

AI評価が AIの信頼性を いかに高めるか

効果的なAI評価は、強固なガバナンス、
コンプライアンス、パフォーマンスを後押し
することができます。

■ ■ ■
The better the prompt. The better the answer.
The better the world works.

EY

Shape the future
with confidence

ACCA

メッセージ



Marie-Laure Delarue
EY Global Vice-Chair — Assurance

人工知能（AI）は転換点を迎えています。ビジネスリーダー、政策立案者、学術関係者、一般の人々は、AIがもたらす変革の可能性を引き出し始めています。同時に、AIの複雑さと大きなリスクをどのように管理していくかという課題にも向き合っています。

EYのチームは、AIの導入を成功させる取り組みの最前線に立っています。AIシステムに対して厳格なAI評価を実施することで、AIが安全かつ効果的に開発・導入されるよう支援します。そうすることで、企業、政府、そして社会全体でAIに対する信頼を築くことができます。

本書では、AI評価が、自主的なものであれ規制に基づくものであれ、慎重かつ独立して実施される場合にはどのような重要な役割を果たし得るのかを考察します。とりわけ、あらゆる業種や地域で、企業・政策立案者・一般の人々が、AIの可能性を最大限に引き出し、そのリスクを最小限に抑えるために不可欠な、信頼と信用の基盤を築く上での役割を果たし得ることについて説明します。

効果的なAI評価は、コーポレートガバナンスを支える上で重要な役割を果たします。AIシステムが意図どおりに機能しているか、適用される法令・規制・基準に準拠しているか、社内のポリシーや倫理基準に沿って管理されているかを確認するのに役立ちます。

私たちは、本書が、ビジネスリーダーや政策立案者にとって建設的で価値ある一助となると考えています。それは、AIガバナンスの重要性と、AIシステムに対するガバナンスを各組織のニーズに合わせて構築し、強固で実効的なものにする上で、AI評価が果たす役割を明らかにするためです。

本報告書の作成に当たり、英国勅許公認会計士協会（ACCA）および国際会計士連盟（IFAC）の専門家の皆さまにご協力いただき、心より感謝申し上げます。今後も、両団体やその他の関係者の皆さまと引き続き連携を深めていくことを楽しみにしており、ビジネスリーダーや政策立案者がAIを活用して大きな進歩と繁栄の未来を築く上での支援を続けてまいります。



Helen Brand

Chief Executive Officer, Association of Chartered Certified Accountants (ACCA)

AIが経済全体に広がる中で、AIが示す内容を信頼できることは重要であるにとどまらず、公益にとって不可欠です。AIから持続可能な長期的価値を生み出していくための取り組みを進める上で、AI評価は重要な役割を果たします。

本書では、AI評価が果たし得る役割を検討します。AI評価が現在どのように理解されているか、堅牢なAI評価を開発する上での課題、そして将来その価値を最大化するために必要な鍵となる要素を取り上げます。

また、本書は、ビジネスリーダーや政策立案者にとって重要な考慮事項を示すとともに、コーポレートガバナンスおよびリスク管理を強化する上でAI評価が果たす重要な役割を明らかにします。さらに、AIへの信頼を築くための自主的な評価の価値や、評価フレームワークにおける目的と構成要素を明確に定義することの重要性についても検討します。さらに、評価を実施する際に、確立された基準を用いることや、評価指標の重要性を強調しています。

本件でEYおよびIFACと協働できることをうれしく思います。本書が、見解やアプローチをさらに発展させようとする方々の議論のきっかけとなることを期待しています。ACCAは今年、Global Policy Prioritiesを刷新しました。これは、スキルギャップの解消や持続可能なビジネスの推進などの領域を対象としています。そして、この分野でのスキル向上の必要性や、AIエコシステムにおいて信頼を築く上でAI評価が果たす役割を踏まえると、AI評価はこれらの優先課題と密接に関連しています。

私たちはこれを長期的な課題と捉えており、この魅力的かつ重要な分野において、政策立案者やその他の関係者と協力していくことを楽しみにしています。



Lee White

Chief Executive Officer, International Federation of Accountants (IFAC)

会計のプロフェッショナルとして、信頼を築くことが私たちの基盤です。今、AIが企業の事業運営において中核的な役割を担うようになる中で、その信頼を構築する上での私たちの役割は、かつてないほど重要になっています。

AIは、スピード、スケール、そして新たな可能性をもたらします。しかし、同時に複雑さも伴います。AIを支えるシステムはしばしば不透明で、その意思決定の過程はたどりにくいものです。

だからこそ、AIシステムの効果的な評価が重要であり、本報告書が今まさに時宜を得ている理由でもあります。本報告書は、この取り組みが単なるチェックリストにとどまってはならないことを示しています。AI評価は、堅牢で、明確かつ有意義なものでなければなりません。適切なスキルと倫理的基盤を備えたプロフェッショナルが主導する必要があります。

テクノロジーがいかに進歩しようとも、AIは自らを省みたり、疑問を抱いたり、「これは正しいのか？」と問いかけることはできません。一方、会計のプロフェッショナルである私たちは、常に一步引いて物事を捉え、批判的に考え、公益に資する役割を担っています。

会計のプロフェッショナルは既に、システムを評価し、データを解釈し、一貫した枠組みを適用し、健全な判断を下すための能力を備えています。AIが仕事の進め方を変えていく中で、私たちもまた進化しなければなりません。テクノロジーを受け入れつつ、私たちの職業を不可欠なものにしている人間的資質、すなわち懐疑心と批判的思考をさらに深めていく必要があります。

テクノロジーが信頼され、人間が進歩の中心であり続ける未来を築いていきましょう。

目次

1. エグゼクティブサマリー	05
2. はじめに	06
3. AI評価に関する公共政策の現状	07
4. AI評価をより効果的にする方法	12
5. ビジネスリーダーの考慮事項	14
6. 政策立案者の考慮事項	14
7. おわりに	15
8. 著者	16
9. 付録	17



1 エグゼクティブサマリー

ますます多くの企業が、戦略目標を達成するために人工知能（AI）を導入しています。こうした導入は企業全体の変革を加速させ、新たなビジネス機会を切り開いています。企業によるAIの導入が進むにつれて、導入するAIシステムの安全性、信頼性、有効性を確保する必要性も高まっています。そのため、AIがイノベーション、生産性、成長を高めるといふ潜在力を十分に発揮するためには、AIシステムに対する信頼が不可欠です。

こうした信頼を築くために、多くのビジネスリーダー、政策立案者、その他のステークホルダーは、AI評価を活用している、または活用を検討しています。AI評価は、「AI監査」あるいは「AI保証」と呼ばれることもあります。これらの評価は、企業が、適切なガバナンスの下で運用され、適用される法令・規制を遵守し、ビジネスリーダーが求める品質基準を満たしたAIシステムを構築・活用する上で役立ちます。

本書では、効果的なAI評価の構成要素を特定し、考察します。そのために、企業や政策立案者がAIへの信頼構築に取り組んでいる主要な国・地域において、自主的なものと規制に基づくもの、両方に関連するAI評価フレームワークを調査しました。

この調査から、企業が単独または組み合わせて活用している3つの新たなタイプのAI評価が明らかになりました。

◆ ガバナンス評価

AIシステムに関する内部ガバナンス体制を評価するためのものです。

◆ 適合性評価

適用される法令・規制・基準に準拠しているかどうかを判断するためのものです。

◆ パフォーマンス評価

あらかじめ定められた品質・性能指標に照らしてAIシステムを測定するためのものです。

また、これらのAI評価の有効性に対する潜在的な課題として、用語の曖昧さ、評価対象・方法論・評価基準の不十分な定義、そしてこれらの評価を実施するための有資格のプロフェッショナルの必要性も特定しています。

これらの課題に対応し、効果的で有用なAI評価を実現するために、本書ではビジネスリーダーおよび政策立案者に向けたいくつかの考慮事項を提示します。

具体的には、ビジネスリーダーには、以下の事項を検討することを提案します。

- コーポレートガバナンスやリスク管理の強化に対してAI評価が果たす役割
- 規制上の要件がない場合であっても、自主的な評価が従業員や顧客の間でAIシステムに対する信頼を築くことができるかどうか。また、自主的な評価がどのような場面で用いられているか
- どのような種類の評価が最も適切か（例：ガバナンス評価、適合性評価、パフォーマンス評価）、そしてその評価を社内で実施すべきか、第三者に委託すべきか

政策立案者に対しては、以下の事項を提案します。

- AIシステムへの信頼を築き、導入の成功を支援し、AIガバナンスに貢献する上で、自主的な（または規制に基づく）AI評価がどのような役割を果たし得るかを検討する
- 評価フレームワークの目的と構成要素を明確に定義し、可能であれば、評価を実施する際に準拠すべき確立された基準、または評価基準を定める
- AI評価の内容およびその限界に関する期待のギャップに対処する
- 高品質で一貫性のある評価を提供できるよう、市場の提供能力を強化するための適切な措置を特定する
- AI評価のコストを削減し、評価の妥当性に対する国境を越えた信頼を高めるため、可能な限り、他の国・地域の基準と整合性および互換性を有する評価基準を支持する

2 はじめに

2022年11月、OpenAIはChatGPTをリリースし、AIの現在の能力と将来の可能性が広く認識されるようになった一方で、AIの開発や導入に関連するリスクに対する懸念も高まりました。それ以来、企業や政策立案者などは、共通かつ根本的な課題、すなわち、目的に適合し、従業員、顧客、市場、そして社会全体から信頼されるAIアプリケーションをどのように開発・導入するかという課題への取り組みを強化してきました。

ChatGPTのような生成AIシステムや、最近ではAIエージェントといったAIがもたらす大きな機会を鑑みれば、AIの開発と導入は今後も拡大し続けるでしょう。例えば、**EY-Parthenon**の推計によると、生成AIだけでも2033年までに世界のGDPを1.7兆米ドルから3.4兆米ドル押し上げる可能性があります¹。しかし、AIに関連する有害なインシデントの増加を考慮すると、導入の成否は、このテクノロジーに対する信頼と確信に懸かっています。実際、**経済協力開発機構(OECD)**の報告によると、有害なインシデントの月間平均発生件数は増加の一途をたどっており、2022年11月の32件から2025年1月には614件へと、約20倍に増加しています²。2025年4月の**EY AIセンチメント指数調査**によると、調査対象となった人々の58%が、組織がAIの不適切な使用に対して責任を果たせていないことを懸念しており、52%が、組織がAIに関する社内ポリシーや規制要件に準拠できていないことを懸念していることが明らかになりました³。

AIの急速な発展の中で、ビジネスリーダー、政策立案者、学術関係者、投資家、保険業界関係者、そしてその他のステークホルダーは、以下のような喫緊かつ根本的な問いを投げかけています。

- AIシステムが信頼性と有効性を備えているかをどのように評価するのか？
- そのリスクをどのように識別し、管理するのか？
- AIシステムが有効性および品質に関する適用される規制やその他の基準を満たしているかをどのように判断するのか？

これらの問いに対応するため、数多くのAIガバナンスの枠組みが登場しています。これらの枠組みの多くは、テクノロジーのガバナンス、適用されるポリシーへの準拠、運用上の整合性、または有効性を検証するために設計された評価を取り入れています⁴。本書では、「AI評価」という用語を、「定義された対象事項(主題)」⁵について、結果、判断、または結論を導き出すための体系的な評価⁶を指すものとして使用します。

AI評価は、規制当局、ビジネスリーダー、投資家、保険会社、消費者など、多様なステークホルダーのニーズや要件に合わせて調整することが可能です。AI評価には、自主的なものと規制に基づくもの、定性的なものと定量的なものがあり、社内または外部の当事者によって実施され、さまざまな報告および開示指標が用いられます。また、AI評価は、特定のユースケースやリスクレベル、あるいはテクノロジーの運用領域に特化させることもできます。

AIシステムの厳格な評価は、開発および導入がガバナンス、コンプライアンス、または有効性に適用される評価基準を満たしていることを検証することで、テクノロジーに対する信頼を高めることができます。

1 [How global business leaders can harness the power of GenAI](#), EY, 1 August, 2024.

2 [AI Incidents and Hazards Monitor](#), OECD, January, 2025.

3 [How a license to lead can transform human potential in an AI world](#), EY, 9 April, 2025; [How responsible AI can unlock your competitive edge](#), EY, 3 June, 2025.

4 コンプライアンスには、適用される法律や規制のガイドライン、社内ポリシー、あるいは基準への準拠が含まれる場合があります。

5 保証業務の文脈において、「主題」とは、保証業務実施者が評価する特定の情報、プロセス、または統制の集合を指します。

6 政策文書や議論において「AI評価」を表す用語の使い方は幅広く、文書間で一貫していません。「保証」、「監査」、「ベンチマークテスト」、「認証」、「適合性評価」、「検証」などの用語が、場合によっては互換的に用いられます。さらに「監査」は、AIの分野では第三者によるあらゆる評価(調査報道、コンプライアンスおよびバイアスの評価、適合性評価など)を指す場合があります。本書では、これらすべての用語を広く「AI評価」の一形態として総称します。



3 AI評価に関する公共政策の現状

このセクションでは関連するAI政策を検討し、企業、AI評価実施者、その他のステークホルダーがこれらの政策を実施する上で直面する課題の一部をまとめます。

政策立案者は、この新たな分野において積極的に活動しており、AI評価に関する政策上の規制に基づく枠組みと自主的な枠組みの両方を策定しています。OECDによると、2025年1月時点で、約70カ国の政策立案者が、法律、規制、自主的な施策、協定を含む1,000件以上のAI関連政策イニシアチブを提案しています⁷。[スタンフォード大学の2025年の報告書](#)によると、39カ国以上が、これらのイニシアチブのうち204件を法制化しています⁸。正確な件数を把握することは難しいものの、提案段階または法制化済みの多数のAI政策イニシアチブにおいて、

AI評価が位置付けられています⁹。2025年7月には、米国のトランプ政権が「AIアクションプラン」を発表し、評価はAIシステムのパフォーマンスと信頼性を測定する上で重要な手段となり得ると指摘しています¹⁰。以下の表では、よく知られているAI評価の政策枠組みのいくつかを取り上げ、政策立案者がどのように多様なアプローチを取っているかを示しています。世界各国の政策イニシアチブに関するより広範なリストは、付録IIに掲載されています。

7 これらのイニシアチブは、多国間機関、各国政府、州や市など、さまざまなレベルで展開されており、それぞれ異なる目的を掲げています。
[National AI policies & strategies](#), OECD, January, 2025.

8 “[The AI Index 2025 Annual Report](#)”, Stanford Human-Centered Artificial Intelligence, April, 2025.

9 “[Global AI Law and Policy Tracker](#)”, IAPP, November, 2024.

10 “[America's AI Action Plan: Winning the Race](#)”, 23 July, 2025, pg. 10.

表1:AI評価を取り入れている公共政策の枠組みの例

枠組み	EU AI法 (EU AI Act)	G7 AI行動規範 (G7 AI Code of Conduct)	英国AI保証ツールキット (UK Toolkit on AI Assurance)	ニューヨーク市地方法 144号(New York City Local Law 144)
政策全体の目的	個人の安全、セキュリティおよび基本的人権の保護	安全でセキュアかつ信頼できるAIの世界的推進	AI保証実務者向けのリソースおよびガイダンスの提供	自動雇用決定ツール(Automatic Employment Decision Tool, AEDT)における潜在的なバイアスからの求職者保護
評価の目的	EU AI法で定められた義務に対するAIシステムの適合性評価	AIシステムの信頼性、安全性およびセキュリティの確保	AIリスクの測定、評価および伝達のための評価の提案	人種、民族、性別などの人口統計データカテゴリーに基づく、自動雇用決定ツール(AEDT)の人々への影響評価
評価の対象事項	AI品質マネジメントシステムおよび技術文書(プロセスおよびガバナンスを含む)	記載なし	手法により異なる(対象:データ、AIモデル、ガバナンスプロセス)	AIシステムの出力
評価の方法論	EU AI法の要件への準拠を示す適合性評価	詳細にわたる評価の記載なし	定義されたAI保証の手法および仕組み	カテゴリー別の選考率またはスコアリング率の算定を含むバイアス監査
評価の実施主体	自己評価または特定のAI用途に対する第三者評価	記載なし	評価類型に応じて検討される複数の選択肢	独立した第三者による評価
評価を表す用語	適合性評価、リスク評価	独立した外部テスト手法、影響およびリスクの評価	AI保証には、コンプライアンス監査、バイアス監査、形式的検証などの用語が含まれる	バイアス監査

現在の政策環境におけるテーマ:

AI評価の3つのカテゴリーの形成

AI評価の目的は、規制や基準への準拠の確認、AIシステムの出力にバイアスがないかどうかの判定、AI出力の正確性の測定など、多岐にわたります。AI評価の目的を明確に定義することは極めて重要です。それによって、評価に求められる要件

や期待事項が方向付けられるためです。以下の通り、AI評価は一般的に3つのカテゴリーに分類でき¹¹、単独または組み合わせて実施されます。

ガバナンス評価

AIシステムのリスク、適合性、信頼性を含め、AIシステムを管理するために適切なコーポレートガバナンス方針、プロセス、および人材が組織内に整備されているかどうかを判断するためのものです。

適合性評価

組織のAIシステムが関連する法律、規制、基準、またはその他の政策要件に準拠しているかどうかを判断するためのものです。

パフォーマンス評価

AIシステムの中核機能の品質（精度、非差別性、信頼性など）を測定するためのものです。これらの評価では、AIシステムの特定の側面を評価するために定量的な指標がしばしば使用されます。



¹¹ これらのカテゴリーは、相互に関連性がないものとして認識するべきではありません。例えば、AIシステムのガバナンスを評価する審査は、ISO/IEC 42001規格に対する組織のAIマネジメントシステムの評価など、適合性評価でもある場合があります。

AI評価に関する政策枠組み間の大きな差異

現在、AI評価に関する政策枠組みについては、規制に基づくものと自主的なものの両方において、そのあらゆる側面で大きな差異が見られます。具体的には、範囲、対象事項、方法論、評価実施者に求められる能力および資格の指定、さらには評価によって提供されることを想定した信頼の水準などに違いがあります。

評価の範囲は狭くもあれば非常に広くもあり、その幅は多岐にわたります。例えば、その範囲には、ニューヨーク市地方方法144号で示されているようなAIシステムの出力におけるバイアス、EUのデジタルサービス法やオーストラリアの保証フレームワークに見られるようなAIシステムを巡る組織のガバナンスおよび統制プロセス、あるいはEU AI法における適合性評価に含まれるようなデータガバナンスの特性などが含まれる場合があります。このような差異は、国・地域ごとの政策全体の目標や目的の違い、あるいは評価の想定するステークホルダーのニーズの違いによって、一部は説明できます。

さらに、AI評価の目的が一致している場合でも、AI評価フレームワークの具体的な要件は、国・地域によって異なる場合があります。例えば、米国のさまざまな州や市では、採用や雇用において使用されるAIシステムのバイアス評価を含む政策を定めています¹²。しかし、これらの評価の具体的な要件は大きく異なります。例えば、ニューヨーク市地方方法144号は、コロラド州やイリノイ州でバイアス評価を義務付ける州法とは、バイアスの測定に関する要件が異なります¹³。

また、AI評価は、必要とされる証拠の範囲や評価実施者に対する要件など、具体的な要件の設計によって、信頼性の程度もさまざまです。第三者が実施する評価は、社内チームが実施する評価よりも妥当性が高いと見なされる場合があります。特に、第三者の評価実施者が、社内チームには遵守義務がない可能性のある専門的責任、倫理、および公開報告に関する基準に準拠している場合、その傾向が強まります¹⁴。

最後に、例えば規制に基づくAI評価（規制準拠の評価）は、米国国立標準技術研究所（NIST）のAIリスクマネジメントフレームワーク¹⁵のようなガバナンス基準に基づく自主的な評価とは、多くの場合、大きく異なるものとなります。

国連の情報環境に関する国際パネル（IPIE）が2024年12月に発表した調査結果が示すように、AI評価のアプローチが多様であるため、一貫した品質と説明責任を確保することが困難となっています¹⁶。



12 “A running list of states and localities that regulate AI in hiring”, HR Dive, 20 May, 2024.

13 トランプ政権は、連邦政府の資金提供や補助金の決定を行う場合を含め、州レベルのAI政策がAIアクションプランとどのように整合しているかを確認するために、各州のAI政策を見直す意向であり、これは今後の州のAI政策の策定に影響を与える可能性があります。
“America’s AI Action Plan: Winning the Race”, 23 July, 2025, p. 3.

Schlemmer, Michael D., Morgan Lewis, “AI in the Workplace: The New Legal Landscape Facing US Employers”, 1 July, 2024.

14 “Enhancing AI Accountability: Effective Policies for Assessing Responsible AI, Business Software Alliance”, 23 October, 2024.

15 “AI Risk Management Framework”, NIST, January 2023.

16 IPIEは、AI評価を「AI監査」と呼んでいます。

“Recommendations for a Global AI Auditing Framework: Summary of Standards and Features”, IPIE, December 2024.

現在のAI評価の有効性に関する課題

国・地域ごとの違いに加え、現在、いくつかの共通要因が、一部のAI評価フレームワークの頑強性や有効性を阻害しており、ひいてはその本来の目的を達成する能力にも影響を及ぼしています。

これらの課題は、主にAI評価における以下のような重要な要素について、明確さや十分な定義が欠けていることに起因しています。

- AI評価の目的
- 評価の対象事項
- 方法論、評価を行う際の基準、証拠、および報告要件
- AI評価実施者に求められる資格、説明責任、および利益相反がないこと

AIテクノロジーの性質もまた、評価を複雑にする要因となります。AIシステムは往々にして複雑であり、より広範な環境に統合され、複数のステークホルダーが関与しています。これらの要因により、評価の対象事項を適切に特定することが困難になる場合があります。さらに、モデルのドリフト(時間の経過に伴うモデル結果の変動)により、評価結果が時代遅れになり、誤解を招く恐れがあります。また、AIシステムの変動性により、再現性の確保が難しくなる可能性があります。最後に、AIテクノロジーの急速な進歩は、性能を評価するための技術基準の策定を追い越してしまう可能性があります。

さらに、曖昧で一貫性を欠き、主観的な用語を使用すると、重要な概念や適切な評価基準について解釈が分かれる可能性があります。その結果、本来の目的を果たさない評価につながる恐れがあります。「公平性」、「信頼性」、「透明性」といった広範な用語は、さらに具体的に定義されない限り曖昧さを生じさせ¹⁷、特定の評価の実現可能性や有用性を制限する可能性があります¹⁸。

最後に、十分に整備されていない基準や方法論は、AI評価の厳密性や比較可能性にとって課題となります。ステークホルダーは、AI評価の基準の設定と適用において、より一層の明確性、一貫性、客観性、および方法論的な厳密性が必要であることに、ますます注目しています。例えば、[国際アルゴリズム監査人協会\(IAAA\)](#)は、専門家を集め、「アルゴリズム監査の基準の基礎を築く」ことを目的として設立されました。ISO/IEC¹⁹、CEN-CENELEC²⁰、NISTなどの標準化機関もこの課題に対し、既存の規格の改訂と新たなAI規格の策定の両方に取

り組んでいます²¹。英国では、規制当局が効果的な「AI保証」エコシステムのロードマップを示し²²、AI評価に関する詳細なガイダンスを提供するイニシアチブを開始しました²³。2025年2月、国連のIPIEは、技術的な考慮事項を示し、評価の範囲、評価者の資格、評価基準、および方法論に関するガイダンスを提供する、[包括的なグローバルな「AI監査」フレームワーク](#)²⁴を発表しました²⁵。

上記の課題に対処することは、AI評価のための一貫性があり効果的なポリシーおよびビジネスフレームワークを構築するために不可欠です。



17 英国の情報コミッショナーオフィス(ICO)が公表した、[AIによる意思決定の説明](#)に関するガイダンスでは、6つの主要な説明のタイプが特定されています。

18 曖昧で主観的な基準は、特定の状況において保証の提供を困難にする可能性があります。

19 [国際標準化機構\(ISO\)](#)と[国際電気標準会議\(IEC\)](#)による共同作業、2025年。

20 [欧州の標準化機関である欧州標準化委員会\(CEN\)](#)と[欧州電気標準化委員会\(CENELEC\)](#)による共同作業(CEN-CENELEC名義)、2025年。

21 例えば、ISO/IECはAIマネジメントシステムに関する新規格[ISO/IEC 42001:2023](#)を策定し、CEN-CENELECは既存のソフトウェア品質に関するISO/IEC規格を基に、AIシステムの品質に関する規格[EN ISO/IEC 25059:2024](#)を公表しました。

22 “[The roadmap to an effective AI assurance ecosystem](#),” 英国科学・イノベーション・技術省, 8 December, 2021.

23 英国の科学・イノベーション・技術省(DSIT)は、AI保証について次のように述べています。「『保証』という用語はもともと会計学に由来するものですが、その後、サイバーセキュリティや品質管理などの分野にも適用されるようになりました。保証とは、システムやプロセス、文書、製品、または組織に関する何らかの事項を測定、評価、伝達するプロセスです。AIの場合、保証とは、AIシステムの信頼性を測定、評価、伝達するものです」

24 “[Towards A Global AI Auditing Framework: Assessment and Recommendations](#)”, IPIE, February 2025.

25 国連のIPIEは、AI監査について次のように述べています。「AIシステムの監査は、個人、コミュニティおよび組織との相互作用を評価し、これらのシステムが適切に開発、導入、運用、管理されているかどうかを判断するのに役立ちます。監査を通じて、AIシステムが公平性、データプライバシー、環境の持続可能性といった重要な社会的、倫理的、法的規範を遵守しているかどうかを確認することができます」



4 AI評価をより効果的にする方法

AI評価フレームワークにおいて、AI評価を効果的にするためには、評価対象、評価の実施方法、評価の実施主体という3つの基本要素を、より明確かつ一貫して定義する必要があります²⁶。

評価対象

AI評価フレームワーク²⁷を効果的なものとするためには、明確に規定されたビジネス上、またはポリシー上の目的を有していることが重要です。目的が明確でなければ、評価によって提供される情報と、AI評価が本来果たすべき目的との間に不整合が生じる恐れがあります。また、評価の目的は、適切な方法論や参照すべき基準を選定する上での指針ともなります。

重要な点として、AI評価フレームワークには、評価の種類（ガバナンス、コンプライアンス、パフォーマンスなど）、対象事項、および評価を実施すべき時期に関するガイダンスを含め、明確かつ十分に定義された範囲を有している必要があります。例えば、評価の対象を、トレーニングデータ、アルゴリズム、セーフガードを含むAIシステム全体とするか、それともその結果のみに限定するかを決定することが重要です。

26 情報技術(IT)、自動車、製薬、サイバーセキュリティなどの分野で確立された評価フレームワークは、AI特有の特性に配慮した必要な調整が行われる限り、AI評価に有益な知見を提供することができます。例えば、IT分野では、組織のITインフラのセキュリティおよび有効性の確保のために、評価(一般に「監査」と呼ばれます)がしばしば用いられており、データの保護、リスク管理および関連する業界の規制への準拠に関する組織の能力を包括的に評価するものです。

27 評価フレームワークには、規制で義務付けられているもの、または自主的に実施されるものが含まれます。

評価の実施方法

明確な方法論(メソドロジー):

対象をどのように評価するかは、方法論と適切な基準により決定されます。同種類のAI評価では、明確に定義された一貫したアプローチを使用することが不可欠です。

例えば、評価によっては明確な意見や結論が含まれるものもあれば、実施された手順の概要のみを提供するものもあります。明確に定義された方法論、基準、証拠、および報告要件が欠如していると、評価結果が損なわれ、評価の利用者との間で誤解が生じる恐れがあります。一貫性があることで、明確な用語と相まって、利用者が評価結果を比較し、その結果がどのように導き出されたかを理解できるようになります。適切な基準(関連性があり、客観的で、測定可能かつ網羅的なもの)は、整合性がとれ、比較可能で、意思決定に役立つ評価結果をもたらします。

方法論は、保証業務の指針となるISAE 3000(改訂版)²⁸などの基準、または形式的検証、レッドチーム活動、品質保証などの他の評価プロセスを参照する場合があります(ISAE 3000(改訂版)の詳細については、付録Iをご参照ください)。評価手法は、評価の対象となるユースケースや状況において許容されるAIシステムの出力のばらつきの範囲など、AIシステムの扱いが難しい特性についても考慮する必要があります。

評価の基準は、ポリシーフレームワーク(政策の枠組み)内で直接定義することも、技術基準を通じて参照することも可能です。評価の基準は、評価結果の理解を容易にするため、評価の利用者にとって適切かつ利用可能なものである必要があります。方法論や基準を選択する際は、評価の目的、対象事項、および求められる信頼の水準と整合していなければなりません。特定の評価には、より適した方法論がある場合があります。

評価の実施主体

効果的なAI評価を行うためには、評価実施者の選択が極めて重要です。これは、評価実施者の客観性、専門知識、そして透明性の高い方法論への準拠が、評価プロセスの妥当性、信頼性、全体的な整合性に直接影響を及ぼすためです。評価実施者を選定する際に考慮すべき主な事項は以下の通りです。

- 能力と資格：妥当性を備えたAI評価には、AIに関する技術的知識と評価手続きに関する能力に加え、倫理的・規制的な枠組みに対する理解を有する専門家が必要です。
- 客観性：評価実施者の客観性(利益相反がないことを証明する能力を含む)は、評価の妥当性に影響し、ステークホルダーの信頼を醸成するのに役立ちます。

- 専門的な説明責任：専門的な説明責任に関する要件は、国際会計士倫理基準審議会(IESBA)の監査専門職のための倫理規程²⁹など、一般に公開され、広く受け入れられている基準やガイドラインに基づく場合があります。これらの基準やガイドラインに準拠する評価実施者は、信頼を醸成し、評価がどのように提供されるかをステークホルダーが理解する一助となります。

28 International Standard on Assurance Engagements (ISAE) 3000 Revised, Assurance Engagements other than Audits or Reviews of Historical Financial Information, December 2013.

29 International Code of Ethics for Professional Accountants, IESBA, 2024.



5 ビジネスリーダーの 考慮事項

- **コーポレートガバナンスやリスク管理の強化に対してAI評価が果たす役割を検討する。**AI評価は、ビジネスリーダーが自社のAIシステムに関連する新たなリスクを特定・管理し、AIシステムが意図した通りに機能しているかどうかを判断するのに役立ちます。
- **規制で義務付けられていない場合であっても、従業員、顧客、その他の重要なステークホルダーの間でAIシステムへの信頼性を構築するため、自主的な評価を実施するかどうかを診断する。**市場の動向、投資家の要求、あるいは内部ガバナンス上の考慮事項から、自発的なAI評価が企業のAIシステムに対する信頼構築に有益となる場合があります。さらに、一部のAIシステムが規制上の義務の対象となる場合、ビジネスリーダーはコンプライアンス状況の測定とモニタリングに役立てるため、AI評価を活用することを選択できます。
- **自主的な評価を実施する場合、最も適切な評価方法を決定する。**ビジネスリーダーは、ガバナンス評価、適合性評価、パフォーマンス評価のいずれを実施するのか、またその評価を社内で行うべきか、第三者に委託すべきかを判断する必要があります。

6 政策立案者の 考慮事項

- **AIシステムへの信頼を築き、導入の成功を支援し、AIガバナンスに貢献する上で、自主的な(または規制に基づく)AI評価がどのような役割を果たし得るかを検討する。**
- **評価フレームワークの目的と構成要素を明確に定義し、可能であれば、評価を実施する際に準拠すべき確立された基準、または評価基準を定める。**
- **AI評価の内容およびその限界に関する期待のギャップに対処する。**この情報は、評価の意義について現実的な期待を設定することで、一般の認識と信頼を高めることができます。
- **高品質で一貫性のある評価を提供できるよう、市場の提供能力を強化する措置を講じる。**政策立案者は、自らの国・地域において、効果的なAI評価を実施するための十分な能力が備わっているかどうかを確認することが望まれます。十分でない場合には、AI評価実施者、専門職団体、その他の関係者と連携し、評価の品質基準の策定や認定研修コースの開発を支援するなどして、能力構築に取り組む必要があります。
- **可能な限り、他の国・地域の基準と整合性および互換性を有する評価基準を支持する。**政策立案者は、評価のコストを削減し、評価の妥当性に対する国境を越えた信頼を高めるため、自国のAI評価基準を国際機関や主要な国・地域が定めた基準と整合させることを検討する必要があります。



7 おわりに

企業がAIシステムの開発と導入を進める中で、AI評価は、AIのメリットを最大限に引き出し、そのリスクを軽減する上で重要な役割を果たすことができます。適切に設計され、かつ資格を有する評価実施者によって実施されれば、AI評価は、この重要なテクノロジーの潜在能力を最大限に引き出すことができます。そうすることで、ビジネスリーダー、政策立案者、そして一般の人々が求めるAIへの信頼を高めることができます。



8 著者

Shawn Maher

EY Global Vice Chair Public Policy
Ernst & Young LLP

Dr. Ansgar Koene

EY Global AI Ethics and Regulatory Leader
Ernst & Young LLP

Anne McCormick

EY Global Digital Technology Public Policy Leader
Ernst & Young LLP

Tate Ryan-Mosley

EY Assistant Director, Technology and Geopolitics
Global Public Policy
Ernst & Young LLP

Richard Jackson

EY Global Assurance AI Leader
Ernst & Young LLP

Cathy Cobey

EY Global Responsible AI Leader, Assurance
Ernst & Young LLP

Narayanan Vaidyanathan

Head of Policy Development
ACCA

謝辞

本書の作成に当たり、以下のEYのメンバーの協力に謝意を表します。

Katie Kummer, Amy King, Marc Ryser, Yvonne Zhu, Julia Tay, Yi Xie, Denise Schalet, Bridget Neill, John Hallmark, Antonio Quintanilla, Andrew Hobbs, Tim Volkmann, Ambrose Murray, Dean Protheroe, Frank de Jonghe, Detmar Ordemann, Brent Simer, Brandon Miller, Kelly Devine, Dan Albanese and Krista Walpole。

また、本書の作成にご貢献いただいた以下のACCA会員各氏に感謝申し上げます。Alex Panait, Andrew Chong, Fahad Riaz and James Best。

お問い合わせ

EYストラテジー・アンド・コンサルティング株式会社
リスク・コンサルティング パートナー

川勝 健司

9 付録

付録I: ケーススタディ:ISAE 3000 (改訂版)のISO 42001への適用

政策立案者は、他の分野で使用されている既存の評価、保証、または認証の枠組み(ISO CASCOツールボックス³⁰、ISO/IEC 17067³¹、ISAE 3000(改訂版)³²、IFRS基準³³など)を、一定の修正を加えることでAIに適用できるかどうかを検討しています。既存の枠組みを活用することで、政策立案者は確立された品質管理および認定プロセスを利用できるようになる可能性があります。

例えば、国際監査・保証基準審議会(IAASB)³⁴が策定したISAE 3000(改訂版)基準は、財務諸表監査の範囲を超える分野における保証業務の要件および方法論を概説し、対象事項を適用される基準と比較するための手順を詳述しています。ISAE 3000(改訂版)は、原則に基づく基準であり、幅広い対象事項に適用可能です。この国際基準は、サステナビリティ、内部統制、規制遵守など、多岐にわたる分野における保証業務の基盤となってきました。ISAE 3000(改訂版)に規定されている保証提供者に対する要件には、以下のものが含まれます。

- 利益相反がないことなど、関連する倫理的要件を遵守していること
- 保証の対象事項および範囲について十分な理解を有していること(「合理的」対「限定的」)
- 適用される基準に照らして対象事項を評価するために必要な証拠を入手すること
- 評価の結果に関する結論を表明すること

保証提供者は、ISAE 3000(改訂版)を用いて、ISO/IEC 42001:2023などの認められた規格に照らしてAIマネジメントシステムを評価することができます。このような業務は、AIマネジメントシステムの国際的に認められた規格への適合性を評価するために用いることができます。

ISO/IEC 42001は、組織内におけるAIマネジメントシステム(AIMS)の確立、実施、維持および継続的改善に関する要求事項を規定しています。本規格は、AIベースの製品やサービスを提供または利用する組織を対象としており、AIシステムの責任ある開発と利用を確保するように設計されています。ISO/IEC 42001は、倫理的配慮、透明性、継続的な学習など、AIがもたらす課題の一部に対応しています。

EU AI法への準拠を支援する適合性評価フレームワークの策定に向けて、CEN-CENELEC JTC21において進行中の作業では、ISO CASCOツールボックスおよびISO/IEC 17067:2013「適合性評価—製品認証の基礎および製品認証スキームのための指針」を主要な参照文書として用いています。これにより、企業は、EU AI法における高リスクAIシステムに関する義務への準拠に向けた準備に当たり、非AIシステムで用いてきた既存の適合性評価手続きを基盤として活用するための手段を得ることになります。

30 Conformity Assessment tools to support public policy, ISO CASCO toolbox - Conformity Assessment tools to support public policy, 2024.

31 ISO/IEC 17067:2013 - Conformity assessment — Fundamentals of product certification and guidelines for product certification schemes, August 2013.

32 International Standard on Assurance Engagements (ISAE) 3000 Revised, Assurance Engagements other than Audits or Reviews of Historical Financial Information, December 2013.

33 International Financial Reporting Standards (IFRS), 2025.

34 International Auditing and Assurance Standards Boards (IAASB), 2025.

35 ISO/IEC 42001:2023 - Information technology — Artificial intelligence — Management system, ISO/IEC 42001:2023 - AI management systems, December 2023.

付録II: AI評価に関連する政策 イニシアチブの例

2022年以降、超国家、国家、地方の各レベルにおいて、各国・地域の政策立案者の活発な動きが見られます。以下の例は、自主的なAI評価および規制に基づくAI評価の両方に関連する、目的およびアプローチの幅広さについて、さらなる知見を提供するものです。

政策イニシアチブ	政策イニシアチブの 現状と目的	地理的範囲	評価を表す 用語	評価の役割と詳細
シンガポールの AI Verify認証	2022年5月に一般公開。この自主的なAIガバナンスのテストフレームワークおよびツールキットは、AIシステムの性能を開発者の主張および国際的に認められたAI倫理原則に照らして検証するよう設計されている。	全世界で一般公開。自主的な利用が可能（制限なし）。シンガポール情報通信メディア開発庁（IMDA）および個人データ保護委員会（PDPC）により公表。	テストおよび保証（外部による検証および第三者によるテストを含む）。	企業が自社のAIシステムの責任ある実装を、国際的に認められた11のAIガバナンス原則に照らして評価するのに役立つAIガバナンスのテストフレームワークである。これらのガバナンス原則（透明性、頑強性、公平性など）は、EUおよびOECDのAIフレームワークと整合している。AI Verifyは、標準化されたテストレポートを通じて、組織が自社のAIシステムの性能をこれらの原則に照らして検証できるよう支援する。

政策イニシアチブ	政策イニシアチブの現状と目的	地理的範囲	評価を表す用語	評価の役割と詳細
EUのデジタル市場法 (DMA)	2022年11月に発効し、2023年5月2日から適用。公正で開かれたデジタル市場の確保を目的とする。	EUにおいて事業を行う大規模デジタル・プラットフォームのうち、DMAのゲートキーパー指定の基準を満たす市場での地位を有するもの。	独立監査。	規制当局（欧州委員会）に対し、デジタル・ゲートキーパー・プラットフォームが中核プラットフォームサービスに適用している消費者プロファイリングのあらゆる手法について、独立監査済みの説明を提供する。
EUのデジタルサービス法 (DSA)	2022年11月に発効。インターネット上の利用者の基本的権利を包括的に保護することを目的とする。	EUにおいて事業を行う大規模デジタル・プラットフォームのうち、DSAの超大規模オンライン・プラットフォームまたは超大規模オンライン検索エンジンの指定の基準に該当するアクティブ利用者数を有するもの。	独立監査、リスク評価。	独立監査では、システミックリスクの年次評価およびリスク軽減措置を含め、DSAの要件への準拠を検証する。
NISTのAIリスクマネジメントフレームワーク (NIST AI RMF)	2023年1月公表。個人、組織および社会に対するAIに関連するリスクをより適切に管理するための自主的なリスクマネジメントフレームワークを提供することを目的とする。	米国NISTは、米国外の利用者向けに指針を示すため、他の国・地域（EU、日本、シンガポールなど）の政策枠組みとのクロスワークを複数実施している。	リスク管理、リスク評価、影響評価、パフォーマンス評価。	個人、組織および社会がAIのリスクを管理し、AIの信頼できる開発および責任ある利用を促進するとともに、AI製品、サービスおよびシステムの評価を促進することを支援するために策定された。
EUのデジタル・オペレーショナル・レジリエンス法 (DORA)	2023年1月に発効し、2025年1月に適用開始。金融機関のITセキュリティを強化し、金融セクターのレジリエンスを確保することを目的とする。	EU域内で事業を行うすべての金融機関。	検証(自主的)。外部監査(自主的)。外部または内部のテスターによるテスト(規制に基づく)。	(自主的)ICTリスクマネジメントフレームワークおよび要件への準拠の検証。ICTの第三者サービス・プロバイダーとの契約上の取り決めの監査。 金融機関のICTツールおよびシステムに対するデジタル・オペレーショナル・レジリエンス・テスト。

政策イニシアチブ	政策イニシアチブの現状と目的	地理的範囲	評価を表す用語	評価の役割と詳細
ドイツ経済監査士協会(IDW)PS 861(AIシステムの監査に関する基準)	本基準の最新版は2023年3月に公表。AIシステムの監査のための自主的なフレームワークの提供を目指す。監査に対する体系的なアプローチを確立することでAIテクノロジーに対する信頼を高め、AIの実装に伴うリスクの管理において組織を支援することを目的とする。	主としてドイツを対象とするが、EU域内およびそれ以外の地域での適用も想定される(例えば、欧州全域またはグローバルに事業を展開する組織に適用される場合)。	自主的な監査、評価基準、適切性監査、AIシステムの有効性監査、合理的保証。	財務監査の対象範囲外におけるAIシステムに対する自主的な監査の要件を明確化し、経済監査士が監査人自身の責任を維持しつつ、これらの業務をどのように計画・実施および報告すべきかについての専門的見解を定める。本基準は、倫理、法規制、トレーサビリティ、ITセキュリティおよびパフォーマンスの要件に基づき、AIシステムに対する相互に関連する評価基準を定める。このAI監査の対象は、AIシステムの記述であり、選定された評価基準への準拠に関する経営陣の説明を含む。AI監査は、適切性監査または有効性監査のいずれかの形式で実施され、いずれも合理的保証を伴う。
ブレッチリー宣言	2023年11月に合意。AIの安全な開発および利用に関する主要な原則およびコミットメントを示す国際的な合意であり、AIシステムに対する強固な安全対策、厳格なテスト、継続的なモニタリングを含む。	署名国28カ国:オーストラリア、ブラジル、カナダ、チリ、中国、フランス、ドイツ、インド、インドネシア、アイルランド、イスラエル、イタリア、日本、ケニア、サウジアラビア、オランダ、ナイジェリア、フィリピン、韓国、ルワンダ、シンガポール、スペイン、スイス、トルコ、ウクライナ、アラブ首長国連邦、英国、米国。EUも署名。	安全性テスト。	企業に対し、フロンティアAIの潜在的に有害な能力を測定、モニタリングおよび軽減するため、安全性テスト、評価、説明責任および透明性の仕組みを含む措置の実施を推奨する。こうした安全性テストおよび説明責任の仕組みの詳細は、本宣言では詳述されていない。

政策イニシアチブ	政策イニシアチブの現状と目的	地理的範囲	評価を表す用語	評価の役割と詳細
ISO/IEC 42001: 2023 AIマネジメントシステム	2023年12月に公表。AIベースの製品やサービスを提供または利用する組織による、AIシステムの責任ある開発と利用を確保することを目的とする。	グローバル。	リスク評価、影響評価、適合性評価、保証、内部監査。	ISO/IEC 42001は、組織内におけるAIマネジメントシステム(AIMS)の確立、実施、維持および継続的改善に関する要求事項を規定する。AIベースの製品やサービスを提供または利用する組織を対象としており、AIシステムの責任ある開発と利用を確保するよう設計されている。
国連のAI監査グローバル基準パネル(IPIE)ーグローバルAI監査フレームワークに関する提言:基準・要素の概要、評価・提言	IPIEは、AI評価に関する報告書を2本公表(2024年12月、2025年2月)。IPIEは、AI監査の基準と方法論を定義し、グローバル基準を確立してAIの公共への影響に焦点を当てた議論を促進することを目指す。	グローバルな範囲。情報環境に関する国際パネル(IPIE)が作成(国連と連携)。	AI監査。	アルゴリズムやAIシステムが期待される結果を生み出しているか、または社会的・技術的に重大な(場合によっては有害となり得る)影響を及ぼしているかを検証する手段としての監査。こうした監査は、AIシステムがAIの責任、説明責任、信頼性または安全性に関する規範や原則と整合しているかを評価するための仕組みと位置付けられている。また、監査は、AIシステムの設計、開発および運用を精査し、多くの場合、その中で使用されるモデルやデータも検証する。監査は、監査対象のAIシステムが、あらかじめ定められた基準に照らしてどのように機能しているかを記述し、その影響を報告するために用いられる。

政策イニシアチブ	政策イニシアチブの現状と目的	地理的範囲	評価を表す用語	評価の役割と詳細
米国商務省国家電気通信情報庁(NTIA)説明責任政策報告書	2024年3月に政策文書として公表されており、法的拘束力はない。信頼できるAIのイノベーションおよび導入を促進し、より広く利用可能な新しい説明責任ツールや情報の必要性を強調するとともに、独立したAIシステム評価のエコシステムを促進することを目的とする。	NTIA(米国の行政府に属する行政機関)によって作成。バイデン政権下で公表。トランプ政権が同様の提言を継続するかどうかは、現時点では不明。	AI説明責任メカニズム、AIシステム保証。	AI監査のより広範な適用を提唱する一方、執行メカニズムの具体化には至っていない。 本報告書は、(将来の連邦政府による)AI政策立案が純粋に自主的なベストプラクティスのみに全面的に依拠すべきではなく、一部のAI説明責任措置を義務付けるべきであると提言している。過去においても、NTIAはAIシステム監査に関する全国的登録制度の創設や、特定のシステムまたはモデルを対象とするリリース前の審査および認証の導入を求めてきた。
コロラド州AI法	2024年5月に可決。2026年2月に施行予定。2025年4月に本法に対する一連の改正案が提出されたが、コロラド州議会が5月7日に閉会するまでに可決されなかった。公共部門を含む分野横断型のAIガバナンス法であり、高リスクAIシステムを対象として、自動化された意思決定システムにおけるバイアスの防止に重点を置いている。	米国コロラド州内のデプロイヤーおよびデベロッパー。	影響評価、リスク評価。	高リスクAIシステムのデベロッパーおよびデプロイヤーに対し、バイアスや差別を含む影響評価およびリスク評価の実施を義務付ける。影響評価には、以下を含めなければならない:1.システムの目的、想定されるユースケースおよび導入状況を開示する記載。2.アルゴリズムによる差別のリスクおよび講じた軽減措置の分析。3.処理されるデータのカテゴリーの記述。4.システムの性能を評価するために用いられる指標および既知の制限事項。5.講じられた透明性措置の記述。6.導入後のモニタリング、および導入に伴い生じる問題に対処するための利用者保護措置の記述。

政策イニシアチブ	政策イニシアチブの現状と目的	地理的範囲	評価を表す用語	評価の役割と詳細
EU AI法に基づく汎用AI(GPAI)の行動規範	EU AI法の一部として採択。策定作業は進行中。関連するEU AI法上の義務は2025年8月2日に適用開始。本行動規範の利用は任意。GPAIモデルの開発者に対する追加のガイダンスを提供し、義務を明確化することを目的とする。GPAIの行動規範に従うことで、利用者はEU AI法の一部の要件への準拠を示すことができる。	本行動規範は、EU域内でAIシステムを開発、流通またはその他の方法で導入するあらゆる企業（EU域外に本社を置く企業を含む）によるEU AI法への準拠を支援する。	リスク評価、システミックリスク評価。	（評価の詳細は現時点では未確定である。しかし、評価の概要はEU AI法において既に大枠で示されており、GPAIのシステミックリスクの評価および管理のための措置、手続きおよび方法の確立〈それらの文書化を含む〉を含む） ³⁶

36 本書の公表時点では、AI法に基づく汎用AI(GPAI)の行動規範はまだ公表されていません。

ACCAについて

私たちACCA（英国勅許公認会計士協会）は、世界的に認知されている職業会計専門家団体として、資格を提供するとともに、会計分野の基準の向上を世界中で推進しています。

1904年に会計専門職への参入機会を広げるために設立されて以来、私たちは長年にわたりインクルージョンを推進してきました。現在では、180カ国で252,500人を超える会員と526,000人の将来の会員から成る多様なコミュニティを誇りを持って支援しています。

私たちの将来を見据えた資格、継続的な学習機会、そして知見は、あらゆる業界の雇用主から高く評価され、信頼されています。これらを通じて、個人は組織や経済がもたらす持続可能な価値を創出し、保護し、報告するために必要なビジネスおよびファイナンスに関する専門知識と倫理的判断力を身に付けることができます。

私たちのパーパスと価値観に基づき、世界が必要とする会計専門職を育成することをビジョンとしています。政策立案者、基準設定機関、ドナーコミュニティ、教育機関、他の会計団体と連携し、すべての人にとって持続可能な未来を推進する専門職を強化し、その発展を支えています。

詳しくはaccaglobal.comをご覧ください。

© ACCA JUNE 2025.

All Rights Reserved.

EY | Building a better working world

EYは、クライアント、EYのメンバー、社会、そして地球のために新たな価値を創出するとともに、資本市場における信頼を確立していくことで、より良い社会の構築を目指しています。

データ、AI、および先進テクノロジーの活用により、EYのチームはクライアントが確信を持って未来を形づくるための支援を行い、現在、そして未来における喫緊の課題への解決策を導き出します。

EYのチームの活動領域は、アシュアランス、コンサルティング、税務、ストラテジー、トランザクションの全領域にわたります。蓄積した業界の知見やグローバルに連携したさまざまな分野にわたるネットワーク、多様なエコシステムパートナーに支えられ、150以上の国と地域でサービスを提供しています。

All in to shape the future with confidence.

EYとは、アーンスト・アンド・ヤング・グローバル・リミテッドのグローバルネットワークであり、単体、もしくは複数のメンバーファームを指し、各メンバーファームは法的に独立した組織です。アーンスト・アンド・ヤング・グローバル・リミテッドは、英国の保証有限責任会社であり、顧客サービスは提供していません。EYによる個人情報の取得・利用の方法や、データ保護に関する法令により個人情報の主体が有する権利については、ey.com/privacyをご確認ください。EYのメンバーファームは、現地の法令により禁止されている場合、法務サービスを提供することはありません。EYについて詳しくは、ey.comをご覧ください。

EYのコンサルティングサービスについて

EYのコンサルティングサービスは、人、テクノロジー、イノベーションの力でビジネスを変革し、より良い社会を構築していきます。私たちは、変革、すなわちトランスフォーメーションの領域で世界トップクラスのコンサルタントになることを目指しています。7万人を超えるEYのコンサルタントは、その多様性とスキルを生かして、人を中心に据え（humans@center）、迅速にテクノロジーを実用化し（technology@speed）、大規模にイノベーションを推進し（innovation@scale）、クライアントのトランスフォーメーションを支援します。これらの変革を推進することにより、人、クライアント、社会にとっての長期的価値を創造していきます。詳しくは、ey.com/ja_jp/services/consultingをご覧ください。

© 2026 EY Strategy and Consulting Co., Ltd.

All Rights Reserved.

ED None

本書は一般的な参考情報の提供のみを目的に作成されており、会計、税務およびその他の専門的なアドバイスを行うものではありません。EYストラテジー・アンド・コンサルティング株式会社および他のEYメンバーファームは、皆様が本書を利用したことにより被ったいかなる損害についても、一切の責任を負いません。具体的なアドバイスが必要な場合は、個別に専門家に相談ください。

本書はAI assessments: enhancing confidence in AIを翻訳したものです。英語版と本書の内容が異なる場合は、英語版が優先するものとします。

ey.com/ja_jp