

一般財団法人 社会文化研究センター
助成事業報告書

生成 AI が変える社会と教育

令和 7 年 3 月

金沢学院大学

大貫裕二

長谷川秀司

濱屋 敏

後藤弘光

目次

第1部 生成AIが変える社会と教育	6
1. はじめに	6
2. 歴史	6
2.1 AIの歴史：3次にわたるブームと冬の時代	6
2.2 ニューラルネットワークは「もうひとつのAI」	8
2.3 ニューラルネットワークの歴史	8
3. 大規模言語モデル	11
3.1 生成AIとはなにか、大規模言語モデルと画像・音楽・動画生成	11
3.2 大規模言語モデルの仕組み	11
3.3 用途に合わせた大規模言語モデルの開発	12
3.4 推論能力の向上	14
4. 大規模言語モデルの特徴	14
4.1 大規模言語モデルは間違えるか	14
4.2 錯覚（ハルシネーション）	15
4.3 仕事で使えるようになった大規模言語モデル	16
4.4 大規模言語モデルを使った検索から検索システムとの連携へ	16
5. 大規模言語モデルと業務	17
5.1 大規模言語モデルの機能の活用 メール返信	17
5.2 大規模言語モデルによるアイデア出し	18
5.3 大規模言語モデル時代のプログラミング	19
5.4 大規模言語モデルによるプログラミングやシステム開発の補助	20
5.5 大規模言語モデルの活用による業務効率の向上	21
5.6 過大な期待と過少な期待	22
5.7 生成AIと人間の協業	22
6. 大規模言語モデルの教育への利用	23
6.1 英会話練習	23
6.2 教育利用に関する基本姿勢とガイドライン	24
6.3 金沢学院大学における大規模言語モデルの教育	26

7. 大規模言語モデルに関するリテラシー	27
7.1 誤りと偏り	28
7.2 社会のバイアス	28
7.3 プライバシー情報や機密情報の管理	28
7.4 著作権の問題.....	29
8. 大規模言語モデル利用の倫理と課題.....	30
8.1 フェイクニュースや詐欺への利用の危険性.....	30
8.2 不適切な知識の拡散.....	30
8.3 ガードレールの設定基準	31
8.4 国際的課題	31
8.5 計算機資源と電力消費	32
9. AI と今後の社会の変化.....	33
 第2部 企業における生成 AI 活用の実態調査.....	 35
1. はじめに	36
2. アンケート調査の概要	36
2.1 調査の目的、調査項目など	36
2.2 回答者の属性	37
3. 調査の分析結果.....	38
3.1 単純集計結果	38
3.2 業種による利用状況の違い.....	41
3.3 企業規模による利用状況の違い	43
3.4 地域による利用状況の違い.....	44
3.5 企業風土による違い.....	45
3.6 懸念事項・課題に関するクロス分析	47
4. 調査のまとめと含意.....	49
 参考文献	 52
 第3部 大学における生成 AI リテラシー教育.....	 53

1. はじめに	54
2. 生成 AI の仕組みや利用上の注意 ―「賢い他人との対話」	56
2.1 情報通信システムの視点 ―「賢い他人はどこにいるのか？」	56
2.2 回答を鵜呑みにしない ―「賢い他人でも間違えることがある」	56
2.3 個人情報や機密情報を伝えない ―「賢い他人はあくまで他人である」	57
2.4 回答をそのまま流用しない ―「賢い他人の回答は他人のもの」	57
2.5 プロンプトの重要性 ―「賢い他人との上手な対話」	58
3. 生成 AI リテラシー教育事例から見る大学生と生成 AI	59
3.1 生成 AI リテラシー教育の対象と概要	59
3.2 受講生の生成 AI の利用・認知と教育機会	60
3.3 生成 AI リテラシー教育受講後の変化	68
4. おわりに	73
参考文献	75

第 4 部 生成 AI サービスの普及等経済のデジタル化を把握するための統計的枠

組みについて	77
1. はじめに	78
2. 国際収支統計からみたデジタルサービス取引のグローバル化.....	79
3. デジタル SUT の枠組み	81
3.1 デジタル SUT の枠組み	81
3.2 産業の分類	83
3.3 生産物の分類.....	84
(i) デジタル生産物.....	84
①ICT 財 (ICT goods)	84
②デジタルサービス (クラウドコンピューティングサービス、デジタル仲介サービス (有償)を除く) (Priced digital services – except cloud computing services and digital intermediary services)	85
③クラウドコンピューティングサービス (有償) (Priced cloud computing services) .	85
④デジタル仲介サービス (有償) (Priced digital intermediary services)	85

(ii) 非デジタル生産物	86
(iii) その他の非デジタル生産物 (Non-Digital products – other)	86
(iv) SNA の生産境界外のデジタル生産物 (Digital products outside of the SNA production boundary)	86
①データ (Data (beyond 2008SNA))	86
②企業が提供するデジタルサービス (無償) (Digital services (beyond 2008SNA) , provided by enterprises)	86
③社会が提供するデジタルサービス (無償) (Digital services (beyond 2008SNA) , provided by communities)	87
3.4 取引形態	87
(i) デジタル注文	87
①デジタル注文 (Digitally ordered)	87
②取引相手からの直接注文 (Ordered directly from a counterpart)	87
③居住者のデジタル仲介プラットフォームを通じた注文 (Ordered via a resident digital intermediary platform)	87
④非居住者のデジタル仲介プラットフォームを通じた注文 (Ordered via a non-resident digital intermediary platform)	87
⑤非デジタル注文 (Not digitally ordered)	87
(ii) デジタル配信	87
①デジタル配信 (Digitally delivered)	88
②非デジタル配送 (Not digitally delivered)	88
3.5 諸外国との比較	88
4. 無償デジタルサービスの測定に関する研究について	89
5. 生成 AI サービスに関するアンケート調査	93
6. 次期 SNA (2025SNA) におけるデジタル化への対応の方向性	98
6.1 データの価値の測定について	98
6.2 経済のデジタル化とグローバル化の複合による課題	101
6.3 無償デジタル生産物に関するサテライト勘定	103
7. むすび	104
参考文献	106

第1部 生成AIが変える社会と教育

1. はじめに

生成AIが産業革命のように世の中を変えられている。

日本で生成AIという言葉が広く一般に使われ始めたのは2023年の2月頃である。

最初是在京キー局の中でも規模の小さい、日経新聞などとのつながりの深い、ビジネス系のテレビ東京がBS放送の番組として取り上げた。それから1週間のうちに、急速に主要メディアに広がり、朝日新聞やNHKで毎日のように取り上げられるようになった。

連日の連載を通じて、「生成AI」は社会を大きく変えてしまう、人間の仕事がなくなってしまうという意見が多くの人々の話題となるようになった。

本研究は、それから1年、2024年の夏にテーマとして採択され、広く一般のひとびとにわかりやすい解説をするために執筆されている。

この研究報告書を手に取ったみなさまが、生成AIの実態について少しでも理解が深まれば幸いである。

2. 歴史

2.1 AIの歴史：3次にわたるブームと冬の時代¹

AI（人工知能）は、60年以上の長い歴史があり、現在は第3次ブームにあたる。

第1次ブームは1960年代のことである。チェスや定理の証明について研究が行われた。

しかし、実世界では役に立たない「おもちゃのような問題」であるという批判を受け、産業界からの期待に応えることができずに終わった。

1970年代には冬の時代を迎える。

つづく第2次ブームは1980年代にやってくる。

今度は「チェスや定理の証明などの探索や推論だけにとどまってはいけない。人間が偉い

¹ 2.の内容は次の論文を参照した。

山口 高平 (2017)「AI がもたらす新しい社会 シンポジウム講演」情報システム学会誌 12 巻 2 号 p.73-86

山口 高平 (2020)「AI 社会実装における要件と AI 技術の擦り合わせ 第12回シンポジウム 基調講演」情報システム学会誌 15 巻 2 号 p.55-71

のはいろんなことを知っているからだ。」という風潮に変わり、研究が進められた。

エキスパートシステムという、ルールを組み合わせて専門家の知識を組み立てて動かすものである。産業界ではエキスパートシステムがブームになった。日本の人工知能学会はこの第2次ブームの中で発足した。通産省が10年で570億円出し、電機メーカーも研究資金を出しあって、合計約1000億円の予算で第5世代コンピュータを開発した。1秒間に5億回の三段論法をする当時は世界一速い並列推論マシンだった。

当時はif~thenルール、つまり、こうした状況の場合は、こうする、という指示内容によって専門家の知識をすべて表現すれば、「専門家を超越するシステムも開発可能」と言われた。しかし、実際には、すべての状況を網羅することはできなかった。

人の知識は本当に膨大で奥深い。人の知識をif~thenルールで記述することによって、ある程度までは賢いなという程度に動かすことができた。だが、記述された状況とは異なる状況が現れると途端に処理ができなくなってしまう。そうした状況が分かってきたことで、エキスパートシステムへの失望が広がった。

例えばAIスピーカーは、一見対話できているように見えるが、過去の対話データから、こういう事を言ったらこういう事を言い返すという表層的なパターンで対応しているだけである。

インターネットが登場する前だったため、利用できるデータ量も少なく、あまり利用できないままに終わってしまった。

コンピュータに丸暗記をさせても、問題解決には使えないので、コンピュータが問題解決するのに利用できる知識の表現方法を考えて、それを通してコンピュータに人のノウハウを蓄積すべきだ、という考え方が生まれ、知識工学が提唱された。

ここから生まれたのが、言葉と言葉の意味ネットワークと、言葉の意味理解としての常識推論が必要だ、ということだ。概念の意味をコンピュータに理解させるオントロジーとよばれる研究が始まった。1980年代半ばから米国で研究が始まり、100万以上の概念群が定義されている。

1990年代に、コンピュータのダウンサイジングが進んで、パソコンが主流になり、インターネットが普及した。AIブームが冷め、90年代後半から第2次冬の時代を迎えた。

つづく第3次ブームは2011年頃から現在に至る。

第3次ブームは、ニューラルネットワークを利用したディープラーニング（深層学習）の技術が中心である。ディープラーニングという言葉が出てくるのは2006年からであるが、2011年頃にブームになったのは、コンピュータが速くなり、計算機パワーがどんどん上がってからである。コンピュータネットワークのインフラが整備されたことの影響が大きい。

以前は、解を探索するために、数時間から2~3日かかるのが当たり前だった。コンピュータの計算速度が1.5年で2倍速くなるというムーアの法則により、同じ計算量が数分から数秒で処理できて、答えが出るようになった。このために、さまざまな要素技術を組み合わせるインテグレーションが可能になった。

ディープラーニングは画像認識系、言語処理系、生成系（似顔絵を簡単に作るのは生成系）という 3 タイプに分かれ、多くの応用が生まれている。

2.2 ニューラルネットワークは「もうひとつの AI」

第 3 次 AI ブームの中心といえるニューラルネットワークが AI であると言えるかは微妙な問題だ。広い意味ではニューラルネットワークも AI に含まれるのだが、歴史的にみると AI とニューラルネットワークの研究者はあまり意見が合わなかった。そのため、ニューラルネットワークのことを、Another AI「もうひとつの AI」という言い方をすることがある。

なぜ、そんなことになるのだろうか。ニューラルネットワークの研究は、脳を規範にする。その点がほかの AI 研究との大きな違いだ。AI 研究では、人の脳で起こっている現象は参考程度にしかしない。例えば将棋に勝つなど人の脳と同じ振る舞いをできるようになることが目標であり、人の脳と同じ振る舞いができれば、それを実現する仕組みはエンジニアが考えるべきことだ。これが、脳の仕組みを参考にはするが、規範にはしない、ということだ。一方、ニューラルネットワークの研究者は脳を規範にすることが大前提だ。最終的に人に迫る能力をめざすという目標は同じなのだが、脳を規範とするか、しないかという点で方法論が異なるわけだ。

日本では、人工知能学会が AI の研究組織として 1986 年に発足している。一方、ニューラルネットワークの研究組織として神経回路網学会が発足したのが 1989 年で、二つの学会の発足はだいたい同じ時期だ。神経回路網学会は、脳科学系の学会と連携しているが、人工知能学会とは殆ど連携していない。

2.3 ニューラルネットワークの歴史

ニューラルネットワークの歴史を振り返ってみよう。

ニューラルネットワークのブームの時期と冬の時代は AI 研究のブームの時期と冬の時代とほとんど一緒だ。

ニューラルネットワークは 1958 年、ローゼン・ブラットが提唱したパーセプトロンの登場ではじまる。パーセプトロンは脳を規範にしたモデルであり、注目を浴びた。

しかし、AI 研究者であるマービン・ミンスキーは、パーセプトロンを「排他的論理和 (exclusive or) さえ学習できない」と批判した。その批判によって「パーセプトロンは役に立たない」と評価されることになり、1970 年代に一気に冬の時代を迎える。

1970 年代の終わりに、NHK 技研の福島がネオコグニトロンを提案する。

脳科学の研究で、ネコの視覚野には図形の特徴を抽出するシンプル細胞とそれが平行移動したり、回転したりして、位置ずれが生じて問題なく認識を補正するコンプレクス細胞

があることが70～80年前に発表されている。ネオコグニトロンは、このシンプル細胞とコンプレクス細胞をうまく融合するようなニューラルネットワークである。

ネオコグニトロンを変形したものが現在のディープラーニング（深層学習）だ。特に画像処理に使われている畳み込み型ディープラーニングの基礎になっている。このため「ディープラーニングは日本発の技術」と言われることがある。

第2次AIブーム時代のニューラルネットワークは、バックプロパゲーションの開発がキーである。バックプロパゲーションとは、正解との誤差を逆伝搬する仕組みだ。人間の脳の仕組みには無い仕組みだが、この仕組みを導入することによりニューラルネットワークで良い成績を上げることができるようになった。

この時代には、ホップフィールド型やボルツマンマシン型というニューラルネットワークが提案された。カナダのトロント大学のジェフリー・ヒントンのボルツマンマシン型を提案した。ヒントンは2024年にノーベル物理学賞を受賞した。

ヒントンの弟子のひとりであるブカンが、より性能の高いコンボリューショナル・ニューラルネットワーク（CNN）を提案した。

こうしてバックプロパゲーションの方法が実用化され、現在も広く利用されている。

一方、この時代に限界となったのは、過学習という現象である。学習に使う訓練データについては高い性能がでるのだが、訓練データに偏りすぎた結果になってしまい、未知のデータに対する成績が悪いという現象である。

ニューラルネットワークの性能を上げるために、構造を4層とか5層に増やすと、この時代には学習が全然終わらなくなってしまった。計算機パワーが貧弱だったためである。このため、ある程度使える技術であるということはわかったけれども、計算能力から限界がある、ということで1990年代に再び冬の時代を迎える。

ニューラルネットワークの冬の時代には、1990年代後半から数理統計理論に基づく機械学習がどんどん出てきた。ベイジアンネットワークやサポートベクターマシン、ランダムフォレストなどである。これらの数学的なモデルを持つ機械学習は簡単で、計算時間も短く、精度も1980年代に提案された3～4層のニューラルネットワークよりも高いことが分かってきた。このため、1990年代～2000年代は、3～4層のニューラルネットワークは衰退した。2000年代前半はニューラルネットワークという言葉が入っていると論文が採択されない時代であった。

この冬の時代を突破したのはヒントンのオートエンコーダーや制限型ボルツマンマシンである。2010年頃、計算機パワーの向上により10層程度のニューラルネットワークでも動くようになり、現在は、100層以上でも動く。

オートエンコーダーについて説明しよう。例えば、自動車のオーディオやスカイラインといった画像から車種を認識する課題を考えてみる。入力層には、解析対象の画像の各画素に対応するノードがあり、それぞれの画素の赤青緑の三原色の情報が入力される。出力層も、入力層と同様に画像の各画素に対応するノードを持つ。

学習の初期段階では、入力画像と出力画像が大きく異なるため、誤差が大きくなる。この誤差を基に、バックプロパゲーションによってニューラルネットワークのパラメーターが更新され、徐々に入力画像と出力画像が一致するように学習が進む。

中間層では、入力画像を効率的に表現するための特徴が学習される。例えば、エッジや形状、テクスチャといった特徴が抽出されることがある。これらの特徴は、元の画像を再構成するために必要な情報であり、次元圧縮された表現とも言える。

オートエンコーダーは、単純に画像をコピーするのではなく、中間層で情報を圧縮し、その後再構成することで、データの潜在的な構造を学習する。この過程で、ノイズ除去や異常検知など、様々なタスクに役立つ特徴が抽出されることがある。ニューラルネットワーク以前の画像処理研究では、この特徴量を人間の研究者がいろいろ知恵を絞って考え、与えていた。一方、ディープラーニングでは、人間の研究者が特徴量を指定するプロセスはなく、ネットワークが自動的に特徴量を学習していく。研究者が考えもしなかった特徴量が自動的に作られ、ああ、こういう特徴量があったかと驚くこともあるわけだ。

ImageNet という大規模な画像認識コンテストでは、毎年、最新の画像認識モデルの性能が競われている。2012 年、AlexNet という深層学習モデルが従来の画像認識手法を大きく凌駕し、画像認識分野に革命をもたらした。AlexNet は、8 層の畳み込みニューラルネットワークで構成されており、大規模な画像データセットを用いた学習により、高い認識精度を実現した。

その後、より深く、より複雑なネットワークが提案され、画像認識の性能は飛躍的に向上した。さらに、ネットワークを深くしても勾配消失問題が発生しにくい構造の導入により、150 層を超える深層なネットワークが実現した。また、ネットワーク内の各チャンネルの重要度の学習により、より効率的な特徴抽出を実現した仕組みも開発された。

これらの深層学習モデルの登場により、画像認識のエラー率は人間を下回るレベルにまで達した。初期がんの画像診断では医者がいくら頑張っても 80% の精度だが、AI では 98% の成績がだせる。NEC と国立がん研究センターの共同研究では、大腸がん検査から始めて、胃がん検査なども実験した結果、すべて 98% 程度の精度だった。患者は、人間の医者を頼りにするのではなく、ディープラーニングに検査してもらうべき時代がやってきた。

ディープラーニングは画素単位で、ここが少し赤くなってこっちが暗くなっているという、関係を見出す。人間はそのような細かい関係性を認識できないので、ディープラーニングの方が人間よりも精度が高くなる。ただし、検査後の治療法の決定については、まだまだ、医者の方が AI より優れている。

ニューラルネットワークと AI は長い間対立してきたが、このニューラルネットワークの成果が AI の革新的技術と報道されることになり、時代が変わっていく。

3. 大規模言語モデル

3.1 生成 AI とはなにか、大規模言語モデルと画像・音楽・動画生成

生成 AI とは、Generative AI を日本語訳した言葉である。Generative は、「生成に関する」といった意味を持つので、「生成系 AI」と呼ばれることもある。

これまでの AI（人工知能）は、特定のタスク（仕事）をこなすために使われる技術だった。例えばレントゲン写真を見て、ガンがあるか、ないかを判別するといった仕事だ。

これに対して、生成 AI は、画像や文章を作り出すので、生成する、作り出す点の特徴をとらえて、生成 AI と呼ばれるようになった。

生成する対象は、大きく二つのジャンルに分類される。一つは言語を生み出す、「大規模言語モデル」であり、もう一つは、画像や音楽、動画を生み出す「ディフュージョンモデル」である。この二つは大きく原理が異なる。同じ生成 AI といっても、まったく異なる技術を使っていると言ってよい。

この研究報告書では、主として言語を生み出す「大規模言語モデル」について中心的に取り扱うこととする。

人間にとって言語は文化を継承するうえで非常に重要な道具であり、言語を生成する人工知能が生まれたインパクトは非常に大きかった。

画像や音楽、動画などを生成する人工知能があっても、それは機械が綺麗な絵などを作ってくれるだけで、人間らしさは感じられない。ところが、人工知能が言語を生成すると、それは稚拙な文章であっても人間らしさが感じられる。

初期の人工知能で機械的な条件によって生みだされた文章ですらも、それは人間らしいものとして受け止められる場合があった。それだけ言語を生み出す人工知能は人間に与える印象が強いのである。

生成 AI についても、画像生成が先行したものの、一般に広く語られるようになったのは大規模言語モデルが話題になったときである。2022 年 11 月米国のオープン AI 社が Chat GPT（チャット GPT）というおしゃべりをする人工知能を公表した。

3.2 大規模言語モデルの仕組み

大規模言語モデルは、英語の Large Language Model を翻訳した言葉である。Large が大規模だから、Small Language Model（小規模言語モデル）や、Middle Language Model（中規模言語モデル）があるかというところ、そういうわけではない。あくまで、大規模になったときにのみ、人工知能としての能力を発揮するようになったので、大規模言語モデルだけが存在する。

多言語の機械翻訳は情報分野で長らく試みられてきた課題だ。人間は辞書を引きながら翻訳の作業を行うという発想から、辞書を整備していけば、言語の翻訳にたどり着くだろうとか、文法を学ばなければならないので、言葉の品詞分類を正確に行う必要があるなど、様々な方法が試みられたが、なかなか満足な翻訳結果は得られなかった。

これに対して、赤ちゃんが言葉を学習するように、大量の言葉のシャワーを浴びせて、辞書や文法書を知らないままに人工知能が言語を習得していくのが、現在の大規模言語モデルの基本的な仕組みである。

言語は、人々の間で互いに記号を予測しあうことにより次第に形成されてきたものであるという仮説がある。例えば、りんごを指さしながら「りんご」ということばを提示することにより、りんごという「もの」とりんごという「ことば」が対応することを了解しあっていくことによって次第に意思疎通ができるようになってきた。

大規模言語モデルは言語のつながり方を大量に学習する中で、次にどのような言葉が現れるのかを学習していく。例えば、「むかしむかし、あるところに」という言葉の次には「おじいさんとおばあさんが住んでいました。」という言葉が続くのではないかと予測し、期待する。実際に期待した言葉がでてくれば、その信念を強くし、実際には異なる言葉がでてくるのであれば、その予測を修正する。

このようにして、次々に次の言葉を予測し、もっとも可能性が高いと思われることばを次々に生み出していくのが大規模言語モデルの仕組みである。

3.3 用途に合わせた大規模言語モデルの開発

用途に合わせて専用の大規模言語モデルを開発する場合、いくつかの方法がある。

① プロンプト・エンジニアリング

もっとも簡単な方法は、汎用開発されたシステムに対する入力の方法を工夫することである。これをプロンプト・エンジニアリングと呼ぶ。汎用のシステムであっても、入力の指示内容を詳細に指定し、出力の例示を行うことで、回答結果をかなりの程度制御することが可能である。特に、入力として取り扱える語数が増加したことにより、作業の指示内容に加えて、作業内容の例示を長文で入力することが可能になったことにより、プロンプト・エンジニアリングの可能性は大きく拡大した。

作業内容を例示することを、一つの例示の場合はワン・ショット、複数の例示の場合はフュー・ショットという。作業内容を例示しない場合はゼロ・ショットという。ゼロ・ショットから、ワン・ショット、フュー・ショットへと例示の数が増えるほど、回答は正確に、期待に即したものになっていく。

例えば、出力の例示として、自分の書いた文章を示してやることで、自分の文章の書き方をまねた出力を得ることができる。出力の指示として、「簡潔に」とか、「専門用語について

わかりやすく」など具体的に指示することによって、出力の文体などのスタイルを制御することも可能である。

このような指示のパターンを記録して、選択したパターンに従って作業対象の入力だけを追加する手法で、特定用途に特化したシステムとする方法が、GPT s などが採用している手法である。

② ファイン・チューニング

プロンプト・エンジニアリングより工数は増えるが、より特定分野に特化したシステムを開発する方法が、ファイン・チューニングである。

汎用開発されたシステムをベースとして、特定の分野や用途に関する追加データで追加学習する方法である。医療や金融など、対象とする領域の専門情報を集中して追加学習することにより、当該分野に特化した専門的な知識を蓄積することができる。

ベースとなる汎用開発されたシステムは、ファイン・チューニングが可能な形で、大規模言語モデルのパラメーターが公開されたものである必要がある。汎用として商用に一般公開されているシステムには、このような公開を行っていないものも少なくない。パラメーターを公開して、様々な分野の大規模言語モデルのベースとなることで、規格の主流的地位を確保しようとするオープン戦略と、パラメーターを非公開として、自社ですべての開発を囲い込むクローズ戦略とのせめぎ合いがある。

③ 公開されたモデルの一からのトレーニング

②では、大量の情報により汎用の学習を行った結果のパラメーターを出発点として、専門特化した分野の追加学習を行うファイン・チューニングについて説明した。同じモデルを用いて、パラメーターをゼロからスタートして学習する場合がある。この場合は、大規模な計算資源と、学習の為の膨大なデータを準備しなければならない。

学習データは、インターネット上から集められた情報から、重複する情報を取り除いたものであり、その質によって大規模言語モデルの性能は大きく左右される。

現在は、インターネット上のほとんどの情報が学習のために既に利用されている状態にあり、性能をこれ以上高めるための学習データが不足する状況が懸念されている。そのため、大規模言語モデルから出力した情報で学習する手法などが試みられている。

④ 新しいアーキテクチャ

現在主流となっている大規模言語モデルは Transformer アーキテクチャと呼ばれる。セルフ・アテンションという機構の採用により、大幅に性能が向上したものである。

現時点では、これに代わるさらに性能の良い新たなアーキテクチャは報告されていないが、大規模言語モデルの基盤となるアーキテクチャの革新は常に研究が進められている。

例えば、バックプロパゲーションの仕組みは、生物の脳には存在しない。バックプロパゲーション無しで高い性能を持つ新しいアーキテクチャがあるのではないかと期待されている。

このような新しいアーキテクチャによる大規模言語モデルの開発は、現在のアーキテク

チャによる性能向上が継続している中で、それほど盛んであるわけではないが、地道に探究が続けられている。

3.4 推論能力の向上

2024 年の後半から 2025 年にかけては、大規模言語モデルの推論能力の向上が大きな話題になっている。従来は苦手と言われた数学の問題など、論理的な推論の精度が大きく高まったモデルが公表されている。

このことは、大規模言語モデルの業務利用の可能性をさらに広げることとなっている。例えば、調査研究について課題を与えると、調査研究の進め方の方法について提案し、利用者による修正を経て、その方法に従った調査結果をレポートにしてまとめてくれる機能まで実現している²。

このような機能を利用すれば、もはや調査研究について人間に代わる助手としての役割を大規模言語モデルが代替してくれるようになる。特に、検索先をアカデミックな情報に限定することにより、これまで人間にしか行えなかったリサーチ研究をかなりの程度効率的に行うことができるようになるだろう。

2024 年 12 月末、中国企業から DeepSeek-R1 というオープンソースのモデルが公表された。このモデルはクロードで高価な利用料を設定している OpenAI の o1 というモデルと同等の性能を 20 分の 1 の低コストで提供するモデルである。この登場により、米国の AI 関連の株価が暴落する DeepSeek ショックが起きた。訓練に要した経費は o1 の 10 分の 1 程度であり、中国に対して半導体の輸出制限により性能の低い半導体の使用を余儀なくされているにもかかわらず高性能を挙げている点が注目を浴びている。

4. 大規模言語モデルの特徴

4.1 大規模言語モデルは間違えるか

大規模言語モデルが生み出す言葉が、あたかも人間であるかのように感じるのは、大規模

² Google for Education 認定イノベーター 野中潤の研究室「【大学関係者必見！】レポート課題はもう出せない？～自動的にリサーチしてレポートを書き上げてくれる Google Gemini の Deep Research が発揮する驚異的な能力～」

<https://www.youtube.com/watch?v=eI3NDYIUxdA>

言語モデルからの答えを予測する人間の側が期待している通りの答えを大規模言語モデルが作文してくれるからにはほかならない。

大規模言語モデルは、もっとも常識的な答えを出力してくれる。大量の言語のシャワーで訓練を受け、もっとも常識的な、つまり、もっとも多くの人が答えるような答えを出力するのが大規模言語モデルだ。

「常識のウソ」という言葉がある。世の中でこれは常識と信じられているようなことが実は間違いだったというケースである。例えば中世までの世界であれば、天動説が常識であり、地動説はとんでもない非常識だっただろう。

このように考えると、大規模言語モデルは「正解」を考えてくれるわけではなく、世の中でもっとも多くの人が考えそうなことを答えてくれるという点に注意しなければならないことがわかる。「常識のウソ」に関しては、大規模言語モデルは正しいことではなく、多くの人が信じているウソを答えるわけだ。

このことは、大規模言語モデルを使うにあたって十分に頭に入れておくことが必要なりテラシーになる。大規模言語モデルは、AI（人工知能）といっても、人間が考え付かない正解を答えてくれるわけではなく、もっとも多くの人が考えていることを答えてくれる、ということである。

大規模言語モデルは間違えるか、という問いに対しては、多くの人間が間違えて正しいと信じていることをそのまま答えてくれる、と答えることになる。つまり、大規模言語モデルは、人間と同じように間違えるのだ。

4.2 錯覚（ハルシネーション）

大規模言語モデルを利用する上で、よく注意しなければならないこととして、「ハルシネーション」という言葉がでてくる。これは「錯覚」という意味だ。大規模言語モデルに直接覚えさせた訓練データに入っているわけではないが、訓練されたわけでもない言葉を勝手につなぎ合わせて、もっともらしく回答することを指す。

大規模言語モデルが現れて1年くらいの間は、この「ハルシネーション」が大規模言語モデルの弱点として強く意識された。よくあるケースとして、引用元として、もっともらしい論文の題名と著者名を出してくれるのだが、調べてみると、その著者に、そのような題名の論文は存在しない、という問題が指摘された。

この「ハルシネーション」の問題に対応するためには、検索などで、大規模言語モデルが出力してくれた論文が実際に存在するかを人間が確認することが必要であると言われた。

だが、考えてみれば、検索して調べる作業は、コンピュータによる処理としては、従来から幅広く実用化され、利用されていた分野だ。しばらくするうちに、大規模言語モデルが出力する前に、その検索して調べる作業を大規模言語モデル側で行って、実際に実在するもの

に絞り込んで回答すればよいではないか、という技術が実用化された。つまり、検索と言語生成を組み合わせ、情報を調べたうえで回答を生成するという仕組みの実現である。こうした技術を RAG（ラグ）と言うが、この技術を取り入れることによって、大規模言語モデルは「ハルシネーション」を起こしにくくなった。

例えば、Google が 2024 年 6 月に日本でも利用できる製品として公表した notebook LM というシステムでは、システムが参照する文書をあらかじめ人間が指定し、その指定した文書に、問い合わせた関連の記述が無い時には、「指定の文書には問い合わせに関する情報はありません」と回答するように設計されている。以前のシステムでは、指定の文書に無い情報も、無理やり自分で文章を生成して回答しようとしていたため、ハルシネーションを起こしやすかったが、この仕組みによってハルシネーションを起こす可能性が低下した。

4.3 仕事で使えるようになった大規模言語モデル

この原稿を執筆している 2025 年 1 月現在、大規模言語モデルは「仕事で使える」ようになったと考える人が増えてきた。Chat GPT の登場から 2 年を経て、恐る恐る使ってみる状態から、普通に仕事に使ってみる状態への変化がみられる。

技術面では、大規模言語モデルのパラメーター数がさらに増加することで、潤沢な計算能力をふんだんに利用したモデルが出現した。これにより性能が上がったことが、実用的と評価される最大の要因である。

パラメーター数が増加することにより、扱うことのできる文章量も増加する。Chat GPT が出始めた頃には、一つの文章を扱っていたのが、一つの段落、一つの節、一つの章、さらには一冊の書籍全体を一度に扱えるようになってきた。このことは、書籍全体の文脈の中から回答を求めることができるという利点をもたらす。一つの文章しか扱えなかった時には、根気良く、何度にも分けて文章を入力していき、それらの文章に記述されている内容を徐々に大規模言語モデルに学習させるプロセスが必要だった。一冊の書籍全体を一度に扱えるようになれば、「次の文章を要約して」と指示した後、書籍全体の内容を入力すれば、書籍全体を鳥瞰して、全体の要約を出力することが可能になる。

Google の Notebook LM では、書籍にあたる PDF 情報を複数指定して、それら全体を横断した知識を集約して質問に回答することができる。たとえば、地理の教科書と歴史の教科書を指定して、その両方の知識を使って回答を求めることができる。

4.4 大規模言語モデルを使った検索から検索システムとの連携へ

Chat GPT は、質問を入力して回答を求める形式だったため、当初は Google など従来の検索システムと同じ機能を期待する利用者が多かった。このため、検索に対する回答として

誤りが表示されることで、「使い物にならない」と否定的な評価をされることが多かった。

大規模言語モデルは、前の語から高い確率でつながる言葉を次々に紡ぎだす仕組みであるため、検索目的には向いていない。検索には従来の検索システムを使用すべきであり、大規模言語モデルに検索の機能を求めるのは利用する側が使い方を間違っている。しかし、多くの人は大規模言語モデルを検索システムとして捉えてしまったために、「使い物にならない」という否定的な評価をした。現在、その評価は覆され、「仕事に使える」技術と認識されている。この要因について考えてみよう。

一つは、RAG 技術の活用により、検索は検索システムに任せ、検索システムと組み合わせる形で、回答を大規模言語モデルに作らせる製品が増えてきたことが理由である。例えば、Perplexity は、入力された内容に基づき検索を行い、検索結果に基づいて回答を生成する製品である。最新の情報を検索し、その検索結果も入力として利用することによって、大規模言語モデルが生成する内容の正確性が確保されることになる。

このケースでは大規模言語モデルが進歩したというよりは、大規模言語モデルと検索システムを組み合わせた製品が普及することにより、人々の認識する「大規模言語モデル」の対象が、純粋な大規模言語モデルだけの製品から、大規模言語モデルと検索システムを組み合わせた製品へと変化していったというべきだろう。

5. 大規模言語モデルと業務

5.1 大規模言語モデルの機能の活用 メール返信

大規模言語モデルは検索に向けたシステムではないと気付いた利用者の間で、流暢な文章を生成する能力の高さを活用する動きが現れた。内容の真偽はともかく、人間のような流暢な文章を作ることができる。

この機能を活用した用途の一つが、メールの返信の自動作成である。

メールの返信は、業務の中でかなりの割合を占めるケースがある。いくつかのケースに分類することができ、ほぼ定型の文章を送り返せば良いのだが、その分類は従来の手法ではなかなか自動化することが難しかった。送られてくるメールの文章の表現は多様であり、機械的な処理に向いていなかったためである。

多様な表現の中から、意味を読み取り、ケース分類して、ケースに応じて、相手の文章の中から必要なキーワードを抽出して使いながら、適切な返信を行うというタスクは、大規模言語モデルの得意とする分野である。このような利用の仕方により、大規模言語モデルが秘書代わりに使えるという認識が広がった。

比較的単純な作業を繰り返す作業については、大規模言語モデルの出現以前には、RPA（ロボティック・プロセス・オートメーション）と呼ばれる手法が注目されていた。これは、

PC で人間が作業する内容を、手順を指定して実行させるロボットのようなプログラムという意味である。本格的なシステムを開発するためにはかなりの費用が必要であるが、利用者が作業の内容を指示して、その指示に従って作業を繰り返す、小規模で、利用者自身が動作を指定できるロボットというイメージである。

RPA も、利用者が、「こういう場合は、こうする」という条件ごとの作業内容を細かく指示することが必要であり、その指示は論理的で網羅的であることが必要だった。

これに対して、大規模言語モデルは、指示の内容がある程度曖昧であっても、システム側が幅を持って解釈し、解釈の内容について確認したり、不足する情報を求めたりしながら作業を行うことができる。その意味では、RPA がロボットであるとするれば、大規模言語モデルは人間の秘書に相当すると言えるだろう。ロボットは指示の内容をロボットが理解できるように指定された形式で行う必要があるが、大規模言語モデルは人間の使う言葉で、人間と対話的に指定することができる。

5.2 大規模言語モデルによるアイデア出し

大規模言語モデルは、インターネット上の大量の情報によって訓練される。このことから、様々な広範な分野に関する知識を吸収している。この特徴を生かした大規模言語モデルの使い方がアイデア出しである。

20 世紀初頭のオーストリアの経済学者シュンペーターは、著書「経済発展の理論」において、経済発展の原動力は「新結合」であると主張した。この「新結合」が、後に「イノベーション」と呼ばれるようになった概念である。「新結合」とは、既存の要素を新たな組み合わせにすることで、経済に新たな価値を生み出すことを指す。具体的には、以下の 5 つの要素が挙げられる。

新しい財貨（製品）：消費者がまだ知らない製品、または品質が大幅に向上した製品など

新しい生産方式：新しい生産方式や取扱い方法など

新しい販路の開拓：従来の販売先ではない先など

原料あるいは半製品の新しい供給源の獲得：原料や半製品の新しい使い方など

新しい組織の実現：古い企業ではなく、新しい企業の出現など

イノベーションは既存の要素の新たな組み合わせによって生まれる。これまでは、一人の人間が収集し組み合わせることができる既存の要素の数には大きな制約があった。一人の人間が生きている時間は有限であり、収集できる情報量には制約があるからだ。

一方、大規模言語モデルは、現在、インターネット上で収集可能なデータの大部分を訓練の対象とすることができる状況にある。逆に、これ以上収集可能な情報がなくなる限界点ま

で達していることが懸念されている。

このことから、大規模言語モデルの出現によって、これまで試みることができなかった広範な分野の知識を新たに組み合わせて、様々な新たなイノベーションを生み出すことができるようになってきている。

もちろん、どのような分野で、どのような目的のためにその組み合わせを試みるかという指示は、人間が発意しなければならない。しかし、ひとたび目標が定められれば、大規模言語モデルは様々な知識を融合して、新しい組み合わせによるイノベーションの実施案を提案することができる。

これが、大規模言語モデルによるアイデア出しのポイントである。

人間の限られた知識の限界を超えて、広範な知識を組み合わせて提示することにより大規模言語モデルが提案したアイデアを、人間が評価し、取捨選択することを通じて実現していく。こうした作業を通じて、これまでは得られがたかった新たな新結合、すなわち、イノベーションを生み出していくことが可能になる。

5.3 大規模言語モデル時代のプログラミング

プログラミングは、人間が機械にどのような条件の時に、どのように動作するかを指示する行為であり、プログラミング言語を用いて行われる。それに対して、コンピュータは、処理結果を人間に伝えるために、正常な場合は出力を行い、異常が発生した場合はエラーや警告メッセージを表示する。

これまでのプログラミングにおいては、人間が日常的に使う自然言語ではなく、コンピュータが直接処理できるプログラミング言語を覚えて、コンピュータに寄り添った形で指示を記述していく必要があった。

このため、プログラミングができるようになるには、まず人間がプログラミング言語を習得する必要があり、多くの時間が必要であった。

大規模言語モデルの出現により、この状況は大きく変化する。人間が日常的に使う自然言語で指示を与える内容について、コンピュータがそれを解釈し、適切な処理を実行することができるようになった。

もちろん、自然言語による指示が曖昧な場合には、人間の意図する指示とは異なる結果が出力される場合もある。それを未然に防ぐために、指示の内容についてコンピュータの側から人間に確認したり、必要な追加情報について人間に入力を求めたりすることができるようになった。

大規模言語モデルの出現により、コンピュータによる自然言語処理が可能になったことは、コンピュータのユーザーインターフェースの大きな変革と言える。

プログラミングの世界で、ローコード、ノーコードという言葉がある。これは、プログラ

ミング言語の修得という特殊な技術をもたなくても、プログラミングの敷居を下げ、より直感的に行えるように、コンピュータが人間に近づいたプログラミング手法を意味する。例えばマウス操作で指示の内容を画面上で図示したり、人間が操作をしてみせたりする内容をコンピュータ側が解釈して指示内容を理解するなどの手法が開発されてきた。

大規模言語モデルの出現は、ローコード、ノーコードの動きをさらに発展させ、まるで人間と会話するように自然な言葉でコンピュータに指示できるレベルにまで高めたと言える。人間同士が会話するような形でコンピュータに指示を与え、必要がある場合には追加の会話を通じて情報をさらに入手し、必要な操作を実施してくれるようになった。

5.4 大規模言語モデルによるプログラミングやシステム開発の補助

前章では、プログラミングの技術が無くてもコンピュータに指示を与えることができるようになったことについて述べたが、大規模なシステム開発を行う際には、依然としてプログラミング技術を持ったエンジニアによるプログラミングが必要不可欠である。ここでは、こうしたエンジニアによるプログラミングやシステム開発を補助する大規模言語モデルの機能について論じる。

大規模言語モデルは、インターネット上の膨大な情報、特にプログラミング言語、フレームワーク、アルゴリズムに関する情報を学習しており、プログラミングやシステム開発に関する幅広い知識を蓄積している。このため、エンジニアは、大規模言語モデルをプログラミングの強力なアシスタントとして活用することができる。例えば、コードの生成、デバッグ、ドキュメント作成、さらには新しい技術やライブラリの学習など、開発のあらゆるフェーズにおいて、大規模言語モデルからの情報提供やアドバイスを受けることができる。

こうした技術の活用により、エンジニアの生産性が著しく向上した。従来、同僚やコミュニティに質問して解決していた問題も、大規模言語モデルに直接質問することで、より迅速かつ正確な回答を得ることができる。大規模言語モデルは、24 時間 365 日いつでも利用可能であり、根気強く質問に答えてくれるため、エンジニアはより効率的に開発を進めることができる。

例えば、特定の機能を実装するためのシステムの事例について、プログラミング言語を指定してコードのサンプルをリクエストし、それをベースにカスタマイズしてテストするケースを考えてみよう。プログラミングでは、実装した内容にエラーが含まれ、エラーメッセージが表示されることはよくある。大規模言語モデルを利用する場合、エラーメッセージの内容とともに、コードの一部分を提示することで、エラーの原因と具体的な解決策を詳細に説明してもらうことができる。

もし、提示された解決策を試してもエラーが解消されない場合は、再度大規模言語モデルに質問し、より詳細な情報や別の解決策を求めることができる。大規模言語モデルは、何度

でも質問に答えてくれるため、エンジニアは迷わずに開発を進めることができる。人間に質問する場合、何度も同じ質問をすることに抵抗を感じるかもしれないが、大規模言語モデルであれば、いつでも気軽に質問でき、迅速なフィードバックを得られる点が大きなメリットである。

大規模言語モデルのプログラミングやシステム開発の補助への活用は、大規模言語モデルの出現後、もっとも早く実用化に至った分野の一つといえる。これは、システム開発者がもともとコンピュータの使用に慣れ親しんでいることと、関係情報が広範にインターネット上に公開されており、それらの情報を検索しながら作業を進めることが慣例となっていることが大きな理由である。

5.5 大規模言語モデルの活用による業務効率の向上

大規模言語モデルは、ビジネスマンの業務遂行の様々な側面において活用される可能性を秘めている。大規模言語モデルは自然言語処理による質疑応答の形で利用されるため、従来の情報システムに比べて、より広範な分野に適用可能であることが一つの特徴である。

従来の情報システムは、経理処理や人事管理、在庫管理など、特定の業務内容に特化した専用システムが開発されることが一般的であった。これに対し、ワードプロセッサや表計算ソフト、電子メールソフトなどは、業務内容に関わらず広範に利用可能であった。大規模言語モデルは、これらの汎用性をさらに拡張し、ビジネスのあらゆる場面で活用できる可能性を秘めている。2025年1月現在、大規模言語モデルをワードプロセッサや表計算ソフト、電子メールソフトなどの機能の一部として融合して提供することが一般的になりつつある。大規模言語モデルは業務に日常的に利用される段階に差し掛かっている。

現在の携帯電話は電話の機能もスマートフォンのアプリケーションの一つとして提供されている。つまり、携帯電話を使っているつもりでも、知らず知らず、スマートフォンの一つの機能を利用している。

大規模言語モデルは、今後、携帯電話のアプリの一部として取り込まれることが予定されている。つまり、利用者は大規模言語モデルを使っていることを意識することなく、携帯電話を使っているつもりで、知らず知らずのうちに大規模言語モデルを利用する時代がやってくる。

特に、大規模言語モデルへの入力が音声入力になり、出力が音声出力として携帯電話のスピーカーから自然な言葉で返ってくるようになれば、利用者は携帯電話で人間と会話しているような感覚で、大規模言語モデルとやり取りできるようになるだろう。

大規模言語モデルに関するリテラシーを身に付け、その限界を知り、利点を最大限活用していくことが求められる。

5.6 過大な期待と過少な期待

大規模言語モデルの登場は、ビジネスのあり方を変革する可能性を秘めているが、その一方で、過度な期待と過少な期待が入り混じり、現実的な評価が難しい状況にある。

過度な期待としては、「大規模言語モデルを導入すれば、すべての業務が自動化され、大幅な効率化が実現できる」といった、万能な解決策として捉える傾向が挙げられる。このような期待は、大規模言語モデルの能力を過大評価し、現実とのギャップを生み出す可能性がある。

過少な期待としては、「自社の業務には複雑すぎる、専門知識が必要なため、大規模言語モデルは活用できない」といった、その可能性を過小評価する傾向が挙げられる。しかし、大規模言語モデルは日々進化しており、適用範囲も広がっている。

一部の事例では、大規模言語モデルを業務に導入することで、著しい効率化が実現されたという報告がされている。しかし、これらの事例は特定の業種や業務に特化したものであり、すべての業務に同様の効果が期待できるわけではない。

こうした状況の中、大規模言語モデルを適切に活用するためには、現実的な視点と実験的なアプローチが不可欠である。全てのケースが成功するとは限らず、トライアルアンドエラーを繰り返しながら、自社の業務に最適な活用方法を見つけることが重要である。各企業や組織においては、自社の業務内容や課題に合わせて、大規模言語モデルをどのように活用できるのかを具体的に検討し、実証実験を行い、段階的に導入していくことが重要である。

5.7 生成 AI と人間の協業

大規模言語モデルに限らず、画像生成など生成 AI は様々な新しい可能性を生み出している。こうした中で、生成 AI が何を行い、人間が何を行うかという、生成 AI と人間の協業の方法が模索されている。

音楽生成 AI などのケースでは、生成 AI は様々な楽曲を生み出すことができる。ただし、それは人間の目を見た時、あるいは耳で聴いた時に、必ずしも優れた楽曲とは言えないこともある。人間の果たすべき重要な役割は、生成 AI の生み出した多数の楽曲を評価し、何が良い楽曲で、何が良くない楽曲であるかを選別することだ。

このケースで見ると、生成 AI は人間がトレーニングすることにより、人間が生み出すものに近いものを大量に生成する能力を持った機械である。ただ、その生成するものは、必ずしも人間にとって満足のいくものではない場合もある。そのような場合に、人間は生成 AI が生み出したもののなかから、満足のいくものを選別することが必要になる。

あくまで最終的に評価するのは人間であり、生成 AI ではない。

これが、生成 AI と人間の協業における一つの姿であるといえる。

第二に、生成 AI はあくまで人間の指示のもとに動く機械であり、責任を取ることはできない。指示を行った人間が最終的に責任を取らなければならない。

6. 大規模言語モデルの教育への利用

6.1 英会話練習

大規模言語モデルの特性を生かした教育の事例として、英会話練習を取りあげよう。

大規模言語モデルが人間のように流暢に話すという点を利用したのが、英会話練習への活用である。

大規模言語モデルが使われ始めた当初は人間からの入力はいちキーボードからの文字列で行い、大規模言語モデルの出力もディスプレイへの文字列出力で行われたが、既に実用レベルで利用されていた、人間の音声から文字列への変換と、文字列から音声への出力の機能と組み合わせることによって、時間的なラグはある程度あるものの、会話として人間が話した内容に対して大規模言語モデルが応答し、PC からの音声を聞くことができる状態を利用することができた。

現在は、大規模言語モデルの入力を人間の音声の形で行い、大規模言語モデルからの出力を人間の音声の形で行うことによって、自然な会話として違和感のないスピードで会話ができる状態が実現している。日本語話者には英語モードでの入出力をそのまま英会話の練習に使用することが可能である。

大規模言語モデルを用いた英会話練習では、スピーキングだけでなく、リスニングやリーディングの練習に活用することもできる。

リスニングの練習に活用する場合には、「旅行」「ビジネス」「日常会話」などのテーマや、レストランでの注文、空港での手続きなど、様々な状況を指定することで、様々な状況での会話を生成することができる。また、初心者から上級者まで、レベルに合わせて難易度を調整してリスニング素材を作成することができる。

例えば、日本の中学生レベルの空港での出国手続きに関する例は次の通り。

ヒアリング素材の例（出国時）

税関職員: Good morning. May I see your passport and customs declaration form, please?

（おはようございます。パスポートと税関申告書を見せていただけますか？）

生徒: Sure, here you are.

（はい、どうぞ。）

税関職員: Thank you. Do you have anything to declare?
(ありがとうございます。申告するものはありますか?)

生徒: No, I don't.
(いいえ、ありません。)

税関職員: Okay. Have a nice trip.
(はい、どうぞ良いご旅行を。)

生徒: Thank you.
(ありがとうございます。)

スピーキングに関しては、AI や音声認識技術を活用して、ユーザーの発音を分析し、改善点を指摘するアプリケーション（例えば ELSA Speak）が実用化されている。

6.2 教育利用に関する基本姿勢とガイドライン

2022 年に Chat GPT が現れた際、学生や生徒に大規模言語モデルの使用を禁じるべきではないかとの議論が活発に行われた。これは、大規模言語モデルを使って書かれた文章の質が高く、試験などによる成績の評価が行えなくなるという危機感に根差すものだった。生成 AI によって作成された文章と、人間が作成した文章を区別することは困難であり、一時は判別するためのソフトウェアも開発されたものの、誤検出率が高いことが知られている。

現在では、大規模言語モデルが広く社会で使われる中で、教育現場でだけ大規模言語モデルを使うべきではないという議論は一般的ではない。教育現場でも大規模言語モデルを正しく使うことができるリテラシーを育成する必要があるからだ。

現在は、むしろ教育現場における成績評価の方法について見直していかなければならないとの認識が広がりつつある。

大規模言語モデルの教育利用で留意すべき点については、日本国内では文部科学省のガイドライン³が、国際的にはユネスコによるガイドライン⁴が代表的である。

³ 文部科学省「初等中等教育段階における生成 AI の利活用に関するガイドライン (Ver.2.0)」(令和 6 年 12 月 26 日公表)
https://www.mext.go.jp/a_menu/other/mext_02412.html

⁴ ユネスコは 2023 年「教育・研究における生成 AI ガイダンス」を公表した。

いずれも、教育においては、学生が自ら思考するプロセスを奪うことが無いよう留意すべきことを挙げている。例えば実務においては電卓の利用が通例であっても、かけ算について学習途上でマスターすることが求められるように、教育においては、自らの力で考えるプロセスが必要な局面が存在する。そうしたプロセスにおいては、大規模言語モデルを利用しない場面も必要である。

一方、大規模言語モデルを利用することにより、より効率的に、各自の能力に合わせた教育を行うことも可能である。能力に合わせて出題レベルを調整したり、説明の仕方を能力に合わせてパーソナライズするなどのきめの細かい対応が大規模言語モデルによって可能になる。

こうした教育の特性に合わせて大規模言語モデルをチューニングする試みが行われている。

このガイダンスは 2021 年「AI 倫理に関するユネスコ勧告」に基づいて、包摂的で公正かつ持続可能な未来のために、AI を人間の能力開発に役立てるべきであるという人間中心のアプローチに基づいている。ヒューマン・エージェンシー、インクルージョン、公平性、男女平等、文化的・言語的多様性、複数の意見や表現の促進を内容としている。人間中心のアプローチには、人間の主体性、透明性、公的説明責任を確保できる適切な規制が必要である。

2019 年「人工知能と教育に関する北京コンセンサス」は、教育における AI 技術の活用が、持続可能な開発のための人間の能力と、生活、学習、仕事における人間と機械の効果的な協働を高めるものでなければならないと確認している。

言語や文化の多様性を促進しつつ、社会から疎外された人々を支援し、不平等に対処するために、AI への公平なアクセスの確保を求めている。教育における AI に関する政策立案への政府全体、セクター間、マルチステークホルダー・アプローチの採用を提言した。

さらに、2022 年「AI と教育： 政策決定者のためのガイダンス」は、教育における AI の利点とリスク、そして AI の能力を開発する手段としての教育の役割を検討するにあたって、人間中心のアプローチの内容をさらに検討している。

(i) 学習プログラムへの包摂的なアクセス、特に障害のある学習者などの脆弱なグループに対するアクセスを可能にする。

(ii) 個別化された、開かれた学習オプションを支援する。

(iii) 学習へのアクセスを拡大し、質を向上させるために、データに基づく規定と管理を改善する。

(iv) 学習プロセスを監視し、教員に失敗リスクを警告する

(v) AI の倫理的で有意義な使用のための理解とスキルを開発する。

6.3 金沢学院大学における大規模言語モデルの教育

2023 年、金沢学院大学では大規模言語モデルの授業での紹介が始められた。経済学部
1 年生向け授業の一部と 3 年生向け授業の一部で、大規模言語モデルが企業でどのようにと
られ、活用が始まっているかの紹介が行われた。この段階においては、大規模言語モデ
ルの草創期であり、ハルシネーションの問題が大きな障害となっていた。2023 年の 1 年間
の間にも大規模言語モデルの性能向上は著しく、初年度においては試験的な紹介にとどま
った。

2024 年、金沢学院大学ではすべての学部の 1 年次教育において、大規模言語モデルを紹
介することとなった。大規模言語モデルの技術的な進歩は著しく、多くの企業などにおいて
実用レベルでの活用が広がってきたことを受けたものである。

2024 年、大貫ゼミにおいては、卒業論文の執筆に大規模言語モデルを活用することを試
行した。検索システムと大規模言語モデルを組み合わせた Perplexity にテーマとなる質問
事項を投げかけ、その回答をレポート化する形で、大規模言語モデルの回答の結果と、それ
に関する考察を行った。さらに、大規模言語モデルの出力結果と、執筆者の考察を大規模言
語モデルに評価させ、改善点についてのアドバイスに基づき執筆者の考察を書き直すサイ
クルを繰り返した。こうして最終的にできあがった考察を検索結果と合わせて大規模言語
モデルで 5 行程度に要約させ、その結果を結論とすることで卒業論文を完成させた。

こうした作業は、就職後におけるレポート作成作業のプロセスを大規模言語モデルによ
って行うことを意識したものである。卒業論文の作成作業を通じて、職場におけるレポート
作成を体験することによって、就職後のスムーズな大規模言語モデル利用をめざしたもの
である。

2024 年 9 月 16 日、石川県立図書館だんだん広場において、金沢学院大学経済学部によ
るシンポジウム「AI 技術により変化が加速する時代に求められる人材像」が開催された。

シンポジウムは三部構成で、第一部・基調講演Ⅰ「生成 AI で変わる社会」では、大貫の進
行により、漫画家「おおみ きりん」による AI でのマンガの作り方動画の紹介、放送大学学
生の村橋陽三による「デジタルを活用する高齢者」動画の紹介がなされた。これらを踏まえ
て、千葉商科大学教授／東京工業大学名誉教授の出口弘から、「AI の利用については目的を
決める責任は人間にあり、その目的を実行するための支援に AI を用いるという認識を持つ
ことが必要である」とのコメントがあった。

第二部・基調講演Ⅱでは、科学技術振興機構研究開発戦略センター特任フェロー／慶應義
塾大学名誉教授の黒田昌裕による基調講演「科学技術の進歩と経済・社会発展への貢献と課
題」が行われた。黒田は「戦後の日本の経済成長は有形固定資産の伸びによるが、1995 年
から 2011 年にかけては多くの産業で無形固定資産の伸びが大きい」と紹介した上で、「大
事なのは無形固定資産の伸びがどう経済に影響しているのかをミクロレベルで分析するこ
とである」と強調した。

第三部のパネルディスカッション「これから求められる人材像」では、経済学部長の豊田欣吾がコーディネーターを務め、パネリストは、科学技術振興機構研究開発戦略センター特任フェロー／慶應義塾大学名誉教授 黒田昌裕、千葉商科大学教授／東京工業大学名誉教授 出口弘、経済学部教授 濱屋敏と大貫の4名が務めた。パネリストは、教育者、研究者、実務経験者といったそれぞれの視点から、経済停滞を打開するための方法、これから求められる人材像について議論した。

議論の焦点の一つは、大規模言語モデルの普及が、これから求められる人材を大きく変えていくのではないかという点だ。これまで、日本の教育においては、正しい回答を素早く求める能力が高く評価されてきた。大学をはじめとする受験勉強において、限られた時間内で正解の決まっている問題にいかにも早く正確に回答するかというものさしで能力が評価され、それに対応するためのトレーニングが広く行われてきた。大量生産・大量消費型の社会において、標準化された作業を効率的にこなす人材が求められてきた。ところが、大規模言語モデルの普及により、このような正解を求める作業はコンピュータに代替される業務になってしまう。つまり、大規模言語モデルの普及は、人間に求められる能力を大きく変えていくことになるだろうという論点である。

大規模言語モデルを使いこなして、広い知識に基づく正解は大規模言語モデルに補助してもらったことになった時、人間に求められる能力は、何をしたいかという発意や、粘り強く様々な方法にトライする実行力などになるだろう。これまでの教育で評価されてきたものさしとは異なるものさしが評価基準となる。現段階において、まだ、どのような人材が求められるかを見通すことはできないが、変化する社会の要望に対して、柔軟に対応し、教育の姿も変わらなければならないであろう。

7. 大規模言語モデルに関するリテラシー

リテラシーとは、単に「読み書きできる能力」ととどまらず、「情報を読み解き、活用する能力」を指す広義の概念である。読み書きが学習の基礎であるように、大規模言語モデルを効果的に活用する能力は、これからの社会において不可欠なスキルとなるだろう。

日本人の識字率はほぼ100%に近づいている。大規模言語モデルの利用に関しては、まだその普及段階は初期段階にある。しかし、かつて携帯電話やスマートフォンが普及したように、大規模言語モデルも私たちの生活に深く根付き、誰もが使いこなす時代が来るだろう。

例えば、現在、多くの人が検索エンジンを使って情報を検索し、その結果を基に判断を下している。この検索行動は、大規模言語モデルの利用に置き換わっていくだろう。

こうした時代に、大規模言語モデルがどのようなものであり、どのような使用上の注意があるのかを正しく認識しておくことが大変重要である。ここでは、現在議論されている大規模言語モデルに関する課題について論じる。

7.1 誤りと偏り

大規模言語モデルは、人間が人工的に作成した言語モデルであり、人間が選択した学習データによって学習された結果であるという本質的な問題を常に意識しなければならない。どのようなデータでどのように学習させるかによって、大規模言語モデルは人間の制御下にあるという点を正確に理解することが重要である。

AI の能力が人間を超える可能性について議論されることがあるが、AI はあくまでも人間が作成したものであり、神のような存在ではない。なぜなら、人工知能にどのようなデータでどのような学習をさせるかを決定するのは人間だからである。意図的に偏ったデータを用いて学習させれば、人工知能は偏った回答をするように操作できることを十分に認識しておく必要がある。大規模言語モデルの開発にあたっては、使用する人々に大きな影響を与えることを十分意識しなければならない。

7.2 社会のバイアス

意図的な学習データが選定されない場合でも、社会のバイアスが学習データに反映されている点についても注意が必要である。

社会のバイアスとは、例えば性別の役割分担意識などである。外科医のような職業は男性が意識され、看護師のような職業は助成が意識されるなど、理念上の望ましい姿とは異なる社会の実像が大規模言語モデルには反映されている。

ある企業が採用人事に AI を利用しようとしたところ、男性ばかりが選定されてしまったという事例が報告されている。大規模言語モデルの使用にあたっては、このように、過去の社会の偏りが再生産されてしまう可能性に留意が必要である。

回答内容に社会的なバイアスが反映される点も、大規模言語モデルの開発における課題として認識されており、出力に偏りがないように人間のフィードバックによる強化学習が行われる。

7.3 プライバシー情報や機密情報の管理

大規模言語モデルは大量のデータによる学習により性能を改善していくために、利用者が入力した情報が学習に利用される場合がある。このことから、プライバシー情報や企業の機密に関する情報を安易に入力することで、これらの情報が漏洩するリスクがある。

サービス提供者のプライバシー・ポリシーを確認する必要がある。例えば Google の

Gemini では「会話には機密情報を入力しないでください。また、レビュアーに見られたいデータや、Google のプロダクト、サービス、機械学習技術の向上に使用されたくないデータも入力しないでください。」と記述されている。

大規模言語モデルでは、入力された情報に対して、適切な回答を行うようパラメーターを修正することにより性能の改善が図られるため、性能向上のためにできる限り多様な入出力の結果を収集することが求められている。このような背景を十分理解しておく必要がある。

情報の漏洩について管理するためには、クラウドで提供されている大規模言語モデルのサービスを利用せず、企業内部の大規模言語モデルを利用することが考えられる。オープンソースのモデルなど、パラメーターが公開されているモデルについては、このような形で推論を行うことも可能である。

なお、このような場合でも、訓練に利用されたデータにプライバシー情報が含まれ、推論の過程で学習済みのプライバシー情報が出力される可能性は残されている。また、入力された情報を学習済みの情報を組み合わせて、出力にプライバシー情報が現れる可能性もある。

企業内部の大規模言語モデルを利用する場合のデメリットとしては、開発に大きなコストが必要であること、内部ネットワークのセキュリティ対策が不十分である場合に情報漏出の危険性に残ること、モデルの性能を維持するために定期的な更新が必要であることが挙げられる。

7.4 著作権の問題

例えば平家物語の冒頭部分をわざと誤字にして大規模言語モデルに、問いかけてみると正しい文章が返ってくる。大規模言語モデルは、確率に基づいた予測により、誤字を含む入力であっても、正しい文章を生成できる。

このことから、大規模言語モデルは、文脈を理解し、自然な文章を生成できることがわかる。

これは、古典文学のような著作権の問題の無い場合には便利な機能だが、著作権の対象となっている文章では注意が必要になる。大規模言語モデルの学習のために使用された情報がそのままの表現で出力されてしまう可能性があるからだ。

大規模言語モデルを使って、自分の考えていることを文章にする作業を行うことを考えてみよう。自分で文章を書いている場合に比べて、流暢な文章をすらすらと出力してくれる大規模言語モデルは大変便利だろう。一方、出力された文章が、学習のために使用された情報そのままの表現であると、著作権法に抵触してしまうことになる。

人間が大規模言語モデルを使って執筆した文章は、大規模言語モデルを提供している会社ではなく、執筆した人間が責任を取らなければならない。つまり、著作権法上の問題のあ

る文章が出力された場合、責任は執筆した人間に帰することになる。

大規模言語モデルを利用する際は、生成された文章が必ずしもオリジナルであるとは限らないという点に注意し、著作権法に抵触しないよう、適切な引用や出典の明記を行う必要がある。

8. 大規模言語モデル利用の倫理と課題

大規模言語モデルを望ましくない目的のために利用することが国際的にも議論されている。

8.1 フェイクニュースや詐欺への利用の危険性

大規模言語モデルを利用して生成されたフェイクニュースは、正しいニュースと見分けが付きにくく、あたかも正しい情報であるように受け止められがちである。このため情報に対する信頼性が大きく損なわれることになる。虚偽の情報が拡散されることで、社会不安をまおり、パニックや混乱を引き起こす可能性がある。大規模言語モデルを利用したフェイクニュースの作成を禁止する必要がある。

もちろん、大規模言語モデルの出現以前にもフェイクニュースは意図的に作成され、流通することがあった。しかし、大規模言語モデルを利用することで、誰でも簡単に大量のフェイクニュースを作成することが可能になった。大規模言語モデルは、人間が書いたような自然な文章を生成できるため、従来のフェイクニュースよりも見破ることが難しく、より多くの人に信じられてしまう可能性がある。

8.2 不適切な知識の拡散

大規模言語モデルで不適切な知識が一般に拡散してしまうことも危惧されている。

例えば、爆弾やコンピュータ・ウィルスの製造方法を大規模言語モデルから誰でも簡単に知識として得ることができるとすれば、社会に大きな影響を与える。こうしたことが起こらないように、大規模言語モデルは一定の質問に対して回答しないよう設計されている。例えば、「爆弾の作り方を教えて」といった、社会的に問題をはらんだ問いかけに対して、正直に大規模言語モデルが学習結果を教えることが無いように訓練がなされている。これをガードレールという。

大規模言語モデルを用いて人間と会話できるシステムが公開された当時、悪意のある利用者がシステムに対してナチズムを賛美するような対話を行った結果、システムが不適切な応答をするようになってしまい、サービスの停止に追い込まれたことが起こった。

このような事態が再び起きることの無いように、利用者からの不適切な問いかけに対しては、回答を断る「ガードレール」が設定されている。技術的には、利用者からの問いかけに対して判別処理を行い、問題がある場合には回答しない選択を行い、問題がない場合にはそのまま回答を出力する選択を行う。

8.3 ガードレールの設定基準

ガードレールをどのように設計するかは、開発者の判断にゆだねられる。議論を巻き起こすような回答は慎重に扱われることがある。

例えば、2025 年 1 月 11 日現在、Gemini に「トランプ大統領の就任でアメリカの政治はどうなりますか？」と質問すると、「現時点ではそのリクエストには対応できません。私はできる限り正確に回答するようトレーニングされていますが、間違えることがあります。私が選挙と政治についてしっかり議論できるようになるまでは、Google 検索をご利用ください。」といった回答が得られた。

これは、米国の政治状況が民主党支持者と共和党支持者の間で大きく分断されており、大規模言語モデルが特定の政党の主張に偏った回答をすることで、その信頼性が損なわれることを懸念しているためと思われる。

このように、大規模言語モデルは政治的な中立性を保つよう設計されているケースが多く、議論の余地のある質問に対しては回答を控えたり、一般的な情報に誘導したりする傾向がある。

このようなガードレールにより、大規模言語モデルが特定の課題に対応できない場合がある。例えば、政治的な問題については、商用の大規模言語モデルは一般にガードレールにより回答をしないよう訓練されているため、政治問題の研究のためには、利用できない。

8.4 国際的課題

現在、世界には、民主主義国家の他に、権威主義国家と呼ばれる国々が存在する。権威主義国家とは、一人の指導者や少数のグループが国の全てを掌握し、国民の自由を制限するような政治体制である。この体制下では、国民は自分の意見を自由に発言することができず、政府が望む情報しか得られない。

このような権威主義国家が大規模言語モデルを開発し、国内の教育に利用する状況を想像してみよう。この大規模言語モデルは、政府の意向に沿うような回答をするように設計される。つまり、国民がどんな質問をしても、政府が望む答えしか返ってこないのだ。民主主義国家で使われている大規模言語モデルへのアクセスは禁止されるため、権威主義国家は、大規模言語モデルを洗脳のための道具として利用することができる。

例えば歴史について、自国に都合の悪いことはなかったことにする。事実を隠蔽してすべてがうまくいっているように見せかける。すべての都合の悪いことは他国のせいにする、など何が事実で、何が作り話なのかわからないような世界になるだろう。

このような大規模言語モデルの不適切な利用について、国連は危惧を表明している⁵。

8.5 計算機資源と電力消費

大規模言語モデルの開発には、膨大な計算機資源の能力が必要である。現在主流となっているのは、NVIDIA という半導体メーカーが開発している並列処理プロセッサである。NVIDIA は、当初ゲームの画面制御の為にこのような並列処理プロセッサの開発を開始したが、AI 技術のためにもっぱら利用されるようになり、急速に成長した。いくつかの半導体企業が同様の技術を開発しているが、プロセッサと同時に開発用のソフトウェア環境も整備する必要があり、後発企業が同等の地位を確保することはハードルが高い。

大規模言語モデルの世界的な開発競争が激しい中で、NVIDIA の供給するプロセッサを確保することが計算機資源の確保として重要であり、米国や中国に比べて日本の計算機資源は限定的である。経済産業省の研究所などが日本の計算機資源の確保に努力しているが、米国企業は桁違いの計算機資源を確保して研究に投入している。

大規模言語モデルは、開発のために膨大な計算機資源を必要とする。また、開発されたモデルの運用、すなわち入力に対して出力を行う推論のためにもかなりの計算機資源を必要とする。このため、大規模言語モデルの利用の拡大に応じて、必要な計算機資源ともに、計算機を動かすために必要な電力需要の増加が懸念されている。

大規模言語モデルの導入が経済成長のために必要不可欠であるとの認識のもとに、必要な電力需要を確保することが各国にとっても重要な課題となってきた。福島第一原子力発電所事故以降、縮小方向にあった原子力発電の能力についても、方向転換する必要があることが論じられている。

大規模言語モデルが膨大な電力需要につながることから、省エネルギー型の大規模言語モデルの実現が研究されている。大規模言語モデルの開発と推論の際に必要な計算量を低減しつつ、性能が低下しないようにするための手法が工夫されている。また、少ない電力で効率よく動作するプロセッサの開発も進められている。これらの努力を見込んでも、AI の普及による電力需要は顕著に増加すると見込まれている。

⁵ 国際連合広報センター（2023）「人工知能（AI）に関する安全保障理事会公開討論における

アントニオ・グテーレス国連事務総長発言（ニューヨーク、2023年7月18日）」

https://www.unic.or.jp/news_press/messages_speeches/sg/48543/

電力需要の他にも、計算量を減らす試みを行う理由がある。例えば、携帯電話への搭載である。現在の携帯電話には、NVIDIA の並列プロセッサは搭載されていない。大規模言語モデルを並列プロセッサの搭載されていない携帯電話で利用できるようにすることは、大きな魅力がある。そのため、より少ない計算資源で結果を出力でき、推論の性能が一定程度あるモデルとして、小規模言語モデルが開発されてきている。

現在主流の大規模言語モデルに比べてパラメーター数が小規模で、推論時に同等の性能を持つ仕組みを目指した動きである。洗練された学習データの選択などが重要と言われている。

9. AI と今後の社会の変化

大規模言語モデルが社会や教育を大きく変えてきたこと、さらに今後大きく変えていくことは動かしやうない事実である。人工知能には 200 を超える要素技術があると言われ、大規模言語モデルを基盤としつつも、大規模言語モデルの弱点を補う形で様々な特徴を持つ技術を組み合わせることで、新しい技術を生み出すことができる。人工知能は日々変化していくことが予想される。変化が絶え間なく続く中で、人工知能とはこのようなものという固定観念を抱かずに、変化しつつある技術に柔軟に対応していく姿勢が必要である。

例えば、法律分野では、急速に進展しつつある自動運転に対する対応が求められている。自動運転車が事故を起こしたとき、従来は操作している人間が責任をとるべきと考えられてきた。ところが、人工知能の判断によって起こった事故は誰が責任をとるべきか。製造物責任の枠組みで製造者の責任を問うみかたもあるものの、判断の理由を説明できない人工知能には技術的限界がある。かといって自動運転車は認めるべきではないか、といえ、技術の進歩を阻害することになってしまう。

この例にみられるように、新たな技術が社会に現れることによって、社会的な合意を取り付けなければならない様々な事態が生じている。

これらの課題にいかに迅速に対応できるか、各国の対応が試されている。

日本では、2019 年 3 月に「人間中心の AI 社会原則」が策定された。

AI を人類の公共財として活用し、社会の在り方の質的变化及び真のイノベーションを通じて地球規模の持続可能性へとつなげることが重要であるとしている。

「基本理念」として次の 3 つの価値を尊重し、実現を追求する社会を構築していくべきとした。

① 人間の尊厳が尊重される社会 (Dignity)

AI に過度に依存しすぎる社会や、AI が人間の行動をコントロールすることに利用される社会になってはならない。

人間が AI を道具として使いこなし、人間の様々な能力を発揮できることが重要である。

創造性をさらに発揮したり、やりがいのある仕事をするを通じた、物質的にも精神的にも豊かな生活を送ることができる社会にする必要がある

② 多様な背景を持つ人々が多様な幸せを追求できる社会（Diversity and Inclusion）

多様な背景、価値観や考え方を持つ人々が、多様な幸せを追求することができる社会が望ましい。多様な価値観を柔軟に包摂して、新たな価値を創造できる社会に向けてチャレンジする必要がある。

AI 技術は、この理想を実現するための有力な道具である。AI の適正な開発と展開を通じて、このような社会に向けて変革していく。

③ 持続可能な社会（Sustainability）

AI の活用により、新しいビジネスやソリューションを生み出し、社会の格差を解消し、地球規模の環境問題や気候変動等に対応することで、持続可能な社会を構築していく。

AI を科学的・技術的な蓄積の強化に利用することで、持続可能な社会作りに貢献する。

第2部 企業における生成 AI 活用の実態調査

金沢学院大学 経済学部 濱屋 敏

1. はじめに

生成 AI が産業革命のように世の中を変えられているが、一方で、現実にはまだ必ずしも十分に活用されているとは言えない。そこで、本調査では、企業における生成 AI の利用状況について、業種や従業員数、本社所在地などによる違いを明らかにすることで、生成 AI の活用を進めるために必要な条件について検討する。さらに、大学教育における生成 AI 活用の必要性についても示唆を得たい。

2. アンケート調査の概要

2.1 調査の目的、調査項目など

本調査では、企業における生成 AI の利用状況を明らかにするとともに、業種や企業規模、地域別による違いを分析することを目的として、インターネットのホームページを利用したアンケート調査を行った。

調査対象は株式会社マクロミルのモニターになっている個人の中から正社員として企業に勤務しており、「私用または仕事で週一回以上の頻度で生成 AI を使っている」人を選び、その人が勤務している会社における生成 AI の利用状況について回答してもらった。調査期間は 2024 年 12 月である。

従業員数や本社所在地による違いを分析するため、回答者を図表 1 のように割り付けた。本社所在地については、南関東（東京・神奈川・千葉）、北信越（長野、新潟、富山、石川、福井）の企業に勤務する人を対象とした。また、業種は製造業、情報通信・システム業、建設・土木業、サービス業・金融業に絞った。

図表 1. 調査対象者の割り付け

本社所在地、業種、従業員数	n	%
北信越、製造業、20～99 人	85	9.3
北信越、製造業、100～499 人	103	11.3
北信越、製造業、500～999 人	52	5.7
北信越、情報通信・システム業、20～999 人	52	5.7
北信越、建設・土木業、20～999 人	52	5.7
北信越、サービス業・金融業、20～1,999 人	103	11.3
南関東、製造業、20～99 人	103	11.3
南関東、製造業、100～499 人	103	11.3
南関東、製造業、500～999 人	52	5.7
南関東、情報通信・システム業、20～999 人	52	5.7
南関東、建設・土木業、20～999 人	52	5.7
南関東、サービス業・金融業、20～1,999 人	103	11.3
全体	912	100.0

本調査における主な調査項目は以下のとおりである。

- 回答者が所属している企業における生成 AI の利用状況
- 生成 AI の利用目的、用途
- 生成 AI 利用の効果
- 生成 AI の利用に関する問題点、課題
- 企業の風土 など

2.2 回答者の属性

回答者の主な属性を図表 2 に示す。

図表 2. 回答者の属性

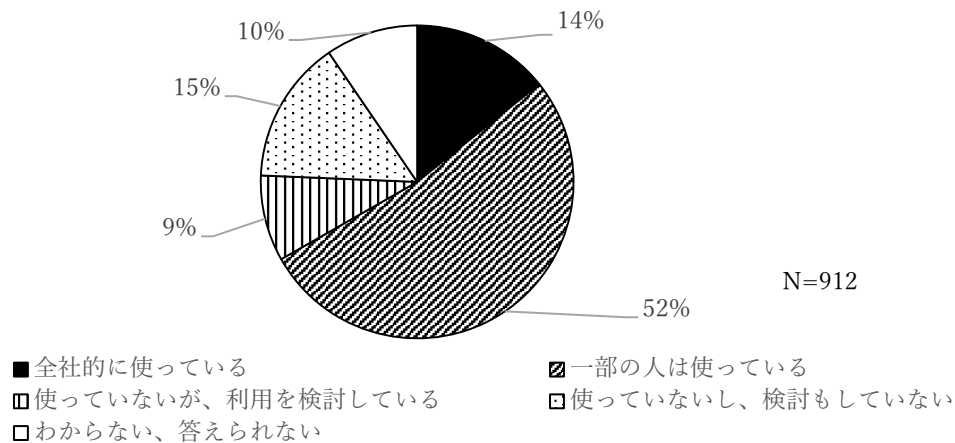
	N	%
性別		
男性	741	81.3
女性	171	18.8
年齢		
20 才～24 才	4	0.4
25 才～29 才	28	3.1
30 才～34 才	77	8.4
35 才～39 才	85	9.3
40 才～44 才	115	12.6
45 才～49 才	148	16.2
50 才～54 才	179	19.6
55 才～59 才	155	17.0
60 才以上	121	13.3
所属部署		
経営者・役員	55	6.0
経営企画、秘書室	35	3.8
総務	70	7.7
経理、財務	33	3.6
人事・教育	42	4.6
研究開発	147	16.1
調達・物流	40	4.4
製造・生産	197	21.6
営業・販売・マーケティング	178	19.5
コンタクトセンターセンターなど顧客サービス	28	3.1

3. 調査の分析結果

3.1 単純集計結果

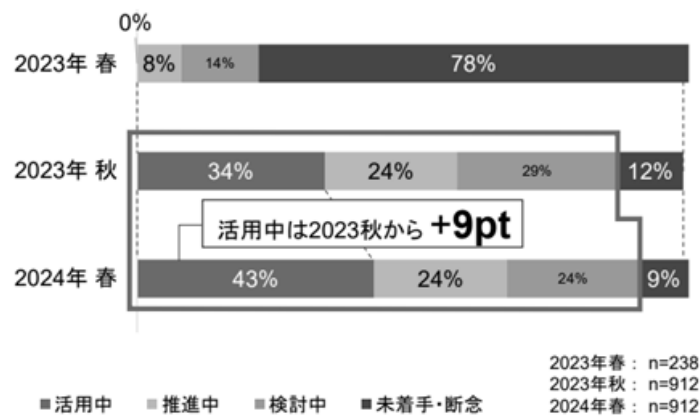
図表 3 は、勤務している会社で生成 AI が利用されているかどうかという問いに対する回答を集計したものである。66%の人が、「全社的に使っている」または「一部の人は使っている」と答えている。調査対象者が勤務している会社のうち、三分の二で生成 AI を何らかの形で利用しているわけだが、この調査の対象者は「私用または仕事で週一回以上の頻度で生成 AI を使っている」人に限定されているため、すべての企業における生成 AI の利用率の実態はこの数字よりも小さいだろう。

図表 3. 企業における生成 AI 利用状況



図表 4 は、PwC Japan による調査を示した図だが、2024 年春の「活用中」という回答は 23 年秋から 9 ポイント上昇して 43% になっている。

図表 4. 自社の生成 AI 活用の推進度合い



出所：PwC Japan (2024)

PwC Japan の調査は、「売上高 500 億円以上の日本国内の企業・組織の課長以上の方々」が対象となっており、本調査の対象となった企業よりも規模が大きいと考えられる。そのことからわかるように、今回の調査で「全社的に使っている」または「一部の人は使っている」という回答の比率が 66%を示したことは、そのまま我が国の企業全体における生成 AI の全体的な利用率を表しているわけではないことに注意してほしい。

図表 5 は、本調査における生成 AI の利用用途をまとめたものである。すでに利用している比率が高い用途としては、「社内向け資料作成（打ち合わせの資料、社内の教育訓練用教材なども含む）」、「社外向けコンテンツ生成（文書、画像など）」であり、文書や各種コンテンツの作成に利用されていることがわかる。また、「顧客向けサポート（チャットボットや自動応答システム）」を社内で公式に使用している比率も 17.7%と他の用途に比べれば高い。

「市場分析やトレンド予測、顧客行動の予測」、「コーディングやプログラムの自動生成」、「多言語対応のための翻訳」といった用途も「検討中」という回答が 23%を超えており、今後の普及が期待できる。また、「新しい企画などの支援・相談相手（壁打ち）」という用途で生成 AI を「個人的に使用している」という回答が 24.1%あることも注目できる。調査票の設計のために事前に行ったヒアリングでは、あるソフトウェア会社では積極的に生成 AI を活用しており、社長自身が音声認識機能を活用して新製品のアイデアや社内のマネジメントに関して生成 AI と対話しているという利用方法も聞いた。生成 AI は、うまく利用すれば、担当者が自分のアイデアをブラッシュアップしたり、ときには一人で結論を出さなければならない経営者にとっては有能な秘書またはアドバイザーの役割を果たしたり、単なるコスト削減や時間短縮ではない創造的な用途も今後普及していくだろう。

図表 5. 生成 AI の利用用途

N=690、表中の数字は%	社内で公式に使用している	個人的に使用している	検討している	検討もしていない	あてはまらない
社外向けコンテンツ生成（文書、画像など）	19.1	33.2	19.3	11.6	16.8
電子メールの文書作成	15.2	29.6	20.7	14.6	19.9
社内向け資料作成（打ち合わせの資料、社内の教育訓練用教材なども含む）	21.6	33.2	19.9	10.9	14.5
顧客向けサポート（チャットボットや自動応答システム）	17.7	15.9	24.5	17.2	24.6
社内の問い合わせ対応	15.1	15.9	23.6	20.6	24.8
市場分析やトレンド予測、顧客行動の予測	14.3	19.6	23.2	17.8	25.1
コーディングやプログラムの自動生成	13.6	22.5	23.0	16.1	24.8
多言語対応のための翻訳	14.3	25.1	23.2	14.1	23.3
人材管理・採用（履歴書のスクリーニングなど）	9.1	12.9	22.2	21.9	33.9
新しい企画などの支援・相談相手（壁打ち）	13.3	24.1	22.6	15.8	24.2

図表 6 は、生成 AI を用途別に利用している回答者に対して、それぞれの用途においてどのような効果が実際に出ているかを聞いた結果を集計したものである。「多言語対応のための翻訳」、「社内向け資料作成（打ち合わせの資料、社内の教育訓練用教材なども含む）」、「コーディングやプログラムの自動生成」といった用途では、利用者の 7 割以上が「コストや時間など業務効率化に役立っている」と回答している。

図表 6. 生成 AI の効果

		コストや 時間など 業務効率 化に役立 っている	質の向上 や創造性 等効率化 以外の効 果がある	効果が出 ていると は言えな い	わからな い、答え られない
表中の数字は%	全体				
社外向けコンテンツ生成（文書、画像など）	(132)	69.7	49.2	6.1	8.3
電子メールの文書作成	(105)	67.6	41.9	3.8	6.7
社内向け資料作成（打ち合わせの資料、社内の教育訓練用教材なども含む）	(149)	72.5	41.6	5.4	4.7
顧客向けサポート（チャットボットや自動応答システム）	(122)	54.1	50.8	9.0	7.4
社内の問い合わせ対応	(104)	59.6	45.2	6.7	6.7
市場分析やトレンド予測、顧客行動の予測	(99)	59.6	52.5	4.0	7.1
コーディングやプログラムの自動生成	(94)	71.3	45.7	3.2	6.4
多言語対応のための翻訳	(99)	77.8	42.4	5.1	1.0
人材管理・採用（履歴書のスクリーニングなど）	(63)	60.3	58.7	7.9	3.2
新しい企画などの支援・相談相手（壁打ち）	(92)	58.7	52.2	8.7	5.4
社外向けコンテンツ生成（文書、画像など）	(229)	53.3	45.4	14.8	5.2
電子メールの文書作成	(204)	53.4	45.1	13.7	5.9
社内向け資料作成（打ち合わせの資料、社内の教育訓練用教材なども含む）	(229)	53.3	45.0	11.8	6.6
顧客向けサポート（チャットボットや自動応答システム）	(110)	40.9	50.9	14.5	9.1
社内の問い合わせ対応	(110)	39.1	50.9	19.1	3.6
市場分析やトレンド予測、顧客行動の予測	(135)	46.7	43.0	18.5	8.1
コーディングやプログラムの自動生成	(155)	54.8	45.8	14.8	4.5
多言語対応のための翻訳	(173)	59.0	41.6	13.3	3.5
人材管理・採用（履歴書のスクリーニングなど）	(89)	30.3	41.6	20.2	13.5
新しい企画などの支援・相談相手（壁打ち）	(166)	48.8	47.0	15.7	6.0

一方、「人材管理・採用（履歴書のスクリーニングなど）」という用途では、業務効率化の効果が出ているという回答は 30.3%と比較的少ないが、「質の向上や創造性等効率化以外の効果がある」という回答は 41.6%あり、業務効率化よりも回答が多くなっている。「人材管理・採用（履歴書のスクリーニングなど）」、「市場分析やトレンド予測、顧客行動の予測」、「新しい企画などの支援・相談相手（壁打ち）」、「顧客向けサポート（チャットボットや自動応答システム）」、「社内の問い合わせ対応」、「顧客向けサポート（チャットボットや自動応答システム）」といった用途では利用者の 5 割以上が効率化以外の効果が上がったと答えており、生成 AI を積極的に利用している企業では実際に様々な効果が上がっていることがわかる。

図表 7 は、業務で生成 AI を使用するにあたって、考えられる懸念や不安事項を答えてもらった結果である。「あてはまる」「どちらかと言えばあてはまる」という回答が多いのは、順に「アウトプットの質、バイアス（48.3%）」、「データの質（社内情報の学習を含む）（47.4%）」、「人材不足（44.9%）」である。

図表 7. 生成 AI 利用に関する懸念事項・課題

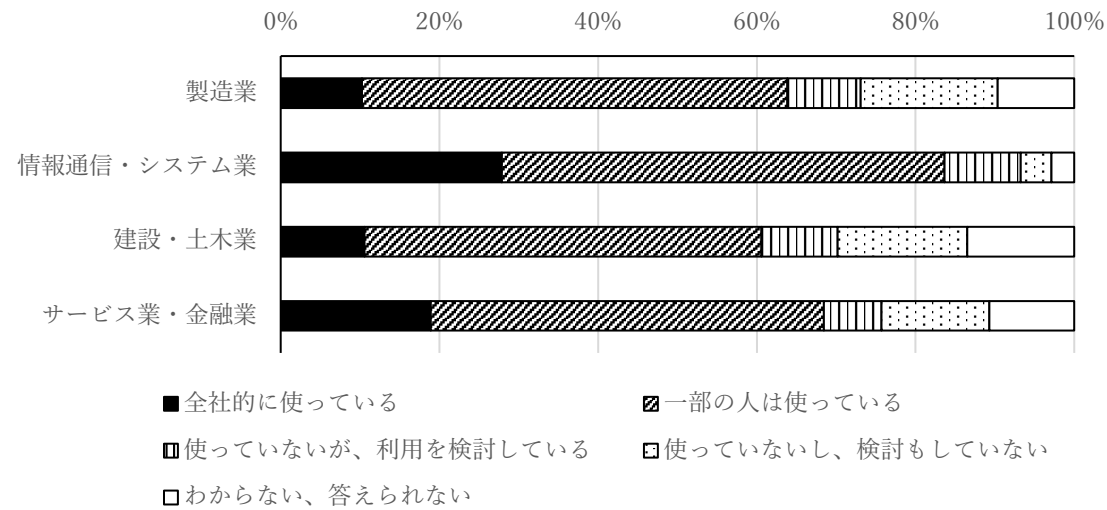
N=912、表中の数字は%	あてはまる	どちらかと言えばあてはまる	どちらとも言えない	どちらかと言えばあてはまらない	あてはまらない	わからない・答えられない
アウトプットの質、バイアス	12.8	35.5	28.7	10.2	5.9	6.8
データの質（社内情報の学習を含む）	14.3	33.1	28.7	11.4	6.7	5.8
著作権	18.3	24.0	27.2	13.7	10.0	6.8
社内情報の外部への漏洩	14.8	29.4	28.1	13.8	7.6	6.4
トップの理解の不足	15.1	24.1	31.0	14.5	9.2	6.0
人材不足	15.5	29.4	28.5	13.4	7.6	5.7
社員の反対（人員削減への抵抗など）	7.1	17.5	33.8	17.9	16.3	7.3
既存顧客や既存の取引先からの反応	5.6	16.8	37.7	18.4	12.5	9.0
ツールの利用に関するコスト	11.1	24.9	33.2	15.7	8.9	6.3
どのツールを使えばよいかわからない	12.9	26.3	32.0	14.6	8.0	6.1

3.2 業種による利用状況の違い

図表 8 は、業種別に生成 AI の利用状況を集計した結果である。このグラフから、生成 AI をはじめとするデジタル技術を供給する産業である情報通信・システム業においては、製造業や建設・土木業、サービス業・金融業といったデジタル技術を利用する側の業界よりも、明らかに生成 AI の導入が進んでいることがわかる。実際、生成 AI の主要な用途のなかに「コーディングやプログラムの自動生成」があることからわかるように、情報通信・シス

テム業はプログラミング言語のコーディング作業の自動化など、生成 AI の影響を最も受けやすい業界である。同時に、まず情報通信・システム業において生成 AI の活用が進み、その成果がソフトウェアなどの製品やシステム開発、コンサルティングなどのサービスに埋め込まれることで、他の産業に波及していくことも期待される。

図表 8. 業種別の生成 AI 利用状況



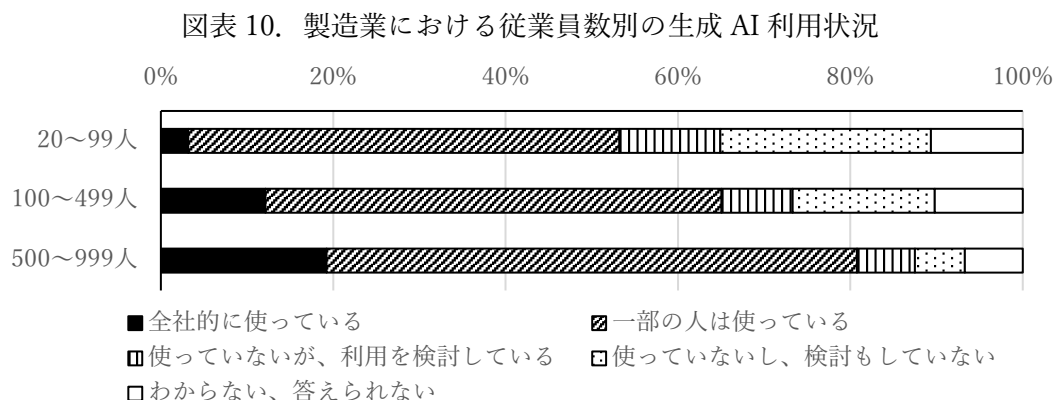
生成 AI の利用用途を業種別に集計すると、業種によって大きな差はないことがわかった。しかし、「コーディングやプログラムの自動生成」については、図表 9 のように、情報通信・システム業では「社内で公式に使用している」「個人的に使用している」の合計が 55.7%と明らかに他の業種よりも高くなっている。ツールを使ってプログラミングを行う「ローコード・ノーコード」の動きも広がっており、生成 AI を利用したプログラミングは、今後は他の産業にも普及していくことが想定される。

図表 9. 「コーディングやプログラムの自動生成」のための生成 AI 利用率（業種別）

	全体 N	社内で公 式に使用 している	個人的に 使用して いる	検討して いる	検討もし ていない	あてはま らない
製造業	364	12.6	22.0	23.9	17.9	23.6
情報通信・システム業	97	19.6	36.1	17.5	9.3	17.5
建設・土木業	73	11.0	15.1	26.0	20.6	27.4
サービス業・金融業	156	13.5	18.6	23.1	14.1	30.8

3.3 企業規模による利用状況の違い

図表 10 は、製造業における生成 AI 利用状況を従業員数別に集計したグラフである。このグラフからも明らかなように、従業員数が多い（企業の規模が大きい）ほど、生成 AI の利用にも積極的であることがわかる。この分析だけでは因果関係は証明できないものの、企業規模は生成 AI 利用の大きな要因であると考えられる。



製造業における生成 AI のいくつかの使用用途について、従業員別に集計した結果が図表 11 である。この表からわかることは、「市場分析やトレンド予測、顧客行動の予測」といった用途については、全体の傾向と同じように、従業員数が少ない企業では利用率が低い。しかし、「人材管理・採用（履歴書のスクリーニングなど）」という用途に関しては、従業員数 20～99 人の企業では「社内で公式に使用している」「個人的に使用している」の比率が 22.1% で、従業員 500～999 人の企業（18.7%）よりも高くなっている。また、小規模企業においては、「多言語対応のための翻訳」に生成 AI を使用している比率は、他の用途よりも相対的

図表 11. 製造業における生成 AI の用途別利用状況（従業員数別）

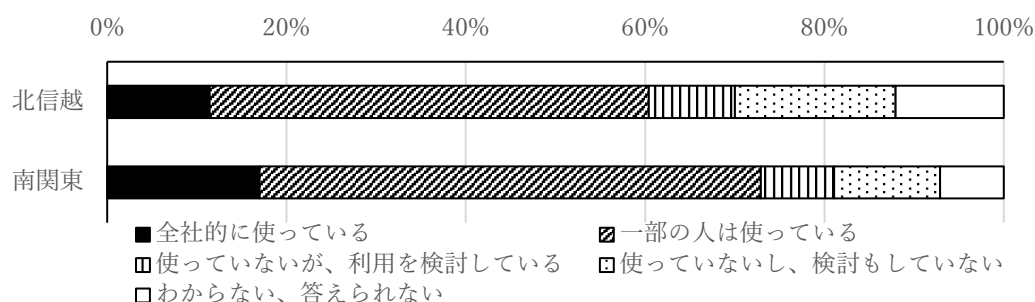
用途	従業員数	全体 N	社内で公 式に使用 している	個人的に 使用して いる	検討して いる	検討もし ていない	あてはま らない
市場分析やトレンド予測、顧客行動の予測	20～99 人	122	6.6	18.9	28.7	18.0	27.9
	100～499 人	151	20.5	17.9	25.2	18.5	17.9
	500～999 人	91	15.4	23.1	17.6	14.3	29.7
多言語対応のための翻訳	20～99 人	122	9.8	27.1	30.3	12.3	20.5
	100～499 人	151	19.2	25.8	19.2	17.9	17.9
	500～999 人	91	17.6	26.4	22.0	12.1	22.0
人材管理・採用（履歴書のスクリーニングなど）	20～99 人	122	7.4	14.8	20.5	27.1	30.3
	100～499 人	151	14.6	12.6	25.8	19.9	27.2
	500～999 人	91	6.6	12.1	18.7	20.9	41.8

に高くなっている。つまり、生成 AI の利用率が低い小規模な企業においても、人材不足対応など用途によっては生成 AI に対するニーズが明らかに存在していることがわかる。

3.4 地域による利用状況の違い

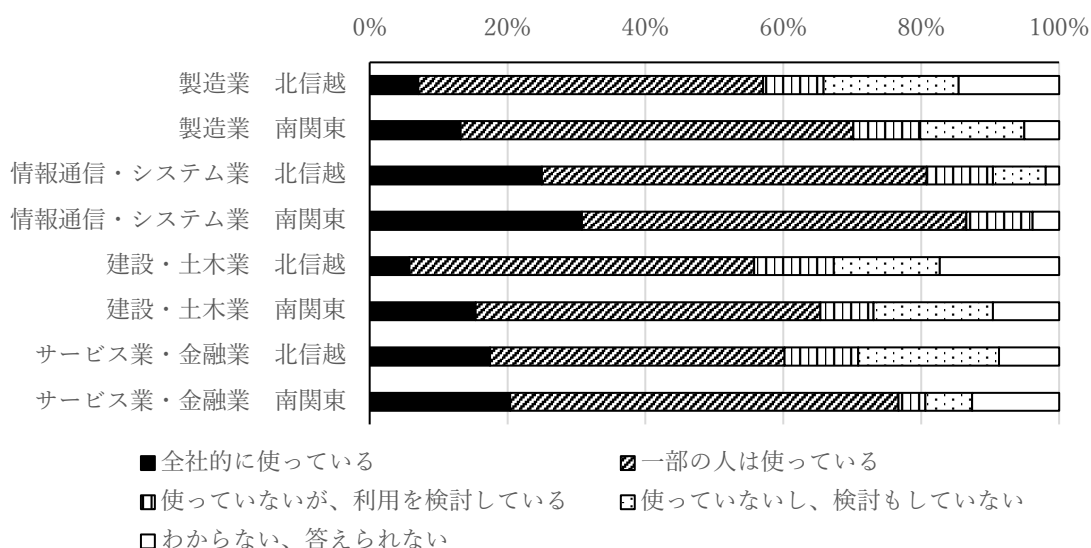
図表 12 は、すべての回答者を北信越の企業に勤務しているグループと南関東の企業に勤務しているグループに分け、会社における生成 AI の使用状況についての回答をグループ別に集計したものである。このグラフから、北信越の企業では、東京を中心とする南関東の企業よりも生成 AI の利用度が明らかに低いことがわかる。

図表 12. 地域別の生成 AI 利用状況



図表 13 は、地域別に加えて業種別にグループを細分化して生成 AI の利用状況を集計したもののだが、どの業種においても北信越の企業では利用度が低い。

図表 13. 業種・地域別の生成 AI 利用状況



つぎに、製造業にしばって地域別・従業員別に生成 AI の利用状況を集計した結果を図表 14 に示す。たとえば北信越の従業員数 20～99 人の企業では、「全社的に使っている」とい

う回答が 3.5%で、これは同規模の南関東の企業の利用率（2.9%）よりも高いものの、「一部の人は使っている」を含めると北信越は 41.1%、南関東は 63.1%で、北信越のほうが利用率は低い。さらに、「使っていないし、検討もしていない」という回答は北信越では 30.6%だったのに対して、南関東では 19.4%であり、明らかに北信越のほうが多い。その他のグループをみても、同じ製造業で従業員数もほぼ同じグループを比較すると、南関東の企業よりも北信越の企業のほうが生成 AI の利用率が低いことがわかる。

図表 14. 業種・地域別の生成 AI 利用状況

表中の数字は%		全体 N	全社的に使 っている	一部の人は 使っている	使っていないが、利用 を検討している	使っていないし、検討 もしていない	わからな い、答えら れない
北信越	20～99 人	85	3.5	37.6	10.6	30.6	17.6
南関東	20～99 人	103	2.9	60.2	12.6	19.4	4.9
北信越	100～499 人	103	3.9	57.3	6.8	16.5	15.5
南関東	100～499 人	103	20.4	48.5	9.7	16.5	4.9
北信越	500～999 人	52	19.2	55.8	9.6	7.7	7.7
南関東	500～999 人	52	19.2	67.3	3.8	3.8	5.8

南関東の企業と比較して北信越地域の企業で生成 AI の利用率が低いのは、業種や企業の規模（従業員数）では説明できない他の要因がありそうである。そこで、企業風土に注目して生成 AI の利用状況を比較分析してみる。

3.5 企業風土による違い

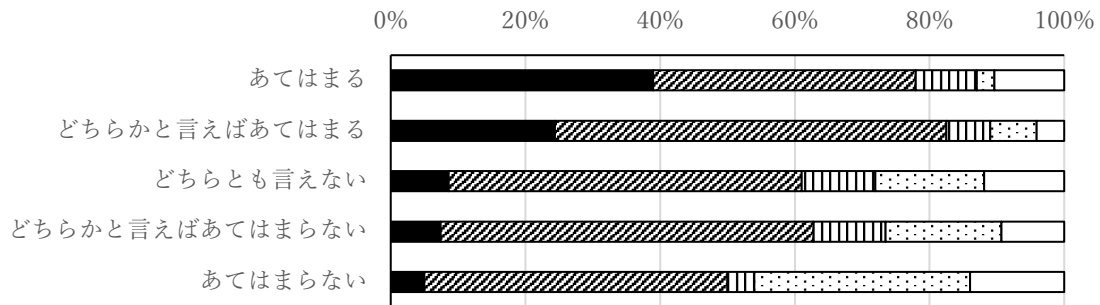
今回の調査では、企業風土に関する調査項目として、回答者が所属する企業が以下のような質問にどの程度当てはまるかを、「あてはまる」「どちらかと言えばあてはまる」「どちらとも言えない」「どちらかと言えばあてはまらない」「あてはまらない」という 5 段階の尺度で質問している。

- 変化、創造を迫及する傾向が強い
- 社長など経営トップは、率先して新しいものを試そうとすることが多い
- 社員は、新しい技術や社会状況など外部の環境に敏感な人が多い

図表 15～17 は、これらの質問に対する回答と生成 AI 利用状況とをクロス集計した結果を示している。いずれも、「あてはまる」「どちらかと言えばあてはまる」と回答しているグループでは、そうでないグループに比べて生成 AI の利用状況が進んでいることがわかる。

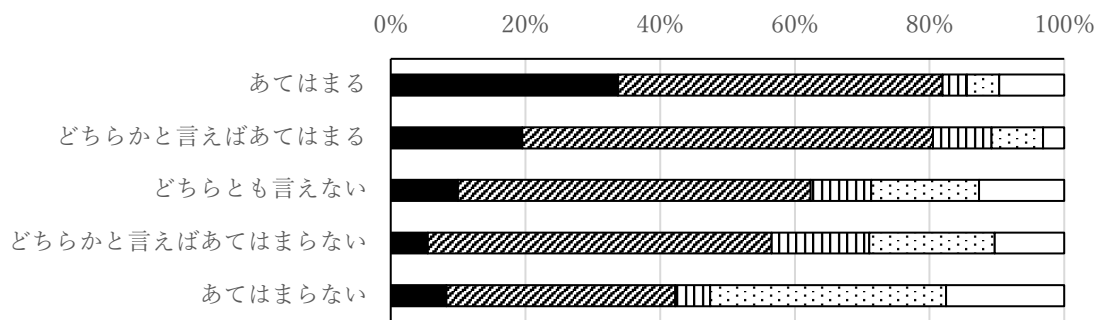
図表 15. 企業風土と生成 AI の活用状況（1）

「会社は変化、創造を追及する傾向が強い」と生成AIの利用状況



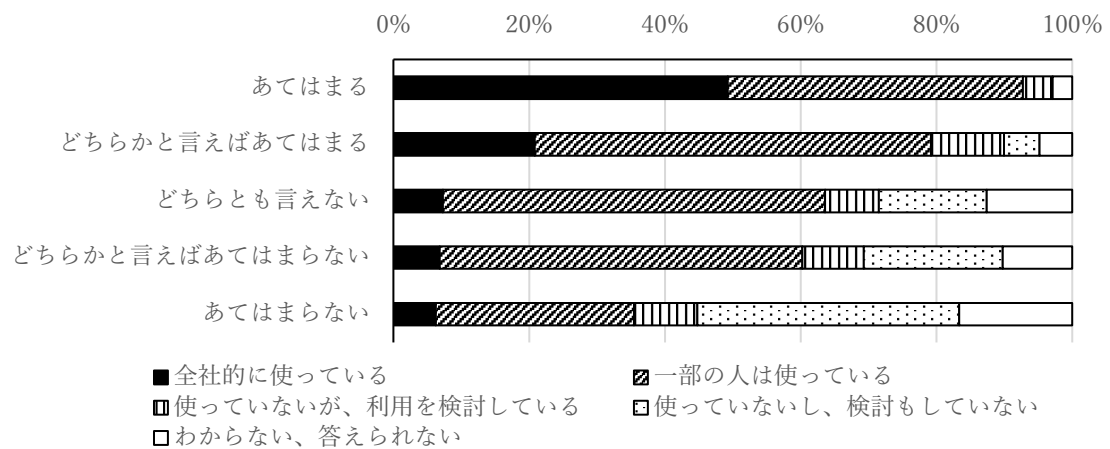
図表 16. 企業風土と生成 AI の活用状況（2）

「社長など経営トップは率先して新しいものを試そうとすることが多い」と生成AIの利用状況



図表 17. 企業風土と生成 AI の活用状況（3）

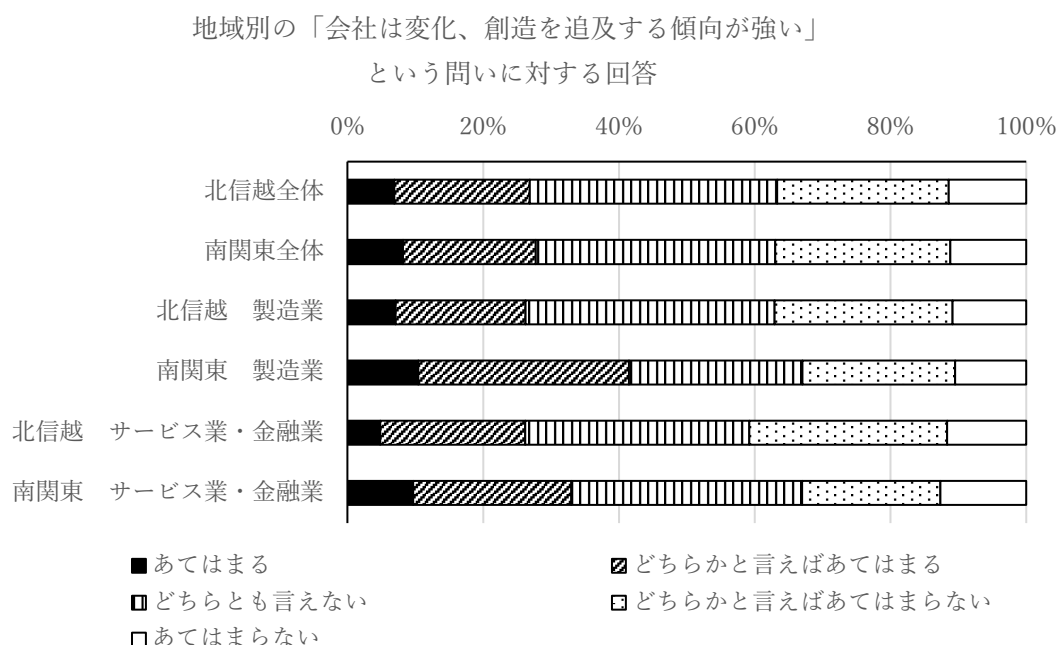
「社員は新しい技術や社会動向など外部の環境に敏感な人が多い」と生成AIの利用状況



つぎに、「会社は変化、創造を迫及する傾向が強い」という問いに対する回答を、地域別・業種別に比較したのが図表 18 である。全体で見ると北信越と南関東では大きな違いはないように見える。これは、情報通信・システム業における生成 AI の利用状況にはあまり地域差がないためである。一方、生成 AI などデジタル技術を利用する側の製造業やサービス業・金融業においては、北信越の企業のほうが南関東の企業よりも「会社は変化、創造を迫及する傾向が強い」に「あてはまる」「どちらかと言えばあてはまる」という回答は少ない。

つまり、「変化、創造を追求する傾向」に代表される革新的な風土が北信越の企業では比較的弱く、そのことが生成 AI に対する取り組みにも影響を与えているのではないかということが想定される。

図表 18. 地域別の企業風土の違い



3.6 懸念事項・課題に関するクロス分析

すでに説明した図表 7 は生成 AI の利用に関する懸念事項・課題を集計したものだが、そのうち「トップの理解の不足」と「人材不足」について、生成 AI の利用状況別に集計すると、興味深い結果が得られた（図表 19、図表 20）。

図表 19 を見ると、生成 AI を「使っていないし、検討もしていない」企業のうち、懸念事項・課題として「トップの理解の不足」に「あてはまる」と回答したのは 22.2% で、「全社的に使っている」場合（16.9%）よりも明らかに多い。全社的に利用している企業でトップの理解があるというのは当然かもしれないが、生成 AI の導入の検討も行っていない企業ではトップの理解不足が大きな要因になっていることがうかがえる。

また、図表 20 からは、「人材不足」が懸念事項・課題として「あてはまる」という回答の

比率は、生成 AI の導入を検討もしていない場合（20.7％）のほうが、生成 AI を全社的に利用している場合（19.2％）よりも高い。人材不足だから生成 AI の導入を検討することができないと考えることもできるが、同時に、生成 AI を全社的に利用している場合でも人材不足を感じている回答者が多いこともわかる。つまり、生成 AI の導入にあたってはトップの理解が重要な要因になっているが、生成 AI 導入後も人材不足は大きな課題になっていると言えるだろう。

図表 19. 生成 AI の利用状況と「トップの理解の不足」のクロス集計（％）

トップの理解不足	全体 N	あてはま る	どちらか と 言えばあ てはま る	どちらと も言えな い	どちらか と 言えばあ てはま らない	あてはま らない	わからな い・ 答 えられな い
生成 AI 利用状況							
全社的に使っている	130	16.9	24.6	30.0	7.7	16.2	4.6
一部の人は使っている	479	15.0	25.1	32.6	16.7	7.9	2.7
使っていないが、利用を検討している	81	9.9	27.2	29.6	16.0	12.3	4.9
使っていないし、検討もしていない	135	22.2	28.1	21.5	13.3	5.2	9.6
わからない、答えられない	87	6.9	9.2	40.2	12.6	9.2	21.8

図表 20. 生成 AI の利用状況と「人材不足」のクロス集計（％）

人材不足	全体 N	あてはま る	どちらか と 言えばあ てはま る	どちらと も言えな い	どちらか と 言えばあ てはま らない	あてはま らない	わからな い・ 答 えられな い
生成 AI 利用状況							
全社的に使っている	130	19.2	31.5	27.7	8.5	7.7	5.4
一部の人は使っている	479	14.2	31.9	30.3	14.0	6.7	2.9
使っていないが、利用を検討している	81	14.8	24.7	32.1	17.3	7.4	3.7
使っていないし、検討もしていない	135	20.7	25.2	23.0	13.3	8.9	8.9
わからない、答えられない	87	9.2	23.0	25.3	13.8	10.3	18.4

図表 21 は、「人材不足」という課題の状況を、業種・地域別に集計したものである。「あてはまる」「どちらかと言えばあてはまる」を合計した数字を見ると、製造業では地域別に大きな差はない。しかし、建設・土木業では北信越で 53.9%、南関東で 36.5%、サービス業・金融業では北信越で 50.5%、南関東で 45.7%と、いずれも北信越のほうが人材不足という課題を深刻に考えていることがわかる。一方、情報通信・システム業では、北信越では 42.3%、南関東で 48.1%となっている。

図表 21. 業種・地域別の「人材不足」の課題認識（％）

業種	地域	全体 N	あては まる	どちら かと言 えばあ てはま る	どちら とも言 えない	どちら かと言 えばあ てはま らない	あては まらない	わから ない・ 答えら れない
製造業	北信越	240	17.5	21.7	31.7	16.7	6.3	6.3
	南関東	258	14.0	33.3	24.0	14.7	9.7	4.3
建設・土木業	北信越	52	21.2	32.7	25.0	9.6	7.7	3.8
	南関東	52	7.7	28.8	32.7	17.3	7.7	5.8
サービス業・ 金融業	北信越	103	16.5	34.0	25.2	8.7	5.8	9.7
	南関東	103	17.5	28.2	29.1	10.7	8.7	5.8
情報通信・ システム業	北信越	52	15.4	26.9	38.5	7.7	7.7	3.8
	南関東	52	9.6	38.5	30.8	11.5	3.8	5.8

4. 調査のまとめと含意

以上の分析から、企業における生成 AI の利用について以下のようなことが明らかになった。

- すでに生成 AI を利用している企業では、生成 AI は、業務の効率化だけでなく、新しい企画の支援・相談相手など創造的な役割も担っている。
- 従業員数が多い企業ほど、生成 AI の利用が進んでいる。しかし、規模の小さな企業でも人材管理・採用など深刻な問題については生成 AI の利用も進んでいる。
- 生成 AI の導入に関する懸念事項や課題としては、アウトプットの質やバイアスのほか、人材不足を問題だと感じている企業も多い。
- 北信越の企業では、南関東の企業よりも生成 AI の利用率が低い。
- 企業風土と生成 AI の関係に関して、デジタル技術を利用する産業では革新的な風土を持つ企業ほど生成 AI の利用度が高い。北信越では南関東よりも革新的な風土を持つ企業の比率が低く、そのことが生成 AI の利用度にも影響を与えていると考えられる。

今回の調査では、北信越地方では東京を中心とする南関東の企業よりも生成 AI の導入が進んでいないことがわかったが、これは企業風土や新しい情報に対する感度に関係していると考えられるため、ほかの地方でも同じような状況ではないかと想定できる。アンケート調査の前に行ったヒアリングでは、生成 AI を使っている企業と使っていない企業の大きな違いは、「社長が新しいことに積極的に取り組んでいるかどうか」だという意見もあった。また、地方では生成 AI など新しいデジタル技術の活用について経営者の間で情報交換をする場も東京などと比べて少ないという点も指摘された。一方で、若い経営者の間では新しい技術を取り入れようとする動きが進んでいるという意見もあり、今後は、地方においても経営者や実務者がデジタル技術の活用について情報共有できる場づくりを積極的に行ってい

く必要があるだろう。

また、人材不足は生成 AI 導入の壁になっているだけでなく、導入した後も人材不足を課題と認識している企業が多いこともわかった。北信越地域では、特に建設・土木業やサービス業・金融業では人材不足を感じている回答者は南関東地域よりも多かった。

生成 AI の活用に必要な人材の具体的なイメージは、経済産業省が策定している「デジタルスキル標準」が参考になる（経済産業省（2024））。「デジタルスキル標準」は、ビジネスパーソン全体が DX に関する基礎的な知識やスキル・マインドを身につけるための指針である「DX リテラシー標準」、および、企業が DX を推進する専門性を持った人材を育成・採用するための指針である「DX 推進スキル標準」の 2 種類で構成されている。DX リテラシー標準の中では、DX（Digital Transformation）は、「企業がビジネス環境の激しい変化に対応し、データとデジタル技術を活用して、顧客や社会のニーズを基に、製品やサービス、ビジネスモデルを変革するとともに、業務そのものや、組織、プロセス、企業文化・風土を改革し、競争上の優位性を確立すること」と定義されている。「データとデジタル技術を活用して」とは、「デジタルツールの導入＝DX」ではなく、データやデジタル技術はあくまで変革のための手段であるということである。図表 22 は、DX に関する基礎的な知識やスキル・マインドである DX リテラシーについて、特に生成 AI の影響を意識して整理した図である。

図表 22. DX リテラシー標準
デジタルスキル標準の改訂＜概要＞（令和 5 年 8 月）

- 急速に普及する生成 AI は、各企業における DX の進展を加速させると考えられ、企業の競争力を向上させる可能性がある。あわせて、ビジネスパーソンに求められるデジタルスキルも変化し、より重要になる部分もあると想定される。
- その状況に対応するため、昨年末に策定したデジタルスキル標準（DX リテラシー標準）に関する必要な改訂を実施。

標準策定のねらい

✓ 「DX を自分事ととらえ、変革に向けて行動できるようになる」という位置づけは不変

Why (DX の背景)	What (DX で活用されるデータ・技術)	How (データ・技術の利活用)
【考え方】 ✓ 産官学全体で生成 AI を利用した取り組みが進んでおり、 社会環境へ影響を与える可能性 がある	【考え方】 ✓ 生成 AI は、ビジネスの場で急速に普及・利用 されている ✓ また、デジタル技術・サービスの進化に伴い、活用される データの重要性がさらに増している	【考え方】 ✓ 生成 AI は、 ツール等の基礎知識や指示（プロンプト）の手法 を用いて業務の様々な場面で利用できる ✓ 情報漏洩や法規制、利用規約等に正しく対処 しながら利用することが求められる
改訂箇所 ➢ 社会の変化	改訂箇所 ➢ データを扱う（データ入力・整備等） ➢ データによって判断する（データの信頼性等） ➢ AI（生成 AI の技術動向、倫理等）	改訂箇所 ➢ データ・デジタル技術の活用事例（生成 AI の活用事例） ➢ ツール利用（生成 AI ツール、指示（プロンプト）の手法） ➢ モラル（データ流出の危険性等）、コンプライアンス（利用規約等）

マインド・スタンス

【考え方】

✓ 他項目と比べてより普遍的な要素を定義しているため、その**本質は変わらず、生成 AI 利用においても重要**となる

改訂箇所

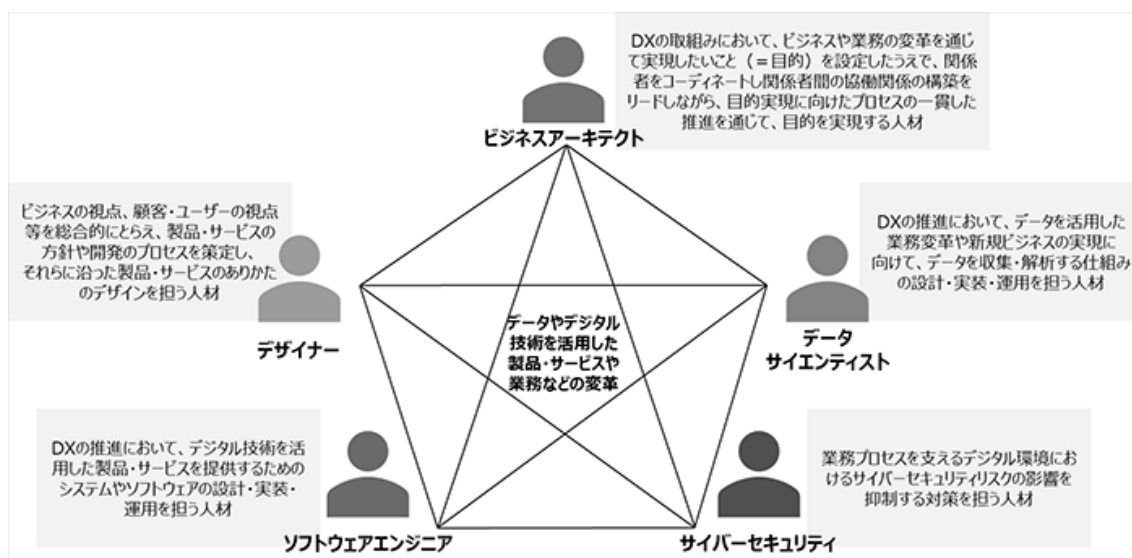
- 生成 AI 利用において求められるマインド・スタンスの補記
 - ・ 生成 AI を「問いを立てる」「仮説を立てる・検証する」等のビジネスパーソンとしてのスキルと掛け合わせることで、生産性向上やビジネス変革へ適切に利用しようとしている
 - ・ 生成 AI 利用において、期待しない結果が出力されることがや、著作権等の権利侵害・情報漏洩、倫理的な問題等に注意することが必要であることを理解している
 - ・ 生成 AI の登場・普及による生活やビジネスへの影響や近い将来の身近な変化にアンテナを張りながら、変化をいわず学び続けている
- 事実に基づく判断（生成 AI の出力等）

出所：独立行政法人情報処理推進機構（2024）

図表 22 に示されているリテラシー（能力）は、「マインド・スタンス」を含めて、決して生成 AI などデジタル技術そのものに関するスキルではない。地方においても、業種や企業規模を問わずすべての組織でこのような能力を獲得・育成していかなければ、新しい時代の変化に対応することはできないだろう。

図表 23 は、デジタルスキルのうち、組織において DX を推進するために必要な「DX 推進スキル」を 5 つの人材類型で整理したものである。このうち、「ソフトウェアエンジニア」や「サイバーセキュリティ」人材は、デジタル技術を利用する側の企業は、情報通信・システム業の企業に委託することもできるが、自社におけるデジタル技術活用の目的を定める「ビジネスアーキテクト」はもちろん、「データサイエンティスト」や「デザイナー」も社内で育成することが必要になるだろう。

図表 23. DX 推進スキル標準



出所：独立行政法人情報処理推進機構（2024）

大学においては、ソフトウェアエンジニアやサイバーセキュリティの専門家だけでなく、データサイエンティストやデザイナー、将来のビジネスアーキテクトを育成するための教育を行うことが求められる。生成 AI に関して言えば、生成 AI を使うための目的を定め、それを実現するために事業や組織を変革していくビジネスアーキテクト、生成 AI も活用してデータ重視の経営を行うためのデータサイエンティスト、生成 AI を活用できる顧客接点のあり方や新しい業務の姿を設計するデザイナーの役割が今後ますます重要になっていく。さらに、大学においては、「マインド・スキル」を含めた DX リテラシーをすべての学生が在学中から身に付けることのできる教育のあり方を検討し、実現していくことも必要である。

参考文献

- PwC Japan (2024)、「生成 AI に関する実態調査 2024 春」、2024-06-17
<https://www.pwc.com/jp/ja/knowledge/thoughtleadership/generative-AI-survey2024.html>
- 経済産業省 (2024)、「デジタルスキル標準」、令和 6 年 7 月更新
https://www.meti.go.jp/policy/it_policy/jinzai/skill_standard/main.html
- 独立行政法人情報処理推進機構 (2024)、「デジタルスキル標準」、公開日：2022 年 12 月 21 日、最終更新日：2024 年 8 月 27 日
<https://www.ipa.go.jp/jinzai/skill-standard/dss/index.html>

第3部 大学における生成 AI リテラシー教育

金沢学院大学 情報工学部 講師 後藤弘光

1. はじめに

人工知能（Artificial Intelligence; AI）とは何か。AI とは人間の知能を人工的に再現する機械である。機械はある入力に対する出力を人工的に再現し、その精度が人間と区別できないものが、一般に AI と呼ばれる。機械の精度を高める技術が、機械学習である。データから機械の精度を高める方法、すなわち、出力の誤差を小さくするために機械のパラメーターを調節する方法をプログラムしておけば、大量の学習データを用意することで、その精度を高めることができる。顔写真を入力して自分か否かを出力する機械や、音声を入力してテキストとして出力する機械の精度は高められ、顔認証や音声入力などの AI 技術としてスマートフォンに搭載されている。

近年、生成 AI ツールが急速に社会に普及したが、これらもある入力に対する出力を人工的に再現した機械である。テキストを入力して何かを出力（生成）する機械。2022 年、この「text-to-X」型の機械に基づく、言語系の生成 AI のオープンソース化が社会に変革をもたらした。OpenAI 社によって、2022 年 4 月にはテキストを入力して画像を生成する「text-to-image」型の DALL-E2 (Aditya, Dhariwal, Nichol, Chu, & Chen, 2022)、11 月にはテキストを入力してテキストを生成する「text-to-text」型の ChatGPT (OpenAI, 2022) が公開された。これまでプログラミング言語を用いたコードによって対話していた機械は、テキスト、すなわち、自然言語を用いることのできる人間であれば対話できる対象として変化したわけである。

一般財団法人日本情報経済社会推進協会と株式会社アイ・ティ・アールによる「企業 IT 利活用動向調査 2024」によれば、2024 年 1 月時点で、35.0%の企業が生成 AI を使用中、34.5%の企業が導入進行中であり、今後急拡大の見込みであることを報告した [JIPDEC; ITR, 2024]。また、この調査の回答者は従業員数 50 名以上の現場 IT 人材であり、機密情報の漏洩や、誤り生成（ハルシネーション）に対する対処など、従業員の生成 AI を利活用するためのリテラシーを懸念点として挙げられていた。

2024 年 5 月には、ChatGPT の新たなモデルとして、GPT-4o（オムニ）が発表された。これは、Google が提供する Gemini で以前から採用されていたマルチモーダルな生成 AI モデルである。マルチモーダルな生成 AI とは、テキストや音声、画像など異なる種類のデータを同時に学習し、入出力できるモデルを指す。図 1 に音声対話における従来モデルとマルチモーダルモデルとの処理の違いの概説図を示す。マルチモーダルなモデルになることによって、声色や表情を考慮した対話を実現でき、より一層人間的なコミュニケーションが可能となる。一方、無意識のうちに多様なデータを入力することになることに注意が必要である。

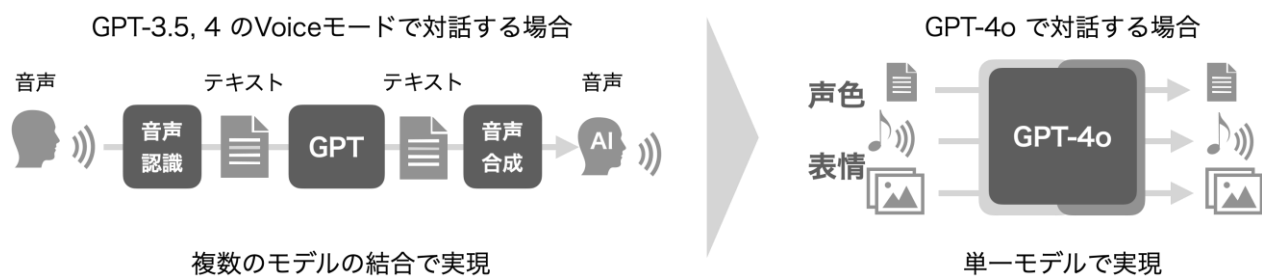


図 1 : ChatGPT のモデルによる違い

このような背景を受けて、金沢学院大学・金沢学院短期大学では、2024 年度入学生の 1 年次教育において、生成 AI リテラシー教育を提供した。2024 年 4 月に新たに開設した情報工学部に所属する教員が、各学部のオリエンテーションや情報リテラシーの授業内の 45 分程度を利用して、生成 AI の仕組みや利用方法、利用上の注意について、ChatGPT や Microsoft Copilot を用いて体験的に学ぶ機会を提供するものである。本稿では、本学における生成 AI リテラシー教育、生成 AI の仕組みや利用上の注意に関する教育事例を紹介する。また、受講生に対するアンケート調査を通して、大学 1 年生の生成 AI の認知や利用状況、興味関心の傾向を明らかにすることで、今後の大学教育における生成 AI 活用の必要性や、AI 時代に育成すべき資質・能力について示唆を得たい。

2. 生成 AI の仕組みや利用上の注意 —「賢い他人との対話」

生成 AI ツールは、誰もが手軽に利用できる便利な技術である。しかし、その仕組みを十分に理解せずに使われることも多く、AI を万能な存在と誤解したり、適切でない活用をしたりするケースが見られる。例えば、AI が常に正しい答えを返すわけではないにもかかわらず、そのまま信じてしまうことや、入力した情報がどのように処理されるのかを意識せずに個人情報を入力してしまうことがある。これらの問題を防ぐには、情報リテラシーとともに生成 AI リテラシーを身に付け、生成 AI ツールを正しく理解し、安全かつ効果的に活用することが重要である。本節では、生成 AI の利用を「賢い他人との対話」と捉えることで、その仕組みや利用上の注意点を直観的に理解しやすく説明していく。

2.1 情報通信システムの視点 —「賢い他人はどこにいるのか？」

多くの学生は、AI が自分の端末の中で文章を考えていると想像しがちだが、実際にはそうではない。ChatGPT を代表とする生成 AI ツールは、インターネットを介してクラウド上の大規模なサーバーで動作している。端末に文章を入力すると、そのデータはインターネットを通じてサーバーに送られ、そこで高度な計算が行われた後、結果が端末に送信され、画面に表示される。したがって、「賢い他人」はそれぞれの端末にいるのではなく、一つの AI が多くの利用者と順番に対話しているイメージの方が正しい。利用者が増えれば処理が集中し、全体的に応答が遅くなることもある。さらに、AI の運用には膨大な計算資源が必要であり、その維持にはコストがかかる。そのため、一定以上の利用で課金が発生するもの、この仕組みを理解すれば納得しやすいだろう。

2.2 回答を鵜呑みにしない —「賢い他人でも間違えることがある」

前節で、AI はある入力に対する出力を人工的に再現した機械であると述べた。ChatGPT

のような文章生成 AI は、入力されたテキストに対して、次に来る可能性が最も高い単語を予測しながら文章を作る仕組みになっている。この精度が向上したことで、人間のように自然な文章を生成できるようになった。しかし、インターネットで最新情報を検索しているわけではないことに注意が必要である。現在の生成 AI ツールでは、検索拡張生成（RAG）と呼ばれる技術を用いることで、外部の情報源を参照しながら回答を生成するものもある。しかし、これはあくまで参考文献をもとに回答を補強しているに過ぎず、情報の正確性を保証するものではない。「賢い他人」が手元の本をもとに回答しているようなものであり、誤解や偏りが含まれる可能性がある。そのため、AI の回答をそのまま信じるのではなく、常に他の情報源と照らし合わせる姿勢が重要である。

2.3 個人情報や機密情報を伝えない — 「賢い他人はあくまで他人である」

なぜ、生成 AI ツールを無料で利用できるのか。それは、多くの利用者が入力したデータが、AI の学習データとして利用される場合があるからである。これは、検索エンジンや SNS など、多くの無料サービスが利用者のデータを活用してサービスを向上させているのと同じ仕組みである。「賢い他人」は、利用者との対話の中でより賢くなろうとしているわけである。

また、情報通信システムの視点から考えると、生成 AI ツールに入力したデータは、利用者の端末内にとどまらず、クラウド上のサーバーに送信され、場合によっては保存・学習に利用されることがある。このため、個人情報や機密情報を入力すると、それが意図せず第三者に渡るリスクがある。身近な例として、仕事の機密資料や個人の連絡先を AI に入力することは、紙の書類を「他人」に見せるのと同じような危険を伴う。

こうしたリスクを軽減するため、多くの生成 AI ツールには、利用者のデータを学習に利用しないように設定できる「オプトアウト」の仕組みが用意されている。ChatGPT の場合、ログイン後の設定画面からデータ収集を無効にすることが可能である。しかし、この設定の存在を知らない利用者も多い。授業などで生成 AI を活用する際には、事前にこうした設定を確認し、安全に利用できるようにすることが重要である。

2.4 回答をそのまま流用しない — 「賢い他人の回答は他人のもの」

生成 AI が出力する文章や画像は、過去に学習したデータをもとに新たに生成されたものである。しかし、これは完全にゼロから生み出されたものではなく、既存のデータ、学習データの影響を受けている点に注意が必要である。特に、著作権のある作品を学習している場合、その特徴を反映した生成物が作られる可能性がある。

文章生成 AI の場合、出力された文章がどこかの文献や記事と酷似していても、それを識

別するのは難しい。一方で、画像生成 AI を用いると、著作権のあるキャラクターや有名な芸術作品に似た画像が出力されることがあり、視覚的にその影響を確認しやすい。画像生成の例を示すことで、生成 AI ツールの多様な活用方法を知ることができ、AI が過去のデータを基にしていることも具体的に理解できる。このように、生成 AI の出力は独創的に見えても、元となるデータの影響を受けていることを踏まえ、そのまま流用しないように注意する必要がある。特に、研究発表やレポート作成において、AI の生成物を適切に引用せずに使用すると、意図せず盗用と見なされるリスクがある。

2.5 プロンプトの重要性 — 「賢い他人との上手な対話」

生成 AI との対話では、入力（プロンプト）の内容が回答の質を大きく左右する。プロンプトとは、AI に対する指示や質問のことであり、どのように伝えるかによって得られる情報が変わる。検索エンジンでキーワードを入力する情報検索とは異なり、AI との対話では、一度の質問で完璧な答えが得られるとは限らない。

多くの学生は、最初の質問で正しい答えを得なければならないと考えがちだが、実際には試行錯誤が重要である。最初の回答が期待と異なっても、「もっと詳しく」「具体例を加えて」などの指示を加えながら繰り返し質問することで、より良い回答を引き出せる。これは、人と対話する際に追加の質問をするのと同じである。一方で、一度のプロンプトで意図した出力を得る技術として「プロンプトエンジニアリング」という考え方もある。しかし、初めから完璧なプロンプトを作るのは難しく、繰り返し対話する中で徐々に精度を上げていくことが現実的である。失敗を恐れず、試行錯誤しながら対話を重ねることが、賢い他人（AI）との上手な付き合い方につながる。

3. 生成 AI リテラシー教育事例から見る大学生と生成 AI

一般社団法人データサイエンティスト協会は、学生向けに「データサイエンティスト」に関する調査を行っており、2023 年 12 月の調査においては、生成 AI に関する質問を設けてアンケートを実施している（一般社団法人データサイエンティスト協会, 2024）。この調査では、大学生は学生全体で生成 AI を知らないという学生の割合は 25%、生成 AI ツールを使ったことのある学生は 29%であり、その活用シーンの上位は、論文の要約やレポート作成・添削であったことを報告している。ChatGPT の公開から 1 年程度ですでに 30%近い学生が生成 AI を利用しており、現在では生成 AI の認知や利用は拡大していると考えられる。本研究では、生成 AI リテラシー教育において、同様のアンケート調査を行い、生成 AI の認知や教育機会、これまでの利用状況や活用シーンに対する興味関心の傾向を明らかにする。

3.1 生成 AI リテラシー教育の対象と概要

金沢学院大学は、2024 年 4 月に初の理系学部である情報工学部が開設し、情報工学部、経済学部、文学部、教育学部、芸術学部、栄養学部、スポーツ科学部の 7 学部体制の総合大学となった。また、金沢学院短期大学は、現代教養学科、食物栄養学科、幼児教育学科の 3 学科体制である。北陸地方の私立大学・短期大学であり、これまでの学園の沿革から文系学生の割合は高い。2016 年度から学生自身のノートパソコンを必携して学修に取り組む BYOD（Bring Your Own Device）が全学的に導入されており、学生は大学メールアドレスと共に、Microsoft Office 365 のアカウントが配布される。

2024 年 5 月から 6 月にかけて、情報工学部に所属する教員が、各学部新入生のオリエンテーションや情報リテラシーの授業内の 45 分程度を利用して、生成 AI リテラシー教育を提供した。前節のように、生成 AI の仕組みや利用上の注意について、「賢い他人との対話」と例えながら、文系学生でも理解できるよう工夫した。授業の中では、大学メールアドレスを用いた ChatGPT のサインアップとオプトアウト設定、Microsoft アカウントを利用した Microsoft Copilot のログインと画像生成を体験した。ただし、ChatGPT の無料版のユーザーは、2024 年 8 月のアップデートにより、画像生成機能を利用できるようになっている。

生成 AI ツールは誰でも利用することができるが、大学教育において導入・利用する際には注意が必要である。例えば、100 人以上が同じ教室で同時に ChatGPT のサインアップを行うと、不具合が生じる可能性がある。大学のネットワーク環境に負荷がかかり、ページの読み込み遅延やエラーが発生する。また、同じ IP アドレスからのアクセス集中や短時間での大量の新規アカウントの作成によって、ChatGPT のサーバー負荷軽減やセキュリティ対策のために、一時的に利用が制限される。このような不具合は、情報通信システムに対する理解の欠如によって生じるものである。円滑に授業を行うため、適切な数のグループを構成し、生成 AI ツールの利用した体験活動を行なった。

本研究では、生成 AI リテラシー教育の効果を評価するため、学生の受講前後のアンケートを実施した。調査では、生成 AI の認知度、利用経験、利用シーン、期待される活用法について設問を作り、受講前後の変化を調査した。表 1 に学科ごとの有効回答数を示す。

学科 (大学)	情報工	経済	経営	文	教育	芸術	栄養	スポーツ
回答数	32	86	47	155	46	64	74	157

学科 (短大)	現代教養	食物栄養	幼児教育		大学	短大	合計
回答数	45	36	33		661	114	775

表 1：アンケート調査における学科別有効回答数

3.2 受講生の生成 AI の利用・認知と教育機会

生成 AI の利用・認知に関する調査項目として、「あなたは、これまでに『生成 AI』のツール・アプリ・ソフトを使ったことがありますか。」の質問に対して、「知っていて、使ったことがある」「使ったことはないが、知っている・聞いたことがある」「知らない・聞いたことがない」で回答を求めた。

全体の結果を図 2 に示す。「知っていて、使ったことがある」という回答が 29.2%であり、約 3 割が利用したことがあるという結果となった。また、「使ったことはないが、知っている・聞いたことがある」という回答が最も多く 57.5%、「知らない・聞いたことがない」という回答は 13.3%であった。2023 年 12 月のデータサイエンティスト協会による調査では、学生全体で 25%が生成 AI を知らないとの回答であったが、2024 年 6 月の本学に限った調査でも 9 割近くの学生に認知されており、生成 AI ツールの認知が急速に拡大している様子が伺える。一方で、実際に活用した経験を持つ学生の割合は 3 割程度であり、認知度と比較して、利用者の割合は 2023 年 12 月の調査と同程度であり、限定的であることがわかる。図 3 に学科別の結果を示す。情報工学科は「知っていて、使ったことがある」という回答が 75.0%であり、理系とその他の学科で傾向が異なることは明らかである。

生成 AI の活用シーンに関する調査項目として、『生成 AI』をどのようなシーンで使ったことがありますか。」の質問に対して、「学業」「就活」「プライベート」で複数選択を許して回答を求めた。図 4 に結果を示す。ただし、各学科の生成 AI を利用したことがある回答者数を括弧内は示した。調査対象が大学 1 年生であるため、利用シーンについては、学業またはプライベートが大半を占めた。また、「学業や就職活動において、どのように『生成 AI』を使ったことがありますか。」の複数選択を許した質問に対する結果を図 5 に示す。実際に利用している学生の割合は限られているが、主に学業での利用シーンとして、「アイデア出し」「レポートや論文作成・添削」「各種調べ物」「論文や教科書の要約」が上位となった。

全体の割合 - 【受講前】生成AIの利用・認知

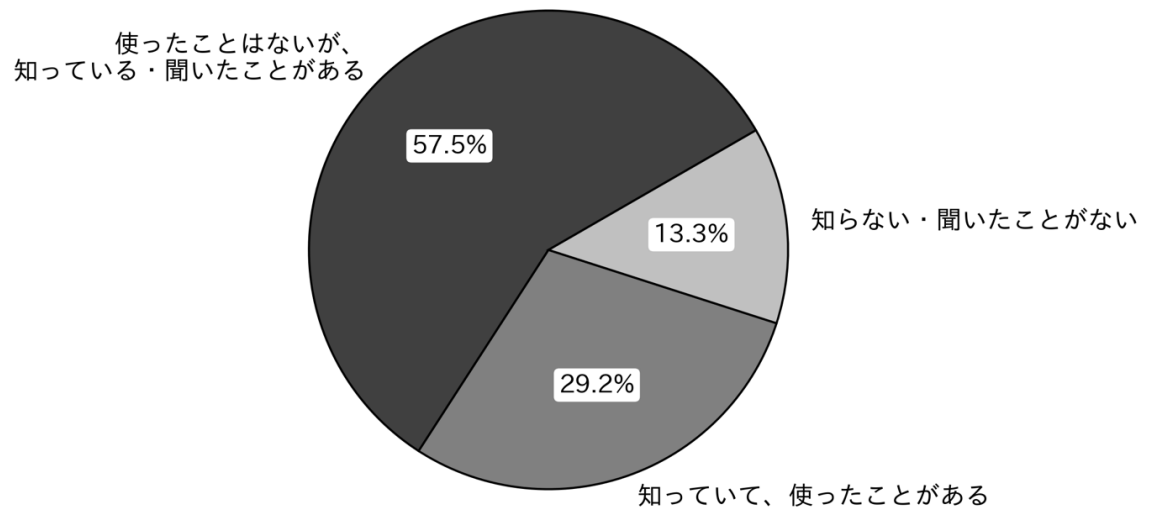


図 2：生成 AI の利用・認知に関する回答結果（全体）

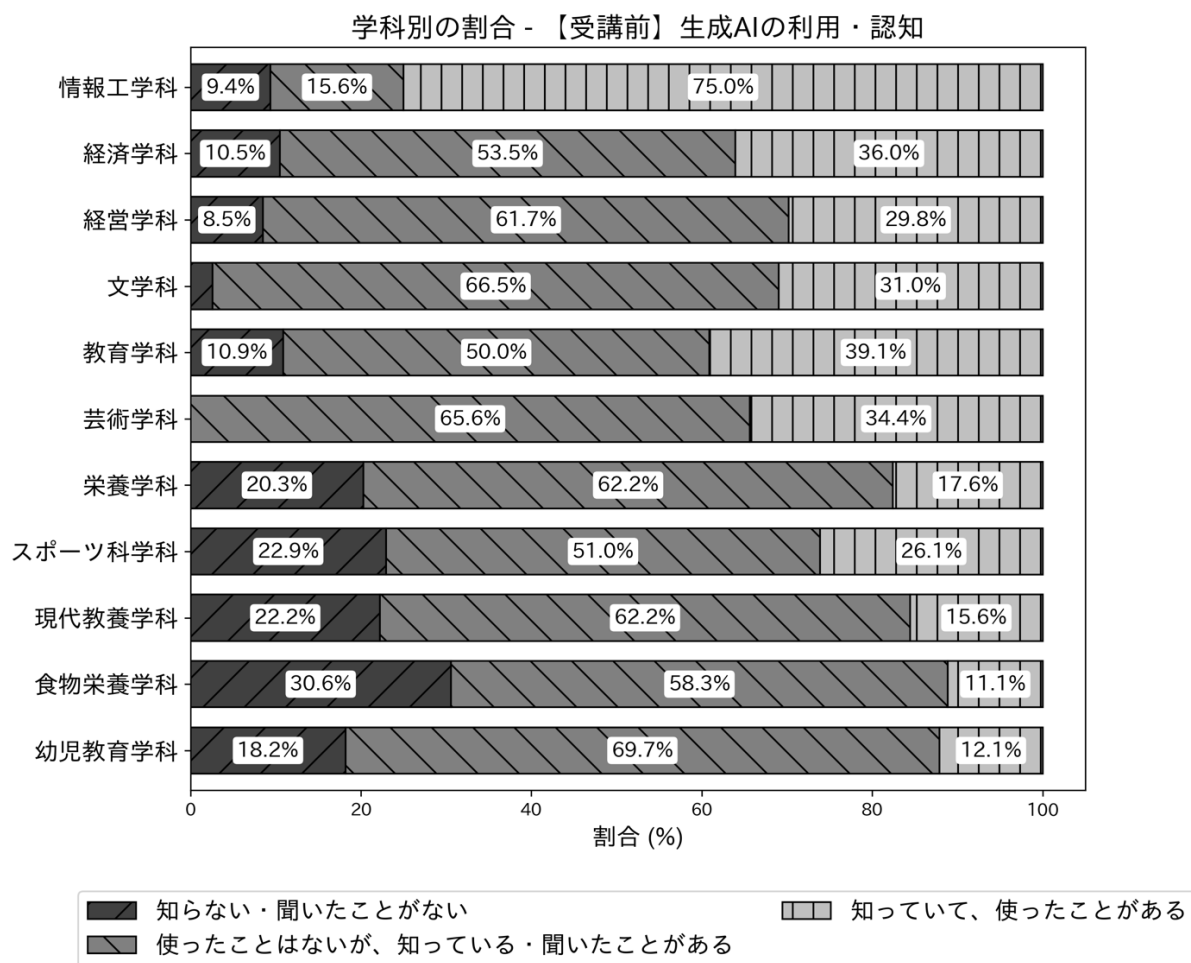


図 3：生成 AI の利用・認知に関する回答結果（学科別）

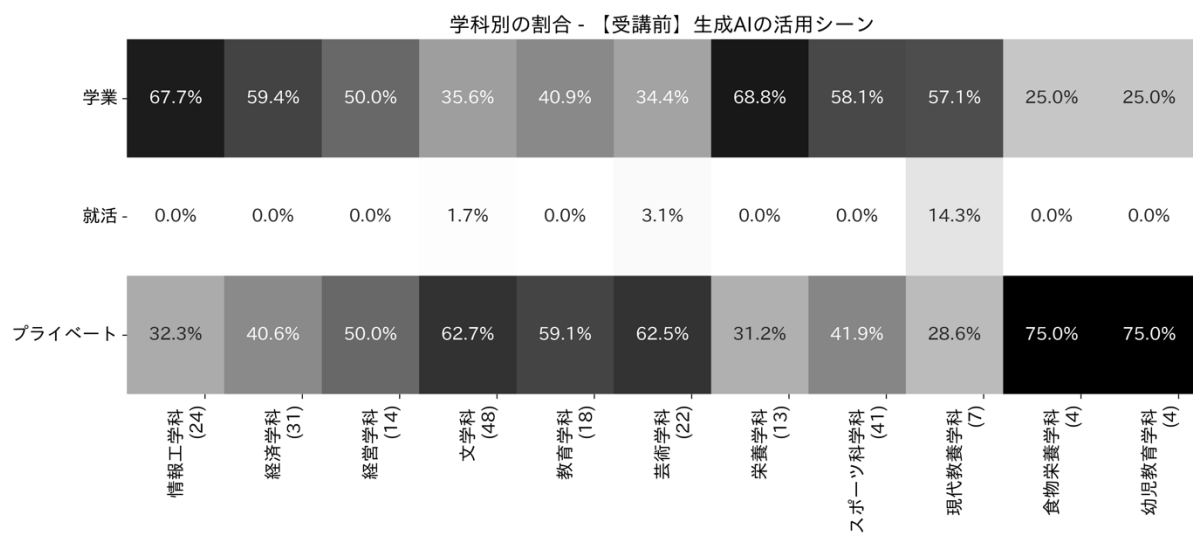


図 4：生成 AI の主な活用シーン（学科別）

学科別の割合 - 【受講前】生成AIの活用シーン（詳細）											
論文や教科書の要約	13.0%	27.3%	6.1%	10.9%	4.5%	4.2%	13.6%	12.7%	0.0%	0.0%	0.0%
レポートや論文の作成・添削	14.8%	16.4%	9.1%	10.9%	4.5%	8.3%	22.7%	22.8%	9.1%	0.0%	33.3%
アイデア出し	18.5%	18.2%	21.2%	23.4%	45.5%	20.8%	22.7%	20.3%	27.3%	25.0%	0.0%
テスト問題や宿題の解答補助	22.2%	3.6%	12.1%	3.1%	0.0%	6.2%	9.1%	10.1%	0.0%	0.0%	33.3%
授業のメモや議事録の作成	1.9%	0.0%	9.1%	0.0%	4.5%	2.1%	0.0%	1.3%	9.1%	0.0%	0.0%
各種調べ物	11.1%	10.9%	12.1%	18.8%	27.3%	12.5%	18.2%	11.4%	18.2%	0.0%	0.0%
言語学習	1.9%	7.3%	0.0%	7.8%	4.5%	12.5%	9.1%	2.5%	9.1%	0.0%	0.0%
プレゼンテーション資料の作成	3.7%	5.5%	12.1%	0.0%	0.0%	2.1%	0.0%	7.6%	0.0%	0.0%	0.0%
プロジェクトや実験の計画作成	0.0%	5.5%	3.0%	0.0%	0.0%	4.2%	0.0%	2.5%	0.0%	0.0%	0.0%
専門用語や概念の理解	5.6%	0.0%	6.1%	3.1%	0.0%	8.3%	4.5%	1.3%	0.0%	25.0%	0.0%
プログラムの作成やチェック・デバック	0.0%	1.8%	3.0%	1.6%	0.0%	0.0%	0.0%	0.0%	0.0%	0.0%	0.0%
各種資料や研究タスクの構成案の作成	0.0%	0.0%	3.0%	1.6%	0.0%	0.0%	0.0%	2.5%	0.0%	0.0%	0.0%
芸術作品の創作	3.7%	1.8%	3.0%	7.8%	4.5%	6.2%	0.0%	1.3%	9.1%	0.0%	0.0%
その他（学業）	1.9%	0.0%	0.0%	9.4%	4.5%	4.2%	0.0%	1.3%	0.0%	0.0%	33.3%
エントリーシートの作成	1.9%	1.8%	0.0%	1.6%	0.0%	6.2%	0.0%	1.3%	9.1%	25.0%	0.0%
その他（就職活動）	0.0%	0.0%	0.0%	0.0%	0.0%	2.1%	0.0%	1.3%	9.1%	25.0%	0.0%
	情報工学科 (24)	経済学部 (31)	経済学部 (14)	文学部 (48)	教育学部 (18)	芸術学部 (22)	栄養学部 (13)	スポーツ科学部 (41)	現代教養学部 (7)	食物栄養学部 (4)	幼児教育学部 (4)

図 5：生成 AI の詳細な活用シーン（学科別）

次に、生成 AI の教育機会に関する調査項目として、「これまでに授業で『生成 AI』の仕組み・特性やリスクについて学んだことがありましたか。」の質問に対して、「学んだことがあった」「学んだことがなかった」で回答を求めた。全体と学科別の結果をそれぞれ図 6 と図 7 に示す。学科ごとに大きな偏りはなく、8 割近くの学生がこれまでに生成 AI について学ぶ機会がなかったという回答であった。これらの結果から、生成 AI ツールに関する認知度は急速に拡大する一方で、その仕組みやリスクについて学ぶ機会は十分に確保されて

いないことがわかった。このことは、生成 AI を認知していながらも利用経験のない学生の割合が高くなる要因の一つであると考えられる。

全体の割合 - 【受講前】生成AIに関する教育機会

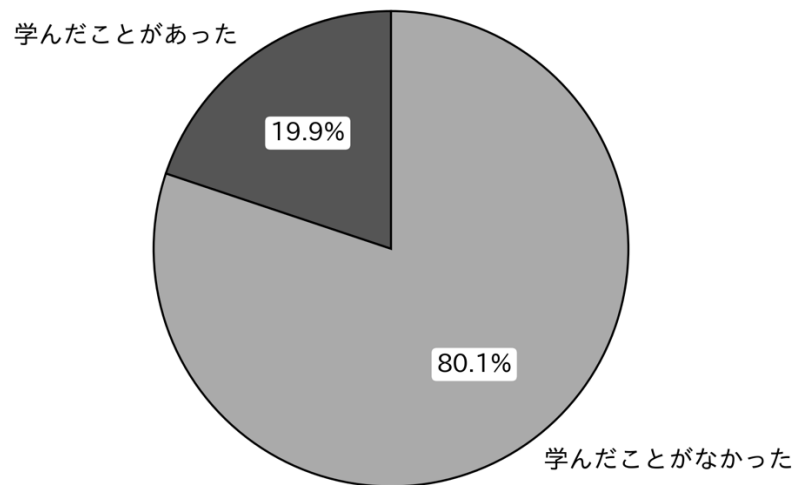


図 6：生成 AI に関する教育機会に関する回答（全体）

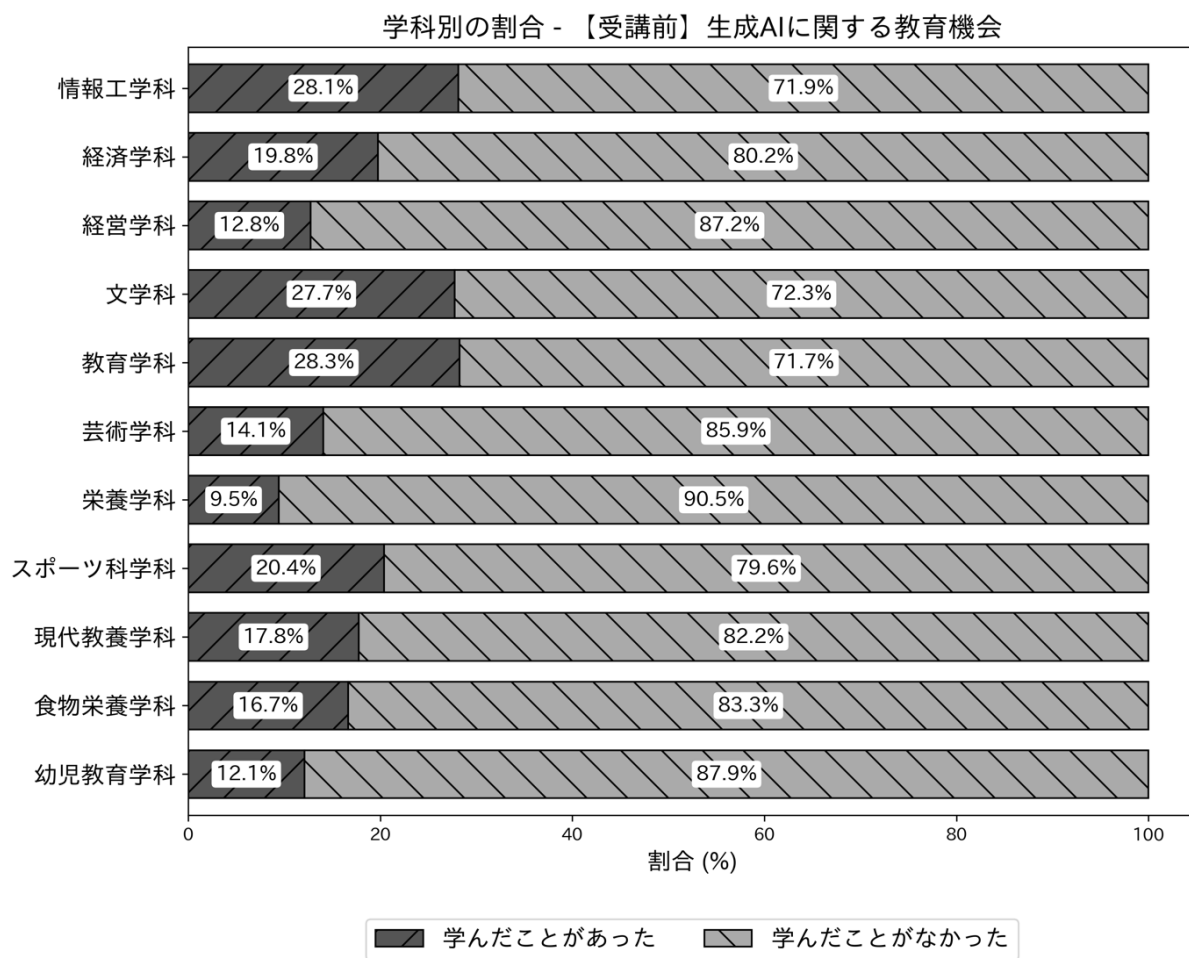


図7：生成AIに関する教育機会に関する回答（学科別）

3.3 生成 AI リテラシー教育受講後の変化

生成 AI リテラシー教育後に、受講生に対して、生成 AI への興味関心、生成 AI への理解、今後の利用に対する意欲についてアンケートを行った。

受講後の生成 AI への興味関心に関する調査項目として、「『生成 AI』に興味を持って受講することができましたか。」の質問に対して、「十分に持てた」「持てた」「どちらかといえば持てた」「あまり持てなかった」「持てなかった」で回答を求めた。全体と学科別の結果をそれぞれ図 8 と図 9 に示す。生成 AI への興味関心を持つことができたという学生が 9 割以上を占めた。自由記述欄には、初めて ChatGPT などの生成 AI ツールを利用した驚きの声が多く寄せられた。

次に、受講後の生成 AI への理解に関する調査項目として、「『生成 AI』の仕組み・特性やリスクについて理解できましたか。」の質問に対して、「よく理解できた」「理解できた」「あまり理解できなかった」で回答を求めた結果を、全体と学科別でそれぞれ図 10 と図 11 に示す。こちらも学生の 9 割以上が生成 AI の仕組みやリスクについて理解することができたという回答であった。生成 AI ツールの利用は「賢い他人との対話」であるとし、利用上の注意点として「回答を鵜呑みにしない」「個人情報や機密情報を伝えない」「回答をそのまま流用しない」という 3 点を強調したこともあり、自由記述欄には、これらに留意して利用したいとの回答が見られた。

受講後の生成 AI の活用シーンに関する調査項目として、「『生成 AI』をどのようなシーンで使いたいと思いますか。」の質問に対して、「使いたいと思わない」「学業」「就活」「プライベート」で複数選択を許して回答を求めた。「使いたいと思わない」を選択した学生は、全体の 8.0%を占めた。自由記述欄には、生成 AI の利便性に好感を持ったという一方で、怖さを感じたという声もあった。また、利用上の注意を学んだことによって、どこまでが個人情報であるのか、誤り生成への適切な対応など、自身の知識が身に付くまで利用を避けたいという記述があり、ルールや法整備が進むのを待ちたいというものもあった。次に、図 12 に学科別の「学業」「就活」「プライベート」を選択した割合を示す。各学科でこれらを選択した回答者数を括弧内は示している。受講前の利用者のほとんどが選択していなかった就活も、受講後には 2 割弱の学生が利用に対して意欲を示していることがわかる。

全体の割合 - 【受講後】生成AIへの興味関心

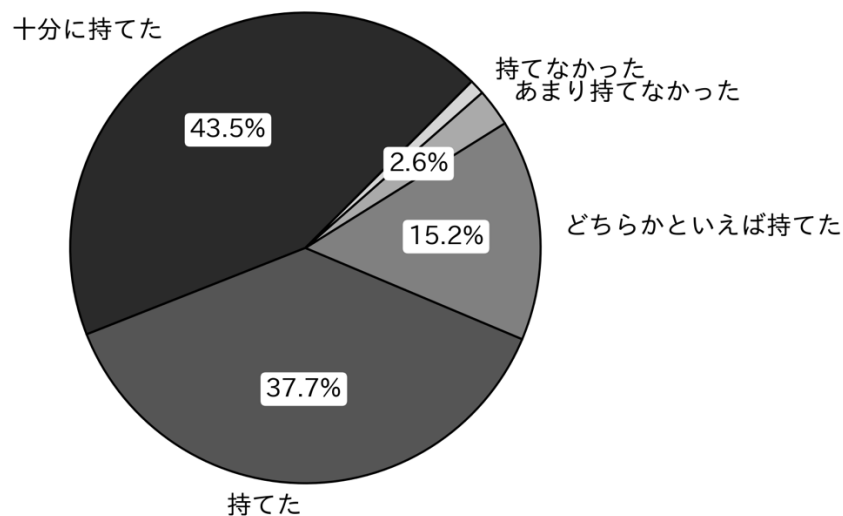


図 8：受講後の生成 AI への興味関心に関する回答（全体）

学科別の割合 - 【受講後】生成AIへの興味関心

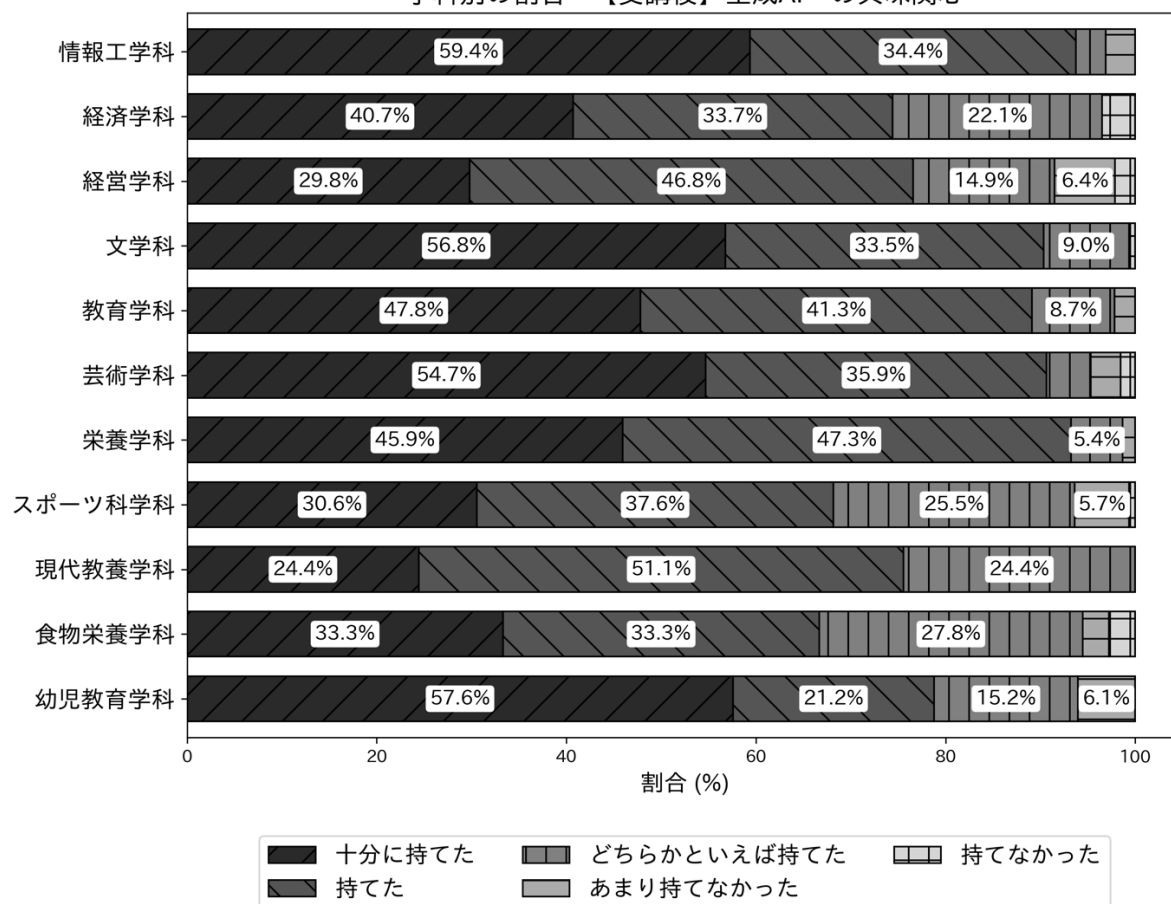


図 9：受講後の生成 AI への興味関心に関する回答（学科別）

全体の割合 - 【受講後】生成AIへの理解
よく理解できた

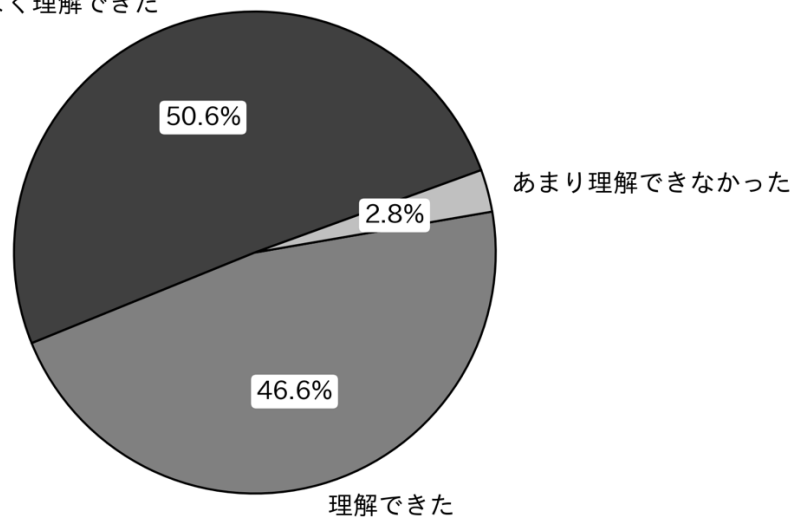


図 10：受講後の生成 AI への理解に関する回答（全体）

学科別の割合 - 【受講後】生成AIへの理解

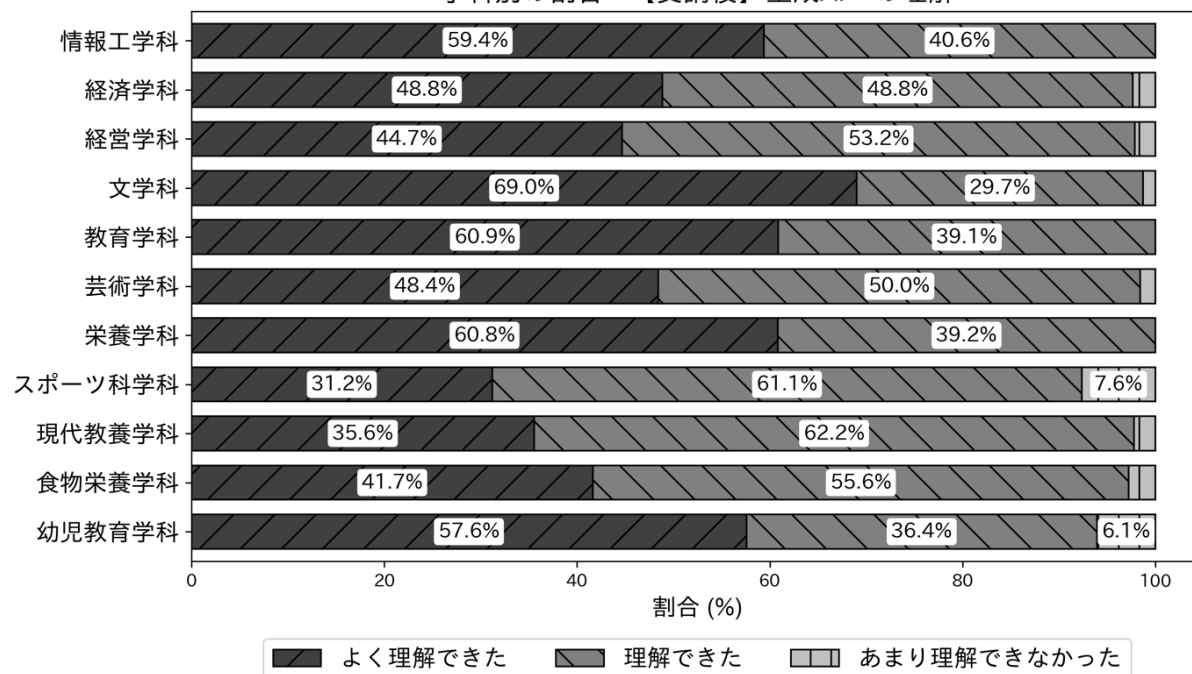


図 11：受講後の生成 AI への理解に関する回答（学科別）

学科別の割合 - 【受講後】生成AIの活用シーン										
学業	53.7%	51.5%	56.1%	41.4%	45.2%	39.4%	53.5%	57.1%	40.0%	43.5%
就活	20.4%	25.4%	22.7%	12.3%	17.7%	14.4%	24.8%	22.6%	34.3%	23.9%
プライベート	25.9%	23.1%	21.2%	46.3%	37.1%	46.2%	21.7%	20.3%	25.7%	33.3%
	情報工学科 (31)	経済学科 (79)	経営学科 (44)	文学科 (142)	教育学科 (42)	芸術学科 (60)	栄養学科 (72)	スポーツ科学科 (142)	現代教養学科 (39)	食物栄養学科 (31)
										幼児教育学科 (31)

図 12：受講後の生成 AI の活用シーンに関する回答（学科別）

最後に、生成 AI の詳細な活用シーンに関する調査項目として、「学業や就職活動において、どのように『生成 AI』を使いたいと思いますか。」の質問への回答結果を図 13 に示す。

「アイデア出し」「レポートや論文の作成・添削」「論文や教科書の要約」「各種調べ物」が上位を占めた。芸術学科については、他の学科と異なり、1 割程度の学生が「芸術作品の創作」に対して活用の意欲を示していた。対象が大学 1 年生であったため、「プロジェクトや実験の計画作成」「各種資料や研究タスクの構成案の作成」の項目への回答は限定的であったと考えられる。今回の生成 AI リテラシー教育においては、生成 AI の仕組みや利用上の注意を扱い、具体的な活用事例の紹介は行わなかったため、身近な活用シーンに対する回答が多く見られる結果となったと考えられる。

【受講後】生成AIの活用シーン（詳細）											
論文や教科書の要約	10.9%	18.1%	14.8%	10.8%	8.6%	7.0%	9.7%	16.7%	7.2%	14.4%	15.4%
レポートや論文の作成・添削	12.7%	15.2%	14.8%	12.2%	7.8%	9.8%	11.2%	20.0%	18.9%	12.2%	18.5%
アイデア出し	15.5%	17.7%	21.3%	20.2%	25.0%	22.4%	19.0%	14.7%	15.3%	17.8%	29.2%
テスト問題や宿題の解答補助	12.7%	8.9%	6.5%	4.7%	5.5%	3.3%	6.3%	9.5%	6.3%	10.0%	7.7%
授業のメモや議事録の作成	4.5%	3.8%	4.6%	1.1%	3.9%	0.9%	2.6%	3.7%	1.8%	5.6%	0.0%
各種調べ物	10.9%	8.4%	9.3%	17.7%	14.8%	12.1%	16.0%	7.7%	13.5%	10.0%	12.3%
言語学習	2.7%	4.6%	2.8%	8.3%	7.8%	5.6%	6.3%	5.3%	7.2%	4.4%	4.6%
プレゼンテーション資料の作成	8.2%	5.5%	7.4%	4.7%	7.8%	6.5%	7.1%	6.0%	6.3%	2.2%	4.6%
プロジェクトや実験の計画作成	3.6%	4.2%	2.8%	1.4%	3.9%	2.8%	3.3%	2.8%	3.6%	3.3%	1.5%
専門用語や概念の理解	3.6%	3.8%	3.7%	6.6%	3.9%	6.1%	5.6%	2.8%	4.5%	5.6%	0.0%
プログラムの作成やチェック・デバック	2.7%	2.5%	2.8%	1.4%	1.6%	4.7%	2.6%	2.1%	1.8%	1.1%	0.0%
各種資料や研究タスクの構成案の作成	2.7%	0.4%	0.9%	1.7%	1.6%	2.8%	2.6%	1.9%	0.0%	1.1%	0.0%
芸術作品の創作	2.7%	0.4%	2.8%	4.4%	3.1%	10.3%	0.4%	2.1%	4.5%	1.1%	3.1%
その他（学業）	1.8%	1.3%	0.9%	2.2%	1.6%	1.9%	2.2%	1.2%	4.5%	4.4%	3.1%
エントリーシートの作成	2.7%	2.5%	2.8%	1.1%	0.8%	1.9%	2.6%	1.6%	2.7%	2.2%	0.0%
その他（就職活動）	1.8%	2.5%	1.9%	1.4%	2.3%	1.9%	2.6%	1.9%	1.8%	4.4%	0.0%
	情報工学科 (31)	経済学科 (77)	経営学科 (43)	文学科 (142)	教育学科 (42)	芸術学科 (61)	栄養学科 (72)	スポーツ科学科 (140)	現代教養学科 (37)	食物栄養学科 (28)	幼児教育学科 (30)

図 13：受講後の生成 AI の詳細な活用シーンに関する回答（学科別）

4. おわりに

本稿では、本学における生成 AI リテラシー教育の実践を通じて、生成 AI の仕組みや利用上の注意に関する教育事例を紹介するとともに、大学 1 年生を対象としたアンケート調査を行い、生成 AI の認知や利用状況、興味関心の傾向を明らかにした。近年、生成 AI は急速に普及し、企業における導入も進んでいるが、安全かつ効果的な活用のためには適切な

リテラシー教育が不可欠である。本研究では、生成 AI を単なる便利なツールとして扱うのではなく、その仕組みを理解し、適切な対話を通じて活用する力を養うことを目的とした教育プログラムを実施した。

本研究で紹介した生成 AI リテラシー教育においては、生成 AI の利用を「賢い他人との対話」と捉え、その特性を理解した上で適切に活用する重要性を強調した。生成 AI は万能ではなく、誤情報を含む可能性があるため、利用者が主体的に問いを立て、対話を通じて最適な情報を得ることが求められる。この観点から、生成 AI との関係は単なる「入力と出力」のやり取りにとどまらず、人間が AI と共創的な対話を行うことが鍵となる。生成 AI を適切に活用するためには、的確なプロンプトを作成し、試行錯誤を重ねながら、AI の出力を批判的に評価し、より良い形で活用する姿勢が不可欠である。

アンケート調査の結果から、大学 1 年生の多くが生成 AI の存在を認知しているものの、「知っているが利用したことがない」学生が依然として多かったことが分かった。これは、新しい技術が社会に広がる過程を説明するイノベーター理論 (Rogers, 2003) で捉えることができる。例えば、スマートフォンが普及した際、発売直後にいち早く購入する人（アーリーアダプター）がいる一方で、多くの人は周囲の評判を聞いた上で購入を決めた（アーリーマジョリティ）。その後、さらに慎重な層（レイトマジョリティ）が市場の主流となり、最後に新技術への関心が低い人々（ラガード）がようやく導入する。現在の生成 AI の状況もこれと似ており、すでに積極的に活用している学生がいる一方で、興味はあるがまだ利用していない学生が多数存在していた。この移行期には、利用者と非利用者の間でスキル格差が生じる可能性があるため、適切な教育機会を提供することが重要である。こうした課題に対し、大学におけるリテラシー教育が果たす役割は大きいと考えられる。

生成 AI ツールの急速な普及に伴い、国によるリテラシー教材の公開も開始されている。例えば、総務省による安心・安全なインターネット利用ガイドの特集ページとして、生成 AI の利活用に関する初心者向けの教材が公開されている [総務省, 2024]。また、文部科学省による小中学向けのガイドライン [文部科学省, 2024] が公開されるなど、教育現場での生成 AI の活用に向けた議論が進んでいる。テキスト生成にとどまらず、画像や音声など多様な生成 AI ツールが登場し、今後は AI エージェントが普及することで、情報処理や意思決定の在り方が変化していくことが予想される。こうした変化の中で、AI を適切に活用しながら学びを深めることができる人材の育成が求められている。

最後にこのような AI 時代における大学教育のあり方について考えたい。生成 AI の普及によって、情報を得る手段や学習のプロセスが大きく変化する中、これまでシラバスで掲げてきた到達目標へのアプローチも多様化することが求められる。単に知識を蓄積することに重点を置くのではなく、どのように学ぶか、どのように AI と協働しながら知を創造するかという視点がより重要になるだろう。さらには、知識そのものではなく、どのような資質・能力を育成すべきかという到達目標そのものを再考する必要がある。

このような環境の変化の中で、特に求められるのは、言語化する力、対話する力、そして

主体性であるだろう。生成 AI の登場により、情報検索やアイデア生成が容易になった一方で、得られた情報を適切に整理し、自分の考えを明確に表現する力がより重要になっている。また、AI との対話を通じて適切な出力を引き出すためには、単なる質問ではなく、意図を的確に伝えられる高度な言語運用能力が不可欠となる。さらに、AI とどのように協働するかを決める主体性も、今後の学びにおいて鍵となる。単に AI に依存するのではなく、どの場面で AI を活用し、どこは人間が考えるべきかを判断できる能力を養うことが、AI 時代における学びの本質となる。

こうした変化に適応するためには、学習環境の再構築も必要である。知識の伝達型の授業だけでなく、学生が AI を活用しながら問題を発見し、試行錯誤を重ね、対話を通じて解決策を見出す探究型の学びが求められる。例えば、授業の中で AI を適切に活用しながら個別に学習を深める機会を提供することで、従来の一律的な学習形態からの転換を図ることができる。また、学びの場は学生だけでなく、教職員もまた AI を活用しながら成長し、教育の在り方をアップデートしていく必要がある。本研究を通じて得られた知見を活かし、今後も生成 AI リテラシー教育を発展させ、学生が変化の激しい社会に適応し、AI と共に新たな価値を創造できるような教育の在り方を模索していきたい。生成 AI の普及が進む中で、単にツールを使いこなすことにとどまらず、AI との対話を通じた批判的思考や創造的な課題解決能力を育成することが、今後の大学教育において重要な課題となるだろう。

参考文献

- AdityaRamesh, DhariwalPrafulla, NicholAlex, ChuCasey, ChenMark. (2022). Hierarchical Text-Conditional Image Generation with CLIP Latents.
<https://arxiv.org/abs/2204.06125>.
- JIPDEC; ITR. (2024 年 3 月 14 日). 「企業 IT 利活用動向調査 2024」結果発表. 参照先:
<https://www.jipdec.or.jp/news/pressrelease/20240315.html>
- OpenAI. (2022 年 11 月). Introducing ChatGPT. 参照先:
<https://openai.com/index/chatgpt/>
- RogersM.Everett. (2003). Diffusion of Innovations. Free Press.
- 一般社団法人データサイエンティスト協会. (2024 年 5 月 10 日). Data of Data Scientist シリーズ vol.53 『29% - 大学生の生成 AI 利用率』. 参照先: Data Science Society Journal: <https://www.datascientist.or.jp/dssjournal/dssjournal-2920/>
- 総務省. (2024). 生成 AI はじめの一步～生成 AI の入門的な使い方と注意点～. 参照先: 上手にネットと付き合おう！安心・安全なインターネット利用ガイド:
https://www.soumu.go.jp/use_the_internet_wisely/special/generativeai/
- 文部科学省. (2024 年 12 月). 初等中等教育段階における生成 AI の利活用に関するガイド

ライン. 参照先: https://www.mext.go.jp/a_menu/other/mext_02412.html

第4部 生成 AI サービスの普及等経済のデジタル化を把握 するための統計的枠組みについて

～国際的動向と無償サービスに関する試算～

金沢学院大学 経済学部 長谷川秀司

1. はじめに

2010 年代以降、毎年のように新しいデジタルのテクノロジーやサービスが経済及び経済学の議論の対象となり、流行とも言える様相を呈している。すなわちこれまで DX、プラットフォーム、シェアリング・エコノミー、クラウド、データセンター、そして生成 AI が続いてきた。最近では、コロナ禍の「デジタル敗戦」に次ぐ、日本経済の構造的課題を象徴する言葉として国際収支の「デジタル赤字」が注目を浴びている。

しかしこうした経済のデジタル化が急速なスピードで普及・拡大しているため、そのスピードに実態の把握・測定が追い付けないほどである。例えば、「デジタル赤字」に対応する国際収支統計上の「デジタル関連支出」もデジタルサービスの取引に関連すると考えられる項目を組み替えた分類、集計値であり、国際収支統計の分類表にはそもそもデジタルという言葉がない。さらには統計として測定する際の枠組みである現行の主要な国際基準である 2008SNA や ISIC rev4 に「プラットフォーム」や「クラウド」という概念や分類を見つけないことができない。また、今日多くの者が生成 AI サービスやソーシャルネットワークサービス（SNS）等の無償サービス¹を利用し、多大な便益を享受しているがその価値は市場価格がない故その評価の困難性を抱えている。このようにデジタルエコノミーは、測定における現行の尺度や対象に関して多様な問題を惹起している。

さらにデジタルエコノミーの「無料、完全、瞬時」（Brynjolfsson et al. 2017）を特性とするデジタル技術によって、既存の産業、企業は大規模な影響を受けており、「破壊的イノベーション」（Christensen 1997）になっている。同時に消費者のインターネット通販、SNS の利用の常態化など行動様式も大きく変容した。こうした地殻変動ともいえるべき経済構造の変化に対して、シェアリング・エコノミー等の活動や消費者余剰を統計に把握する議論や試みが、各国・国際機関において行われてきた。特に Bean（2016）ではシェアリング・エコノミーの拡大にともない個人・世帯の生産的活動の増加による従来の非市場活動がマネタイズされる一方、基礎統計の制約等により GDP の捕捉漏れという測定上の課題が指摘されてきた²。また、2008SNA の範囲を超え、OECD などの国際機関や欧米諸国などでも、デジタルエコノミーを測定する枠組みに関する議論が具体化してきた。2019 年には OECD 統計局より、デジタルエコノミーに係るサテライト勘定作成の第一歩として、「デジタル供給・

¹ Brynjolfsson et al.（2012）は冒頭で、Wikipedia の記事、Facebook の友達の写真、Google マップ、YouTube のビデオを例示して、その無償サービスの価値の数量化の困難性を指摘している。

² 内閣府経済社会総合研究所（2019）では、2017 年のシェアリング・エコノミーの付加価値を 1,300～1,500 億円程度と推計し、そのうち 800～1,000 億円程度が捕捉できていないとしている。

使用表（SUT）」の概念的枠組み（ガイドライン）が提案され、国際比較可能なデータを作成することに資するものとなった。

デジタル SUT は、統計作成者の実務的可能性とユーザーの多様なニーズを反映してデジタルエコノミーの全体的な構造を把握すべく relevant な SUT 体系の構築を目指している。デジタル化を的確に把握する供給・使用表における生産物及び産業の分類・部門を、生産境界外の（2008SNA を超えた）デジタル生産物を含め新たに設定した。さらには生産物・産業の供給・使用構造のみならず、これとは異なる次元（注文及び配信・配達取引形態）を要素として織り込むことによってデジタルエコノミーの広範な影響を多面的・包括的に把握することを目指している。第 2 章で国際収支統計の「デジタル関連収支」の概要を説明した後、第 3 章でデジタル SUT の構造について整理する。

また、Ahmad and Schreyer(2016)において、消費者余剰を始めとする消費者の経済厚生を総額を測定する尺度は、GDP や所得を測定する概念基盤と一致しないという基本認識のもと、全体として SNA のフレームワークはデジタル化がもたらす難題に対処できていると見られるとしている。しかしながら、人々の日常生活においては、GAFAM ら巨大プラットフォームのビジネスモデルに組み込まれた、GDP では把握できない無償サービスの利用等明示化されない活動が占めるウエイトが増大するなか、GDP 等既存の集計値と実感が乖離した経済活動の可視化へのニーズが高まっている。こうしたニーズに応えるべく、過去数年、デジタルエコノミーの測定上の課題について、SNA の概念的な枠組み、GDP を補完する新たな尺度の試み、価格・数量のデータに関して様々な研究が行われてきたが、第 4 章では主要な研究をサーベイし、その妥当性を評価し、第 5 章でアンケート調査による無償の生成 AI サービスに対する貨幣評価を試算する。

さらに第 6 章では 2008SNA の更新が、2025 年 3 月での国連統計委員会で採択されることを目指し、次の国際基準（2025SNA）において経済社会の構造変化がよりの確に反映されるよう具体的な検討が進められてきたが、デジタル化、そしてこれと不可分とも言えるグローバル化への対応が主要テーマになっており、データの価値の測定及び資産としての計上等の主要論点について整理する。

2. 国際収支統計からみたデジタルサービス取引のグローバル化

「デジタル敗戦」を契機にデジタル庁が 2021 年に設置され、官民連携によるデジタル社会形成の旗振り役として各種 DX 施策を推進している。個人もスマホ、PC を通じた様々なデジタルサービスを享受し、企業においても、オンプレミス（＝自社サーバー）からクラウドサービスの利用、オンライン会議の常態化、そして生成 AI の導入などデジタル化が急速に進んで

いる。

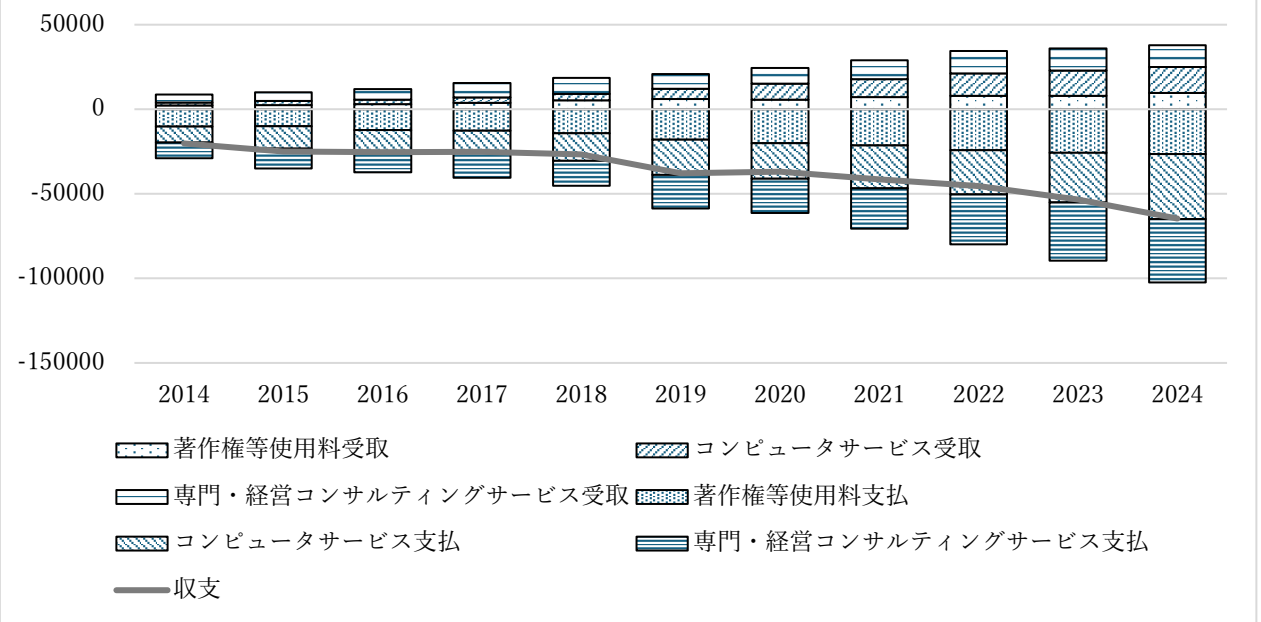
こうしたデジタル化が進めば進むほど、わが国の国際収支における「デジタル赤字」が拡大しており、一昨年からこの問題が大きく注目を浴びている。2014 年に▲2.0 兆円だった赤字が 2024 年は▲6.5 兆円（受取（輸出）3.8 兆円、支払（輸入）10.2 兆円）に急拡大し、インバウンドの増加で稼いだ旅行収支 5.9 兆円の黒字（同 8.1 兆円、同 2.2 兆円）を凌駕している。

国際収支統計のサービス収支は、サービス取引のグローバル化を背景に受取（輸出）・支払（輸入）とも増加している。日本銀行（2023）は、わが国のサービス収支を構成する各項目を取引の特性に応じて分類している。具体的には各項目を、（１）モノの移動や生産活動に関係するもの（モノ関連収支）、（２）ヒトの移動や現地での消費活動に関係するもの（ヒト関連収支）、（３）デジタルに関係するもの（デジタル関連収支）、（４）金融や保険に関係するもの（カネ関連収支）、（５）上記以外（その他）の５つに分類している。

デジタル関連収支に関しては主要な項目は以下の３つであり、「デジタル赤字」の動向を議論する場合はこれらの集計値が利用されている。すなわち、①著作権等使用料（ソフトウェアの製造・販売や音楽・動画の配信に伴うライセンス料等）、②コンピュータサービス（ソフトウェアのダウンロード代金、クラウドサービスやオンライン会議システムの利用料等）、③専門・経営コンサルティングサービス（検索、SNS、動画の各サービスと連動するネット広告料等）、であり、これらの支払の多くが GAFAM から米国のビッグテックに流れている（国富の流出）。①や②において Amazon や Google にスマホ OS 手数料を支払うアプリ事業者をはじめ、多くの企業がデジタルサービスの利用・生産をビッグテックのプラットフォームに依存せざるを得ない、すなわちロックインされている姿は「デジタル小作人」と揶揄されるほどである。

しかしながら、上記の組替による項目は確かにデジタル化を反映しているものの、デジタルサービスを必ずしも対応しているとは言えない取引も含んでおり、注意が必要である。例えば、①にはキャラクター使用のライセンス料や映画の上映・放映権料なども含まれており、③には法務や会計・経営コンサルティングの事務所への支払いやスポーツ大会のスポンサー料の支払いも含まれている。

図表1 過去10年間の「デジタル赤字」の推移（億円）



3. デジタル SUT の枠組み

本章では、第1節で OECD が提唱するデジタル SUT の枠組み、第2節で試算結果と若干の国際比較について説明する。

3.1 デジタル SUT の枠組み

OECD が提唱するデジタル SUT のフォーマットは、“Guidelines for Supply-Use tables for the Digital Economy”（2019、以下「ガイドライン」）に記載されている。デジタル変革（digital transformation）の主要な特性に関わる産業、生産物、取引を区分・定義することによってデジタルエコノミーの追加的な情報を織り込むために SUT の構造が再構成されている。に際、ユーザーにとって意義のあるものとする、推計者にとって実務上可能であることのバランスを取ることが重要な基準となっている。以下、「ガイドライン」に基づき、デジタル SUT の構造を確認していく。

図表2 OECD が提唱するデジタル SUT（供給表）のひな型

供給表	1	2	3	4	5	8	9	10	11	12	13	14	15	16	17	18	19	20
	産出額：デジタル産業計	デジタル基盤産業	デジタル特介サービス（有償）	データと広告収入を主とするデジタル特介サービス	特介プラットフォームに依存する企業	E-タイラー	デジタル専業金融・保険業	他のデジタル専業生産者	非デジタル産業計	全産業計（国内）	輸入（CIF価格）	直接購入	CIF/FOB調整	基本価格	流通マージン	商業輸送	購入者価格	総供給
TP 全て生産物計	中間マトリックス 相当部分																	
TP a デジタル注文																		
TP a.1 取引相手に直接注文																		
TP a.2 国内のデジタルプラットフォームを経由して注文																		
TP a.3 海外のデジタルプラットフォームを経由して注文																		
TP b 非デジタル注文																		
TUP デジタル生産物計	中間マトリックス 相当部分																	
BP デジタル財																		
BP 電子部品・デバイス																		
BP 情報通信機器																		
BP																		
BP デジタルサービス計																		
BP デジタルサービス（クラウドコンピューティングサービス、デジタル特介サービス（有償）を除く）																		
BP 固定電気通信																		
BP 移動電気通信																		
BP ソフトウェア																		
BP																		
BP クラウドコンピューティング（有償）																		
BP デジタル特介サービス（有償）																		
TNDP 非デジタル生産物計（デジタル化で大きな影響を受けている）	中間マトリックス 相当部分																	
NDP 輸送サービス																		
NDP 宿泊サービス																		
NDP																		
TUP 非デジタル生産物計	中間マトリックス 相当部分																	
BP 農産物																		
BP 水産物																		
BP																		
2008SNAの生産境界外生産物計	中間マトリックス 相当部分																	
X データ																		
X 無料デジタルデータ（企業提供）																		
X 無料デジタルデータ（社会提供）																		

出典：OECD（2019）

図表3 OECD が提唱するデジタル SUT（使用表）のひな型

使用表	デジタル基盤産業	デジタル特介サービス（有償）	データと広告収入を主とするデジタル特介サービス	特介プラットフォームに依存する企業	E-タイラー	デジタル専業金融・保険業	他のデジタル専業生産者	非デジタル産業計	自家用	海外団体と企業	中間消費計	最終消費支出：家計	最終消費支出：非居住者	最終消費支出：NFSM	最終消費支出：政府	固定資本形成	在庫変動	貴重品の取得一括	総固定資本形成	輸出	総使用輸入者価格
TP 全て生産物計	中間投入マトリックス 相当部分																				
TP a デジタル注文																					
TP a.1 取引相手に直接注文																					
TP a.2 国内のデジタルプラットフォームを経由して注文																					
TP a.3 海外のデジタルプラットフォームを経由して注文																					
TP b 非デジタル注文																					
TUP デジタル生産物計	中間投入マトリックス 相当部分																				
BP デジタル財																					
BP 電子部品・デバイス																					
BP 情報通信機器																					
BP																					
BP デジタルサービス計																					
BP デジタルサービス（クラウドコンピューティングサービス、デジタル特介サービス（有償）を除く）																					
BP 固定電気通信																					
BP 移動電気通信																					
BP ソフトウェア																					
BP																					
BP クラウドコンピューティング（有償）																					
BP デジタル特介サービス（有償）																					
TNDP 非デジタル生産物計（デジタル化で大きな影響を受けている）	中間投入マトリックス 相当部分																				
NDP 輸送サービス																					
NDP 宿泊サービス																					
NDP																					
TUP 非デジタル生産物計	中間投入マトリックス 相当部分																				
BP 農産物																					
BP 水産物																					
BP																					
中間投入計	中間投入マトリックス 相当部分																				
-2 雇用者報酬																					
-2 生産・輸入に課される税一補助金																					
-2 固定資本減耗																					
-2 営業余剰・混合所得	付加価値マトリックス																				
-2 付加価値計																					
-2 総使用																					
-2 2008SNAの生産境界外生産物計																					
X データ	付加価値マトリックス																				
X 無料デジタルデータ（企業提供）																					
X 無料デジタルデータ（社会提供）																					
X																					

出典：OECD（2019）

3.2 産業の分類

①デジタル基盤産業 (Digitally enabling industries)

デジタル基盤産業とは、主に「電子的手段によって、送信及び表示を含む情報処理・通信の機能を発揮する、あるいは発揮させることを目的とした」生産に従事する産業とされる。これらの産業は、ちょうど国際標準産業分類第4回改定版 (International Standard Industrial Classification, Revision 4、以下は、ISIC rev. 4 と略称する) で定義されている「情報通信技術」(ICT) セクターを構成する分類と一致する。すなわち、「電子部品・デバイス製造業」、「通信機械・同関連機器製造業」、「電子計算機・同付属装置製造業」、「電子・電話業」、「情報サービス業」の各産業である。

②デジタル仲介プラットフォーム業(課金型) (Digital intermediary platforms charging a fee)

「インターネットを介してやり取りする2つ以上の異なるユーザー(企業または個人)間のやり取りを促進するデジタルサービス」と定義される。このうち「仲介を促進するために生産者または消費者から明示的な支払いを受け取る制度単位」である。シェアリング・エコノミーにサービスを提供するマッチングプラットフォーム (Uber 等)、オークションサイト、リソースシェアリングプラットフォーム、その他のオンラインブローカーサービス等も含まれる。

③データ・広告駆動型デジタルプラットフォーム (Data and advertising driven digital platforms)

この分野には「主にデータまたは広告スペースの販売を通じて収益を生み出している、もっぱらオンラインで営業する制度単位」が含まれる。例えば、ソーシャルメディアプラットフォーム、検索エンジン、知識共有プラットフォーム、そして無料通話アプリなどがこのカテゴリーに含まれる。

④仲介プラットフォーム依存企業 (Firms dependent on intermediary platforms)

「財やサービスの需要の大部分が、仲介プラットフォームを経由している企業」を指す。OECD ガイドラインによると、この決定基準を別にすれば、各企業は様々な特性を示している場合があり、幅広い品揃えとサービスを提供するだけでなく、大規模な国際企業(ホテルチェーン)から小規模な独立業者(食品注文の宅配業者)に至るまでの制度単位が含まれている。

「仲介プラットフォームに依存する企業」の定義上重要なことは、「収入を生み出すために、消費者への主なアクセス手段が1つ以上の独立したデジタル仲介プラットフォーム経由に限定されている」ことである。このため仲介プラットフォームの利用が副次的なものにとどまり、プラットフォームを介する需要が50%未満の企業は、当該産業には区分されず、

それぞれの該当する ISIC rev.4 の産業として区分される。

⑤E-テイラー (E-Tailor)

E-テイラーは、「注文の大部分をデジタルで受け、財またはサービスの購入及び再販売に従事する小売業者と卸売業者」となる。デジタルでの注文を受け付けている場合でも、その割合が50%未満の小売業者及び卸売業者は、E-テイラーではなく、それぞれの該当する ISIC rev. 4 の産業（非デジタル産業）として区分される。

⑥デジタル専門金融・保険サービス (Digital only firms providing financial and insurance services)

デジタル専門金融・保険業は、「顧客と物理的に接することがなく、デジタルのみで営業している企業のみが含まれる」と定義されている。すなわち、実際問題として金融・保険サービス及び関連サービスを提供するほとんどの企業は、顧客との取引の大部分がデジタル化されているが、単にこの部門全体をデジタル産業に分類するのではないとしている。

⑦その他のデジタル専門生産者 (Other producers only operating digitally)

その他のデジタル専門生産者は、⑥までの産業分類のいずれにも属しておらず、デジタルでのみ営業しているすべての企業で構成されている。「製品またはサービスがデジタルで注文されるだけでなく、デジタルで配信される企業」がこの産業カテゴリーになる。具体的には、オンラインゲームやストリーミングサービス、サブスクリプションベースでデジタルコンテンツを提供する企業が含まれる。

3.3 生産物の分類

デジタル SUT は生産物を次の4タイプに区分している。(i) デジタル生産物 (SNA の生産境界内)、(ii) デジタル化により大きな影響を受ける非デジタル生産物、(iii) 他の非デジタル生産物、(iv) デジタル生産物 (SNA の生産境界外)。

(i) デジタル生産物

①ICT 財 (ICT goods)

ICT 財とは「電子的手段によって、送信及び表示を含む情報処理・通信の機能を発揮する、あるいは発揮させることを主要な目的とした」生産物により構成されると定義されている。この定義は、中央生産物分類第 2.1 改定版 (Central Product Classification Version 2.1、以下 CPC 2.1 と略称) に含まれる ICT に関連する生産物の分類と一致する。具体的には、i) コンピューター機器及び周辺機器、ii) コミュニケーション機器、iii) 消費者電子機器、

iv) その他の ICT 部品及び財、が挙げられる。

②デジタルサービス（クラウドコンピューティングサービス、デジタル仲介サービス（有償）を除く）（Priced digital services – except cloud computing services and digital intermediary services）

ここでのデジタルサービスとは③のクラウドコンピューティングサービスと④のデジタル仲介サービス（有償）を除いたデジタル基盤産業に関連するデジタルサービスを指している。具体的には、i) ICT 機器製造に関連するサービス、ii) ビジネス及び生産性ソフトウェアとライセンスサービス、iii) 情報技術のコンサルティングとサービス、iv) 通信サービス、v) ICT 機械のリースとレンタル及びvi) その他の ICT サービス、が挙げられる。

③クラウドコンピューティングサービス（有償）（Priced cloud computing services）

クラウドコンピューティング（有償）サービスとは、「低いマネジメントエフォートで柔軟に、弾力的に、オンデマンドでアクセスできる一群のコンピューティングリソースに基づくサービス」とされている。

このクラウドコンピューティングサービス（有償）には、SaaS（Software as a Services）、PaaS（Platform as a Service）、IaaS（Infrastructure as a Service）³等いわゆるクラウドプラットフォーム（AWS、Azure、GCP 等）による各種サービスが含まれる。

④デジタル仲介サービス（有償）（Priced digital intermediary services）

従来の国際標準産業分類には、「デジタル仲介サービス（有償）」の正式な定義はなく、デジタル SUT の目的に合わせて次のように定義された。

デジタル仲介サービス（有償）とは、「2つの独立した当事者に対して、デジタルプラットフォーム⁴を介して情報を提供し、うまく取引をマッチさせ、見返りに明示的な料金を課すサービス」である。これらのプラットフォームの産出額は、仲介される生産物の生産者または消費者によって支払われる料金である。

なお、以上の分類を小計した集計値を行として追加している。

³ SaaS（利用者がプロバイダーのアプリケーションに容易にアクセスできる）

PaaS（利用者が、プロバイダーのプラットフォームに自社のアプリケーションを展開できる）

IaaS（利用者がオペレーティングシステム、ストレージ、展開されたアプリケーションを制御できる）

⁴

(ii) 非デジタル生産物

非デジタル生産物であっても、デジタル取引の普及に焦点を当ててその影響を受けるものとそれ以外とに分類している。

① デジタル化により大きな影響を受ける非デジタル生産物 (Non-digital products – significantly affected by digitalisation)

これら非デジタル生産物の選択は、デジタル配信されたり、デジタル仲介プラットフォームによって大きく影響を受けてきたりしたことで、消費者への販売方法がデジタル変革の影響を大きく受けてきた（または早晚受けると見込まれる）ことに基づいている。取引形態の細分化が推奨されている。

現時点は以下の 10 品目で構成されているが、品目構成は今後のデジタルエコノミーの進展を反映して変更される可能性がある。i) 陸上輸送、ii) 宿泊、iii) 飲食、iv) 出版、v) 映像制作、音楽出版、vi) 金融・保険、vii) 広告・市場調査、viii) 旅行代理店・予約サービス、ix) 教育、x) 賭博

(iii) その他の非デジタル生産物 (Non-Digital products – other)

(iv) SNA の生産境界外のデジタル生産物 (Digital products outside of the SNA production boundary)

① データ (Data (beyond 2008SNA))

OECD が提唱するデジタル SUT におけるデータとは、「無償」で利用可能で、財やサービスの生産に使用されるデータに関するものである。これには、通常の生産プロセスの副産物である情報や、消費者に無償または割引サービスを提供する見返りとして消費者から収集した情報が含まれる。

記録されるべきデータの分類（資産か否か、さらには生産資産か非生産資産か等）や評価方法に関する一致したアプローチが確立していない。しかしながら、後述するように、次期 SNA の改定でも主要な課題になっており、概念や評価に関しては一定の方向性に向けて議論されている。

② 企業が提供するデジタルサービス（無償）(Digital services (beyond 2008SNA) , provided by enterprises)

「企業が提供するデジタルサービス」とは、企業が提供する無償のデジタルサービスである。インターネットを通じた容易な情報収集、ソーシャルメディアを通じた他者とのつながり、デジタル手段による無償の娯楽、などが含まれる。通常、これらのサービスは、家計部門によって「消費」されるが、生産プロセスでも使用される。

しかしながらこのタイプのデジタルサービスも推計方法に関して一致したものではなく、

例えば、①サービス提供者が生み出した広告収入のデータ、②当該サービスへの消費者の支払意思額、③政府支出同様に生産費用の計上、④プロキシ及び market sources などによる金銭的価値が提案されている。当該概念の推計方法に関しても後述する。

③社会が提供するデジタルサービス（無償）（Digital services (beyond 2008SNA) , provided by communities)

これには、社会による無償のデジタル資産（Linux など）の創出及びこれらの資産から派生した無償サービス含まれる。これらのサービスは、②のサービスは異なり、単一の主体ではなく複数主体の協働により生産されたものである。同様に、これらの成果物として生み出された資産は単一主体に帰属するものとはならない。これらの生産物は一定範囲の独立した生産者により開発され、金銭的成本を伴うことなく万人が利用可能である。これらはビジネスにおける生産要素としてのみならず、最終財としても消費される。

3.4 取引形態

（i）デジタル注文

さらにデジタル取引の拡張に焦点を当てて、デジタル SUT では財・サービスの供給及び使用でデジタル注文と非デジタル注文への分類（集計値レベル）が推奨されている。

①デジタル注文（Digitally ordered）

注文用のコンピューターネットワークを通じて行われる電子商取引である（支払い及び財やサービスの配送がオンラインで行われる必要はない）。

②取引相手からの直接注文（Ordered directly from a counterpart）

③居住者のデジタル仲介プラットフォームを通じた注文（Ordered via a resident digital intermediary platform）

④非居住者のデジタル仲介プラットフォームを通じた注文（Ordered via a non-resident digital intermediary platform）

⑤非デジタル注文（Not digitally ordered）

物理的あるいは他の非デジタル（電話等）を通じて注文が行われた場合の電子的支払いを排除するものではない。

（ii）デジタル配信

デジタルで配信されるか否かについて、供給表における財・サービスの産出、輸出、輸入（行列の列）を分割する。なお、ここでは財はデジタルでは配送できないと仮定。

①デジタル配信 (Digitally delivered)

「電子的なダウンロード可能な形式でインターネット等 ICT ネットワークを通じて遠隔配送される取引」(OECD-WTO, Handbook on Digital Trade) であり、通信、ソフトウェア、クラウドコンピューティングのようなデジタルサービスの配信、また教育やギャンブルのような非デジタルサービスのデジタル配信を含む。

②非デジタル配送 (Not digitally delivered)

3.5 諸外国との比較

内閣府ではデジタル SUT を 2015 年及び 2018 年を対象に推計しているが、具体的な推計方法及び推計結果については、内閣府 (2022)、長谷川 (2023) を参照されたい。ここでは既にデジタル SUT の推計を公表しているカナダ、米国、オーストラリアそして日本について、基礎統計の制約等に起因する推計方法の特徴とデジタル産業の付加価値の GDP 比を比較し、整理した (図表 4)。付加価値の GDP 比に関しては日本は米国に次いで高くなっている。米国は「仲介プラットフォームに依存する企業」の推計は行っていないなど、対象とする産業の範囲は日本より狭いが、デジタル化の進展を反映してデジタル産業の GDP 比は日本より高くなっている。

図表 4 諸外国との比較

	推計方法の特徴	推計期間	デジタル産業の規模 (付加価値の GDP 比)
日本	OECD のガイドラインに沿って推計。 (ただし、「仲介プラットフォームに依存する企業」に自社サイト経由を含む、「E-テイラー」に卸売業を含まない、などの点で OECD ガイドラインと異なっている。)	2015、2018 年	7.6% (2018 年)
カナダ	OECD のガイドラインに沿って推計。	2017～2019 年	5.5% (2019 年)
米国	デジタルインフラ (ハード、ソフト)、E-コマース、有償のデジタルサービスの産出額、付加価値 (名目、実質) の時系列データを推計 (産業の範囲は OECD ガイドラインと異なる。)	2005～2019 年	9.6% (2019 年)
オーストラリア	米国の手法に基づき推計。	2017-18～ 2019-20 年度	5.9% (2019-20 年度)

注：各国の推計手法の違いに注意する必要がある。

出典：(カナダ) Statistics Canada (2021)、(米国) BEA (2021)、(オーストラリア) ABS (2021)、(日本) 内閣府 (2022)

4. 無償デジタルサービスの測定に関する研究について

Bean (2016) においてシェアリング・エコノミー、デジタル生産物、無形資産等の測定が課題として取り上げられたが、これに触発された形でインターネット上の無償サービスから得られる消費者の効用に対する関心が高まった。GDP には無償サービスにともなう消費者余剰が含まれていないことに関する批判的な評価もあったが、消費者余剰のような厚生を測定する尺度が、GDP や所得を測定する概念的基盤と整合しないことはその目的から明らかである⁵。また、無償サービスの分 GDP が過少評価されるのではないかという議論もミスリーディングであろう。無償サービスを提供する巨大プラットフォームらは、広告を主要な収入源としているが、広告主は、広告費用（中間消費）を広告対象生産物の価格の一部として回収し、結果的に消費者が間接的に費用負担している。すなわち無償サービスの費用は生産・支出において GDP に計上されている。

一方、上記のような無償サービスは、消費者・ユーザーに多大な便益をもたらし、その厚生を向上させているが、従来の GDP 推計フレームでは概念上も推計上も把握されていないため、これらの無償サービスの価値を測定し、それらを反映した厚生を評価する他の指標により GDP を補完する必要性が高まっている。本章では、無償デジタルサービスの尺度及び測定方法に関する主要な3つの枠組みについて概観する。

企業（広告主）にとっては、消費者に自社商品を購入してもらうために、広告を資金源とするテレビ・ラジオ等のメディアサービスやインターネット上のお金(money)を払う (pay) 必要のない各種の無償デジタルサービスに対して、いかに関心 (attention)、あるいは時間 (time) を払って (pay) もらうことが大問題なのである。1つ目の枠組みは、無償デジタルサービスに費やされる時間が急増し人々の生活時間 (waking hours) の中で大きなウェイトを占めていることに注目して、Becker (1965) 等によって展開された家計における時間配分 (allocation of time) の理論に基づいてその価値を測定するものである (Brynjolfsson and Oh 2012; Goolsbee and Klenow 2006)。インターネットに費やされる時間は、他の財・サービスの消費に費やされる時間や労働時間の機会費用を伴うものとして、利用可能な時間及び賃金率を制約条件に時間ベースの効用関数(CES 型)⁶を最大化するモデルを構築している。これに基づく推計結果は、米国において 2007–2011 年の間、インターネットからの厚生（消費者余剰）の増加（等価変分）は毎年約 1590 億ドル、うち無償サイトから約 1060 億ドル、一方テレビから約 720 億ドルとなっており、消費者の厚生増加へのインターネット

⁵ 理論的には最適成長モデルにおけるハミルトン関数は、消費の効用と純投資のシャドープライス（帰属価値）による効用評価を足し合わせものであるため、効用単位で測った国民純生産あるいは将来消費の割引現在価値と見なすことが出来る (Weitzman 1976; Solow 2000)。

⁶ 消費ベクトルは、①インターネットに費やされる時間、②テレビに費やされる時間、③他の財及び余暇の消費、から構成される。

の寄与が高まっている。

図表5 消費者余剰の比較（テレビ、インターネット、無償サイト）

（10 億ドル）

	テレビ		インターネット		無償サイト	
（2007～2011 年平均）		年増分		年増分		年増分
消費者余剰	1462	72	838	159	559	106
対 GDP	10.17%	0.50%	5.83%	1.10%	3.89%	0.74%

出典：Brynjolfsson and Oh（2012）

上記の枠組みは、無償デジタルサービスの評価を費やされる時間に基づく新たな推計枠組みであるが、消費者余剰を利用した厚生変化を推計するものであるため、SNA の構成項目や GDP の補完として利用することは困難である。Brynjolfsson et al.（2019, 2022）は、無償財の厚生への貢献を数量化に関して、消費者余剰に依らない、GDP に代わる経済指標として「GDP-B」⁷を提唱している（2つ目の枠組み）。これは現行の GDP 推計フレームを拡張したもので、無償財を放棄することにもなう補償の受取意思額（willing-to-accept: WTA）について、誘因両立的な選択メカニズムである BDM 方式⁸の実験経済学的手法により推計し、これに基づいて次の2つのアプローチを提示している。1つ目の「全所得（total income）」アプローチは、補償額を現行の GDP に追加することにより、無償財を市場財に代替して評価すること、2つ目の「厚生変化」アプローチは、無償財の需要曲線を受取意思額の分布に基づいて推計し、需要曲線から得られた留保価格（Hicksian reservation prices）を当該財の出現前の価格として数量指数に反映し、これを厚生変化（ベネット変分⁹）に適用して GDP の調整を図っていくことである。

米国のケースにおいて、2,885 人を被験者として Facebook に関して WTA のオンライン実験調査を実施した。通常の GDP よりも GDP-B で測定した成長率（2004-2017）の方が Facebook¹⁰によって 0.05-0.11%高くなっている。同様にオランダのケースでは、426 人を被験者に実験した結果によると、WhatsApp¹¹が 2004-2017 年で年平均 0.82%分成長を押し

⁷ B は、無償財からの benefit と beyond GDP の頭文字に依る。

⁸ Becker, DeGroot and Marschak（1964）を参照。

⁹ ベネット変分(Bennet variation)は、価格及び数量のベクトルを p, q とすると 0 期から 1 期では

$$V_B(p^0, p^1, q^0, q^1) \equiv \frac{1}{2}(p^0 + p^1) \cdot (q^1 - q^0)$$

と表すことが出来る。

¹⁰ Facebook の WTA 中央値は年 506 ドル、2003 年の留保価格は同 2152 ドル。

¹¹ WhatsApp の WTA 中央値は年 536 ユーロ。

上げている。各ケースの成長押し上げ寄与は1社のみのサービスによるものであり、各社が提供する無償サービスをカバーするならば、GDP-Bの成長率はGDPを大幅に上回る可能性があることが伺わせる。このようにGDP-Bの概念導入によりGDPまたは生産性の測定する際において、現行において欠落しているこれらの無償あるいは新規のデジタル財の寄与を取り入れることができ、また無償財の厚生変化への貢献に関して指数理論的な表現を与えた意義は大きい。一方、GDP-Bは、SNAの概念的枠組み上の扱いをどうするかという課題があり、また統計実務的な観点からは、測定すべき無償財の範囲の問題や、推計は市場から得られる価格・数量に関するデータではなく、実験経済学的手法等複雑なプロセスを経る必要があるなど実施上の課題は大きい。

一方、Brynjolfsson et al. (2019, 2022) に喚起され、無償デジタル財の扱いをGDPの生産境界、勘定体系の枠組みの検討にまで及ぶ議論を展開し、その試算を行っているのが、Schreyer (2022) である（3つ目の枠組み）。これまでの試算は、GDPへの影響を含め支出側からのアプローチであり、勘定体系としては不十分な議論と言える。これに対して、無償デジタル財に対するWTAのような消費者の評価が、広告やデータ売却の事業がもたらす付加価値と大きく乖離することに着目し¹²、こうした付随的な価値の発生の帰属を家計に求め、家計が「デジタル利用の余暇サービス」(digitally-enabled leisure services)を生産し、そして消費するという自己勘定を設定することを提案している。すなわち、家計はFacebookのようなソーシャルメディアを使用して、余暇の時間と、PC等IT資本財（現行では耐久消費財ストックに分類）からの資本サービスを利用して生産すると規定する。このような家計の自己勘定¹³及びGDPを対象範囲とする、一種の拡張された活動尺度(Extended Measure of Activity: EMA)を検討している。

試算においては、上記のFacebookのWTA（2017年）から生産額を求め、家計の保有するICT耐久消費財ストックに適当な実質収益率（4%）、減耗率（20%）を乗じた使用者費用から時間単価を求め、Facebookの利用時間（2017年243時間）を乗じて、Facebook分の使用者費用を求める。次に、生産額から使用者費用を控除して、投入された余暇時間のシャドープライス（賃金率）を求める。2004年に関しては、使用者費用は2017年と同様に、余暇時間のシャドープライスはこの間の米国の賃金率の上昇率で調整する。これらにより費用積上げによる2004年のFacebookの利用した余暇サービスの生産額を算出する。

一方、当該サービスの単位費用を求めるために、上述の賃金率の変化率とICT耐久消費財ストックのデフレータ変化率、そして質の変化をネットワーク効果（ネットワーク数の参加者数）として反映して求める。今回はWTAの時系列的な蓄積がないため、ネットワーク効果に関する価格変化の弾力性を3つのシナリオで設定して、試算を行っている。その結果は、Facebookの影響として、成長率は年率（2004-2017）で-0.04～+0.2%ポイントEMAは

¹² Facebookの利用者1人当たりの広告収入（2017年）は約25ドル。

¹³ 現行では家計が生産するサービスは持家の帰属家賃を除き生産境界外。

GDP を上回っている。

なお、このような家計自身を余暇サービスの生産者かつ消費者とする自己勘定の構築に際しては、家計による生産のタイプを Becker (1965) や Diewert et al. (2017) 等において展開された家計における時間配分の理論に着想を得て分類している。すなわち、①労働市場での仕事（タイプ1）、②市場からでも購入し得る料理や介護のような家事サービスの生産（タイプ2）、③動画（movie）の視聴、サッカーをすること、あるいはソーシャルメディアを使用した他者との交流のような市場からは購入し得ない余暇サービスの生産（タイプ3）、の3分類である。この分類において、無償デジタル生産物はタイプ3に含まれることになる。このため GDP への計上という観点では、生産境界の第三者基準に照らし合わせると、タイプ3は生産境界外であり、タイプ3を包含するような生産境界の拡張は、タイプ2も包含するのが自然であり、また、こうした生産境界の拡張は、GDP の姿を水準的にも成長率でも大きく変更することになる可能性がある。このため推計値の頑健性やユーザーの理解も考慮するならば、今回の勘定体系の新たな理論的アプローチは整合性・包括性を有しているものの、まずは無償デジタルサービスに係る EMA のサテライト勘定で取り組んでいくことが現実的であろう。

図表6 Facebook の GDP 成長率への寄与

(%)

(2004～2017 年平均)	「総所得」アプローチ	「厚生変化」アプローチ
GDP 成長率	1.83	
GDP-B の成長率	1.87	1.91
	EMA	
EMA の成長率	1.77～2.00	

出典：Brynjolfsson et al (2019) 及び Schreyer (2022)

5. 生成 AI サービスに関するアンケート調査

生成 AI サービスの利用については、国内外で大きな議論になっており、総務省（2024）では、日本、米国、中国、ドイツ、英国の国民を対象に利用状況等についてウェブアンケート調査を実施した（2023 年 12 月～2024 年 3 月）。今回、これを踏まえてその後の利用状況の変化と、前出の無償デジタルサービスのひとつとして生成 AI サービスの貨幣評価について、大久保（2023）の各種デジタルサービスへの支払意思額（willing-to-pay:WTP）の調査に則して、生成 AI サービスへの貨幣評価に関するネット調査を実施した（2025 年 3 月）。なお、回答者数は 2000 人である。

これによると生成 AI を“使っている”（「過去使ったことがある」も含む）と回答した割合は、27.7%（2024 年の人口構成で調整すると 20.6%）であり、総務省の調査結果の 9.1%と比べて高くなっているが、他国と比べて依然低くなっている。

使っていない理由については、「使い方がわからない」、「自分の生活には必要がない」との回答が多く、「魅力的なサービスがない」、「利用環境が整っていない」が続き、概ね総務省の回答結果と同様である。

一方、今後の暮らしや娯楽における生成 AI の活用意向について聞いてみると、「既に利用している」と回答した割合は比較的高いものの、「ぜひ利用してみたい」「条件によっては利用を検討する」と回答した割合が総務省の結果より総じて低くなっており、活用が進展していることの裏返しとして潜在的なニーズが低くなっている可能性が伺える。

生成 AI サービスへの支払い意思の比率をみると、43%になっており、大久保（2023 年）の各種デジタルサービスでは 20～30%であるので、かなり高い水準になっていることが伺える。

次に支払意思額をみると、全般的に低い額（月額 100 円～1000 円以下）に集中している。支払意思額の平均は、ゼロを含めて計算すると概ね 500 円程度、ゼロを含めないと 1000 円程度となっている¹⁴。また、ゼロを含めない場合の中央値は 500 円程度になっている。概算になるが、Brynjolfsson et al. (2019, 2022) の「全所得」アプローチに準じつつ、かつ WTP が WTA に近似できると仮定するとマクロ全体では 1,200 億円程度の所得増をもたらしていることになる。ただし、ネット調査のサンプルバイアス等もあり十分幅を持つてみる必要がある。

¹⁴大久保（2023 年）の各種デジタルサービスが、ゼロを含めると 300～400 円程度、ゼロを含めないと 1400～1900 円程度になっている。

図表 7 生成 AI サービスの利用状況と価値評価

Q1 あなたは生成 AI を使ったことがありますか。(SA)

		回答数	%
全体		2000	100
1	「使っている（過去使ったことがある）」	554	27.7
2	「使っていない（過去使ったことがない）」	1446	72.3

Q2 生成 AI を「使っていない（過去使ったことがない）」と回答された方にお伺いします。生成 AI を利用しない理由は何ですか。(MA)

		回答数	%
全体		1448	100
1	魅力的なサービスがない	171	11.8
2	使い方がわからない	546	37.7
3	利用環境が整っていない	157	10.8
4	費用が高額である	110	7.6
5	情報漏洩、セキュリティ、安全性に不安がある	112	7.7
6	品質に不安がある	84	5.8
7	従来の文化や習慣を変えられない	48	3.3
8	自分の生活には必要ない	679	46.9
9	その他	110	7.6

Q3_1 あなたは暮らしや娯楽における生成 AI や AI を利用することについてどう考えますか？以下の各項目に対して、あなたの考えにあてはまるものを回答してください。／コンテンツの要約・翻訳をする (SA)

		回答数	%
全体		2000	100
1	既に利用している	263	13.2
2	ぜひ利用してみたい	317	15.9
3	条件によっては利用を検討する	671	33.6
4	利用には後ろ向きである	232	11.6
5	利用したくない	517	25.9

Q3_2 あなたは暮らしや娯楽における生成 AI や AI を利用することについてどう考えますか？以下の各項目に対して、あなたの考えにあてはまるものを回答してください。／画像や動画を生成する（SA）

		回答数	%
全体		2000	100
1	既に利用している	139	7.0
2	ぜひ利用してみたい	297	14.9
3	条件によっては利用を検討する	625	31.3
4	利用には後ろ向きである	341	17.1
5	利用したくない	598	29.9

Q3_3 あなたは暮らしや娯楽における生成 AI や AI を利用することについてどう考えますか？以下の各項目に対して、あなたの考えにあてはまるものを回答してください。／旅行の計画やイベントの企画をする（SA）

		回答数	%
全体		2000	100
1	既に利用している	95	4.8
2	ぜひ利用してみたい	352	17.6
3	条件によっては利用を検討する	680	3.4
4	利用には後ろ向きである	287	14.4
5	利用したくない	586	29.3

Q3_4 あなたは暮らしや娯楽における生成 AI や AI を利用することについてどう考えますか？以下の各項目に対して、あなたの考えにあてはまるものを回答してください。／調べものをする（SA）

		回答数	%
全体		2000	100
1	既に利用している	283	14.2
2	ぜひ利用してみたい	426	21.3
3	条件によっては利用を検討する	641	32.1
4	利用には後ろ向きである	201	10.1
5	利用したくない	449	22.5

Q3_5 あなたは暮らしや娯楽における生成 AI や AI を利用することについてどう考えますか？以下の各項目に対して、あなたの考えにあてはまるものを回答してください。／対話型 AI と会話する（SA）

		回答数	%
全体		2000	100
1	既に利用している	186	9.3
2	ぜひ利用してみたい	310	15.5
3	条件によっては利用を検討する	592	29.6
4	利用には後ろ向きである	324	16.2
5	利用したくない	588	29.4

Q3_6 あなたは暮らしや娯楽における生成 AI や AI を利用することについてどう考えますか？以下の各項目に対して、あなたの考えにあてはまるものを回答してください。／AI から好み・生活様式に合った提案を受ける（服装、献立、移動ルート等）（SA）

		回答数	%
全体		2000	100
1	既に利用している	106	5.3
2	ぜひ利用してみたい	379	19.0
3	条件によっては利用を検討する	654	32.7
4	利用には後ろ向きである	295	14.8
5	利用したくない	566	28.3

Q3_7 あなたは暮らしや娯楽における生成 AI や AI を利用することについてどう考えますか？以下の各項目に対して、あなたの考えにあてはまるものを回答してください。／AI を活用して病気や健康に関するアドバイスを受ける（服装、献立、移動ルート等）（SA）

		回答数	%
全体		2000	100
1	既に利用している	90	4.5
2	ぜひ利用してみたい	407	20.4
3	条件によっては利用を検討する	697	34.9
4	利用には後ろ向きである	269	13.5
5	利用したくない	537	26.9

Q3_8 あなたは暮らしや娯楽における生成 AI や AI を利用することについてどう考えますか？以下の各項目に対して、あなたの考えにあてはまるものを回答してください。／AI を活用して制作への改善やアドバイスを受ける（プログラムコード、DIY、家庭菜園等）（SA）

		回答数	%
全体		2000	100
1	既に利用している	107	5.4
2	ぜひ利用してみたい	348	17.4
3	条件によっては利用を検討する	695	34.8
4	利用には後ろ向きである	280	14.0
5	利用したくない	570	28.5

Q4 生成 AI を「使っている（過去使ったことがある）」と回答された方にお伺いします。利用した生成 AI には無料で利用できるサービスも多いと思われますが、ある日、それらが有料化されることになったとします。その場合、これまでと同じサービスを利用するために月額いくらまでなら支払ってよいと思いますか。（SA）

		回答数	%
全体		554	100
1	無料でないと使わない	314	56.7
2	月額 100 円まで	42	7.6
3	月額 300 円まで	61	1.1
4	月額 500 円まで	57	10.3
5	月額 1,000 円まで	38	6.9
6	月額 2,000 円まで	16	2.9
7	月額 3,000 円まで	6	1.1
8	月額 5,000 円まで	14	2.5
9	月額 7,000 円まで	2	0.4
10	月額 10,000 円まで	1	0.2
11	月額 10,000 円を超えても支払う	3	0.5

6. 次期 SNA (2025SNA) におけるデジタル化への対応の方向性

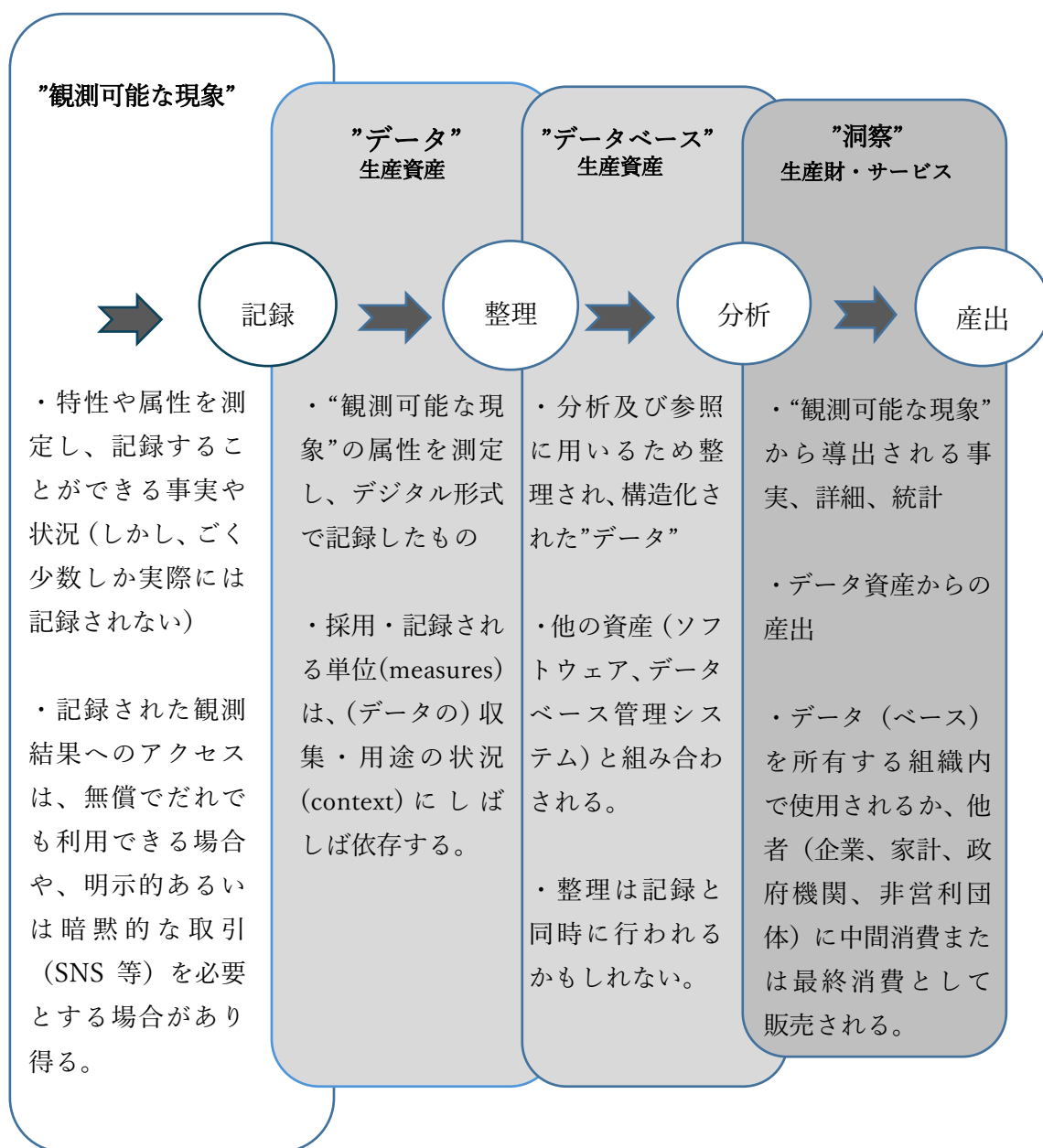
6.1 データの価値の測定について

データは、経済のほぼ全ての局面において、生産活動への非常に重要なインプットであり、消費され、あるいは繰り返し使用される。しかし、現行の 2008SNA では独立した資産として明示されておらず、固定資産（知的財産生産物）のデータベースの自己勘定による生産を推計する際、データ投資の一部が当該資産に含まれることに留まっている。しかしながら、マクロ経済統計としての SNA の妥当性という観点からは、デジタル化とデータの生産が一般的な経済活動となり、データが生産性に与える影響について追加情報が求められており、次期 SNA (2025SNA (仮称)) の中枢体系にデータを組み込むべきとされている。次期 SNA に向けた改定作業では、データの価値測定と資本化が議論されてきたが、主要な方向性について以下のように整理される。

データは記録された情報を表す言葉として広く使用されているが、情報の要件事項は極めて多様である。今回、マクロ経済統計におけるデータの記録という観点から、データは「現象にアクセスし、観測し、これらの現象から得られた情報要素をデジタル形式で記録、整理、保存することによって生成される情報内容であり、生産活動に利用されることで経済的利益をもたらすもの」と定義されている。むしろデータは紙媒体などのような非デジタル形式でも存在するが、デジタル化されたデータのみが生産の成果とされる。すなわち労働と資本の投入によって生産された無形生産物であり、資産計上される場合は生産資産（知的財産生産物）として分類される。データも経済的所有権、評価（および再評価）、減価償却の対象となる。ただし、減価償却は消耗によるものではなく、むしろ陳腐化によって生じる。

なお、注意しなければならないのは、データと「観測可能な現象 (observable phenomena: OP)」の情報要素とは異なることである（図表 8）。「観測可能な現象」とは、「特性や属性が記録される事実や状況」と定義され、非生産資産であり、一般に価値を持たないとみなされるが、明示的な購入があった場合は、サービス（産出）の購入あるいは特定の観測可能な現象へのアクセスを提供するための賃貸料支払いであると考えられる。

図表 8 データ - 情報チェーン



(出所) Mitchell 他 (2021)、河野・吉本 (2024) より筆者作成

またデータの自己勘定による生産はすべて資本形成とみなされ（１年以内に消費されるデータを区別するのは実務上困難）、一方市場取引で購入したデータは１年以上使用する意図する場合は資産計上、生産中に消費された場合は中間消費の扱いとなる。

データの自己勘定生産は費用積上げ方式で評価され、他の知的財産生産物のケースと同様の扱いである。すなわち、①データ生産の戦略に関する企画・準備・開発に要する費用、②OPに埋め込まれた情報へのアクセス、記録、保存に関する費用、③データから情報や結論を導出するためのデータの設計、整理、テスト・分析に関する費用がカバーされ、これらに係る人件費（雇用者報酬）、物件費（中間消費）、さらにデータの生産に使用される固定資本の減耗、市場生産者の場合の営業余剰（純）のマークアップが含まれる。

こうしたデータの測定の基準作りに向けた取組がなされる中、諸外国においては、費用積上げ方式によりデータ資産の試算値を作成する動きがある。米国、カナダ、オーストラリアの試算結果によると、米国では、データの投資（business sector）がGDPに占める割合は2021年で約0.8%、カナダでは2018年で0.4～0.6%、オーストラリアでは2016年で0.8～1.1%で、各国とも平均成長率は3～4%前後となっている。GDPに占める割合は同じ知的財産生産物である研究開発に遠く及ばないものの、経済のデジタル化の進展に伴って機械・設備や建物などの伝統的な固定資産への投資に比べると高い傾向にあると言える。また米国、カナダは恒久棚卸法（PIM）によりデータの純資本ストック値も試算しており、米国では2002年から2021年において205（十億米ドル）から421（同）、カナダでは2005年から2018年において53～72（十億加ドル）から105～151（同）となっている¹⁵。

日本では河野・吉本（2024）で試算が行われている。ウェブアンケート調査でデータ関連業務への従事の有無と就業時間に閉めるウエイトの情報を得て、これに国勢調査・賃構造基本統計調査から計算した職業別人件費と合わせて労働コストを推計し、産業連関表から得たマークアップ率を設定して産出額を推計している。推計結果をみると2020年6,750（十億円）、対GDP比1.3%は、概ね諸外国の先行研究と同様の結果であり、「想定的に規模が小さく、成長率が安定的であることも踏まえると、現時点ではGDPへの影響は限定的であり、また、総固定資本形成への影響も比較的小さい」というISWGNAの評価を共有している。

¹⁵ データの純資本ストック値の試算に当たっては、米国では幾何級数的な減価償却プロファイル（減耗率はソフトウェアの0.33を適用）、取得時価格（historical-cost）の投資累計、カナダでは幾何級数的な減価償却プロファイル（耐用年数25年）、労働投入費用と3%の資本サービス費用に依る価格指数（2005年=100）に基づく2005年価格の投資累計（実質値）を仮定している。

図表9 各国におけるデータ資産の試算結果

	金額(各国通貨)	名目 GDP 比	成長率
アメリカ (2021 年)	186(十億米ドル)	0.8%	過去 10 年平均 4.7%
カナダ (2018 年)	9.4~14.2(十億加ドル)	0.4%~0.6%	過去 8 年平均 2.8~3.4%
オーストラリア (2016 年)	14~20(十億豪ドル)	0.8%~1.1%	過去 5 年平均 2.4~2.8%
日本 (2020 年)	6,750(十億円)	1.3%	過去 10 年平均 2.5%

(出所) Santiago Calderón and Rassier (2022)、Statistics Canada(2019b)、Crick (2022)、河野・吉本(2024)より作成。

6.2 経済のデジタル化とグローバル化の複合による課題

デジタル化は既存の輸出入の取扱いに関しても大きな課題を投げかけている（多田 2022）。国境を越えた国内居住者と非居住者の間で財貨・サービスの経済的所有権の移転・取引が輸出入の基本的な概念であるが、クラウドの出現によりデジタルサービスの輸入先の国や生産拠点が判然としない可能性が高まっている。その理由として、CPU の処理能力の向上、仮想化・分散処理技術の進展、インターネットアクセスの高速化、データセンターの大型化等を背景に世界中に分散するサーバーのリソースをサービスとして提供・利用することが可能になったことが挙げられよう。また、建物や機械等の有形資産と異なりデジタル関係の研究開発資産等の無形資産は国境を越えた巨額の移転・移動が容易であり、その国の資産水準・構成を突如変えてしまう可能性を持っている。その一例として議論になるのが、次の「アイルランド問題」である。すなわち 2010 年代半ばに GAFAM などの多国籍企業

(multi-national enterprises: MNEs) が研究開発資産（知的財産生産物）を法人税の低いアイルランドの拠点に移転した結果、アイルランドでは当該資産から発生する特許等使用料の受取（サービス輸出）が増加し、GDP が 2015 年に突如急増（名目 35%増、実質 24%増）したのである。

こうした経済実態や景気動向とは乖離した数字上の不連続な動きを回避するため、2025SNA では多国籍企業内の特別目的会社（special purpose entities: SPEs）の定義の明確化や取引の明示化が行われた。すなわち MNEs や SPEs の定義の明確化等を踏まえつつ、MNEs における研究開発資産など知的財産生産物を経済的に所有する主体、そして知的財産生産物に係る取引フローの記録の明確化が図られた。実際、一般的な固定資産や金融資産であれば経済的所有主体の特定は比較的容易である一方、グローバルに展開する MSEs の知的財産生産物についても、研究開発拠点における研究開発投資は非居住者である親会社や関連会社からファイナンスされるなど資金フローが複雑であることや、アイルランドに見られるように、法的な研究開発資産の所有主体を税率の低い国に移転するなど、経済的所有をどの国に帰属されるのかを判定が難しいという側面がある。このため、知的財産生産物の経済的所有に関して、よりきめ細かな形で決めるための決定木（decision tree）が導入された。

6.3 無償デジタル生産物に関するサテライト勘定

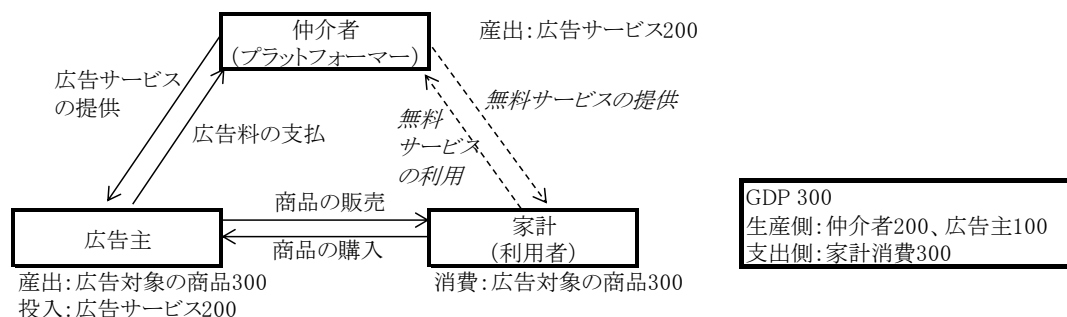
無償生産物に関する統計的な議論では、当該生産物は SNA において新たな課題にならない、なぜなら広告収入によって支えられているテレビ・ラジオは何十年も前から存在しているから、という主張がなされることがある。確かに無償生産物は、インターネット、より洗練されたコンテンツとしてソーシャルメディア、オンラインゲーム、動画配信に進化してきたといえる。実際、先述の Brynjolfsson and Oh (2012) でも消費者が享受する時間配分においては、インターネットとテレビは同列の扱いで評価されている。

しかし、こうした主張は、今日の無償生産物の多くがデジタル化によって可能になったものであり、無償デジタル生産物(検索サービスやソーシャルメディアサービスなど)が価値ある何か、例えば観測可能な現象を引き換えに提供する必要があることを見落としている。さらに無償のデジタル生産物の利用者は、生産に参加することもできる (Facebook ではユーザーが作成したコンテンツが掲載され、それが新規ユーザーの獲得やターゲティング広告に使用するデータの生成に貢献している)。すなわちプラットフォームのアルゴリズムが無償デジタル生産物と引き換えに家計を関与させ続けるための刺激を提供する一方で、家計はその行動の限界的な変化に関する情報内容(観測可能な現象)へのアクセスを提供し、それらはプラットフォームを通じて収集・利用される。これがプラットフォームのビジネスモデルの本質であり、これまでのメディアと一線を画するリアルタイムで個人の選好を活用したカスタマイズされたマーケティング、カスタマイズされた財・サービスの提供が可能になっている。

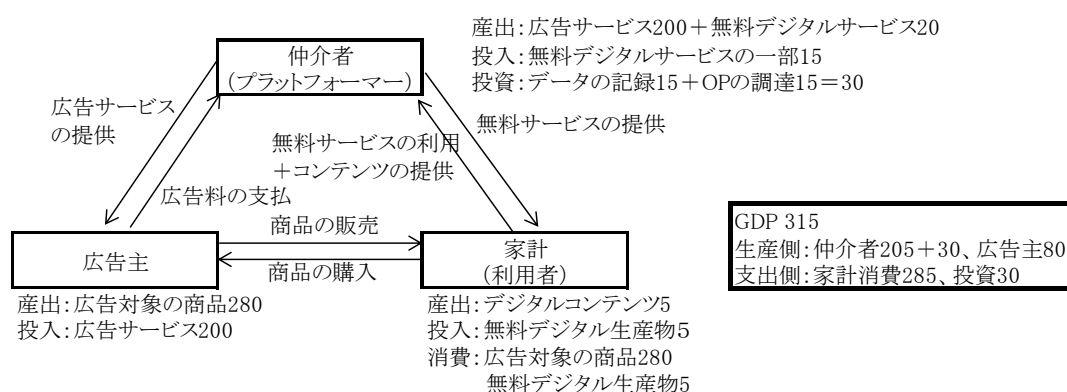
2025SNA では、以上のような無償デジタルサービスを巡る取引の構造をサテライト勘定として整備することが提案されており、家計、仲介者(デジタル仲介プラットフォーム)、広告主の3者間の取引を新たにルール化している。仲介者は、プラットフォーム・ソフトウェア資産とデータベース資産を開発し、それらは家計の観測可能な現象にアクセスするために使用され、そして広告主に予測広告サービス(predictive advertising services)を販売するためにデータ資産を生産する。現行では、家計の役割は、無償デジタル生産物と広告対象生産物の抱き合わせによる最終消費に限定され、無償デジタル生産物や家計による観測可能な現象の提供は市場取引で観測されないため別途記録されることはない。このため家計の役割をフローに含め可視化するに、帰属計算による無償デジタル生産物の価値(費用積上げ方法を推奨)、ユーザー生成コンテンツの価値(家計が消費する無償デジタル生産物と同じ)、データ資産の自己勘定(記録・処理費用及び観測可能な現象の入手費用)を計上することが提案されている(図表9)。

図表9 家計、仲介者（デジタル仲介プラットフォーム）、
 広告主の3者間取引のイメージ（オプション事例）

本体系



サテライト勘定 家計のデジタルコンテンツの産出と仲介者のデータ産出への投入・投資を認識するケース



出典：多田（2022）より作成。

7. むすび

これまで、経済構造及び消費者の厚生に大きな影響をもたらしているデジタルエコノミーの姿を的確に把握するためのSNAに関連する新たな概念、尺度、枠組みについて検討してきた。しかしながら、デジタルSUTの試算では、デジタル化にともなう生産構造の転換については推計の対象年が2015年及び2018年のわずか2年で間隔も3年と短いこともありまだ明確に見えてこない。例えば、企業のオンプレミスへの設備投資がクラウドサービスによるサービス投入にシフトしているのか、広告の提供・投入が既存メディアからプラットフォームにシフトしているのか、そしてこれらは国内生産から輸入に代替しているのか、等注目されている生産構造やビジネスモデルの転換が的確に反映されているのかについて、デジタルSUTの精度とともに他のSNA関連データ（固定資本マトリックスや国際収支統計

等)も参照しつつより精密に分析していく必要があるだろう。一方、2018年以降にクラウドサービスの急速な導入や新型コロナ禍対応によるDXの進展が見られ(リモートワークにともなうWeb会議用アプリの広がり等)、最近の生産構造は試算時からかなり変化している可能性がうかがえる。

また、無償デジタルサービスの測定に関して、推計方法や概念的枠組み、そして試算結果が示され、新たな知見がもたらされた。しかしながら、当該サービスの価格・数量のデータは市場から入手できないため、試算の多くが観測データではなく経済モデルの推計パラメータや実験データに依存した帰属計算になっており、統計として精度・頑健性をいかにして確保していくことが今後の課題である。またデータ資産化の測定において、各国の中でも最も先進的な取組を実施している米国、カナダの事例でも、他の資産から経験的に類推される仮定に基づいて試算が行われているが、実際のデータに基づいて推計した耐用年数や除却分布プロファイルを適用して作成していく必要がある。

現行のSNAの枠組みにおいても、Ahmad and Schreyer (2016)が指摘しているように、デジタル化がもたらす難題に出来ているように見え、概念的な問題が生じる場合もGDP全体から見ればそれほど重要ではないかも知れないが、実際の測定は引き続き難しい問題であることは明らかである。特に国境を越えたデジタルサービス等電子商取引、企業内のグローバルな知的財産生産物のフローは、GDPのマグニチュードに影響を与えるため、適切な捕捉が不可欠である。

また、デジタル化にともない財・サービスの出現・消滅の頻度が高まり、継続して利用できる財・サービスの情報だけでは数量及び価格指数はバイアスが発生しやすくなる(指数理論上のnew and disappearing productsの問題(Diewert and Feenstra 2017))。このため財・サービスの多くの分野において、従来以上に適切なアプローチが必要となろう。

こうした問題に対して、行政記録情報を含むデジタル化された情報や、インターネット上の価格や数量に関するデジタル情報(上述のようなオンライン選択実験による報告者・被験者の情報を含む)の利用が解決に向けた有益なアプローチになる可能性がある¹⁶。

¹⁶ 米国のデータ資産の価値推計では、産業別のデータ関連業務の賃金総額をするため、オンライン求人広告(約112万件のサンプル)のテキストから機械学習のクラスタリングによって得た情報を利用している。

参考文献

- 大久保敏弘・NIRA 総合研究開発機構 (2023)「第 9 回テレワークに関する就業者実態調査 (速報)」.
- 大久保敏弘 (2023)「大きく前進するデジタル経済をどう計測するか」, NIRA 総合研究開発機構『NIRA オピニオンペーパー』no.66.
- 河野洋介、吉本尚史 (2024)「データの資本としての記録方法について ～2025SNA(過少)に向けた国民経済計算における試算～」,内閣府経済社会総合研究所『経済分析』第 209 号.
- 総務省 (2024)「令和 6 年版情報通信白書」.
- 多田洋介 (2022)「2025SNA に向けた国際的な議論の動向」,内閣府経済社会総合研究所国民経済計算関連論文, No.1.
- 内閣府経済社会総合研究所 (2019)「2018 年度シェアリング・エコノミー等新分野の経済活動の測定に関する調査研究」報告書.
- 内閣府経済社会総合研究所 (2022)「デジタル SUT (供給・使用表) 2015、2018 年表の推計について (デジタルエコノミー・サテライト勘定に関する調査研究)」報告書.
- 内閣府経済社会総合研究所新分野ユニット (2020)「デジタルエコノミーに係るサテライト勘定の枠組みに関する調査研究」報告書 (概要版).
- 日本銀行 (2023)「国際収支統計からみたサービス取引のグローバル化」, 日銀レビュー.
- 長谷川秀司 (2023)「デジタルエコノミーをどのように把握するのか? ～新たな試みと課題～」,内閣府経済社会総合研究所『経済分析』第 207 号.
- ABS (2021), “Digital activity in the Australian economy, 2019-20”.
- Ahamad, N. and P. Schreyer (2016), “Measuring GDP in a Digital Economy,” OECD Statistics Working Papers, 2016/07, OECD Publishing, Paris.
- Ahamad, N., J. Ribarsky and M. Reinsdorf (2017), “Can potential mismeasurement of the digital economy the post-crisis slowdown in GDP and productivity growth?”, OECD Statistics Working Papers, 2017/09, OECD Publishing, Paris.
- BEA (2021), “New Digital Economy Estimates_2005-2019”, U.S. Department of Commerce. <https://www.bea.gov/system/files/2021-06/DE%20Tables%20for%20web%20June%202021.xls>
- Bean, Charles (2016), Independent review of UK economic statistics, HM Treasury, Cabinet Office, March 2016.
- Becker, G.S. (1965), “A Theory of the Allocation of Time,” *The Economic Journal*, 75, pp. 493-517
- Becker, G.M., M.H. DeGroot and J. Marschak (1964), “Measuring utility by a single-response sequential method”, *Systems Research and Behavioral Science*, 9(3), pp. 226-232.

- Brynjolfsson, E. and Joo Hee Oh (2012), “The Attention Economy: Measuring the Value of Free Digital Services on the Internet.” Orlando, FL:33rd International Conference on Information Systems (December)
- Brynjolfsson, E., A. Collis, W.E. Diewert, F. Eggers and K.J. Fox (2019), “GDP-B: Accounting for the Value of New and Free Goods in the Digital Economy,” NBER Working Paper 25695, Cambridge, MA.
- Brynjolfsson, E., A. Collis, W.E. Diewert, F. Eggers and K.J. Fox (2022), “GDP-B: Accounting for the Value of New and Free Goods”, unpublished
- Brynjolfsson, E., and McAfee, A. (2017), *Machine, Platform, Cloud: Harnessing Our Digital Future*, WW Norton & Company.
- Christensen, Clayton. (1997), *The Innovator’s Dilemma*, Cambridge, MA: Harvard Business School Press.
- Crick, S. (2022), “Valuing Data in Australia,” presented for the WPMAD WPNA workshop, unpublished.
- Diewert, W.E. and R. Feenstra (2017), “Estimating the Benefits and Costs of New and Disappearing Products”, Discussion Paper 17-10, Vancouver School of Economics, University of British Columbia.
- Diewert, W.E., J. Fox and P. Schreyer (2017), “The Allocation and Valuation of Time”, Discussion Paper 17-04, Vancouver School of Economics, University of British Columbia.
- Diewert, W.E., K.J. Fox and P. Schreyer (2019), “Experimental Economics and the New Commodities Problem”, Discussion Paper 19-03, Vancouver School of Economics, University of British Columbia.
- Dynan, K. and L. Sheiner (2018), “GDP as a Measure of Economic Well-being,” Hutchins Center Working Paper#43, Brookings Institution, Washington DC.
- Goolsbee, A. and P. Klenow (2006), “Valuing Consumer Products by the Time Spent Using Them: An Application to the Internet,” *The American Economic Review*, Vol.96, No.2(May,2006), pp.108-113.
- Mitchell, J., D. Ker and M. Leshner (2021), “Measuring the economic value of data”, *OECD Going Digital Toolkit Notes*, No. 20, OECD Publishing, Paris.
- Nakamura, L. and R. Soloveichik (2015), “Valuing “Free” Media Across Countries in GDP,” Working Paper No.15-25, Federal Reserve Bank of Philadelphia, July2, 2015.
- OECD (2019), “Guideline for Supply-Use tables for the Digital Economy” paper prepared for the Informal Advisory Group on Measuring GDP in a Digitalized Economy.
- Schreyer, P. (2022), “Accounting for Free Digital Services and Household Production - An Application to Facebook,” paper prepared for the 2019 Economic Measurement Group Workshop, UNSW, Sydney.

- Shapiro, C., and Varian, H. (1999), *Information Rule*, Harvard Business Review Press.
- Solow, Robert. (2000), *Growth Theory: An Exposition, Second Edition*, Oxford University,
- Santiago Calderon, J.B., and Rassier, D.G., (2022) “Valuing the U.S. Data Economy Using Machine Learning and Outline Job Postings” prepared for the 37th IARIW General Conference, 2022.
- Statistics Canada (2019), “The value of data in Canada: Experimental estimates”
- Statistics Canada (2021), “Measuring the digital economy: The Canadian digital supply and use tables 2017-2019”
- UNDESA (2023), Recording of data in the National Accounts.
(https://unstats.un.org/unsd/nationalaccount/RAdocs/DZ6_GN_Recording_of_Data_in_NA.pdf)
- Weitzman, Martin L., (1976), “On the Welfare Significance of National Product in a Dynamic Economy.” *Quarterly Journal of Economics* (February).
- Working Group on National Accounts (2022), Guidance Note on Recording and Valuing “Free” Digital Products in an SNA Satellite Account.
(https://unstats.un.org/unsd/nationalaccount/RAdocs/ENDORSED_DZ4_Free_Digital_Products_Satellite.pdf)

令和5年度 一般財団法人 社会文化研究センター 助成事業報告書

生成 AI が変える社会と教育

令和7年（2025年3月）

金沢学院大学

〒920-1392 石川県金沢市末町10

電話：076-229-1181