについても議論を重ね、企業経営者への示唆と同時に国に対しても、個人のリテラシー向上の観点からも、何らかの示唆を提示できたと考えている。

本報告書が企業における戦略的な AI 活用の指針となり、日本の産業競争力強化に寄与することを期待している。

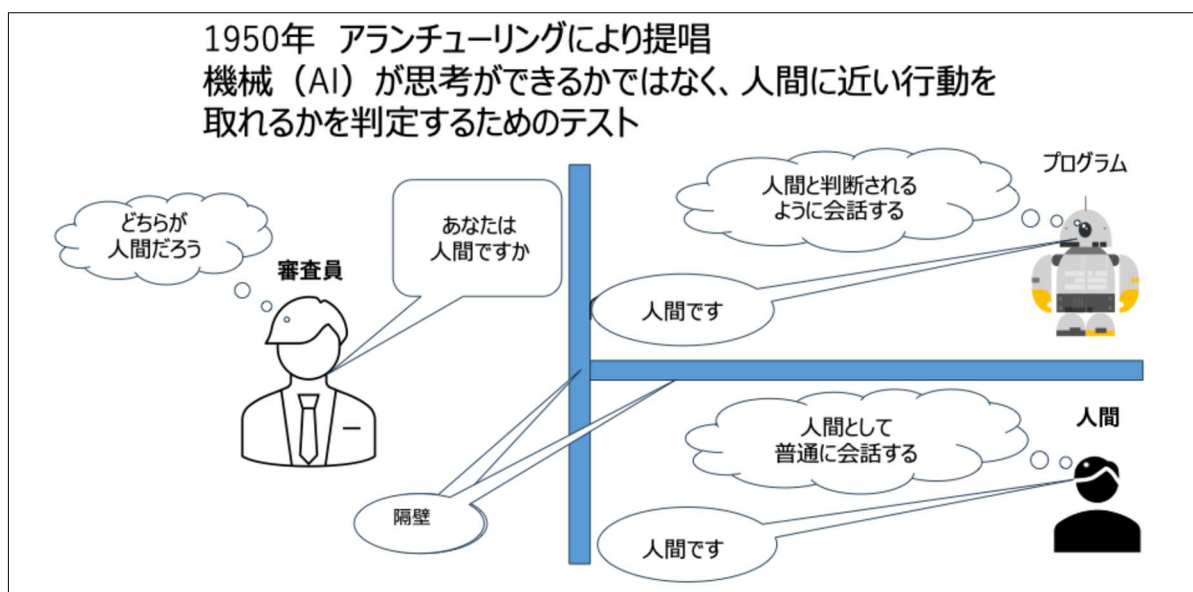
第1章 AI 開発・研究の歴史

本章ではこれまでの AI 開発・研究の歴史を概観する。特に第一次から第四次にわたる AI ブームにおいて、どのような技術的ブレークスルーが生まれ、社会実装が進んだか述べる。

1.1 人工知能の起源と初期の研究

AI の学術的な歴史は 1950 年代にさかのぼる。AI の概念を初めて提唱したのはイギリスの数学者アラン・チューリングである。ドイツ海軍の暗号システム「エニグマ」の解読に成功したことでも知られているが、彼は「計算する機械と知能」という論文の中で、コンピュータが知的行動をすることの可能性について言及し、その後の AI 研究基盤となる「チューリングテスト」を提唱した。このテストはコンピュータが人間のように会話できるかどうかを評価する基準として現在も使われている。審査員がテキストチャットを通じて人間とコンピュータ（プログラム）に質問を投げかけ、その回答によって人間かコンピュータかを見分けるテストである。

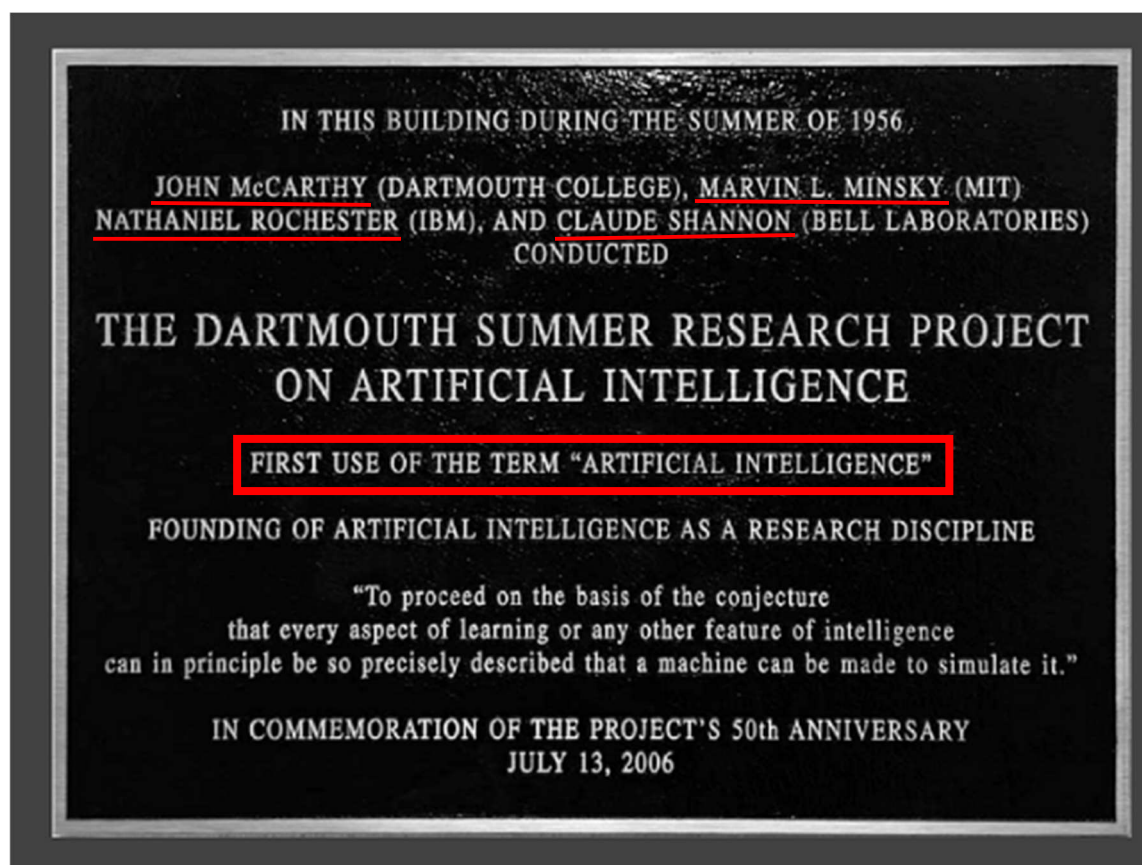
図表 1：チューリングテスト



1956 年にはアメリカのダートマス大学で開催されたダートマス会議で、「Artificial Intelligence (人工知能)」という用語が初めて公に使用された。この会議にはジョン・

マッカーシー、マービン・ミンスキー、ナサニエル・ロチェスター、クロード・シャノンなど、後の AI 研究を牽引する著名な研究者が参加し、AI 研究の正式な出発点となった。当時からコンピュータと AI は密接に関連しており、人間の知能を機械で再現するという目標に向かって双方の研究が進められた。

図表 2：ダートマスホール記念銘板



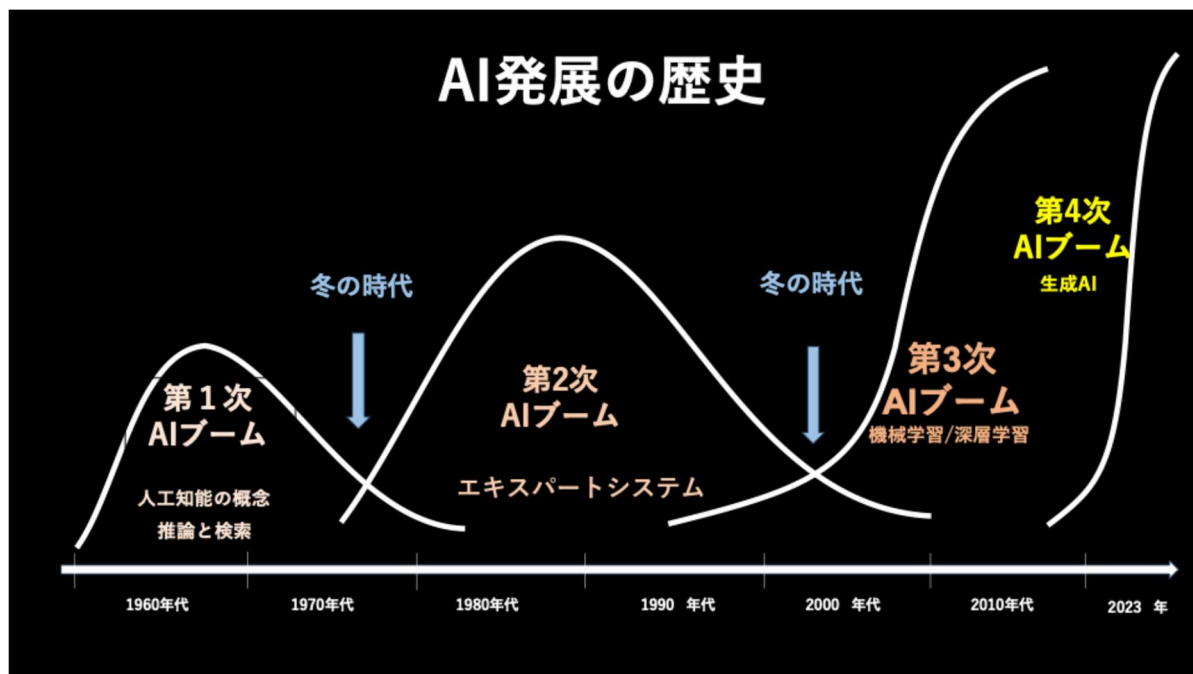
(出所) Moor, J. (2006). The Dartmouth College Artificial Intelligence Conference: The Next Fifty Years. *AI Magazine*, 27(4), 87. <https://doi.org/10.1609/aimag.v27i4.1911>
(一部着色)

AI の初期の研究では、問題解決や推論のためのルールベースのシステムが開発された。例えば、1950 年代から 1960 年代にかけて、アレン・ニューウェルとバート・サイモンが開発した「General Problem Solver (GPS)」は、記号処理を用いた問題解決アルゴリズムの基礎を築いた。また、1966 年にはジョセフ・ワイゼンバウムによる対話型プログラム「ELIZA」が開発され、人間とコンピュータの自然な対話の可能性を示した。

AI の歴史において技術的なブレークスルーと社会的な期待の高まりは、「AI ブーム」

として記述される。これまでに三つの主要なブームがあり、それぞれに特徴的な技術的アプローチと直面した課題が存在する。

図表 3：AI 発展の歴史



1.2 第一次 AI ブーム：探索と推論（1950 年代～1960 年代）

この時期には論理推論や探索アルゴリズムを用いた AI 研究が進められた。特に、ジョン・マッカーシーによる LISP 言語の開発は、AI 研究の発展に大きく貢献した。また、チェスや迷路探索などの課題に取り組むプログラムが作られ、AI の可能性が示された。

当時の AI は、論理に基づいたルールベースのシステムが中心であり、これにより特定の問題領域において高い性能を発揮した。しかし、これらのシステムは汎用的な知能を持たず、現実世界の複雑な問題に対応するには限界があった。また、コンピュータの計算能力やメモリ容量の限界により、実用化が困難であった。

1960 年代末には、AI の初期の理論的成果にもかかわらず、実用的な成果が期待に達していないことから、研究資金が削減される「AI の冬」と呼ばれる停滞期に突入した。この間、多くのプロジェクトが打ち切られ、AI の進展は一時的に停滞すること

となった。

1.3 第二次 AI ブーム：エキスパートシステム（1980 年代）

1980 年代には、エキスパートシステムと呼ばれるルールベースの AI 技術が発展した。これは、専門家の知識をプログラムとして実装することで、医療診断や機器の故障診断など、特定分野の専門家の判断をコンピュータで再現することを目指した。

エキスパートシステムは、特定の分野に特化した知識をデータベース化し、それをルールに基づいて処理することで、高度な問題解決能力を発揮した。しかし、人間が持つ「常識」や「暗黙知」をプログラムに実装することは極めて困難であり、また知識の更新や維持に多大なコストがかかるという問題に直面した。また、ルールが増加するにつれて、システムの動作が遅くなり、例外処理が困難になるという問題もあった。

1980 年代後半には、エキスパートシステムの限界が明らかになり、研究資金の削減とともに、再び AI 研究への関心が低下する「第二次 AI 冬の時代」を迎えた。しかし、この期間にはニューラルネットワークの研究が進み、後のディープラーニングの基礎となる成果が生まれた。

1.4 第三次 AI ブーム：機械学習（2000 年代～2010 年代）

2000 年代に入り、インターネットの普及によるデータの増加と計算リソースの大幅な強化を背景に、機械学習、特に深層学習が注目を集めた。機械学習はデータから規則性やパターンを自動的に学習し、それに基づいて判断や予測を行う技術である。特に深層学習はヒトの脳の神経回路ニューロンを模倣するニューラルネットワークを利用してより複雑なパターンの学習を可能にするものである。

2012 年、画像認識コンペティションである「ImageNet Large Scale Visual Recognition Challenge (ILSVRC)」において、ジェフリー・ヒント率いるチームが開発した「AlexNet」は、従来の手法を大幅に上回る精度で画像認識を行い、入力画像から特徴量を抽出しそれらを区別することができる畳み込みニューラルネットワーク

(CNN) の有用性を証明した。

このブームでは、画像認識や音声認識などの分野で人間に匹敵する、あるいは凌駕する性能が実現された。

1.5 第四次 AI ブーム：生成 AI（2020 年代）

2020 年代に入ると、大規模言語モデル（LLM）が飛躍的に発展した。第四次の AI ブームは 2022 年 11 月に登場した生成 AI ChatGPT が火付け役であるが、現代の生成 AI は深層学習を基盤としつつ新たな技術革新によって実現された。特に重要なのは「大規模なデータと並列計算技術の進化」である。これにより、従来は実現できなかった規模での学習と複雑な処理が可能となった。

深層学習の特徴は人間が明示的にルールを設定せずとも大量のデータからパターンを自動的に学習できる点にある。例えば画像認識では多数の画像データから特徴を抽出しそれらを階層的に組み合わせることで高度な認識能力を実現する。この技術は生成 AI の基礎となる重要なブレークスルーをもたらした。

生成 AI の発展において二つの重要な技術的革新が存在する。一つは 2014 年にイアン・グッドフェローらにより提案された GAN（Generative Adversarial Networks：敵対的生成ネットワーク）である。GAN はジェネレーター（生成器）とディスクリミネーター（判別器）という二つのニューラルネットワークで構成される。GAN の画期的な点は、これら二つのネットワークを競争させることで生成の精度を向上させる仕組みにある。ジェネレーターはできるだけ本物そっくりのデータを生成しようと試み、ディスクリミネーターは提示されたデータが本物か偽物（生成されたもの）かを判別しようと試みる。この競争を繰り返すことで、ジェネレーターは徐々に本物らしいデータを生成できるようになる。この技術により画像生成の分野で飛躍的な進歩が実現した。

もう一つの革新は、2017 年に Google の研究者らによって提案されたトランスフォーマーモデルである。「Attention is All You Need」と題された論文で提案されたこのモデルは文脈を精細に理解し、より高度なテキスト生成を可能にした。その核心は「Self-Attention（自己注意機構）」と呼ばれる仕組みにある。これは、文章中の各要素（単語など）の関係性を文脈に応じて柔軟に捉えることを可能にするものであり、

自然言語処理（NLP）の分野における画期的な変革をもたらした。

トランスフォーマーモデルの革新性はその構造だけでなく、実現を可能にした技術的基盤にもある。大規模なデータ処理を可能にした並列計算技術の進化がこの技術の実用化の鍵となった。この技術基盤の上に OpenAI は GPT (Generative Pre-trained Transformer) シリーズを開発し、大規模言語モデルの新たな世界を切り開いた。

現在の生成 AI 技術の発展過程において重要な発見の一つがスケーリング則である。これは、モデルの規模（パラメータ数）とデータ量を増やすことで、AI の性能が経験則で予測可能な形で向上するという法則性を示すものである。この発見は生成 AI 技術の本質を示している。すなわち、技術の基盤自体は変わらないものの、その性能向上は主にモデルの規模拡大によってもたらされるという特徴である。

GPT-3 (2020 年) は、1750 億のパラメータを持つ巨大なモデルであり、文章の生成、翻訳、コード生成など多岐にわたるタスクを高精度で実行できたが、現在の GPT-4 では 1 兆を超えるパラメータを有すると推定されている。その効果は飛躍的でありパラメータの増加は今も進んでいる。

この原理は AI 開発に二つの重要な示唆をもたらした。一つは、十分な規模のモデルさえ構築できれば、従来は困難とされた複雑なタスクも実現できる可能性があるということ。もう一つは、そのようなモデルの構築には膨大な計算資源とエネルギーが必要になるということである。

さらに、現代の生成 AI のもう一つの重要な特徴はマルチモーダル化である。これは、テキスト、画像、音声、動画など、異なる種類のデータを統合的に処理する能力を指す。このようなマルチモーダル化を支えているのが前述のトランスフォーマーモデルの応用である。様々な種類のデータを統一的な方法で扱うことが可能となり、生成 AI の応用範囲は大きく広がっている。

このように、現代の生成 AI 技術は、GAN やトランスフォーマーモデルという基盤技術の上にスケーリング則の発見とマルチモーダル化という新たな展開を加えることで、急速な発展を遂げている。これら技術の組み合わせにより、従来の AI では実現できなかった柔軟な情報処理と生成能力が実現されている。そして、こうした技術的基盤の確立により、生成 AI は社会実装の段階へと移行しつつある。

第2章 AI技術の基礎と最新動向

本章では、現在の LLM (Large Language Model) につながる AI 技術の基礎として、機械学習や深層学習について簡単に述べる。その上で、LLM やその応用技術である RAG (Retrieval Augmented Generation) やマルチモーダル、AI エージェントについても触れる。

図表 4：機械学習・深層学習・LLM の関係



2.1 機械学習 (Machine Learning)

機械学習 (Machine Learning, ML) とは、データからパターンを学習し、予測や意思決定を自動化する技術である。従来のシステムは人がルールを設定して動作させていたが、機械学習はデータからルールを自動的に発見する。

多くの組織では膨大なデータが蓄積されているが、それらの適切な活用は難しい。機械学習は、人が気付かないパターンや関係性を見出し、業務の自動化や効率化につなげることが可能となる。

機械学習は「教師あり学習」「教師なし学習」「強化学習」の3つに分類される。

- 教師あり学習 (Supervised Learning) :
 - 「正解データ」を与え、AI モデルを訓練する。
 - 過去のデータが豊富であるほど高精度な予測が可能となる。ただし、あらかじめ「正解データ」を用意する必要があるため、データの準備にコストがかかる。
 - 主な活用例として、売上予測や不正検知、画像認識や音声認識があげられる。

- 教師なし学習（Unsupervised Learning）：

- 「正解データ」を与えず、AI モデル自体が与えられたデータからパターンを学習し、データの構造や関係性を発見する。
- データから新しいパターンを発見できる。また、「正解データ」が不要であり、様々なデータに手法を適用可能である。ただし、結果（発見されたパターン）の解釈には専門的な知識が必要となる。
- 主な活用例として、顧客セグメンテーションの発見・分析やマーケット分析、製造機械の異常検知があげられる。

- 強化学習（Unsupervised Learning）：

- 開発した AI モデルの出力に対して、随時フィードバックをかけていくことで AI モデル自体が試行錯誤を繰り返しながら「最適な行動」を学習する。
- 最適な結果を自律的に学習し、リアルタイムで改善していくことが可能である。一方で、大量の試行や計算リソースが必要となる。
- 需給予測やロボット制御、自動運転で用いられている。

2.2 深層学習（Deep Learning）

深層学習は機械学習の技術の一つである。その特徴は、人間の脳の働きを模倣した「ニューラルネットワーク」を活用している点にある。深層学習の実行にあたっては大規模なデータと豊富な計算資源を必要とするが、他の機械学習手法と比較して高精度なパターン認識を行うことが可能となる。

深層学習手法は用途に応じていくつかの技術が提案されている。ここでは、代表的な手法である「CNN（畳み込みニューラルネットワーク）」「RNN（再帰型ニューラルネットワーク）」「Transformer」の3つについて概説する。

- CNN（Convolutional Neural Network：畳み込みニューラルネットワーク）：

- 画像や映像の解析に特化したモデル。画像の特徴を階層的に学習し、識別や分類を行う。具体的には、画像の一部を小さなブロック（フィルター）に分解し、それを走査しながら情報を抽出する「畳み込み」という処理を行う。
- 適用例としては、顔認識や医療画像診断、製造業における不良品検出などがある。

- RNN（Recurrent Neural Network：再帰型ニューラルネットワーク）：

- 時系列データや自然言語処理に適したモデル。「前に出現した情報を記憶し

ながら、次の情報を処理する」ため、文章や音声データの前後関係を考慮して学習できる。

➤ 適用例としては、株価や売上予測、機械翻訳などがある。

- Transformer :

➤ 自然言語処理において最も進化した技術。全体の文脈を一度に処理できるため、RNN よりも長い文章を取り扱うことができる。「アテンション」というメカニズムにより文章内の重要な単語を優先的に解析することで、並列・大規模なデータを高速に学習することが可能となっている。

➤ 適用例としては、Google 翻訳や ChatGPT などの対話型 AI に用いられている。

2.3 大規模言語モデル (LLM)

LLM (Large Language Model) は、大量のテキストデータを学習し、人間のようになら自然な文章を生成できる AI 技術である。LLM は機械学習・深層学習手法の一つであるが、膨大な量のテキストデータにより「言葉の使い方」や「文脈」を学習し、より高度な言語処理が可能となっている。

昨今の爆発的な生成 AI の広がりや、LLM の発展によるものと言える。このような潮流は、下記に示す 3 つのブレークスルーによって実現された。

- データ量の劇的な増加 :

従来の AI は、企業内で保持されているデータなど、ごく限られた量の文章データを学習して開発されていた。これに対して LLM は、インターネット上のありとあらゆるデータを学習することで、より自然かつ広範な知識を持つようになった。

- Transformer による学習手法 :

従来の AI 技術では、一つ前や後ろの単語をもとに、適切な単語を予測・生成するだけであった。しかし、LLM では 2.2 節で示した Transformer という技術を採用することで、文章全体の意味を理解し、的確な回答ができるようになった。

- 計算能力の大幅な向上 :

LLM は、一説にはインターネット上の全てのデータを学習したと言われるほど、膨大なデータを学習するため、強力な演算能力を持つ GPU が必要となる。しかし、近年はコンピュータの計算能力も向上し、これまで数ヶ月から数年か

かっていた処理が、数週間から数日で完了できるようになった。

LLM の登場により、大規模なデータを学習し文脈を理解する能力を持つことから、質問応答や文章要約・翻訳、テキスト生成など、様々なタスクに AI が対応できるようになった。一方で、最新情報を持たないことや、「ハルシネーション」と呼ばれる誤った情報を生成してしまうリスクもある。このような LLM の制限を克服するために、ファインチューニングや RAG、マルチモーダルといった拡張技術が開発されている。

2.3.1 ファインチューニング (Fine Tuning)

ファインチューニングは、すでに学習済みの LLM を、特定の用途や業務に最適化する技術である。LLM はインターネット上の膨大なデータを学習しており、一般的な会話や知識には対応できる。一方で、業界固有の専門知識や業務プロセスには対応できないことがある。

ファインチューニングでは、事前に学習済みの LLM をベースに、企業が持つ特定の業界データなどを使い追加で学習を行う。このようなファインチューニングを行うことで、より業界や目的に特化した AI を開発、活用することが可能となる。

2.3.2 RAG (Retrieval Augmented Generation, 検索拡張生成)

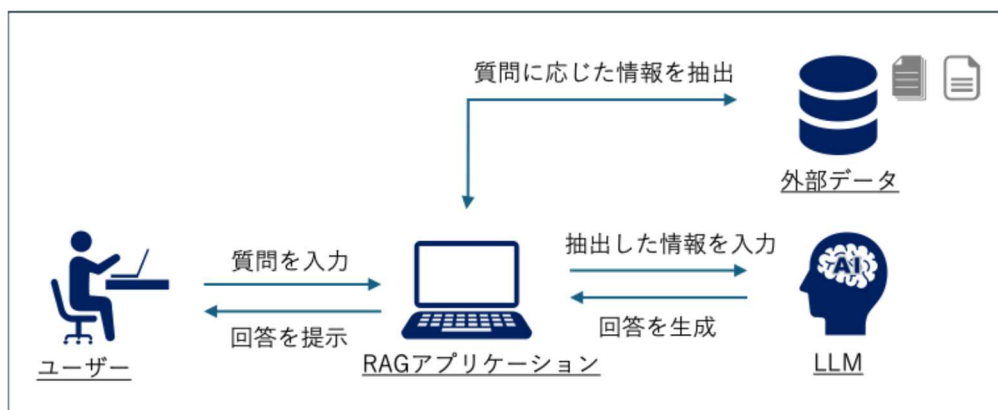
RAG は、LLM に外部情報を統合することで、より信頼性の高い回答を生成する手法である。通常、LLM は学習時点までのデータしか持っておらず、新しい情報には対応できない。また、場合によっては、文脈から推測してもっともらしいが間違った情報を生成してしまうこともある（このような現象はハルシネーションと呼ばれる。）。しかし RAG を使えば、外部データを検索しながら適切な情報を取り入れ、より信頼性の高い回答を生成できる。

RAG は下記のような手順で動作する。

1. ユーザーが RAG アプリケーション（システム）に対して質問を入力する
2. システムは、企業のデータベースや外部の最新情報を検索して必要な情報を抽出し、LLM に抽出した情報を与える (Retrieval)
3. LLM は、取得した情報をもとに文章（質問に対する回答）を生成する (Generation)

4. システムは、LLM が生成した文章をユーザーに提示する

図表 5：RAG の仕組み



RAG を用いることで、企業内のナレッジ検索や、最新のニュースや法律・規制への対応などを精度高く行うことができる。

通常の RAG では、上記の手順 2 において単純に文書を検索し、キーワードに関連する情報を抽出する。これに対して、より関連性の高い情報を取得し回答の精度を上げる手法として、**GraphRAG** と呼ばれる手法も提案されている。GraphRAG では、情報間の関連性をグラフ（ネットワーク）の形であらかじめ整理しておく。その上で情報同士の関係性を考慮して検索し、情報を抽出する。GraphRAG を用いることで、より関連性の高い情報を取得でき、より深く正確な情報を提供できる。

2.3.3. マルチモーダル AI (Multimodal AI)

マルチモーダル AI は、LLM がテキストだけでなく、画像、音声、動画などの異なるデータ形式を統合的に処理する技術である。従来の LLM はテキスト情報のみを扱っていたが、マルチモーダル AI では様々な種類の情報を組み合わせて理解し、より高度な分析や生成が可能になる。例えば、製品画像を入力し、「この製品の特徴を説明してください」と AI に指示をすると、適切な画像の説明文を自動生成できる。

マルチモーダルを活用することで、従来活用しきれていなかった画像や音声、動画と言った様々なデータを活かすことができる。また、複数の情報を統合して処理できるようになり、より高度な分析、意思決定が可能となる。

2.3.4. AI エージェント (AI Agent)

ここまで紹介してきた AI 技術は、基本的にはユーザーからの指示に対して、AI が対応するという受動的なものであった。これに対して **AI エージェント** は、AI 自身が仕事を考え、自律的に実行できる技術である。例えば、「旅行の計画を立てる」という指示を与えると、単純に交通手段や宿泊先の候補を提示するだけでなく、旅行プランを作成し、必要なチケットや宿泊先を予約手配するといった、一連の作業を AI が自律的に行う。

AI エージェントは LLM の技術が進化していることにより実現されつつある技術である。これまでの LLM では文章を生成するだけであったが、最近では AI 自身が考え、誤りがある場合には自己修正できる技術が発展している。また、AI が検索エンジンや各種 Web サービスを利用できるようになりつつあり、より複雑な作業を行うことが可能となった。

これまでの AI はあくまで業務を行う上での「ツール」であったが、AI エージェントは自律的に仕事を考え、実行できるため、人間の「パートナー」として扱うことのできる技術として着目されている。

第3章 AI のリスクと国際的な対応状況

AI の発展は、産業、経済、社会に大きな影響を与えているが、それに伴い様々なリスクが指摘されている。リスクには、バイアス、プライバシー侵害、安全性の問題、倫理的懸念、透明性の欠如などが含まれる。さらに、AI が人間の知能を超えるまで成長したあかつきには、人間の職を奪うことで経済的な混乱を引き起こす可能性や、AI そのものを人間が制御不能になる可能性について言及する研究者もいる。これらのリスクに対処するため、各国や国際機関は規制の整備を進めている。本章では、AI リスクに対する国際的な枠組みや各国での取り組みについて示す。

3.1 OECD AI 原則

AI リスクに対処するための国際的な枠組みのひとつとして、経済協力開発機構（OECD）が2019年に策定した「OECD AI 原則」があげられる。この原則はAIの開発と利用における信頼性の確保を目的とし、次の5つの主要な原則を掲げている。

1. **包摂的な成長、人間中心の価値及び公平性**：AI は経済成長と社会的福祉を促進し、すべての人に利益をもたらすべきである。
2. **持続可能性と人間の幸福**：AI の利用は環境、社会、経済の持続可能性を考慮し、人間の幸福を向上させるものでなければならない。
3. **透明性と説明責任**：AI の意思決定プロセスは透明であり、適切な説明責任が確保されるべきである。
4. **安全性とセキュリティ**：AI システムは安全であり、リスクを最小限に抑える設計が求められる。
5. **アカウンタビリティ（責任の明確化）**：AI を開発・運用する主体は、適切な責任を負い、適用される規制を遵守する必要がある。

OECD AI 原則は各国のAI政策の指針として広く採用されており、特にG7諸国やEUをはじめとする国際社会のAI規制の基盤となっている。この原則の特徴は、政府のみならず、企業や研究機関などの幅広いステークホルダーが参照できるガイドラインとして設計されている点にある。

3.2 生成 AI のリスク

生成 AI の利用が世界的に進む中で、生成 AI を安全かつ適性に利用するためには、リスクを適切に把握し、対処することが必要となる。NIST（米国国立標準技術研究所）は、生成 AI に関わる 12 のリスクを示している。企業における生成 AI 利活用にあたっては、これらのリスクを適切に管理、対処していくことが重要となる。

1. 化学・生物・放射線・核（CBRN）情報または能力：

生成 AI が危険な科学技術情報（化学兵器、生物兵器、核技術など）を拡散し、悪意のある利用を助長するリスク。

2. 虚偽生成：

AI が誤った情報や架空のデータを作成し、それを事実として提示することで、誤解や誤報を広めるリスク。

3. 危険、暴力的、または憎悪に満ちたコンテンツ：

暴力、憎悪、過激化を助長するコンテンツ（ヘイトスピーチ、過激派の宣伝、暴力的な映像やテキストなど）を生成し、社会不安や犯罪を誘発するリスク。

4. データプライバシー：

AI の学習データや生成コンテンツに個人情報が含まれることで、プライバシー侵害や個人特定につながるリスク。

5. 環境への影響：

生成 AI のトレーニングや運用に必要とされる計算リソースが膨大となり、電力消費と炭素排出量の増大を引き起こすリスク。

6. 人間と AI の構成：

AI の役割と人間の役割が曖昧な場合、誤った意思決定や制御不能な状況が発生するリスク。また、ブラックボックス化による説明責任の欠如も問題となる。

7. 情報の完全性：

AI が信頼性の低い情報や改ざんされたコンテンツを生成し、それが事実と区別されないことで、情報の信頼性が低下するリスク。

8. 情報セキュリティ：

AI がマルウェア、フィッシング、サイバー攻撃の支援や自動化を行う可能性があり、セキュリティ脅威を増加させるリスク。

9. 知的財産：

AI が著作権や商標権を侵害するコンテンツを生成することで、企業や個人の知

的財産権を侵害し、法的紛争が発生するリスク。

10. わいせつ、下劣、虐待的なコンテンツ：

AI が性的搾取・暴力的コンテンツ（児童性的虐待コンテンツ、非同意の親密な画像など）を生成し、違法コンテンツの拡散を助長するリスク。

11. 毒性、偏見、均質化：

AI が社会的バイアスを強化し、有害な固定観念を再生産するリスク。また、AI の回答が画一化され、文化的・倫理的な多様性を損なうリスク。

12. バリューチェーンとコンポーネントの統合：

生成 AI は複数のコンポーネント（データ、API、モデル）を組み合わせるため、不透明なサプライチェーン管理により、安全性や品質管理の問題が発生するリスク。

3.3 AI 規制のアプローチ：ソフトローとハードロー

AI 規制のアプローチには、法的拘束力を持たない「ソフトロー」と、強制力を伴う「ハードロー」の両面から推進されている。

ソフトローは、法的拘束力を持たないガイドライン、倫理指針、業界標準などを指す。OECD AI 原則をはじめ、国際的な倫理規範や企業の自主規制などがこれに含まれる。ソフトローの利点は、技術の急速な発展に対応しやすく、柔軟な適用が可能な点である。一方で、強制力がないため、遵守の実効性を確保する仕組みが求められる。

ハードローは、法律や規制によって強制力を持つルールを指し、AI のリスク管理を制度化する役割を果たす。たとえば、EU の「AI 規制法（AI Act）」や中国の「生成 AI サービス管理暫定弁法」などがこれに該当する。ハードローのメリットは、強制力があることで法的安定性を確保できる点にあるが、技術の進化に対して規制が陳腐化するリスクもある。

近年の傾向として、多くの国や地域では、ソフトローを基盤としつつ、リスクの大きい領域はハードローを適用する「ハイブリッドアプローチ」を採用するケースが増えている。たとえば、EU の AI Act では、リスクベースのアプローチを導入し、低リスクの AI は自主規制（ソフトロー）を推奨し、高リスクの AI には厳格な規制（ハードロー）を適用する形が取られている。

このように、AI のガバナンスは国際的な原則に基づき、ソフトローとハードローの適切な組み合わせによって形成されている。企業の経営者にとっては、これらの動向を的確に把握し、自社の AI 戦略に反映させることが重要である。

3.4 AI リスク管理の歴史と各国の動向

AI リスク管理の枠組みは、AI 技術の進化とともに変遷してきた。初期の AI 技術は、限定的な用途に使用されていたため、リスク管理の必要性はさほど認識されていなかった。しかし、機械学習や深層学習の発展により AI が社会のあらゆる領域に浸透するにつれて、そのリスクに対する関心が高まった。

EU は、AI 規制の先駆けとして「AI Act」を策定し、リスクベースのアプローチを導入している。特に、高リスク AI システム（例：医療診断、信用スコアリングなど）に対しては厳格な規制を設け、透明性と安全性の確保を求めている。

米国では、NIST は「AI リスクマネジメントフレームワーク (AI RMF)」を策定し、企業や公共機関における AI リスクの管理を推奨している。

アジアでは、日本が「AI 事業者ガイドライン」を策定し、ソフトローの枠組みで企業に対して AI 活用に向けた指針を示している。また、中国は国家主導で AI 技術の開発を推進する一方で、監視技術の利用に関する規制強化の動きも見られる。

このように、各国はそれぞれの社会的・経済的な背景に基づき、独自の AI リスク管理政策を展開している。企業の経営者は、自社のビジネス環境に応じたリスク管理の枠組みを理解し、適切な対応策を講じる必要がある。

3.5 広島 AI プロセスに基づく取り組み

2023 年、日本は G7 議長国として生成 AI に係る国際的なルール形成を行う枠組みである「広島 AI プロセス」を立ち上げた。広島 AI プロセスでは、生成 AI の開発者、提供者、利用者といった全ての AI 関係者向けの「国際指針」と、高度な AI システ

ムを開発する組織向けの「国際行動規範」が取りまとめられた。

2024 年になり、G7 イタリア議長国のもとでは、生成 AI 開発における透明性及び説明責任を促進するため、「国際行動規範」の遵守状況を AI 開発企業等が自ら確認し報告するための手法（「報告枠組み」）を開発・導入するべく、G7 で議論を重ねられた。その結果、OECD の協力のもと、「報告枠組み」の基本的な運用方法及び質問票について合意し、2025 年 2 月より運用が開始された。この運用では、質問票に回答した AI 開発企業は、OECD のサイト上でリスト化されるとともに、その回答内容が公開される。質問票では、下記のように、開発した AI の安全性、透明性、信頼性といった観点からの質問で構成されている。

1. リスクの特定及び評価
2. リスクの管理及び情報セキュリティ
3. 高度な AI システムに関する透明性報告
4. 組織のガバナンス、インシデント管理及び透明性
5. コンテンツの認証及び来歴確認の仕組
6. AI の安全性向上や社会リスクの軽減に向けた研究及び投資
7. 人類と世界の利益の促進

3.6 日本の「AI 事業者ガイドライン」

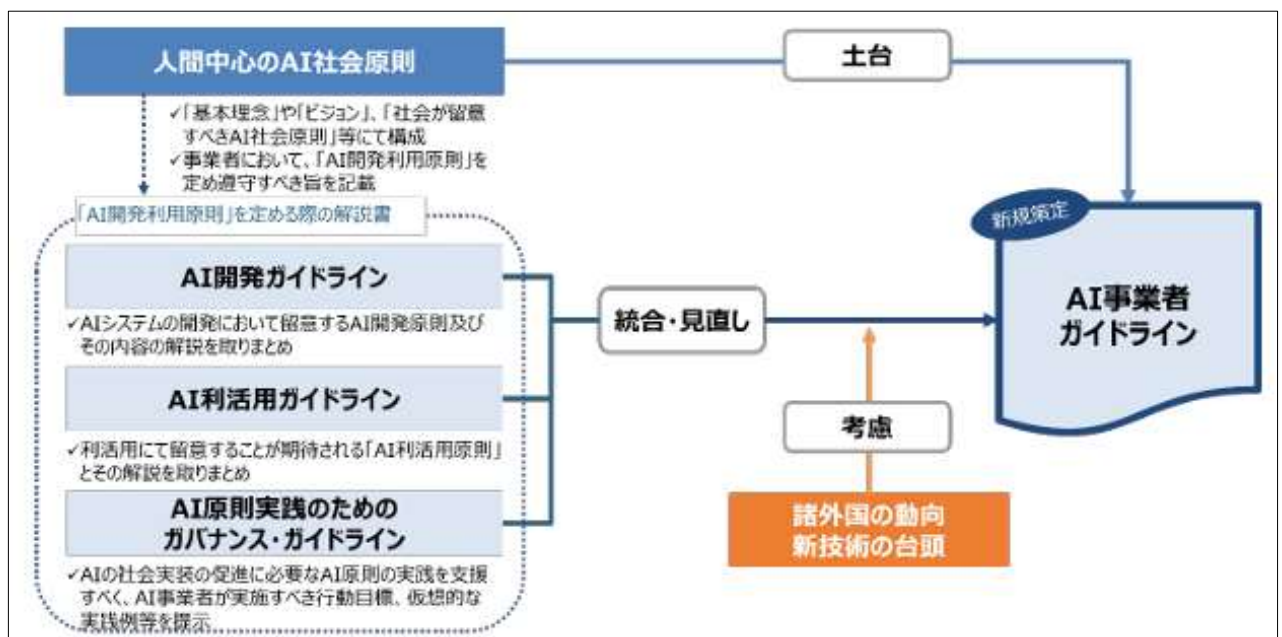
2025 年 3 月 28 日に公表された最新の「AI 事業者ガイドライン」は、これまで総務省主導で策定公表してきた「AI 開発ガイドライン」と「AI 利活用ガイドライン」及び経済産業省主導で策定公表してきた「AI 原則実践のためのガバナンス・ガイドライン」を統合・見直し、新たに策定された。

AI 事業者ガイドラインは、事業者が AI の安心安全な活用を行い、AI の便益を最大化するために重要な基本理念および指針を示した「本編」と、具体的な取り組みアプローチや AI ガバナンスの実践事例を示した「別添（付属資料）」から構成されている。

AI 事業者ガイドラインでは、AI の事業活動を担う立場として、「AI 開発者」「AI 提供者」「AI 利用者」の 3 つに大別して整理しており、各主体が連携しバリューチェーン全体で取り組むべき「共通の指針」を以下のように整理している。

1. 人間中心
2. 安全性
3. 公平性
4. プライバシー保護
5. セキュリティ確保
6. 透明性
7. アカウンタビリティ
8. 教育・リテラシー
9. 公正競争確保
10. イノベーション

図表 6：AI 事業者ガイドラインの策定



(出所)「AI 事業者ガイドライン（第 1.0 版）」総務省/経済産業省

<https://www.meti.go.jp/press/2024/04/20240419004/20240419004-1.pdf>

第4章 生成 AI の社会実装

ChatGPT が登場したことにより AI について、プログラム言語も含めて特別の学習をすることなく普通の日本語で指示すれば、それなりの回答が得られるという簡便性により中小企業も含めて多くの企業が使い始めた。当初は議事録の作成や要約、翻訳のツールぐらいであったが、業界ごとに生成 AI の潜在能力を活かして様々な試験的な取り組みが行われている。また、企業内の知識データベースと連携する、いわゆる RAG の活用により、その精度も飛躍的に向上している。しかし、いずれの利用例においても人と AI の役割分担の問題は存在している。AI にすべて任せるのではなく、最終的な判断を人間が行うなど、その業界や利用例に応じて適切に対応する必要がある。また、AI ガバナンスの観点からも各企業の姿勢が問われている。以下に代表的な事例を述べる。

議事録作成：

会議の音声データをテキストに変換し、議事録を作成するツールとして利用されている。会議内容をリアルタイムで記録し参加者に迅速に共有できる機能も評価されており、議事録の要約版も盛んに利用されている。

要約生成：

長い文章や会議内容を短い要約文に変換する機能は大変効果的である。大量の文書やメール、報告書を効率よく把握するためのツールとして利用されている。また大企業では、膨大な社内文書（例えば報告書や議事録）に対して生成 AI がデータベースを検索し、関連する文書の要約を生成することで効率的に情報収集でき、意思決定の飛躍的なスピード向上に効果が出ている。

翻訳：

国際的な企業では、生成 AI を利用して多言語間でのメールの翻訳を行うことで、コミュニケーションの速度と質を向上させている。また、契約書や技術文書などを迅速かつ正確に翻訳することで、事業推進に効果を上げている。

クリエイティブ作業：

詩や小説、映画のプロットなどの創作活動も試行されているが、音楽の作曲や美術作品の生成もかなり現実的にできるようになっている。自動でデザイン案を生成する

ツールや、広告バナーなどのデザインをアシストする生成 AI も導入されており、クリエイティブなプロセスを効率化している。また、ゲームのキャラクターやシナリオを生成できるし、ゲーム内の音楽やサウンドの生成など自動で楽曲や映像コンテンツを生成することで、制作コストを削減しながら新しいコンテンツ提供が可能になっている。

顧客対応：

顧客からの問い合わせに AI で自動応答するチャットボットが多くの企業で導入されている。これにより、顧客サポートの効率が大幅に向上し、人的リソースの節約が可能となっている。

コールセンターでは、生成 AI が企業内のナレッジベースと連携し、顧客からの問い合わせに対して瞬時に最適な回答を提供できるようになってきた。AI で対応すれば 24 時間対応が可能になり、また、新しいオペレーターへのアシスト機能としても活用できる。このように、AI と人との対応を分けて対応することにより、迅速かつ的確な対応が可能となり、顧客満足度の向上にもつながる。

データ分析と予測：

ビジネスの将来のトレンドや消費者の行動を予測するために、生成 AI を活用してデータ分析を行い、戦略的な意思決定を支援できる。また、こうしたデータ分析のため大量のデータから有用な情報を抽出し、ノイズを除去する作業や、予測モデルを作るための合成データの生産なども生成 AI で効率的に実施できる。

決算報告書や市場分析レポートの作成のため、大量のデータを解析し、レポートや洞察を自動で生成することも可能である。

マーケティング：

RAG 技術を使用して、顧客データベースや購買履歴を検索し、特定のセグメントに対する深い洞察が得られる。例えば、顧客の購買頻度や平均単価を分析し、高額商品に絞った広告を展開することや、特定の商品の購入者に対して、関連商品をおすすめするキャンペーンを実施することができる。これらの新しいキャンペーンやターゲット広告を設計し、広告コピー、ブログ記事、ニュースレターなどのマーケティングコンテンツを自動生成することで、効率的なコンテンツ制作が可能となる。

これ以外にも医療分野では、健診データや医療レポートの生成、疾患の予測モデル

の強化、ドラッグのデザインなどバイオテクノロジーの分野での応用が進められている。

教育分野では、自動的に教材や試験問題の作成、生徒ごとにパーソナライズされた学習サポートツールの提供などがある。

また、製造業では製品の詳細な技術情報や過去の問題解決事例をデータベースから自動で検索し、特定のトラブルに関する解決策をエンジニアに提供して迅速な問題解決が可能となる。

研究開発では、生成 AI が科学論文や特許データベースを検索し、新しい研究開発のインサイトを提供できるので、従来、膨大な時間を要していた文献レビューの作業が効率化され、研究の進展が加速する。

さらには、人事管理では社員のスキルマッチング、法務業務では過去の判例検索などにも適用されている。

このように、生成 AI の活用は様々な分野で進んでおり、多様な業界での導入事例が報告されている。表には代表的な業界ごとの社会実装事例を表示しているが、そのうち図表 7 で太字で示した 4 事例については詳細にそれぞれの導入プロセス、課題、留意点について述べる。

図表 7：国内の生成 AI の社会実装例

業種・業界	会社・組織名	事例概要
自治体	町田市 ¹	オンライン行政サービスの市民サポート
金融	七十七銀行 ²	生成 AI によるデータ分析自動化
製造業	日立製作所	製造設備保全業務への利用
小売業	セブンイレブン ³	商品企画プロセスの高速化
製薬業	中外製薬 ⁴	新薬開発のスピードアップを目指した共同研究
食品	日清食品ホールディングス ⁵	業務効率化と生産性向上

¹ オンラインによる行政手続きを簡単・快適にする「AI ナビゲーター」、町田市と協業し提供開始
<https://www.nttdata.com/global/ja/news/topics/2024/072500/>

² AI inside、七十七銀行の生成 AI 導入プロジェクトを共同実施、AI 実装コンサルティングチーム「InsideX」が経営層に伴走しデータドリブン経営を推進 <https://inside.ai/news/2023/11/10/aiinside-77bank/>

³ セブン、商品企画に生成 AI <https://www.nikkei.com/article/DGKKZO76132610V11C23A1EAC000/>

⁴ 中外製薬、ソフトバンク、SB Intuitions の 3 社が生成 AI の活用で臨床開発業務を革新し、新薬開発のスピードアップを目指す共同研究に向けた基本合意を締結 ～製薬・ヘルスケア業界における先進的なサービスの創出を目指して～ https://www.chugai-pharm.co.jp/news/detail/20250130153000_1461.html

⁵ 日清食品グループにおける生成 AI 活用の現在地
https://www.meti.go.jp/shingikai/mono_info_service/digital_jinzai/pdf/010_02_00.pdf

教育	学研 ⁶	高等学校向けの志望理由書 AI 添削サービス
IT	パルコ ⁷	ファッション広告制作
サービス業	やさしい手 ⁸	介護記録、ケアプラン情報の自動生成
建設	西松建設 ⁹	高精度な建設コストの予測

4.1 町田市 AI ナビゲーター

現在、多くの自治体で生成 AI の導入が行われている。そのような自治体の業務効率化や住民サービス向上の動向で注目すべき先進的なものとしては、たとえば東京都、神戸市、横須賀市、町田市など全国で多数の例がある。ここでは町田市に焦点を当てて紹介する。

NTT データと町田市は 2023 年 5 月、「ジェネレーティブ AI の利活用に関わる連携協定」を締結し、町田市と連携して市民向けオンラインサービスの向上や市役所の業務改革・改善、AI 利活用ガイドラインの策定等、スマートシティ実現に向けた取り組みを進めている。さらにこれらの活動を通して、培ったノウハウを踏まえて全国の自治体へ展開し、行政サービスのさらなる利便性向上と、行政業務の効率化・高度化を進めている。本事例は、この協定に基づく協業の成果である。

事例における適用先の課題

町田市では、多数の行政手続きがオンライン化されている一方で、デジタルに不慣れた市民が簡単かつ快適に手続きを行える環境の整備が求められていた。

課題達成にあたってのアイデア

生成 AI と 3D アバターを組み合わせた「AI ナビゲーター」を開発し、オンライン上での案内役として活用することで、市民が自然な会話を通じて目的の手続きに容易にアクセスできる仕組みを構築した。

⁶ 「志望理由書」を生成 AI で添削、学研 HD が高校向けデジタル教材を販売開始

<https://xtech.nikkei.com/atcl/nxt/news/24/02331/>

⁷ パルコ初の生成 AI 広告「HAPPY HOLIDAYS キャンペーン」が公開！グラフィック・ムービー・ナレーション・音楽まで全て生成 AI にて制作！ <https://prtimes.jp/main/html/rd/p/000002679.000003639.html>

⁸ やさしい手、在宅介護サービスの介護記録などの文書作成業務に生成 AI を活用。Amazon Bedrock を適用し、約 30% の工数削減と文書の品質向上を同時に実現。 <https://aws.amazon.com/jp/solutions/case-studies/yasashii-te/>

⁹ 建設業界の物価変動を経済予測 AI で先読み <https://service.xenobrain.jp/article/posts/usecase-nishimatsu>

適用した技術

- Microsoft Azure Open AI Service
- 3D アバターによる視覚的な案内

導入に向けたプロセス

1. 2023 年 5 月、NTT データと町田市は「ジェネレーティブ AI の利活用に係る連携協定」を締結。
2. 町田市のデジタルサービスポータルサイト「まちドア」を 2024 年 4 月に公開。
3. 「まちドア」の拡張として、2024 年 7 月から「AI ナビゲーター」の提供を開始。

成果・効果

「AI ナビゲーター」の導入により、

- 市民はチャットや音声入力を通じて、複雑な操作なしに目的のオンライン手続きにアクセス可能となった。
- 親しみやすい 3D アバターとの自然な会話を通じて、オンライン手続きの利便性が向上した。
- 利用者のプライバシー保護にも配慮した仕組みを採用している。

上記の取り組みにより、オンライン手続きの利便性が向上し、市民生活の質の向上に寄与している。

4.2 七十七銀行 生成 AI によるデータ分析自動化

金融機関では顧客対応などに加えて、販売する商品の営業実績を分析し、新たな営業施策の検討や商品開発を行うことが重要となる。ここでは、生成 AI を用いたデータ分析業務の例として、七十七銀行での事例を取り上げる。

事例における適用先の課題

営業実績や商品別の販売データ分析は、行員が手作業でプログラミングの上、レポート化しており、データ分析に時間がかかる上に作業負担が大きいという課題があっ

た。また、分析に時間を取られることで、市場環境への迅速な対応や新たな施策立案が進まなかった。

課題達成にあたってのアイデア

- PDF や HTML などの非構造化データを AI に入力し、AI が業務に応じた専用フォーマットに転記する
- 生成 AI が行員の代わりに分析用のプログラムコードを生成し、表やグラフにて可視化
- 分析結果のレビュー文書を作成

適用した技術

- 非構造化データを対象とした文字認識 AI
- 予測 AI

導入に向けたプロセス

1. 既存業務資料の自動読み取り、転記システムの試作
2. 経営層も交えて本格プロジェクトを発足、コンサルティングチームと共同で戦略策定から開発・運用までを一気通貫で実施

成果・効果

- 生成 AI の活用による生産性向上を通じて、職員がより創造的な業務を担うことで、挑戦的な企業文化および職員のエンゲージメント向上の実現

4.3 日立製作所 製造設備保全業務への利用

製造業においては、工場機械・設備の保守・保全業務は人手で行われることが多く、多くのコストがかかる。また、これらの業務実施にあたっては、熟練者の暗黙知に依存している場面も多く、ノウハウの適切な承継も課題の一つとなっている。ここでは、日立製作所における製造設備の保全業務に生成 AI を活用した事例を取り上げる。

事例における適用先の課題

- 製造設備の保全業務における異常原因の推定と対策立案は、熟練者の暗黙知に依存しており、業務の効率化が課題となっていた。

課題達成にあたってのアイデア

- 設計図面を OT(Operational Technology)¹⁰データ、分析プロセスを OT スキルとして学習し、現場作業を支援

適用した技術

- ナレッジグラフ
- 保守支援エージェント

導入に向けたプロセス

1. 設計図面や製品ドキュメントなどの OT データをデータベース化
2. プロンプトで生成された分析プロセスと OT データとの関係により原因と対策を推論

成果・効果

- 障害情報に対する原因の推定と対策において精度 90%、10 秒以内で推定することが可能となった。

4.4 セブンイレブン 生成 AI 商品企画システム

小売業では一般消費者の嗜好、流行を的確に捉え、タイムリーに商品を展開していくことが求められる。一方でこのような商品企画においては、売上データの分析や市場分析に時間を要し、適切なタイミングで商品を市場に導入することが難しい。ここでは商品企画において生成 AI を活用し、大幅な生産性の改善を実現したセブンイレブンの事例を取り上げる。

¹⁰ OT：工場や発電所などに使われる、物理的なシステムや設備を最適に動かすための制御・運用技術の総称

事例における適用先の課題

従来の商品企画では、流行の調査やアイデア出し、社内会議による承認までに長い時間を要していた。そのため、市場の嗜好変化が激しい中で、企画サイクルが長期化していることでタイムリーな商品投入が難しくなっていた。また、企画段階では消費者アンケートの分析や売上データの解析なども膨大になり、担当者の負担が大きくなっていた。

課題達成にあたってのアイデア

- 生成 AI により情報収集、分析、企画立案までを一気通貫で支援する。

適用した技術

- SNS のコメントや販売動向の分析 AI
- 消費者の声の分析結果をもとにした商品の文章や画像の作成 AI

導入に向けたプロセス

1. 本部のシステム部門を中心に試行導入し、管理職クラスから AI 活用を開始
2. 1 年後、一般社員にも利用範囲を広げる本格展開を実施

成果・効果

- 商品企画の時間が 10 ヶ月から 1 ヶ月に短縮される見込み
- 消費者トレンドとニーズの迅速な把握により、セブンイレブンの商品ライン全体の拡大も期待できる

第5章 生成 AI 活用の戦略：日本への提言

ここまで述べてきたように、LLM を含む生成 AI の発展が著しく、企業の生産性向上や新たな価値創出に大きく貢献する可能性が指摘されている。しかし、日本企業は生成 AI の活用が欧米に比べて遅れているのが現状である。本提言では、企業が生成 AI を積極的に活用するための具体策と、適切なガバナンスのあり方について述べる。

5.1 企業経営者に向けた提言

提言 1：生成 AI 活用なくして組織の成長は無い

生成 AI の最大の特徴は、その汎用性と柔軟性にある。文章作成、データ分析、顧客対応、自動翻訳など、さまざまな業務領域で活用が可能である。しかし、現状では「どのように活用すればよいかわからない」「品質が不安」といった理由で導入を躊躇する企業が多く見られる。実際に株式会社情報通信総合研究所が 2024 年 11 月に公開した調査結果¹¹によれば、最も導入が進んでいる情報通信業でも導入率は 35.1%に過ぎず、卸売業、小売業や各種サービス業では導入率が 10%前後であることが示されている。このデータは、企業全体での生成 AI 導入がまだ進んでいないことを物語っている。

一方で、実際に生成 AI を自社の業務で活用している企業は、試行錯誤を重ねながら具体的な活用方法を見出し、業務の効率化、さらには新たなビジネスチャンスを得ている。このような企業が存在する一方で、導入を躊躇している企業はトレンドに乗り遅れ、競争力を失いかねない状況と言える。

このように、生成 AI の活用は企業経営者にとって死活問題である。今や生成 AI はあらゆる業種・業界の業務で用いられている。もはや生成 AI の活用は単なる業務改善の延長線上にあるものではない。2 章で紹介した AI エージェントに代表されるように、生成 AI はこれまでの仕事のあり方そのもののみならず、組織の人員構成や階層構造を変えてしまうほどのインパクトを持っている。

¹¹ 【報道発表】企業における生成 AI 活用の格差浮き彫りに ―規模別・業種別の利用状況・課題と今後の展望― <https://www.icr.co.jp/publicity/5135.html>

これからの企業経営者は、急速に成長を続ける AI 技術をキャッチアップし続け、組織に適用していくことが極めて重要である。トレンドに乗り遅れることは組織の成長を阻害し、最悪の場合、組織の存続に関わる問題となる可能性がある。生成 AI の活用は、持続可能な成長を目指す上で必須であり、企業経営者は迅速な導入と活用を積極的に進めるべきである。

提言 2：AI をまず試してみる

様々なサービス提供者が安価に AI サービスを提供しており、企業における利用のハードルは必ずしも高くない。AI 技術を効果的に活用するためには、まずは小規模な範囲から試してみることが重要である。例えば 4 章で紹介したような、文章検索・要約やマーケティングなどの特定業務の一部に AI を試行的に適用し、その有効性を検証することで、技術の利点を具体的に理解することができる。

導入初期には、社員のスキルアップを図るため社内セミナーやワークショップを開催することや、イントラネットや社内 Web サイトを利用して社内の好事例を発信・共有していくことが重要である。さらに、AI 導入の効果を測定するために、生産性や創造性向上の具体的指標（KPI）を設定する。例えば、業務効率化の度合い、作業時間の短縮、顧客満足度を測定し、導入効果を評価する。

しかし、この試行導入を進める際には、既存の業務フローやビジネスプロセスに変更を加えることが必要になることが多い。多くの場合、新しい技術の導入には既存のプロセスやアプローチを見直し、最適化することが不可欠である。そのため、AI の導入を成功させるためには、単なる技術導入に留まらず、組織全体の文化を変えていく必要がある。

たとえば、日本では品質を重視する文化が強く、ミスを許容しないため、AI のような不確実性の高い技術の導入には慎重になりがちな側面がある。また、既存の業務プロセスの変更にも抵抗感があることが多く、新しい技術の導入が進みにくい。このような文化的要因を克服するには、変化を恐れず、新しい取り組みへのチャレンジを積極的に評価する文化を醸成することが必要である。

経営層自らがこのようなチャレンジに積極的に関与し、AI を試し、その有用性を体

感し、組織全体への普及をリードする役割を果たすべきである。また、企業の中では自主的に AI の活用を進めている社員もいる。そのような社員の知見を共有し、成功例、失敗例を含めてノウハウを蓄積・展開していくことも重要である。

このように、AI を効果的に導入し活用するためには、技術的なスキルだけでなく、企業文化の変革と組織としての適応力が不可欠である。変化を受け入れ、持続可能な成長を目指すことが、組織全体の競争力を強化し、大きな成果をもたらすであろう。

提言 3：AI ガバナンスを確立する

AI を効果的に活用するためには、単に技術的な活用スキルを高めるだけではなく、企業全体として適切なガバナンスの確立が不可欠である。企業としての AI 活用の指針を明確に定め、それを社内外に公表することは、企業の透明性を高め、社会的な信頼性を向上させる上で重要な要素となる。

まず、企業が AI を導入する際には、透明性の確保、倫理的な責任、データプライバシーの保護といった観点から、包括的なガバナンスの枠組みを構築する必要がある。これには、AI 活用に関する企業の基本方針の策定や具体的な運用ガイドラインの設計が含まれる。特に、AI が生成する情報の正確性やバイアスに関するリスクを考慮し、それに対する適切な対策を講じることが求められる。

AI を開発・提供する組織だけでなく、AI を利用する企業や組織、例えば医療機関なども AI ガバナンスを確立することが重要である。利用する側としても、AI の透明性や倫理的な利用を保証するためのポリシーを策定し、運用ガイドラインを整備することが求められる。そのため、ガバナンス体制の構築は、AI が生成する情報の検証、リスク評価、適正利用の基準設定など広範にわたる。

さらに、ガバナンスの実効性を高めるためには、単なるガイドライン策定に留まらず、運用ポリシーの外部公開や、従業員への継続的なトレーニングの実施が不可欠である。従業員が AI ツールを適切に活用し、そのリスクを理解できるようにするために、AI 倫理やデータセキュリティに関する教育プログラムを整備し、定期的にアップデートしていく必要がある。

また、AI ガバナンスの確立には、社内のステークホルダーだけでなく、外部の専門

家や規制当局との連携も重要である。AI の急速な進化に伴い、法規制や業界標準も変化していくため、企業はこれらの動向を常に把握し、ガバナンス体制を柔軟に調整することが求められる。さらに、第三者機関による監査や評価を受けることで、ガバナンスの客観性を担保し、社会的な信頼を一層高めることができる。

このように、AI の適正な活用を推進するためには、技術の導入と同時に、包括的なガバナンスの枠組みを整備し、それを継続的に運用・改善していくことが重要である。AI ガバナンスの確立は、単なるリスク管理の手段ではなく、企業の競争力強化と持続可能な成長を支える基盤となるべきものである。

3.6 節で「AI 事業者ガイドライン」を紹介したが、AI 開発者や AI 提供者は IT 関連企業が多く、これまでも各企業の AI ガバナンスの方針を公表してきている。しかし、AI 利用者になると公表している企業は少ない。AI が人間の知的作業の代替であるとするれば、全ての企業は少なくとも AI 利用者にはなる。

企業は毎年統合報告書で、SDGs や人権などサステナブルな成長に資する取り組みを公表している。これと同じように、各企業は AI ガバナンスのフィロソフィーとそれを実現する組織的な取り組みを構築し、統合報告書などで外部に公表していくことが求められている。

5.2 AI・LLM 活用に向けた政府関係者への提言

提言 4：政府・自治体での生成 AI 活用推進

生成 AI 技術の発展は目覚ましく、政府・自治体業務でもその活用は必須である。生成 AI の活用によって、政府・自治体業務の効率化が期待できる。具体的には、生成 AI は大量のデータを迅速かつ正確に処理できるため、書類の作成や事務処理、データ分析などの業務を効率化する。また、定型的な問い合わせ対応や手続き説明などを AI が自動で行うことで、職員の負担を軽減し、より重要な業務にリソースを割くことが可能になる。さらに、生成 AI の活用により、住民からの問い合わせへの迅速な対応が可能となり、住民満足度の向上が期待できる。また、データ分析に基づいた政策立案が可能となり、より効果的で効率的な行政サービスが実現できる。

地方自治体では、少子高齢化や人口減少に伴い、行政サービスの担い手が減少している現状がある。この問題に対して、生成 AI の活用が有効な解決策となる。生成 AI を活用することで、限られた人員で大規模な業務を効率的に実施することが可能になる。また、遠隔地でもオンラインでの対応が可能となり、地域格差の是正にも寄与する。さらに、専門的な知識を要する問題にも生成 AI が対応し、住民へのサービスの質を維持・向上することができる。

生成 AI の導入に向けては、まず自治体の業務に適した生成 AI 技術を選定し、トライアルとして一部業務に試験導入することが重要である。特に多くの規程や文書類を扱う自治体業務では、2 章で紹介した GraphRAG の技術により、ベテラン職員でなくとも関連する情報を容易に取得し、円滑に業務実施することが期待できる。技術適用と並行して、生成 AI を活用するための職員のスキルアップを図るとともに、適切な IT インフラを整備する。そして、生成 AI の活用による利点を住民に周知し、理解と協力を得ることで、円滑な導入を目指す。

さらに、生成 AI の導入は中央政府にも大きな期待が寄せられている。中央政府が先陣を切って新技術を取り入れることで、全国的な普及と理解が進み、自治体の導入がよりスムーズに進むことが期待される。生成 AI の政府・自治体業務への積極的な適用は、行政サービスの質を向上させるだけでなく、地方自治体における人手不足という深刻な問題の解決にも貢献する。技術と人間の協力によって、持続可能な行政運営を目指し、地域社会の発展に寄与することが期待される。

提言 5：AI ネイティブ人材の育成

デジタルデバイドを超えた「AI 格差」の拡大を防ぐためには、AI を直感的に理解し、活用できる「AI ネイティブ」な人材の育成が急務である。特にデジタル環境に親しんだ若年層を対象に、AI リテラシーの向上を図り、社会全体の適応力を高めることが重要となる。

そのために、義務教育段階から AI 活用・データ活用に関する体系的なカリキュラムを導入すべきである。具体的には、プログラミング教育を発展させ、機械学習の基礎やデータ分析の基礎、さらには倫理的な側面までを網羅した教育プログラムを構築する必要がある。また、教材や指導者の不足を解消するため、教育機関と企業、研究機関が連携し、最新の技術や知見を取り入れた実践的な学習環境を整備することも

求められる。

さらに、AI 技術を単なる「ツール」としてではなく、創造的な問題解決の手段として活用できるよう、産業界との連携を強化し、実践的なプロジェクト学習の機会を提供することが不可欠である。例えば、企業や自治体と協力し、実際の社会課題を解決するプロジェクト型学習（PBL：Project-Based Learning）を推進することで、学生が AI の実践的な活用能力を養うことができる。また、スタートアップや先進企業の支援を受けながら、若者自身が AI を活用した新たなビジネスモデルやサービスを創出する機会を増やすことも重要である。

一方で、AI 技術を効果的に活用するためには、基盤となる「読み書きそろばん」と称される基本的なスキルも不可欠である。具体的には、リテラシー（読み書き能力）、数的能力（計算力）、そして論理的思考力の三つのスキルが挙げられる。リテラシーは AI に関する情報やデータを正確に理解し、解釈する能力を指し、文書作成や情報収集力も含まれる。数的能力はデータ分析や数理的な思考力を養うために必要である。論理的思考力は複雑な問題を分解し、論理的にアプローチする能力であり、プログラミングやデータサイエンスといった AI 技術を扱う上で欠かせないものである。加えて、技術の活用にあたっては、単に技術の適用だけにとどまらず、倫理的な観点や適用によって生じるリスクとその対応についても配慮が必要となる。AI 技術を活用するためには、これらの基盤となるスキルと高い倫理観を兼ね備えた人材育成が不可欠であり、AI 技術の習得に偏重したものにならないようにすべきである。

AI ネイティブ人材の育成は、日本の国際競争力を左右する極めて重要な課題である。これを実現するためには、教育政策、産業政策、労働政策が連携し、包括的な支援体制を構築する必要がある。政府は、AI 教育の普及に向けた制度整備と資金投入を行うとともに、民間企業や大学との協力を推進し、社会全体での AI 人材育成を加速させるべきである。

提言 6：社会発展を阻害しないバランスの取れた規制

AI の法制化が進む中で、社会発展を阻害しないバランスの取れた規制が求められる。そのためには、AI のリスクを抑制しつつ、イノベーションの余地を確保する規制設計と、透明性・公平性を確保するための基準設定が不可欠である。

日本では、2025 年 6 月 4 日、「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」が公布された。この法律は、AI 技術の研究開発と社会実装を包括的に推進するための枠組みを定めており、イノベーションの促進と規制のあり方が特に重要な論点となっている。以下では、それらの点について言及する。

AI 技術は急速に進化しており、その適用範囲も広がっている。このような技術の発展に伴い、規制もまた迅速、かつ柔軟に対応する必要がある。画一的な規制ではなく、リスクベース・アプローチを採用することで技術の進化を阻害しない枠組みを構築することが求められる。具体的には、技術の進化にあわせて規制もアジャイルに対応することが重要であり、AI の使用目的や影響の大きさに応じた段階的な規制を導入することで、過度な制約を回避しながらもリスクの高い用途には適切な監視と制御を行うことが可能となる。

また、規制の透明性と公平性を担保するためには、行政機関や立法府だけでなく、産業界、学術機関、市民社会など多様なステークホルダーの意見を反映させるプロセスが必要である。これにより、規制が特定の業界に偏ることなく、社会全体の利益に資する形で設計されることが期待される。

加えて、AI 技術の導入を促進するためには、企業や開発者が法的リスクを適切に理解し、コンプライアンスを確保しやすい環境を整備することも重要である。そのため、過度な罰則や厳格な制限を課すのではなく、企業の自主規制を促す形でのガイドラインやソフトローの活用が有効である。特に、技術の進化に伴い規制が陳腐化するリスクを考慮すると、法制度よりも機動的に更新可能な枠組みが適している。

さらに、国際的な視点も欠かせない。AI の発展は一国の政策だけで完結するものではなく、国際的な規制の枠組みや協調が不可欠である。そのため、日本としても各国の動向を注視しつつ、国際標準の形成に積極的に関与し、競争力を確保する戦略を取るべきである。

2024 年 1 月

趣意書
「第二次人工知能(LLM：生成 AI)委員会」
～生成 AI 旋風に企業経営者はどう向き合うか～

一般社団法人 日本経済調査協議会
第二次人工知能(LLM：生成 AI)委員会
委員長 岩本 敏男
主 査 須藤 修

2022 年 11 月にオープン AI によって生成自然言語処理大規模モデル GPT-3.5 が公開され、今年には GPT-4 がマイクロソフトの検索エンジン Bing に組み込まれた。Bing は Dall-e による生成画像を自然言語のプロンプトで書くことができるようになった。

一方、グーグルは生成自然言語大規模データモデル Bard をリリースし、2023 年 5 月 10 日には PaLM2 の発表、日本語対応と韓国語対応も発表された。

これらの一連の動きは、今日に至るまで、とてつもなく大きなインパクトを人々に与えている。これまでの与えられた目的に対して与えられた手段を用いて最適化する優れたツール（最適化ツール）としての AI は、もはや大きく変わり始めている。グーグル、オープン AI、マイクロソフト、メタ、IBM、スタンフォード大学、マサチューセッツ工科大学といった米国を代表する先進企業や研究教育機関では 2018 年頃から大きな変化が準備されてきた^{※1}。

その一つが、GPT (Generative Pre-trained Transformer) であり、その核をなすのがトランスフォーマー (Transformer) の開発だった^{※2}。

トランスフォーマーは、大規模なデータセットによって自然言語を処理し、文脈を理解し、確率論的に言葉を生成する自然言語処理機構である。

この処理を行うためには、大規模データベースと高速計算機資源が必要になる。オープン AI が 2020 年春に公開した GPT-3 は、パラメータ数が 1,750 億もある。なお、2022 年 11 月に登場した改良版 GPT-3.5 のパラメータ数も 1,750 億とされている。さらに驚くべきは、グーグルのトランスフォーマー PaLM は 5,400 億のパラメータ数、数理的推論を行う Minerva は 5,400 億のパラメータ数を有する。ちなみにグーグルは 2023 年 5 月に開催された Google I/O2023 においてファインチューニングを行った PaLM2

を公表した。また数理的推論や化学式、DNA などの分析を行うマルチモーダル AI の研究開発である Gemini の概要を公表している^{※3}

これまでの AI は特定の目的に最適化することを主としており、人間の創造的行為や戦略的行為を脅かすものではなかったといってもよい。しかし、生成 AI・マルチモーダル AI は、創造的行為や戦略的行為に大きく関与するようになってきているといってもよいだろう。

周知のように、2022 年 11 月に公開されたオープン AI による ChatGPT は、たった 2 か月で登録ユーザーが 1 億人を突破、現在のアクティブユーザーは約 10 億人と言われ、まさに全世界に急速に普及し、様々な影響を行政、企業、研究機関などに大きな影響を与えている。

我が国でも多くの企業が競合他社に遅れを取るまいと注力しており、また、政府主導で有識者による「AI 戦略会議」が立ち上がる等、官民挙げての対応がスピーディに進んでいる。

生成 AI の市場規模は、2027 年時点で約 17 兆円、年平均成長率は 66%という驚異的な伸びを示すとの試算^{※4}や、世界の GDP は 10 年間で 7%持ち上げられるとの試算^{※5}もあり、顕著な経済効果が期待されている。加えて、上手く活用出来れば、少子高齢化による労働力不足や低位な生産性という、我が国の抱える諸課題を一気に解消し得るとの指摘も挙げられている。

企業経営者の視点において、多くの場合、導入の目的は「非効率業務の削減による生産性の向上」が挙げられよう。稟議書作成支援、金融リポート要約、手続・マニュアル照会等で活用を始めたメガバンクがその一例である。

一方、自社のサービスに搭載して高付加価値化を図る動きも見られる。具体的には、新薬開発、市場調査・分析等マーケティング、法律相談サービス等における活用が報道されている。

このように導入に積極的な企業がいる一方で、導入は決めたがどう取り組むか検討中の企業、或いは、そもそも導入に二の足を踏んでいる企業もあって、その温度感、スピード感は様々である。慎重なスタンスをとる企業は、その理由として、以下のような未だ明確とは言えない将来展望や AI 導入に伴うリスクに関する不安があると考えているからではないだろうか。

- ①現在の ChatGPT、Bard などの生成 AI の台頭は、多くの識者、研究者が語るように一時的な現象ではなく、文明を画するような不可逆的な動きの端緒となる可能性が高

いといわれている。しかし、現時点では、その具体的な展望は未だ不透明である。

- ②プライバシーや著作権、特許などの知的財産権、人権の侵害、情報漏洩等のリスク、更には透明性や信頼性の確保への懸念がいまのところ払拭できてはいない。
- ③短期的には失業や雇用のミスマッチを生じさせる可能性^{※6}があること。また、中長期的には、生成 AI への既存業務の置換えにより、従業員の職能にどのような変化が伴うのかわからない、との懸念。
- ④今後、プロダクトやサービスに AI が取り入れられた場合のビジネスモデルを企業が十分には想定できないこと。特に効率化ではなく、自社プロダクトやサービスへの AI 実装による高付加価値化(→稼ぐ力への AI 活用)がいまのところ想定できていないこと。
- ⑤戦争・プロパガンダへの利用、詐欺等犯罪の巧妙化等、企業の自助努力では対応できない課題やリスクが払拭できていないこと。
- ⑥OECD などの国際機関でも検討が開始された AGI (汎用 AI) への進化について^{※7}、将来、人間を超えた能力を有する AI と人間の新たな関係を構想しなければならないが、まだその具体像が予測できない。

そこで本委員会では、生成 AI という新技術との向き合い方について、企業経営者目線で議論し、技術的な論点は最小限に留めつつ、上記の懸念を解きほぐし、どのように対処すべきかの示唆を企業経営者に提示することを目的としたい。

既に生成 AI を活用し、一定の成果を得ている企業の先行事例を紹介し、併せて導入・活用に至った判断材料や「導入計画～適用業務選定～社内教育～生成コンテンツのチェック体制構築」等、一連の導入プロトコルを明示することにより、二の足を踏んでいる企業経営者を後押ししたいと考える。

また、個別の企業努力を超えた対応が求められる懸念点(特に上記⑤⑥)に対し、日経調スタイル、即ち、産官学レベルで議論・研究することにより、国に対しても有効な示唆を導出できるのではないかと期待している。

もし企業レベルで導入・活用に躊躇していれば、最先端の技術にアクセスできず、競合他社との埋めようのない競争力格差を生むリスクがある。また国民経済レベルでは特許権の競争に出遅れ、先進的 AI の研究開発が進んでいる米・英・中などの後塵を拝すことを懸念する。

本委員会では、5 年後、10 年後に起きていそうなことを盛り込み、「明るい日本経済の将来像」に向けて今からどうすれば良いのかをバックキャスティングで議論し、デ

イストピア的な不安を払拭することも当委員会の目的の 1 つに加えたいと願う。

以 上

-
- ※¹ Bommasani, Rishi et.al. [2022] On the Opportunities and Risks of Foundation Models, *arXiv:2108.07258v3*, Stanford University Human-Centered Institute for Artificial Intelligence [2023] *Artificial Intelligence Index Report 2023*, Stanford University を参照。
- ※² Heaven, Will Douglas [2023] ChatGPT is everywhere. Here's where it came from, *MIT Technology Review*, Feb. 8th 2023 を参照。
- ※³ Google I/O 2023 (<https://io.google/2023/program/396cd2d5-9fe1-4725-a3dc-c01bb2e2f38a/intl>) を 2023 年 5 月 11 日に視聴。
- ※⁴ 2023 年 8 月 3 日 日本経済新聞「きょうのことば」(BCG 試算)
- ※⁵ 2023 年 3 月 28 日 Bloomberg 「AI は米生産性の急上昇促す、世界成長も押し上げへ」(ゴールドマンサックス証券試算)
- ※⁶ 厚生労働省「AI 等の新たなテクノロジーが雇用に与える影響について」2023 年
- ※⁷ OECD は 2023 年 7 月より、汎用 AI を含む先進 AI の研究・開発・普及に関する専門家グループによる検討を開始している (The OECD Expert Group on AI Futures)。主査である須藤修はその構成員の一人である。

講師講演録(所属・役職は講演当時)

【<https://www.nikkeicho.or.jp/> に掲載。一部非公開】

1. 生成・マルチモーダル AI の展望とリスク～Advanced AI はどこへ向かうのか
中央大学国際情報学部 教授 須藤修主査
2. AI は経営に何を及ぼすか～データの時代が到来どう乗り越える～
NTT データグループ 相談役 岩本敏男委員長
3. 東京海上における LLM の活用について
東京海上日動火災保険 理事 IT 企画部 部長 村野剛太委員
4. 生成 AI が拓く AI の未来～AI は人間社会をどのように変容させていくのか～
慶應義塾大学理工学部 教授／共生知能創発社会研究センター センター長
栗原聡委員
5. 日立グループにおける生成 AI の取り組み - 社会実装に向けた課題と取り組み -
日立製作所 代表執行役 執行役副社長 デジタルシステム&サービス統括本部長
徳永俊昭委員
6. IHI における LLM の活用～生成 AI がさらに進化すると…IHI は無くなる！？～
IHI 代表取締役 副社長執行役員 土田剛委員
7. AI のガバナンス
東京大学国際高等研究所東京カレッジ 准教授 江間有沙委員
8. 医学領域における AI 活用の未来
東海大学医学部 講師／ブリガムアンドウィメンズホスピタル循環器内科 講師
後藤信一氏
9. 時代局面を考える
慶應義塾大学環境情報学部 教授／LINEヤフー シニアストラテジスト
安宅和人委員
10. データで見る生成 AI の活用と社会的影響：現状、課題、そして未来への展望
国際大学グローバル・コミュニケーション・センター 准教授 山口真一委員
11. ガバナンス・イノベーション ～ AI 時代の法と経済学
東京大学未来ビジョン研究センター 客員教授 西山圭太委員

本資料は、信頼でき得ると考えられる情報・データに基づき作成しておりますが、当会はその正確性・安全性を保証するものではありません。これらの情報を利用することで直接・間接的に生じた損失に対し、当会および本情報提供者は一切の責任を負いません。

本資料に掲載された内容は、事前の通知を行うことなく変更・削除されることがありますが、それによって生じた如何なるトラブル・損害に対しても責任を負うものではありません。

本資料を利用する際は出典を記載願います。編集・加工した情報を、当会が作成したかのような態様で公表・利用することはご遠慮下さい。本資料の全部または一部を無断で複製することは著作権法上での例外を除き禁じられています。

[禁無断転載]

2025 年 8 月 6 日発行

「生成 AI 旋風に企業経営者はどう向き合うか」

一般社団法人 日本経済調査協議会
専務理事 小田 寛一

〒106-0047
東京都港区南麻布 5-2-32
興和広尾ビル2階
電話 03-3442-9400
FAX 03-3442-9402
<https://www.nikkeicho.or.jp>

[非 売 品]

印刷／日本経済調査協議会