

Adobe

AIの転換点

組織で責任を持ってAIを導入する方法



目次

I. 概要：責任あるAIの導入は現在の必須事項	3
II. 枠組みの概要：スケーラブルで倫理的なAIの未来を構築	4
1. 評価：組織の準備状況を評価し、責任を持って構築されたAIテクノロジーを選択する..	5
1.1 組織の準備状況を評価する.....	5
1.2 責任を持って構築されたAIテクノロジーを選択する	6
2. 試験運用：影響の大きいユースケースを特定し、試験運用する.....	8
2.1 優先度の高いユースケースを特定し、準備する	8
2.2 ビジネスと責任あるAIの基準に基づく試験運用	8
3. 導入：組織全体に責任を持ってAIを統合する	9
3.1 トレーニングと教育を組織的に行う.....	10
3.2 責任を念頭に置いて展開する	10
4. 監視：継続的な管理と改善	11
4.1 ビジネスと責任あるAIのベンチマークに対するパフォーマンスを監視する	11
4.2 継続的なリスク管理	12
III. ベストプラクティスを組織に浸透させる	16
1. 従業員利用に関する指針：	16
1.1 データの機密性：.....	16
1.2 AI利用の透明性：.....	16
1.3 アカウント管理方針：.....	16
2. ベンダー評価：質問サンプル	17
3. AIガバナンスの方策	18
IV. 責任ある実装が責任あるイノベーションを実現.....	19

このページに戻るには、このメニューア
イコンをクリックしてください

I. 概要：責任あるAIの導入は現在の必須事項

AIが産業界全体の変革の推進力となるにつれ、経営陣は急速なイノベーション、競争上の脅威への対応、業務効率の向上を迫られるというかつてないプレッシャーに直面しています。しかし、企業へのAI導入競争は新たなリスクをもたらします。AIを、慎重な監視体制なしに急速に導入すると、規制上の失策、業務の混乱、長期的な評判の低下につながる可能性があります。スピードを求めることと責任を果たすことのバランスを取ることは、もはやトレードオフではなく、戦略上不可欠な要素です。

AIを管理し、そのリスクを管理することは、非常に困難な作業のように感じられます。新しいガイドライン、枠組み、ポリシーの登場や、国際法、連邦法、地方自治体法の法整備の進展などにより複雑性が増し、企業がどこから着手すべきか不明瞭になり、複雑性が当初から関係者の課題になることがよくあります。アドビでは、説明責任、実行責任、透明性に関する[AI倫理原則](#)に基づく責任あるイノベーションの経験から、これらの課題に対処するための深いインサイトを得ることができました。

責任あるAIのイノベーションへの道のりは困難に見えるかもしれませんが、アドビの経験から、適切なツール、戦略、考え方に基づけば成功は手の届くところにあることは確実です。

企業がAI戦略を策定する際に直面する重要な意思決定のひとつは、AIソリューションを構築するか、購入するか、カスタマイズするか、あるいはその3つを組み合わせるかということです。次に紹介するアプローチでは、AIソリューションの購入を検討している企業に注目し、AIソリューションを外部から調達しようとしている企業において、既存の価値観やビジネス慣行を基盤として構築することを目指しています。AIガバナンスに関する独自調査と専門家へのインタビューを基に作成されたこの枠組みは、実行可能な進むべき道筋を示し、組織が現在の状況を評価し、責任あるAIの原則を企業全体に浸透させるためのベストプラクティスを提供します。この枠組みには、従業員向けの生成AI利用ガイドラインの策定、詳細なアンケートによるベンダーの評価、進化する状況に対応するためのAIガバナンスプロセスの更新など、実用的なステップが含まれています。

AIの導入状況がどの段階にあっても、この枠組みは、組織のAI対応能力の評価や既存の戦略の改善など、人間の創意工夫と最先端のAIガバナンスを融合し、責任ある規模拡大を実現する実証済みのアプローチを提供します。このロードマップに従うことで、組織はAIソリューションを効果的に評価、試験運用、導入、監視することができ、信頼を醸成し、リスクを軽減し、持続的な事業価値を生み出す回復力のある基盤を構築することができます。

II. 枠組みの概要: スケーラブルで倫理的なAIの未来を構築

生成AIの実装に成功するには重複しているので削除、単なるチェックリスト的な行動以上のものがが必要です。持続可能なイノベーションと倫理的なAI利用のための基盤を構築するには、各段階が次の段階の基盤となる戦略的かつ段階的なアプローチが必要です。この枠組みは、相互連携する一連の構成要素からなり、組織の準備状況の評価、効果的な拡張、AIシステムの継続的なモニタリングなど、あらゆる段階で責任あるAIの利用を統合するように設計されています。

この枠組みでは、AIの導入をプロセス主導の取り組みとして捉えるのではなく、組織のニーズと調和しながら進化するシステムの構築に重点を置いています。この枠組みでは、人間による管理と高度なAIテクノロジーの間のバランスを重視し、組織がAIの潜在能力を活用しながら、倫理的、規制上、業務上の目標にも対応することを目指しています。

この枠組みの各段階（準備状況の評価、責任ある試験運用、導入の拡大、継続的なモニタリングなど）は、あらゆる段階において互いに補強し総合的な基盤として、長期的な成功を支えています。各段階に責任あるAIの利用を組み込むことで、企業は信頼性、透明性、説明責任を促進しながら、AI導入の複雑性を乗り越えることができます。

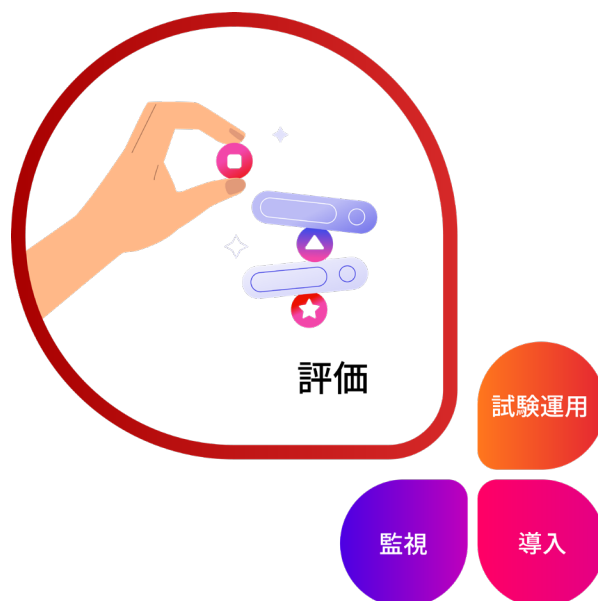


専門知識を基に調査で裏付け

アドビは、独立系調査会社に依頼して生成AIの導入に関する調査を実施し、様々な業界の200人以上のIT、組織、コンプライアンスのリーダーからインサイトを収集しました。この調査では、AI導入における現在の実務、課題、成功戦略が浮き彫りになりました。さらに、アドビは業界の専門家と詳細なインタビューを行い、[欧州連合のAI法](#)、[NISTのAIリスク管理フレームワークシンガポールのAI Verify](#)、[IEEE標準規格7000](#)、[ISO 42001](#)などのグローバルスタンダードに関する調査を行いました。これらの取り組みにより、AIの導入状況に関わらず、業界や組織の規模を問わず、この枠組みが適用可能であることが裏付けられました。

1. 評価:組織の準備状況进行评估し、責任を持って構築されたAIテクノロジーを選択する

AIを責任を持って導入する過程は、それを導く人々から始まります。評価段階では、AIが戦略的な優先事項にどのように適合するかを評価するために必要なツール、データ、インサイトを意思決定者に提供します。この段階では、組織のテクノロジーインフラ、ガバナンスの枠組み、AIリテラシーなどを部門横断的なチームのリーダーたちが検証し、全体的な準備態勢を判断します。



1.1 組織の準備状況进行评估する

多くの組織でAIの導入が始まっている一方で、調査対象の組織のうち、責任あるAIの優先事項を詳細に定めている組織は21%にとどまり、78%はまだ進行中または計画段階にあることが明らかになっており、準備態勢に対する明確なアプローチの必要性が浮き彫りになっています。責任あるAIの基盤を構築するには、IT、コンプライアンス、リスク管理、戦略などのリーダーが不可欠です。まず、AI導入に影響を与える可能性のある課題を特定するために、組織のガバナンス体制とAIリテラシーを包括的に見直すことから始めます。

AIの導入準備状況进行评估するには、**包括的なアプローチ**を採用する必要があり、**トップダウン型のリーダーシップ主導の取り組み**と、AIに日常的に関わる従業員からの**ボトムアップ型のフィードバック**を融合させることが求められます。

導入準備のためのアクション:

準備状況の包括的な監査を実施する — 組織のテクノロジーインフラ、ガバナンス基準、AI関連のポリシー、責任あるイノベーションの枠組み、コンプライアンスの実践などを評価し、強みと改善すべき領域を特定します。戦略目標と、AIを責任を持って導入する上での要求の整合性をとります。

主要課題を特定し、協調して対処する — IT、法務、コンプライアンス、事業部門などを含む部門横断的なチームを編成し、セキュリティ、プライバシー、法務、コンプライアンス、透明性基準における追加のAIポリシーの必要性を文書化し、実行可能な次のステップの優先順位を決定します。

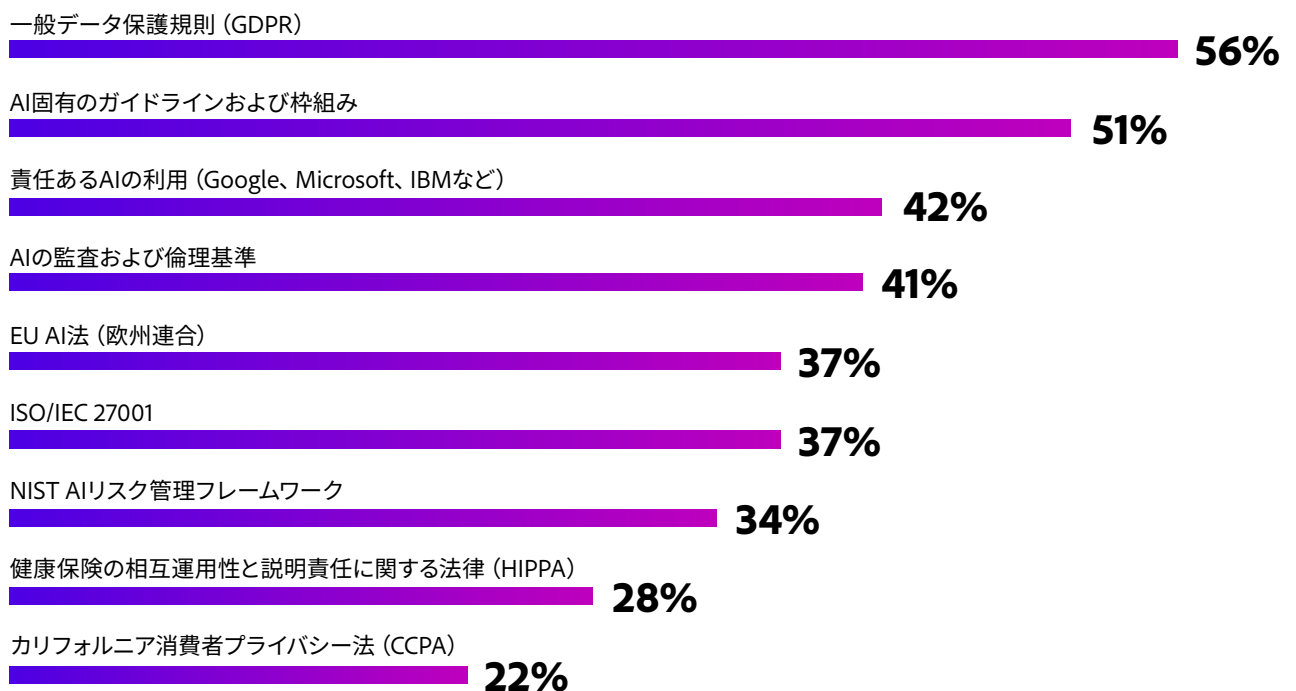
ガバナンスチームを編成し、権限を付与する — AIガバナンスである、社内のAIに関する責任ある基準と外部規制の枠組みが守られているか監督するチームを任命します。このチームには、リスクを積極的に管理し、進化する要件に適應するための権限とリソースが与えられます。

1.2 責任を持って構築されたAIテクノロジーを選択する

まず、自社の既存のガバナンス基準を徹底的に見直すことから始めます。これらの基準は、すでに**プライバシー、セキュリティ、アクセシビリティ、法的考慮事項**などの主要分野を網羅している可能性が高いでしょう。一般データ保護規則 (GDPR) やAI固有の枠組みなどのグローバルなベンチマークは、多くの組織におけるコンプライアンスの維持やリスク管理の一部になっています。さらに、**地域政策や業界固有の基準 (AI監査や責任基準など)** もガバナンス基準に組み込むべきです。

AIテクノロジーに必要なセキュリティおよびプライバシー基準または認証情報

回答者全体に占める割合が高い順に並び替え



責任あるAIへの期待とガバナンスの枠組みを策定したら、次に、責任を持って構築されたAIテクノロジーの選択基準を定めます。これらの選択基準は、既存の基準を統合して定め、生成AIに関連する独自の要素、例えば、出所の透明性、出力の正確性、トレーニングデータのライセンス、偏りの是正、文化的なローカライゼーションなどに焦点を当てます。

調査結果によると、企業が生成AIテクノロジーを評価する際に最も重視する基準には、次のようなものがあります。

- | | |
|----------------------|-----------------|
| 1. トレーニングデータの評価(72%) | 4. ソースの透明性(55%) |
| 2. AI利用に関する開示(63%) | 5. 偏りの緩和(50%) |
| 3. 有害性の緩和(60%) | |

これらの要素を確認することで、選択したAIテクノロジーがビジネスニーズと倫理的責任の両方を満たし、組織の長期的な成功を支えることができます。

組織は、AIソリューションを、戦略的なビジネス目標と責任あるAIの原則の両方と整合性をとることを目的とした選択基準を策定すべきです。これらの基準では、次のような点を重視します。

透明性 AIプロセスが説明可能であり、追跡可能であることを保証します。	文化的ローカライゼーション 様々な文化的かつ地域的文脈を尊重するようにAIシステムを適応させます。
正確性 データの正確性と予測の信頼性について高い基準を維持します。	偏りの緩和 公平かつ公正なAI結果を提供するために、積極的に偏りを低減します。

評価と選択プロセスの各段階を文書化することで、適応性と説明責任が補強され、AIの進歩や規制の変化に対応できる柔軟なガバナンスモデルを構築できます。次に、評価段階における手順の概要を示します。

評価

ステップ1: 組織の準備状況进行评估する

- AIを含むテクノロジーの責任ある利用に関する企業基準を定義し、周知します。
- CIOおよび／または全社横断的な委員会が、責任あるAIの導入から最も利点を得られる分野を特定するために、現在のシステムとビジネスプロセスをレビューします。
- 責任あるAIの導入を検討する追加のユースケースについて、社内のビジネスおよび機能部門のリーダーから意見を収集します。

ステップ2: 責任を持って開発されたAIテクノロジーを選択する

- AIの考慮事項として、プライバシー、セキュリティ、アクセシビリティ、法律に関する既存のガバナンス基準を見直します。
- 透明性、正確性、偏り、文化的なローカライゼーション、コンプライアンスに重点を置いて、責任あるAIへの期待に応えるために、既に確立されている基準を統合する選択基準を策定します。
- 策定された基準とビジネスニーズで評価して、最も適したAIテクノロジーを選択し、意思決定プロセスを文書化します。

2. 試験運用:影響の大きいユースケースを特定し、試験運用する

試験運用段階は、AIの実験と実際の運用を橋渡しする過程です。この段階では、主要な関係者が、ビジネスの目標と責任あるAIの目標にテクノロジーがどの程度適合しているかという観点から、テクノロジーのパフォーマンスを評価します。技術的な実現可能性のテストにとどまらず、主要なリーダーや関係者がテクノロジーと有意義な形で直接関わることを重視します。人々がAIと協働し、その拡張性について十分な情報を得た上で意思決定を行い、倫理的、運用上、規制上の基準を満たしていることを確認できるようにします。

試験運用の導入により、AIシステムを実際の状況下でストレステストすることができ、説明責任の評価や透明性に関する文書化が必要な箇所の確認や、期待される新しい機能のパフォーマンスの把握に役立ちます。インサイトを文書化し、実行可能な学習内容を収集することで、AIを責任を持って拡大するためのロードマップを策定し、短期的および長期的な目標をサポートする基盤を構築することができます。



2.1 優先度の高いユースケースを特定し、準備する

説得力のあるAIのビジネスケースを構築するには、AIの潜在的可能性についての包括的な見解を提供するために、主要な関係者である第一線で働く従業員に参与してもらうことが重要です。テクノロジーと直接的に関わる人々を早い段階から参与させることで、マーケティングコンテンツの制作、コーディング、ワークフローの自動化、データ管理など、AIが具体的な利点をもたらす影響度の高いユースケースを特定することができます。

具体的であること — 役割ではなくプロセスに焦点を当てる: 特定の役割 (例: 開発者向けAI) に焦点を当てたユースケースを構築するのではなく、AIによって合理化や改善が可能なプロセス、例えば「ルーチ的なコードレビューとエラー検出の自動化を目的としたAIによるコーディング支援」などに焦点を当てます。

利用状況とコスト削減に関する測定可能な指標を設定する: ROIは重要ですが、AIの試験運用では、生産性、市場投入までの時間、従業員の満足度、顧客体験の向上など、より幅広いメリットも重視すべきです。この指標は「体験利益率 (Return on Experience)」と呼ばれます。

目先の利益にとどまらない影響力を高める: AIの取り組みを長期的な変革の推進力として位置づけます。ユースケースは、直近の業務上のニーズへの対応だけでなく、デジタル変革や競争上の差別化といった戦略目標にも対応させるべきです。

2.2 ビジネスと責任あるAIの基準に基づく試験運用

ビジネスパフォーマンスと責任あるAIの基準という2つの観点から試験運用プログラムを評価することで、AIの取り組みが業務目標と責任あるAIのベンチマークの両方を満たしていることを確認できます。調査対象企業の半数以上 (54%) が、優先度の高いユースケースについて許容可能なリスクレベルを設定しています。企業は、これらの評価を体系的に文書化し、今後のAIプロジェクトに役立つ知見を記録しておくべきです。このような体系的なアプローチは、拡張可能なAI導入のための強固な基盤を構築します。

試験運用プログラム実施時のアクション:

ビジネスと責任あるAIのベンチマークを設定する:運用目標(生産性、コスト削減など)と責任あるAIの指標(透明性、公平性など)の両方を定義します。

リスクの閾値を設定する:リスクパラメータを設定し、AI関連のリスクを効果的に管理、軽減するための継続的な評価の枠組みを策定します。

学習内容を記録し共有する:透明性を確保し、今後の拡張の取り組みを導くために、試験運用結果を文書化する標準プロセスを構築します。

試験運用

ステップ1:優先すべきユースケースを特定する

- 既存のビジネスにおいて優先すべきユースケースから、AIの倫理と責任が重要な2〜3の試験運用プロジェクトを特定します。
- これらのユースケースについては、ビジネスと責任あるAIのパフォーマンスの両方を追跡するための指標と閾値を設定します。

ステップ2:ビジネスおよび責任あるAIの基準に照らして試験運用を実施する

- 必要に応じて、テクノロジー、ビジネス、責任に関する追加の検証およびテストを行いながら試験運用を実施します。
- ビジネスおよび責任あるAIへの期待について、事前に定義した指標および閾値に照らして試験運用の結果を評価し、学習した内容を今後の評価およびテストのアプローチのために文書化します。
- 試験運用の結果とインサイトに基づいて調達／導入を進めます。

3. 導入:組織全体に責任を持ってAIを統合する

導入段階では、試験運用から組織全体への統合に移行します。この段階では、実験的なアプリケーションから全面稼働するAIシステムへの移行に重点が置かれます。試験運用から得られた教訓を現実世界の業務に組み込みながら、責任を持ってAIを展開します。

この段階で従業員は、既存のワークフローにおけるAIの役割に積極的に関与します。従業員は、試験運用段階での実践的な経験を基に、AIの導入を推進する能力を身に付けています。



3.1 トレーニングと教育を組織的に行う

AIの導入を効果的に拡大するには、AIの能力と、その使用に伴う倫理的責任の両方を理解している知識豊富な人材が必要です。それぞれの役割や部門の従業員がAIツールを活用するには、カスタマイズされたトレーニングプログラムが役立ちます。多くの組織（89%）がトレーニングの重要性を認識しており、そのうちの3分の2近くが責任あるAIのガイドラインを導入しています。トレーニングでは、技術的能力、説明責任、透明性、規制遵守の原則を総合的に教育する必要があります。

トレーニングと利用開始時のアクション:

トレーニングとガバナンスの整合性をとる: 責任あるAIのガイドラインをトレーニング教材に組み込み、従業員のコンプライアンス、リスク管理、透明性の要件の認識を高めます。

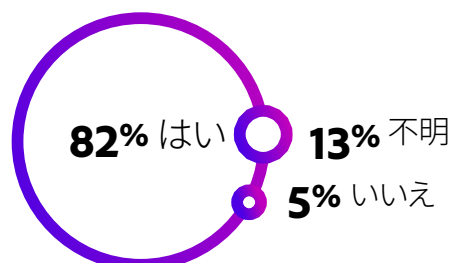
役割に応じてトレーニングをカスタマイズする: 特定の機能のニーズに対応するカスタマイズされたトレーニングモジュールを開発し、ビジネスと責任あるAIの両方に関するベストプラクティスを含めます。

3.2 責任を念頭に置いて展開する

AIを大規模に導入するには、責任ある利用に関するガバナンスの枠組みを構築する必要があります。企業は、AIの取り組みと既存のガバナンス方針の整合性をとると同時に、進化する規制、運用、責任あるAIの基準を満たすために、継続的に方針を改善していくべきです。

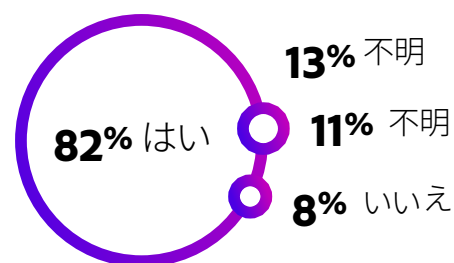
責任あるAIの考慮事項を、広範な展開のための基準に組み込む計画がある

回答者全体に占める割合が高い順に並び替え



テクノロジーのガバナンスの取り組みに責任あるAIの考慮事項を含める

回答者全体に占める割合が高い順に並び替え



責任ある展開のためのアクション:

説明責任の文化を醸成する: 従業員に、業務上のワークフローと関係者の信頼に対するAIのより広範な影響を理解させ、あらゆるレベルで責任感を浸透させます。

継続的にトレーニングを進化させる: AIのガバナンスの枠組みが進化するにつれ、新しいベストプラクティスと規制の変更を反映してトレーニングプログラムを更新します。

導入

ステップ1: トレーニングと教育を組織的に行う

- 責任あるAIの原則に則り、AIをいつ、どのように使用するかに関するユースケースに特化したガイドラインを策定します。
- ソリューションを展開する際に、関連グループに、ベストプラクティスを含めた包括的なトレーニングを行います。
- 組織全体で成功を分かち合い、共有します。

ステップ2: 責任を念頭に置いて展開する

- 各テクノロジーについて、導入の主要要件（ビジネスへの影響、統合の容易性、リスクの軽減など）を明確化します。
- 必要に応じて、ビジネスリーダーと緊密に連携し、トレードオフについて調整します。
- AIと責任に関する主要な考慮事項を、既存のガバナンスの枠組み（アクセス、制御、役割など）に組み込みます。

4. 監視: 継続的な管理と改善

AIシステムを本格的に展開するようになると、継続的な管理と改善が不可欠になります。監視段階では、AIシステムの効果、コンプライアンス、組織目標との整合性などを維持するために、リアルタイムの追跡、厳格なパフォーマンス評価、積極的なリスク管理アプローチに重点を置きます。責任あるAIの指標を組み込み、体系的な評価プロセスを確立することで、信頼を維持し、業務上の価値を育みながら、進化する規制上の要求や新たなリスクに対応することができます。

4.1 ビジネスと責任あるAIのベンチマークに対するパフォーマンスを監視する

テクノロジーの成果をモニタリングするクラス最高の手法では、自動化されたパフォーマンス追跡と人間の専門知識を組み合わせています。多くの組織がAIシステムのパフォーマンスを評価するためにリアルタイムのモニタリングツールを導入していますが、人間による管理を組み込むことで、これらのツールの有効性は向上します。人間のチームは、データの分析、リスクの特定、必要な調整に関して情報に基づいた意思決定を行うのに適しています。

69%の組織がリアルタイムのモニタリングツールを使用していますが、これらは人間の判断と組み合わせることで、はるかに効果を発揮します。多くの組織は技術的な指標を優先しており、72%が正確性、69%がROIに重点を置いています。しかし、責任ある拡張には倫理的な側面への配慮も必要です。人間による管理を組み込むことで透明性と予測可能性を確保し、社内外の信頼を気づくことができます。49%の企業がこの

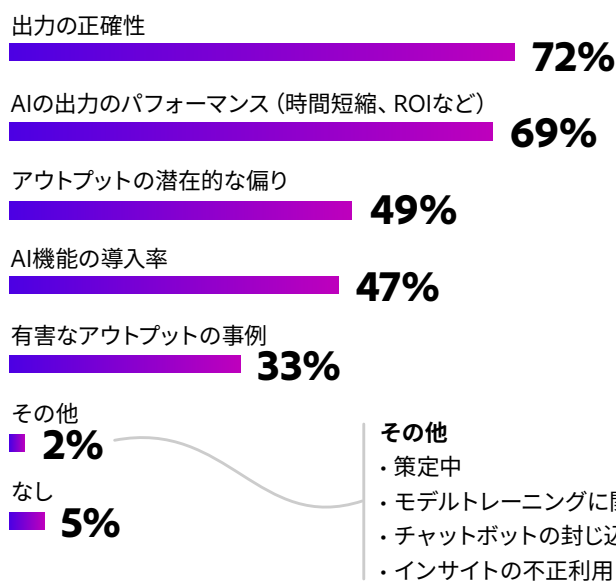


指標を追跡し、33%が有害なアウトプットを監視しており、積極的に偏りを検出しています。積極的な監視を一貫して行わなければ、AIシステムは整合性と信頼性を損なう可能性があります。技術的リスクと倫理的リスクの両方に対処するためにパフォーマンスモニタリングを改善することで、自社のブランドを保護し、ユーザーの信頼を獲得し、AIを責任を持って拡大するための回復力のある基盤を構築することができます。

テクノロジーと人との連携により、データの不正確さ、新たな偏り、コンプライアンス違反などの潜在的なリスクを早期に特定することができます。

テクノロジーのパフォーマンスと有効性を追跡するためのAI固有の考慮事項

回答者全体に占める割合



ほぼ3分の1の組織

テクノロジーのパフォーマンス指標とビジネス成果の継続的な改善をサポートするスタッフを配置していない

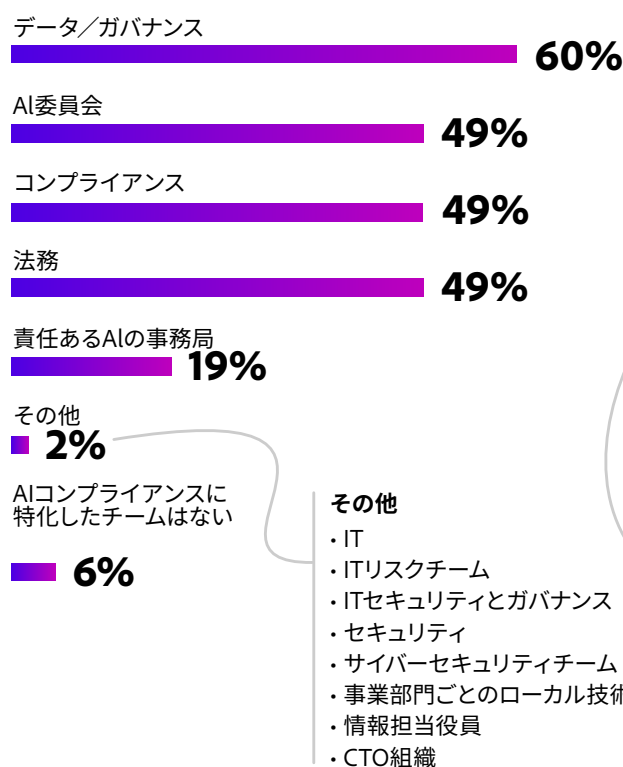
4.2 継続的なリスク管理

AIのリスク管理は、AIシステムとともに進化していくべき継続的なプロセスです。AIのリスク管理に対する体系化された部門横断的なアプローチを確立することで、ビジネスリスクと風評リスクの両方に積極的に取り組むことができます。予定されたレビューには、データサイエンティスト、ビジネスリーダー、法務／コンプライアンス担当役員など、組織全体から関係者が参加し、技術的なパフォーマンスと責任あるAIの目標の両方を徹底的に評価することが求められます。

AIのリスク管理は、テクノロジーの進歩とともに進化する継続的なプロセスです。60%の組織でデータおよびガバナンスチームが関与し、49%がAI委員会、コンプライアンス、法務チームが関与しており、部門横断的な連携の必要性が浮き彫りになっています。この積極的なアプローチにより、社内の価値観と外部の期待の両方に沿った対応が可能になります。「なぜ」という疑問は、規制の変化に対応できる回復力の構築に関係します。68%の企業がリスク管理における責任あるAIを重視していることから、包括的な文書化と継続的なリスク評価は不可欠です。

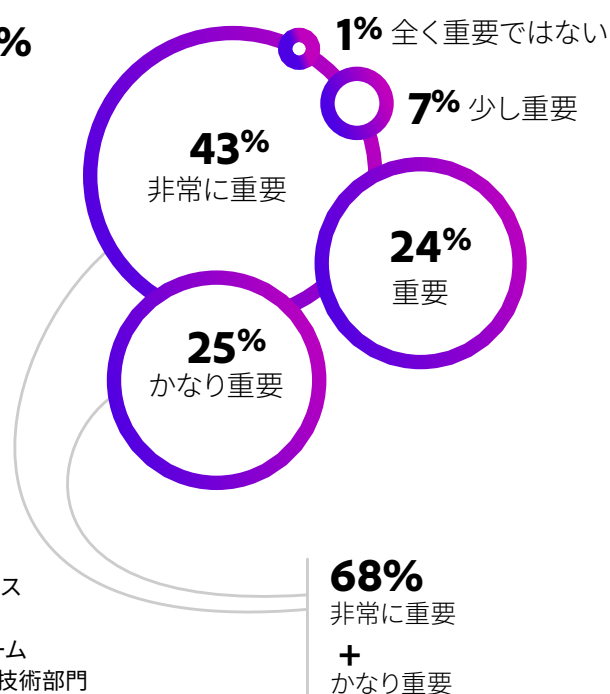
テクノロジーのパフォーマンスと有効性を追跡するためのAI固有の考慮事項

回答者全体に占める割合



コンプライアンスを維持するチームにとって重要な、責任ある倫理的なAI使用の考慮事項

AIの規制と基準の監視に重点的に取り組むチームが存在するとして回答者に占める割合



厳格なパフォーマンス指標を設定し、継続的なリスク管理の文化を育むことで、AIの導入をビジネス目標と責任あるAIの優先事項の両方に沿ったものとすることができます。部門横断的なレビューと徹底した文書化と組み合わせた積極的な監視により、自信と説明責任を持ってAI主導の変革をリードする体制を整えることができます。

リスク管理における主なアクション:

部門横断的なリスクレビューの確立: データサイエンティスト、コンプライアンス担当者、法律の専門家を交えて定期的にリスク評価を実施し、リアルタイムのデータに基づいて新たなリスクを特定します。

一貫した追跡と報告: 従業員やエンドユーザーから継続的にフィードバックを収集し、ユーザビリティの課題、偏り、予期せぬ動作などを検出します。

監視

ステップ1: パフォーマンスを監視する

- ビジネス目標や責任目標を含めた、より長期的なAIパフォーマンス指標を定義し、追跡します。
- 組織の責任原則を守りながら、ビジネスパフォーマンスを継続的に改善するために、継続的なレビューを実施し、調査結果について検討します。

ステップ2: 責任を念頭に置いて展開する

- 進化するAIの規制や基準 (例: シンガポールのAI-Verify、米国議会の提案) を追跡する役割を定めて権限を与えます。企業の基準を規制や基準の変化に応じて更新します。
 - AIの利用に伴うリスクを継続的に特定し、軽減するためのプロセスを構築します。
 - 企業基準の遵守状況に関する文書を更新します。
-

アドビのユーザー事例

生成AIの社内利用

アドビは、生成AIを人間の創造性を高める力を備えた変革テクノロジーと捉えており、人間の創造性を置き換えるものではないと考えています。アドビは、説明責任、実行責任、透明性に関する独自のAI倫理原則に沿った、生成AIテクノロジーの責任ある探究を奨励しています。

2023年6月、アドビは、従業員がアドビ社内で生成AIの研究と利用を安全かつ責任を持って機敏に行うことを支援するために、CIOとCHROがスポンサーを務める部門横断的な社内ワーキンググループを設置しました。このグループは、社内のリーダーや専門家と協力して、生成AI利用に関する仮説の全体像を把握し、適切なガイドラインを策定し、実験を合理化することで、従業員による草の根的な試行錯誤を慎重に導くことに重点的に取り組んでいます。この取り組みでは、アドビ全社における生成AIのユースケースを代表する4つのペルソナに基づくワーキンググループを結成し、倫理、セキュリティ、プライバシー、その他の法的考慮事項の進化を考慮した、受け入れプロセス、生成AIのリスク許容の枠組み、ユースケースレビューの設計図を策定しました。特定のユースケースに基づく承認済みの生成AIツールおよびモデルの一覧、および生成AIの従業員利用に関するガイドラインも利用可能です。2024年3月に発表されたベンダー生成AIガイドラインには、生成AIの利用方法と選択した製品の機能に関するトレーニングセッションが含まれています。

この取り組みの実施により、より迅速な実験と適用拡大のプロセスが可能な範囲で合理化され、全社的に生成AIの状況を評価することが可能になりました。アドビは、ビジネス全体で学習とインサイトの共有を促進し、共同探求のエコシステムの構築を継続しています。この構築プログラムは、アドビ製品への生成AIの拡張、生成AIテクノロジーやモデルの普及、進化する法的および規制上の指針とともに、進化し続けています。実験のレビューは監視されています。チームは、スケールアップのために本番環境に移行する実験も含め、承認後の実験の追跡を行っています。

III. ベストプラクティスを組織に浸透させる

AIシステムを責任を持って導入、監視、最適化するためには、いくつかの業務領域に注力すべきです。具体的には、包括的な従業員向けガイダンスの提供、ベンダーの厳格な評価、堅牢なAIガバナンスの確立などです。そうすることで、AIの取り組みを、進化する規制基準を満たすものにするだけでなく、信頼性、透明性、説明責任を、既存のガバナンスやリスク管理業務を基盤にして促進する業務慣行に沿ったものにすることができます。

このセクションでは、これらのベストプラクティスを日常業務に浸透させるための実践的なステップを概説します。

1. 従業員利用に関する指針：

AIの利用ガイドラインを組織の特定のニーズやリスクに合わせて調整することは、責任ある導入を確実に行うために不可欠です。これらのガイドラインは、従業員が規制基準やガバナンスプロトコルについて理解するのに役立ち、AIテクノロジーと、データセキュリティ、透明性、説明責任の取り組みとの整合性についての説明を提供します。

1.1 データの機密性：

データ処理をローカルで実行すべきか、あるいは厳格なアクセス制御下で実行すべきかを明確に指定し、不正アクセスを防止します。これには、機密性の高い出力を生成または操作する可能性のあるプロンプトの使用を抑えることも含まれます。このガイドラインは、企業の秘密情報を保護し、データプライバシー規制へのコンプライアンスを維持します。

1.2 AI利用の透明性：

社内文書、顧客向けインターフェイス、外部コミュニケーションの作成などにAIが使用されている場合は、AIの関与を速やかに開示します。この慣行は、説明責任を育み、AI生成コンテンツの真正性と信頼性を維持し、企業のブランドと評判を維持します。

1.3 アカウント管理方針：

アカウント登録を必要とする生成AIツールの使用について、明確な方針を策定します。これには、組織メールアカウントの使用の可否、業務目的での使用を許可するツール、業務目的への個人アカウントの使用の非推奨などが含まれます。これにより、不正使用を防止し、組織のより広範な情報セキュリティプログラムとの整合性を維持することができます。

これらのガイドラインを組織の状況に合わせて調整することで、従業員は生成AIツールを自信と責任を持って使用できるようになり、イノベーションと誠実さが両立する環境の実現に寄与することができます。

2. ベンダー評価:質問サンプル

AIベンダーを評価するには、有益な質問が必要であり、その回答が何を意味するのかを理解する必要があります。そうすることで、責任あるAIの基準、法的基準、規制基準に確実に準拠することができます。以下の質問は、基本評価を行うことを目的としており、組織が潜在的なパートナーシップについて情報に基づいた意思決定を行い、AI導入に伴うリスクを軽減することを目指しています。

テーマ	ベンダーへの質問	根拠	アドビの例
データの由来と利用	「AIシステムの開発とトレーニングにどのような種類のデータを使用しましたか？」	AIモデルのトレーニングに使用したデータのソース、性質、範囲に関するインサイトを提供します。この質問により、ベンダーのデータ処理が、購入者の責任あるAIの基準に合致し、法的目標に準拠していることを確認できます。	Firefly: アドビは、Fireflyの基礎モデルのトレーニングに使用するデータセットに、企業ユーザーのコンテンツ (Fireflyの入力および出力を含む) を含めていません。
知的財産のコンプライアンス	「著作権、知的財産、またはライセンス制限がある可能性があるデータセットを使用しましたか？」	この質問により、すべてのデータソースが合法であり、承認されたメカニズムを通じて取得されていることを確認し、潜在的な法的紛争を回避できます。	Firefly: 顧客は、アドビの認証情報でログインした状態で、 stock.adobe.com/Dashboard/LicenseHistory にアクセスして、いつでもライセンス履歴情報を確認できます。
トレーニングデータとロジックの透明性	「AIシステムの開発に使用したトレーニングデータとロジックについて、詳細な説明を入手できますか？」	こうした側面の透明性は、潜在的な偏りを特定し、モデルの推論プロセスを理解するのに役立ちます。これは、信頼性と公平性を評価する上で非常に重要です。	Adobe Experience Platform AIアシスタント: アドビは、Azure OpenAIサービスのトレーニングや調整に顧客データを使用していません。
出力の明確性	「AIシステムの出力について、平易な言葉で説明できますか？」	技術分野以外のレビュー担当者が出力について理解できることで、効果的な意思決定と責任ある利用が可能になります。	Firefly: アドビは、特定のFirefly生成アセットに対して自動的に コンテンツ認証情報 を生成し、そのアセットが生成AIを使用して生成されたものであることを示す透明性の確保に役立てています。
人間による監視	「AIシステムに人間によるレビューが含まれる場合、人間の関与の範囲と性質はどのようなものですか？」	自動化プロセスと人間の判断のバランスを把握することで、システムの運用動態を評価し、品質と責任基準を維持するために人間の介入が必要な領域を特定することができます。	Acrobat AIアシスタント: アドビは、この情報にアクセスできる者を、アドビの生成AIサービスの開発に直接関わる少数の訓練を受けたアドビ社員に厳しく制限しています。
公平性と偏りの評価	「AIシステムの偏りはどのように評価しますか？ その結果はどのようなものですか？」	この質問は、ベンダーが公平なAIの利用にどれだけ真剣に取り組んでいるか、また、デモグラフィックグループによっては偏った影響を及ぼす可能性を検知し、緩和する方法をどれだけ保持しているかを明らかにします。	Acrobat AIアシスタント: Adobeチームは、生成AI製品における偏った有害な結果を低減するためのテストを実施しています。 ビジネス向け生成AIソリューション概要 をご覧ください。
リスクの軽減	「潜在的に有害なアウトプットの評価を行っていますか？ これらのリスクを軽減するためにどのような対策を講じてしますか？」	この質問は、潜在的な悪影響を特定し、対処するためのベンダーの積極的な取り組みを評価します。つまり、AIシステムが安全かつ責任を持って動作するようテストされているか確認します。	Adobe Experience Platform AIアシスタント: アドビは、社内で開発したコンテンツフィルターを使用して、 (a) Adobe Experience PlatformのAIアシスタントに入力されたプロンプトがアドビの生成AIユーザーガイドラインに準拠しているかどうかを判断し、 (b) 生成された応答がこれらのガイドラインに違反する場合 (ハイトスピーチや冒涔的な表現など)、フィルタリングします。

組織は、これらの的を絞った質問をすることで、AIベンダーを評価し、AIの責任基準と業務上の優先事項に沿って、情報に基づいた選択を行うことができます。このアプローチにより、リスクを管理し、AIベンダーとの関係が、透明性、コンプライアンス、責任あるイノベーションの基盤の上に構築することができます。

3. AIガバナンスの方策

強固なAIガバナンスを実装することで、AIシステムが組織の価値観や規制基準に沿ったかたちで開発、導入、監視されていることを確認できます。EUのAI法やシンガポールのAI Verifyなどなど規制基準には多くの枠組みが存在します。米国では、企業は州の包括的なプライバシー法と、将来の規制の基盤となる可能性が高いNISTのAIリスク管理フレームワークに従うべきです。

以下のガバナンスの方策は、組織がAIリスクを管理し、透明性、説明責任、セキュリティを強化するのに役立ちます。

AIインベントリ	AIシステムのインベントリを作成し、リスクプロファイルと戦略的優先度に従って分類し、一元的なリポジトリとして利用します。	AIのユースケースを文書化し、関連リスクを分類しましたか？
フィードバックの仕組み	エンドユーザー、顧客、一般市民など、AI開発チーム外部の人々からのインサイトを収集するための強固なフィードバックチャネルを構築します。	エンドユーザー、顧客、一般市民からのインサイトを収集するために、どのようなフィードバックチャネルを使用していますか？
システム制限に関する文書化	AIモデルの知識不足に関する情報や、出力を確実に利用できる状況など、AIシステムの制限事項を文書化します。	AIのユースケースにおける既知または想定される制限事項を文書化していますか？
コンテンツの由来	AI関連データの由来、履歴、変更を追跡し、検証します。これには、トレーニングデータソースの追跡、使用したアルゴリズム、変換などが含まれます。	AI関連データの由来、履歴、変更をどのように追跡していますか？これには、生成から最終使用までのデータソースや変換などを含みます。
AIのテストとレッドチームの編成	トレーニングデータの意図せぬ露出、リバースエンジニアリングに対する脆弱性、モデル抽出に伴うリスクなどのリスクを評価します。	どのようなテストプロトコルを実施していますか？それらはAIのユースケース固有のリスクにどのように対処していますか？
セキュアなソフトウェア開発	AIシステムは、組織のセキュアなソフトウェア開発ライフサイクルに統合され、コーディングと展開に関する確立されたベストプラクティスに従うべきです。	AIのユースケースは、既存のセキュアなソフトウェア開発プロトコルとどのように統合されていますか？
トレーニング	トレーニングでは、関連方針、手順、コンプライアンス要件などをカバーし、関係者がAIリスクを効果的に管理し、組織の基準に従って行動するための知識を習得できるようにすべきです。	AIのガバナンスとリスク管理に関するトレーニングに参加したことがありますか？

これらのベストプラクティスをそれぞれ既存のプロセスに組み込むことで、AIテクノロジーの責任ある導入、管理、最適化を実現できます。充実した従業員ガイダンスの提供、ベンダーの厳格な評価、包括的なAIガバナンスの仕組みの確立といった活動を統合することで、ガバナンス基準や規制要件に沿った、弾力性のある利用方法の基礎を構築し、今日の課題に対応だけでなく、将来の成功に向けてAIを責任を持って拡大していくための体制を整えることができます。

IV. 責任ある実装が責任あるイノベーションを実現

企業において、AIの潜在能力を最大限に引き出すには、責任を持って構築されたテクノロジーを導入し、明確な利用ガイドラインを策定し、専用のトレーニングを開発し、強力なガバナンスを展開する必要があります。このアプローチは、ビジネス価値を推進し、AIの取り組みが規制当局の期待に応え、責任ある導入基準を満たすことを保証し、組織全体に責任あるAIの文化を浸透させます。

一貫した管理と適応により、AIの取り組みを軌道に乗せることができます。パフォーマンス指標を定義し、定期的に評価を行い、リスクを積極的に管理することで、規制の変化に先手を打ち、AIプロジェクトの整合性を維持することができます。

AIとともに未来へ向かう上で、この枠組みは、責任あるAIの分野において企業が主導的な役割を果たすことを可能にします。影響、統合、整合性に重点を置くこのアプローチは、持続可能なイノベーションと永続的な成功への道筋を確立します。

