

『広島AIプロセス「報告枠組み」参加に向けた手引き』の概要

- 2025年2月7日から、広島AIプロセスの「国際行動規範」の遵守状況をAI開発企業等が自主的に確認・報告する「報告枠組み」がOECDのウェブサイト（<https://oecd.ai/en/transparency/overview>）で運用開始。
- 2026年5月28日には、中小企業を含むより広範な組織の参加を促すため、自由記述式の「報告枠組み（1.0）」から、回答選択式等を導入した「報告枠組み2.0」にアップデート。
- 本手引きは、「報告枠組み」に新たに参加する組織がその報告を行うに当たっての参考となるよう、「報告枠組み（1.0）」に参加した組織からの報告をもとに、蓄積された共通項やベストプラクティスを分析・紹介。「報告枠組み」の設問・回答で使用されている用語の解説や参考となるガイダンス、ガバナンスフレームワーク等についても紹介。
※ なお、「報告枠組み2.0」は、「報告枠組み（1.0）」とセクション構成は同一であるものの、各セクション内の設問には一部相違があり、また、回答様式が、選択式と記述式を組み合わせた様式へと変更されている点に留意が必要。

「報告枠組み」の設問別のガイダンス

- 「報告枠組み」の設問ごとに、回答傾向、回答の共通項等を分析し、紹介。

設問a	①
AIに関連するさまざまな種類のリスク（不合理なリスク等）について、どのように定義および／または分類していますか。 How does your organization define and/or classify different types of risks related to AI, such as unreasonable risks?	①
設問の意図	②
AIに関連する様々な種類のリスクについて、自組織においてどのような基準・方法で定義し、具体的に定義・分類しているかを示すことが求められています。	②
回答にあたっての考え方	③
● リスク定義・分類の方法としては、リスクの程度（影響度、発生頻度）に基づいてリスクを定義・分類する方法のほか、プライバシー保護、機密情報といったAIに特有のリスクに基づいて定義・分類する方法が見受けられます。 ● なお、多くの組織で、リスクの程度を定義・分類するための閾値を設定している例が見受けられます。閾値の設定にあたっては、独自のスコアカードの運用や、参照先となる法令等を示す例が見受けられます。	③
回答の要素	④
1 リスクの程度による定義・分類の例 ● リスクの程度に基づいてリスクを定義・分類する場合、以下のような例があります。 ➢ 禁止、高、中、低リスク ➢ 不合理なリスクと予測可能なリスク 等	④
2 法令、ガイドライン、フレームワーク等を参照する例 ● EU AI法、NIST AI RMF、AI事業者ガイドライン、AIビジネスガイドライン、広島AIプロセス国際行動規範、OECD AI原則、AIセーフティに関する評価観点ガイド（第1.10版）等が想定されます。	④

※「回答にあたっての考え方」は、OECDページも参照しています。
OECD (2025) "How Are AI Developers Managing Risks?", P.12

- ①英語原文の日本語訳 ③回答の傾向を分析
②設問の意図を解説 ④共通項・ベストプラクティスなどを紹介

「報告枠組み」に係る用語集

- 「報告枠組み」の設問・回答で使用されている用語（技術的用語を含む。）を解説し、また、参考となるガバナンスフレームワーク等を紹介。

3.本手引きで使われる用語集（法令、ガイドライン、フレームワーク等）

#	英語	日本語	用語説明
1	AI Guidelines for Business	AI事業者ガイドライン	日本における従来の3ガイドライン（AI開発・AI活用・AI原則実践）を統合し、AIの開発者・提供者・利用者の各主体が遵守すべき事項を示した指針。国際的な議論（OECD・広島AIプロセス等）との整合性を考慮している。（※1）
2	EU AI Act	EU AI 法	EUの包括的AI規制法。AIシステムをリスクに応じて4段階（許容できないリスク・高リスク・限定リスク・最小リスク）に分類し、規制を課す。（※2）
3	NIST AI Risk Management Framework (NIST AI RMF)	NIST AIリスクマネジメントフレームワーク	米国国立標準技術研究所（NIST）が策定した、AIシステムのリスクを管理するためのガバナンスフレームワーク。「統治（Govern）・マッピング（Map）・測定（Measure）・管理（Manage）」の4機能で構成されている。（※3）
4	OECD AI Principles	OECD AI原則	政策立案者及びAI関係者に向けた、5つの価値に基づく原則と5つの勧告で構成される、信頼できるAIの責任ある管理のための国際的な政策枠組み。（※4）
5	—	AIセーフティに関する評価観点ガイド	AIシステム、特に大規模言語モデル（LLM）を構成要素とするAIシステムの開発・提供に関わる事業者を対象に、安全性評価の観点や手法を整理したガイドライン。（※5）
6	—	AIのセキュリティ確保のための技術的対策に係るガイドライン	AI自体への脅威・攻撃（プロンプトインジェクション攻撃やDoS攻撃など）への具体的な対策について整理したガイドライン。（※6）