

日本の AI 反転攻勢：人工知能戦略専門調査会（第 1 回）と新国家戦略の黎明に関する分析レポート

Gemini

第 1 章 戦略的概観と創設時の負託：日本の AI 戦略を駆動するエンジンの設計

日本の人工知能（AI）政策は、2025 年 9 月 1 日の「人工知能関連技術の研究開発及び活用の推進に関する法律」（以下、AI 法）の全面施行をもって、新たな実行フェーズへと移行した。この新時代の幕開けを象徴するのが、同年 9 月 19 日に開催された「人工知能戦略専門調査会」の第 1 回会合である。本章では、この専門調査会が単なる諮問機関ではなく、日本の国家 AI 戦略を技術的・知的に実行するために設計された中核エンジンであることを、その構造、構成員、そして負託された任務の分析を通じて明らかにする。

1.1 AI 戦略本部と専門調査会の役割

人工知能戦略専門調査会は、内閣総理大臣を本部長とする「人工知能戦略本部」の下に設置された専門機関である¹。この階層構造は、AI 戦略がトップダウンの強力な政治的コミットメントに基づいていることを明確に示している。専門調査会の具体的な任務は、AI 法に規定された専門的事項の調査であり、これには国家戦略の青写真となる「人工知能基本計画」の策定や、具体的な「指針」の整備が含まれる³。また、前身である「AI 戦略会議」の検討事項を引き継ぐことで、政策の継続性を担保している点も重要である⁵。

運営規則に目を向けると、ウェブ会議システムの活用が認められており、議事概要と配布資料は原則として公開される方針が示されている³。これは政策決定プロセスの透明性を確保するための重要な措置である。ただし、座長の判断により議事の一部または全部を非公開とするこ

とも可能であり³、この規定の運用は今後の透明性を測る上での試金石となるだろう。

委員は学識経験者の中から内閣総理大臣によって任命され、必要に応じてワーキンググループを設置できる権限を持つ³。この構造は、高度な専門的知見に基づくトップレベルの指針策定と、特定の技術的課題に対する迅速かつ深い調査の両方を可能にするよう設計されている。公開された委員名簿⁶を分析すれば、法律、倫理、コンピュータサイエンスといった分野の専門性がどのように組み合わせられているか、政府の重点領域を推し量ることが可能である。

1.2 第1回会合：国家のアジェンダ設定

記念すべき第1回会合は、令和7年9月19日（金）の16時30分から18時00分にかけて開催された⁶。AI法が9月1日に全面施行されてからわずか18日後という迅速な開催は、政府の強い危機感と実行への意欲を物語っている。

配布資料から明らかになった議題は、極めて戦略的に構成された3つの柱に集約される⁶。

1. **国家の青写真**：「人工知能基本計画の骨子（たたき台）」
2. **規制のガードレール**：「AI法に基づく適正性確保に関する指針の整備について」
3. **喫緊の脅威への対応**：「AI法に基づく調査研究等について【報告】」

この3本柱のアジェンダは、単なる議題の羅列ではない。それは、まず国家としての壮大な戦略（基本計画）を定義し、次に行動規範（指針）を確立し、そして最後に社会が直面する最も目に見えるリスク（調査報告）に即座に対処するという、極めて構造化された政策実行アプローチの現れである。

この構造は、従来の省庁縦割りの緩慢な政策形成プロセスとの決別を意図している。AIという指数関数的に進化する技術に対し、従来の5カ年計画のような硬直的なアプローチが機能不全に陥ることを政府は認識している。その認識が、後述する基本計画の「毎年改定」方針³に結実している。そして、その毎年改定という「アジャイル・ガバナンス」を実際に機能させるための組織的装置こそが、この専門調査会の構造なのである。ワーキンググループの設置権限³や前身組織からの継続性⁵は、単なる運営上の細則ではなく、迅速で柔軟な政策サイクルを実現するための必須の設計思想と言える。このアジャイルな政策サイクルは、企業に対して、これまでの静的なコンプライアンス対応から、継続的な政策動向の監視と戦略的適応へと、新たな対応を迫ることになるだろう。

さらに、会合のタイミングと提示された資料の性質は、政府の姿勢が「議論」から「実行」へと明確にシフトしたことを示している。AI法の施行、総理大臣主導の戦略本部の始動、そして専門調査会の迅速な開催という一連の流れの中で、提示されたのが抽象的な論点ではなく、具

体的な「たたき台」や「報告」であったという事実は重要である⁶。これは、抽象的な議論の時代が終わり、国家戦略を実行に移す段階に入ったという政府の断固たる意思表示に他ならない。

表 1：第 1 回人工知能戦略専門調査会 主要配布資料と目的

議題分類	主要資料名	主要目的	導入された主要概念	典拠
国家 AI の青写真	人工知能基本計画の骨子（たたき台）	日本の AI 開発・利活用に関する国家ビジョン、目標、戦略的枠組みの定義。	「反転攻勢」、イノベーション促進とリスク対応の両立、「3つの原則と 4つの方針」	³
規制のガードレール	AI 法に基づく適正性確保に関する指針の整備について	AI システムの信頼性を確保し、開発者から利用者までが遵守すべき行動規範の確立。	広島 AI プロセス・OECD AI 原則との整合性、利用者を含む広範な対象範囲	³
能動的リスク評価	AI 法に基づく調査研究等について【報告】	顕在化している AI リスク（性的ディープフェイク、雇用 AI）の実態を把握し、今後の対策の基礎情報を提示。	具体的リスク領域の特定、実態調査に基づく現状評価	³

第2章 人工知能基本計画の解体：日本の「反転攻勢」の設計図

第1回会合の中心的な議題は、「人工知能基本計画の骨子（たたき台）」であった。この文書は単なる政策計画ではない。それは、日本の国家としての野心、危機感、そして未来への処方箋を物語る、戦略的ナラティブである。本章では、この基本計画の骨子を詳細に分析し、その根底にある思想と具体的な行動計画を解き明かす。

2.1 「反転攻勢」のナラティブ：国家再興のツールとしてのAI

本計画の最も際立った特徴は、AIを「人口減少」「国内への投資不足」「賃金停滞」といった、日本が長年抱える深刻かつ構造的な社会経済課題を解決するための切り札として明確に位置づけている点である³。これにより、AIは単なる一技術から、国家再興を担う戦略的ツールへとその意味合いを昇華させられている。

特に「反転攻勢」という言葉の選択は、極めて意図的かつ強力である³。この言葉は、日本がこれまで世界のAI開発競争において守勢に立たされ、あるいは周回遅れであったという現状認識を暗に認めつつ、今こそ状況を打開し、主導的地位を取り戻すための決定的行動を開始するという強い意志を示すものだ。このナラティブは、官民双方の危機感を煽り、行動を喚起することを狙いとしている。

その根底にあるのは、日本の個人や企業の生成AI利活用が他国に比べて低水準にあり、「『AIを使わない』ことが最大のリスクである」という厳しい現状認識である³。これは、リスク回避を優先する従来の姿勢から、行動しないことによる機会損失をこそ最大のリスクと捉える、政策思想上の重大な転換点（ピボット）を示している。このナラティブの転換は、一種の社会工学的な試みと言える。AIに対する漠然とした不安（雇用の喪失、倫理的リスク等）が社会に蔓延すれば、AIの導入は遅々として進まない。政府が「使わないことがリスク」と断言し、AIを人口減少という誰もが認める国難の解決策と結びつけることで³、AIを「脅威」から「機会」へと国民の認識を転換させ、変革への広範な合意形成を図ろうとしているのである。

2.2 二軸の哲学：イノベーション促進とリスク対応の両立

計画は、「イノベーション促進とリスク対応の両立」という基本原則の上に構築されている³。これは、米国の市場主導型アプローチと EU の規制先行型モデルとの間で、日本独自の「第三の道」を模索する試みである。

このバランスを支える思想的支柱が、「人間中心の A I 社会原則」である³。経済効率や技術的進歩が、人間の尊厳や基本的な価値を損なうことがあってはならないという理念が、計画の冒頭に明記されることが検討されており、あらゆる施策の基盤となることが担保される⁹。また、リスク認識も進化しており、技術的な不具合だけでなく、「差別や偏見の助長」といった社会的な危害にも明確に言及している³。

2.3 行動の枠組み：「3 つの原則と 4 つの方針」

この戦略は、9 月 12 日に開催された AI 戦略本部の初会合で総理大臣が指示した、「3 つの原則」と「4 つの方針」によって具体化される⁵。

4 つの方針とは、(1) A I を使う、(2) A I を創る、(3) A I の信頼性を高める、(4) A I と協働する、である³。この 4 つの方針は、相互に関連し、好循環を生み出すように設計されている。まず、社会全体で AI を広く使うことで、膨大なデータと多様なニーズが生まれ、それが国内での AI 開発、すなわち AI を創る活動を活性化させる。そして、AI が社会に浸透するにつれて、その信頼性を高めるための技術的・制度的取り組みが不可欠となり、最終的に人間が AI と効果的に協働できる社会が実現される、というシナリオである。

さらに計画は、国内に「AI エコシステム」を構築し、魅力的な待遇や生活環境の整備を含めた包括的な施策によって国内外から AI 開発者を確保することの重要性を指摘している¹⁰。また、医療・ヘルスケア、金融、防災といった重要分野で「国産 AI」の開発・導入を促進する方針も打ち出されている¹¹。

2.4 アジャイル・ガバナンスの実践：年次更新サイクル

本計画の最も革新的な側面の一つが、技術動向を踏まえて「当面は毎年変更を行う」というコミットメントである³。これは、前述したアジャイル・ガバナンスの理念を政策文書として正式に規定するものであり、日本の政策決定プロセスにおける画期的な試みと言える。この計画は、「本年冬まで」に閣議決定される予定であり、骨子策定から公式政策化までのタイムラインも極めて迅速である¹²。

しかし、このアジャイルなアプローチは、新たな形態の規制リスクとビジネスチャンスをもたらす。企業の AI 関連投資（データセンター建設、人材再教育、独自モデル開発など）は、通常、複数年の計画で実行される。国家戦略が毎年変更される可能性があるということは、企業にとって大きな政策不確実性をもたらす。例えば、ある年の重点分野が翌年には変更され、補助金の対象や研究開発の優先順位が変わる可能性がある。これにより、企業は静的なコンプライアンス体制から、常に政策動向を監視し、迅速に対応する動的な体制への移行を余儀なくされる。その結果として、日本の流動的な AI 政策をリアルタイムで分析し、戦略的助言を提供する政策アナリスト、コンサルタント、法律専門家といった新たな市場が生まれることは想像に難くない。

第 3 章 信頼のアーキテクチャ：AI 法適正性確保指針の分析

本章では、AI システムの「適正性」を確保するために提案された指針を分析する。これは、基本計画で示された高次の理念を、開発者や利用者のための具体的な行動規範へと落とし込むための主要なツールであり、日本の「ソフトロー」あるいは共同規制的アプローチの中核をなすものである。

3.1 範囲と目的：信頼できる環境の構築

指針が目指すのは、事業者や国民が「AI を信頼して利活用できる環境を構築」することである³。ここで極めて重要なのは、指針の対象が「AI を開発・提供する者」のみならず、「国民を含む利用者」にまで及んでいる点である³。この広範なスコープは、AI ガバナンスが一部の専門家や事業者だけの責任ではなく、社会全体で共有されるべき責務であるという思想を示唆している。

基本計画と同様に、この指針も固定的なものではなく、技術の進展に合わせて柔軟に見直されることが前提となっている³。

3.2 国際的整合性と基本原則

指針は、既存の国際的な枠組み、すなわち「広島 AI プロセス」および「OECD AI 原則」を明確に参照し、その上に構築されている³。これは、日本の AI ガバナンス・アプローチが国際的に相互運用可能であることを保証すると同時に、日本がグローバルな AI ガバナンス議論において責任ある主導的役割を果たすという意味を示す戦略的な動きである。

指針案には、「信頼できる AI の責任ある管理運営原則」や「信頼できる AI のための国家政策と国際協力」といった主要原則が含まれており³、これらは AI 倫理とガバナンスに関する世界標準の議論と軌を一にしている。

3.3 「ソフトロー」によるアプローチ

配布資料は、これらのルールが罰則を伴う厳格な法律（ハードロー）ではなく、あくまで「指針」であることを示している。これは、イノベーションを阻害することを避け、自主的な遵守と業界主導のベストプラクティスを重視する日本全体の AI 規制アプローチと一致している¹。このアプローチは、より規範的でリスクベースの規制を志向する EU の AI 法とは対照的であり、ハードローとソフトローの中間的な性質を持つと評されている²。

この国際整合性の強調は、単なる技術的な互換性の確保にとどまらない、高度な外交戦略でもある。自国の国内指針を広島 AI プロセスや OECD 原則に明確に根拠づけることで、日本はこれらの国際規範の形成における自国の中心的な役割を内外にアピールしている。これは、世界の AI ガバナンスが米国のイノベーション優先モデルと EU の規制優先モデルに二極化する中で、多くの国が模索している「第三の道」の具体的な実装例を提示する試みである。これにより、日本はグローバルなデジタル経済のルール形成におけるソフトパワーを強化することができる。

また、指針の対象を一般市民にまで広げたことは、将来的な国家規模での AI リテラシー向上政策を予見させる。一般市民を対象とする規制文書は異例であり、その存在自体が、今後政府が学校教育のカリキュラム改訂、大規模な広報キャンペーン、企業研修プログラムなどを通じて、国民全体の AI 理解度を引き上げるための施策を展開する必要性を示唆している。これは、教育産業や企業研修サービスにとって新たな事業機会の創出につながる可能性がある。

表 2 : AI ガバナンス原則の比較分析

ガバナンス原則	日本の適正性確保 指針（アプローチ	OECD AI 原則 / 広 島 AI プロセス（ア	整合性・差異の分 析
---------	----------------------	-------------------------------	---------------

	要約)	プローチ要約)	
透明性・説明可能性	文脈に応じた柔軟な説明責任を重視しつつ、必要性を指摘。	AI システムの動作に関する透明性を要求。	高い整合性。日本の指針は、EU の解釈よりも実用的で、過度に規範的でないアプローチを示唆。
公平性・非差別	差別や偏見の助長といった社会的リスクの緩和に焦点を当てる。	AI システムは法の支配、人権、民主的価値を尊重し、公平・公正であるべき。	理念において完全に整合。日本は特に社会実装時の具体的な危害の防止を重視。
アカウンタビリティ	開発者から利用者に至るまでの広範な関係者の責任を問う。	AI システムとその結果に対するアカウンタビリティのメカニズムが必要。	高い整合性。日本の指針は、利用者を含むエコシステム全体での責任共有を強調する点が特徴的。
安全性・セキュリティ	技術的リスクと社会的リスクの両方を含む、広範なリスクへの対応を求める。	AI システムはライフサイクル全体を通じて堅牢、安全、セキュアでなければならない。	高い整合性。
人間による監督	「人間中心の AI 社会原則」を最上位に置き、人間の介在を担保。	AI システムは常に人間による実効的な監督が可能であるべき。	理念において完全に整合。日本の戦略全体を貫く中核思想。

第 4 章 新たなリスクとの対峙：ディープフェイク、雇用 AI、そして将来の優先事項

本章では、専門調査会が最初に取り組んだ具体的なリスク評価報告書を検証する。政府が初期の調査対象として何を選んだかは、その喫緊の懸念事項を明らかにし、将来の調査や特定分野への規制導入の可能性を示唆するロードマップとなるため、極めて重要である。

4.1 緊急かつ具体的な脅威の優先順位付け

専門調査会の初期調査は、社会的影響が大きく、既に問題が顕在化している2つの特定の領域に焦点を当てた。すなわち、**(1) 性的なディープフェイクを生成するAI**と**(2) 雇用（採用・人事評価等）におけるAI活用**である³。

この選択は、国民が直接的な危害を感じやすく、既に社会的な懸念を生んでいる問題から着手するという、極めてプラグマティックなアプローチを示している。

4.2 性的ディープフェイクに関する調査結果：明白かつ現在の危険

調査により、技術の進歩によってディープフェイクの生成が容易になり、ダウンロード可能なアプリが少なくとも1万種類以上公開され、中には中高生でも利用可能なものがあるという深刻な実態が確認された³。

報告書のトーンは強い懸念を示しており、この分野が単なる指針の範囲を超え、将来的に何らかの法的措置の対象となる可能性が高いことを示唆している。

4.3 雇用AIに関する調査結果：監視と継続調査

採用や人事におけるAI活用については、報告書は「現段階で大きな問題は確認されていない」と結論づけている³。

しかし同時に、今後の活用拡大が予想されるため、引き続き実態把握に努めるとしており、一種の「監視下」に置く姿勢を明確にした³。これは、HRテック関連企業に対し、そのアルゴリズムや運用方法が政府の監視対象であることを明確に伝えるシグナルとなる。

この初期テーマの選定は、国民の信頼を醸成するための戦略的なコミュニケーションと解釈できる。ディープフェイクによる名誉毀損や、アルゴリズムによる不公正な解雇といったリスクは、一般市民にとって最も身近で直感的に理解しやすい AI の脅威である。政府がこれらの「生活に直結する」問題に真正面から取り組む姿勢を見せることで、AI の負の側面を管理する能力があることを示し、国民からの信頼、すなわち社会的な許容性（ソーシャル・ライセンス）を獲得しようとしている。この信頼なくして、基本計画が掲げる「AI を使う」という目標の達成はあり得ない。国民がリスク管理を信用しなければ、社会全体での AI 導入に対する抵抗が強まるからである。

一方で、雇用 AI に関する「現段階で大きな問題はない」という調査結果は、企業に危険な安心感を与えかねない。政府が現時点で規制に踏み込んでいないからといって、これが不透明でバイアスのあるアルゴリズムを自由に導入してよいという青信号を意味するわけではない。むしろ、これは規制導入前の猶予期間と捉えるべきである。HR 分野での AI 活用が拡大すれば、統計的に見ても、採用におけるシステミックなバイアスなどの問題が顕在化することは避けられない。政府が「引き続き実態把握に努める」と明言している以上³、賢明な企業は、この猶予期間を利用して、自社の HR 関連 AI システムに公平性、透明性、監査可能性を確保するベストプラクティスを自主的に導入し、将来の規制に先手を打つべきであろう。

第 5 章 専門家たちの対話：多角的な評価と産業界の反応

本章では、当初の問いにあった「評価と評判」について、政府外部からの分析やコメントを統合し、日本の新 AI 戦略が、それに最も影響を受けるであろうステークホルダーたちにどのように解釈されているかを明らかにする。ここで見えてくるのは、専門家コミュニティ内での深い関与と、一般の主要メディアにおける驚くべきほどの沈黙という、二極化した反応である。

5.1 専門家の視点：慎重な楽観論と原則への着目

最も詳細な分析は、専門的なテクノロジーブログや法律関連のポータルサイトから発信されている³。

これらの論者は、「イノベーション促進とリスク対応の両立」という基本原則を評価し、特に政府が「AI を使わないことが最大のリスク」と問題を提起した点を、重要な視点の転換として称賛している³。また、AI の社会実装が単なる技術的な課題ではなく、「法制度や社会全体の対話」を必要とする複雑なプロセスであるという認識が示されたことにも、強い支持が寄せ

られている³。

会合自体は非公開であるものの、配布資料や議事概要が原則公開されるという透明性確保の姿勢も、肯定的に受け止められている³。

5.2 メディアの空白：雄弁な沈黙

朝日、読売、毎日、日経、NHK といった日本の主要報道機関を対象に本件に関する報道を調査したところ、該当する記事は一件も確認できなかった⁶。

この主流メディアにおける報道の欠如は、それ自体が極めて重要なデータポイントである。これは、国家の AI 戦略が、その初期段階において、社会全体を巻き込む変革プロジェクトではなく、一部の専門家や官僚が関わる技術的な政策課題として扱われていることを示唆している。この状況は、政策立案者と一般市民との間に「言説のギャップ（discourse gap）」を生み出し、戦略の成功に不可欠な社会全体の合意形成を阻害する潜在的なリスクをはらんでいる。

5.3 産業界・法曹界の動向

商事法務のような法律情報ポータルは、専門調査会の資料を積極的に追跡・配信しており、法務・コンプライアンス部門が高い関心を持って動向を注視していることがわかる⁷。国際社会経済研究所のような研究機関は、日本の動きを米国や欧州の政策と比較する国際的な文脈の中に位置づけ、国内に「AI エコシステム」を構築するという戦略目標に注目している¹⁰。

この「言説のギャップ」は、日本の AI 戦略にとって、技術的な課題を除けば最大の脅威と言えるかもしれない。「反転攻勢」は、専門家だけが戦うのであれば成功しない。AI が公共サービスや教育といった国民生活の根幹に関わる領域に導入される際、その背景にある戦略や理念が広く国民に理解されていなければ、政府は強い政治的逆風に直面することになるだろう。AI の大量導入は社会の広範な信頼と理解を前提とするが、その信頼を醸成する主要な担い手である主流メディアが沈黙している現状⁶は、国民が置き去りにされ、AI 戦略の具体的な影響が生活に及び始めたときに、恐怖や誤情報に基づく反発が広がるリスクを高めている。

一方で、現在の専門家による論評の状況は、政府のメッセージが現時点でどの層に届いているかを明確に示している。それは、企業の法務・コンプライアンス部門と業界の戦略担当者であ

る。法律情報ポータル¹⁵やアナリストのブログ³での活発な議論は、ビジネス界が政府のシグナルを確かに受信していることを示している。しかし、この B2B セクターへの情報伝達の成功は、皮肉にも、より広範な国民とのコミュニケーションの失敗を浮き彫りにしている。政府は「産業界」と対話することには成功しているが、まだ「国民」と対話するには至っていないのである。

第 6 章 戦略的インプリケーションと将来への提言

本章では、これまでの分析を統合し、本レポートの読者である意思決定者層に向けた、高次の戦略的かつ実用的な洞察を提示する。日本の新 AI 戦略の全体的な整合性と可能性を評価し、主要な課題と機会を特定した上で、将来を見据えた提言を行う。

6.1 総括的評価：整合的で野心的、しかし困難を伴う戦略

日本は、整合性が取れ、野心的で、かつよく構造化された国家 AI 戦略を打ち出した。この戦略は最高レベルの政治的支援を受け、アジャイル・ガバナンスという実用的なモデルに基づいて設計されている。

戦略は、日本の課題（AI 導入の遅れ）を的確に特定し、それに対処するための健全な枠組み（利用、創造、信頼の好循環）を提案している。そして、「世界で最も AI を開発・活用しやすい国」を目指すという目標¹³は、イノベーションとリスクの「両立」が効果的に管理されれば、十分に達成可能な野心的なものである。

6.2 目前に迫る主要課題

- **言説と導入のギャップ**：国民や中小企業の AI 導入に対するためらいを克服するには、政策文書以上のもの、すなわち大規模なコミュニケーションと教育キャンペーンが必要である。
- **人材と計算資源**：計画は人材確保の必要性を認識しているが¹⁰、世界のトップレベルの AI 研究者をめぐる獲得競争や、最先端の計算資源へのアクセス確保は、依然として巨大な障壁である。

- 「アジャイル・ガバナンス」の実行：政府が、官僚的な惰性や特定の利益団体の影響を受けることなく、基本計画を毎年真に意味のある形で更新し続けられるかどうか、この新しいモデルの真価を問う試金石となる。

6.3 ステークホルダーへの戦略的提言

6.3.1 企業に向けて

- **政策インテリジェンス機能の確立**：AI 基本計画を静的な文書として扱ってはならない。専門調査会による年次更新を継続的に監視するためのリソース（内部または外部）に投資すべきである。
- **「信頼できる AI」原則の自主的導入**：厳格な規制が導入されるのを待つのではなく、特に人事や金融などのハイリスク領域において、公平性、透明性、説明責任の枠組みを今すぐ導入し始めるべきである。これは将来的に競争優位性となる。
- **研究開発と国家優先事項の整合**：基本計画は、「国産 AI」開発の重点分野（医療、防災など）を示唆している¹¹。自社の研究開発をこれらの分野に整合させることで、将来的な補助金や官民連携の機会を獲得できる可能性が高まる。

6.3.2 投資家に向けて

- **戦略の「実現要因」への投資**：AI リテラシー教育、企業向け研修サービス、AI コンプライアンスを支援する規制テクノロジー（RegTech）、そして専門的な政策アドバイザー企業といった分野に投資機会を見出すべきである。
- **「AI 政策への俊敏性」に基づく企業評価**：投資先を評価する際、新たなデューデリジェンス項目として「その企業が、日本の変化し続ける AI 規制環境にどれだけ迅速に適応できる体制を備えているか」を加えるべきである。

6.3.3 政策立案者に向けて

- **国民的 AI リテラシー・キャンペーンの開始**：メディア分析で明らかになった「言説のギ

ャップ」を埋めるため、広範なコミュニケーション戦略に直ちに着手すべきである。

- 中小企業への具体的な支援策の提供：「AI を使う」という目標は、大企業だけに浸透しても意味がない。中小企業が AI ツールを導入するのを支援するための、的を絞った、利用しやすいプログラムを開発する必要がある。

引用文献

1. 国内初の「AI 法案」が閣議決定～罰則はなく自主性重視、イノベーション促進の方向, 10 月 27, 2025 にアクセス、<https://business.ntt-west.co.jp/bizclip/articles/bcl00071-160.html>
2. 中小企業への影響は？ 9 月全面施行の「AI 新法」を読み解く, 10 月 27, 2025 にアクセス、https://chusho-dx.bcnretail.com/dx_learn/detail/20251010_178648.html
3. 内閣府「人工知能戦略専門調査会（第 1 回）」を読み解く：国の AI 戦略策定の全体像と今後の展望, 10 月 27, 2025 にアクセス、<https://oregin-ai.hatenablog.com/entry/2025/10/02/210307>
4. 内閣府「第 1 回人工知能戦略本部」関連資料を読み解く：日本の AI 国家戦略の全体像と推進体制, 10 月 27, 2025 にアクセス、<https://oregin-ai.hatenablog.com/entry/2025/09/14/195637>
5. AI 法 全面施行 一次なるフェーズへ - 内閣府, 10 月 27, 2025 にアクセス、https://www.cao.go.jp/press/new_wave/20251003.html
6. 人工知能戦略専門調査会（第 1 回） - 科学技術・イノベーション ..., 10 月 27, 2025 にアクセス、https://www8.cao.go.jp/cstp/ai/ai_expert_panel/1kai/1kai.html
7. SH5603 内閣府、人工知能戦略専門調査会の第 1 回会合を開催 中崎尚 (2025/10/21), 10 月 27, 2025 にアクセス、<https://portal.shojihomu.jp/archives/76822>
8. 中小企業サイバーセキュリティ関連ニュースクリップ, 10 月 27, 2025 にアクセス、<https://www.cybersecurity.metro.tokyo.lg.jp/security/KnowLedge/506/>
9. 第 217 回国会 内閣委員会 第 15 号（令和 7 年 4 月 18 日（金曜日）） - 衆議院, 10 月 27, 2025 にアクセス、https://www.shugiin.go.jp/internet/itdb_kaigiroku.nsf/html/kaigiroku/000221720250418015.htm
10. 「米国 AI 行動計画と日米欧 AI 政策比較」 - 国際社会経済研究所, 10 月 27, 2025 にアクセス、https://www.i-ise.com/jp/information/report/pdf/rep_it_202603a_2509.pdf
11. 世界で「規制を緩く」競争勃発..日本で利活用推進の AI 法施行、欧米も規制強化から一転、開発優先に【やさしく解説】AI 法とは | JBpress (ジェイビープレス), 10 月 27, 2025 にアクセス、<https://jbpress.ismedia.jp/articles/-/90796?page=2>
12. 城内内閣府特命担当大臣記者会見要旨 令和 7 年 8 月 1 日, 10 月 27, 2025 にアクセス、https://www.cao.go.jp/minister/2411_m_kiuchi/kaiken/20250801kaiken.html
13. <AI Update> 日本の「AI 法」案の概要と実務上のポイント（速報）, 10 月 27,

2025 にアクセス、https://www.noandt.com/wp-content/uploads/2025/03/technology_no59.pdf

14. 日本における AI 法 / AI 戦略動向「人工知能戦略本部」の始動と国家の基本方針「AI 基本計画」 - note, 10 月 27, 2025 にアクセス、
https://note.com/gifted_viola8806/n/n78d6ad0ece95
15. 内閣府、人工知能戦略専門調査会（第 1 回）資料（2025/09/19） | 商事法務ポータル NEWS, 10 月 27, 2025 にアクセス、
<https://wp.shojihomu.co.jp/archives/145706>