

調達チェックシート						(参考) 要求事項の参考										AWSサンプル回答
分類	評価 観点 #	評価観点	評価・選定時の 項目の分類	要求 事項 #	要求事項	対策例の対象AI					対策例	対策例詳細	裏付けとなる情報の例	(参考) 「DS-120 デジタル・ガバメント推進標準ガイドライン実践ガイドブック」の「各種テンプレート」との対応関係	(参考) 対策例・対策例詳細の参考文献	
						入力		出力								
						テキストを含む	音声を含む	テキストを含む	画像を含む	音声を含む						
	4	情報セキュリティシナシデント・生成AIシステム特有のリスク発生時の対応	基本項目	4	情報セキュリティシナシデント・生成AIシステム特有のリスク（事業者の責任の範囲に属するものに限る。）対応体制・手順を整備していること（開発・運用するサービスにおいて、ユーザーからインシデント報告を受け付け、対応の協力をすることを含む。）	●	●	●	●	●	情報セキュリティシナシデント・生成AIシステム特有のリスク発生時に備え以下のような体制をあらかじめ整備する ・連絡受付窓口の設置 ・プロジェクトに対し責任ある役職者のアサイン ・対応担当者個々の役割分担 ・対応アプローチ・プロセス ・影響を受けるステークホルダーに対して通知するプロセス等	生成AIシステムの開発・運用時に発生する情報セキュリティシナシデント・生成AIシステム特有のリスク発生時の対応体制や方針及び計画を策定する	・AIガバナンス関連ルール（情報セキュリティシナシデント・生成AIシステム特有のリスク発生時の対応体制・手順が整備されていることが分かる資料）等	「調達仕様書標準テンプレート」 6.作業の実施に当たっての遵守事項 6.1.機密保持、資料の取扱い 「要件定義書標準テンプレート」 3.非機能要件定義 3.10. 情報セキュリティに関する事項	-	実際のシステム運用における具体的なインシデント対応体制・手順の整備は事業者の責任範囲となります： ただし、AWS では以下のリソースを提供しています： 1. インシデント検知・制御 - Amazon Bedrock Guardrailsによる不適切な入出力の制御 - CloudTrailによるAPI呼び出しの記録 - CloudWatchによる異常検知 2. セキュリティシナシデント発生時のAWS側の対応 - AWSセキュリティ対応チームによる24時間365日の監視 - AWS セキュリティ通報を通じて脆弱性情報の提供 3. Amazon Bedrock のガードレールを通じて生成AI特有のリスクへの対応 - ハルシネーション対策機能の提供 - プロンプトインジェクション対策の実装 - データの漏洩防止機能の提供 事業者は、これらのAWS提供のリソースを活用しつつ、具体的な運用体制や手順を整備することが可能です。
						●	●	●	●	●	情報セキュリティシナシデント・生成AIシステム特有のリスク発生に関する教育や訓練を実施する	生成AIシステムの開発・運用時に発生する情報セキュリティシナシデント・生成AIシステム特有のリスク発生時の対応方針及び計画について、適宜実践的な予定演習を実施する	・AIガバナンス関連ルール（情報セキュリティシナシデント・生成AIシステム特有のリスク発生時の対応体制・手順が整備されていることが分かる資料）等	-	(参考) ・Amazon Bedrock アブレーションで責任ある AI のコアディメンションに対応するための考慮事項：https://aws.amazon.com/jp/blogs/news/considerations-for-addressing-the-core-dimensions-of-responsible-ai-for-amazon-bedrock-applications/ ・AWS CloudTrailを使用した API コールのログ記録：https://docs.aws.amazon.com/ja_jp/fis/latest/userguide/logging-using-cloudtrail.html ・CloudWatch 異常検出の使用：https://docs.aws.amazon.com/ja_jp/AmazonCloudWatch/latest/monitoring/CloudWatch_Anomaly_Detection.html ・AWS Security Incident Response：https://aws.amazon.com/jp/security-incident-response/ ・セキュリティ通報：https://aws.amazon.com/jp/security/security-bulletins/ ・Amazon Bedrock のガードレール：https://aws.amazon.com/jp/bedrock/guardrails/	
	5	関係者への生成AIに関する教育・リテラシー向上	基本項目	5	生成AIシステムの開発・運用に従事する者又は組織について、生成AIに関するリテラシー向上の取組を実施していること	●	●	●	●	●	生成AI開発・運用に従事する者が適切なリテラシー（生成AI関連の技術知識やリスク、倫理等）を具備できるよう、役割に応じた教育や訓練を実施する	生成AIシステムの開発・運用時に発生する情報セキュリティシナシデント・生成AIシステム特有のリスク発生時の対応方針及び計画について、適宜実践的な予定演習を実施する	・AIガバナンス関連ルール（教育・訓練等、AIリテラシーの向上の取組を実施していることが分かる資料）等	「調達仕様書標準テンプレート」 5.作業の実施体制・方法に関する事項 5.2.作業要員に求める資格等の要件	総務省、経産省「AI事業者ガイドライン 第1.2版」	本項目における開発・運用要員への教育・訓練は、主として事業者が実施すべき内容となります： ただし、AWSでは、事業者の教育・訓練をサポートするため、以下のリソースを提供しています。これらのリソースは、事業者が自社の教育・訓練プログラムを構築する際の参考として活用可能です。特に、AWS認定試験は、AIおよび機械学習に関する知識・スキルを客観的に評価する基準として活用できます。 1. AWS認定試験による知識・スキルの体系的な習得 - AWS Certified Machine Learning - Specialty - AWS Certified AI Practitioner 2. AWS公式のトレーニングリソース - Amazon Bedrockのデモリソース - AWS Skill Builderでの生成AI関連コンテンツ - AWS Well-Architectedフレームワークにおける生成AI関連のベストプラクティス 3. 責任あるAIの理解促進 - AI Service cardsを通じて技術特性やリスクの解説 - Guardrailsの設定・運用に関するガイドライン (参考) ・AWS 認定：https://aws.amazon.com/jp/certification/ ・Amazon Bedrock Documentation：https://docs.aws.amazon.com/ja_jp/bedrock/ ・AWS Skill Builder で 生成 AI を勉強する 4 ステップ：https://aws.amazon.com/jp/blogs/news/skill-builder-generative-ai-4step/ ・Amazon Bedrock を使用した Well-Architected な生成 AI ソリューションによる運用上の優秀性の実現：https://aws.amazon.com/jp/blogs/news/achieve-operational-excellence-with-well-architected-generative-ai-solutions-using-amazon-bedrock/ ・AI/ML について学ぶ：https://aws.amazon.com/jp/training/learn-about/machine-learning/ ・AWS の AI サービスのカードを詳しく見る：https://aws.amazon.com/jp/ai/responsible-ai/resources/#service ・Amazon Bedrock のガードレール：https://aws.amazon.com/jp/bedrock/guardrails/
開発・運用工 程要件	6	データの取扱い	基本項目	6	生成AIシステムへの入出力又は処理されるデータの取扱いを適切に管理していること	●	●	●	●	●	プロジェクトの目的に照らし、必要な範囲でインプット（プロンプト、学習用の生データ等）の利用目的や利用条件を定める		事業者によるインプット（プロンプト、学習用の生データ等）の利用目的を定めていることがわかる資料等	「要件定義書標準テンプレート」 3.非機能要件定義 3.2.システム方式に関する事項 3.3.システム規模に関する事項 3.10.情報セキュリティに関する事項	経産省「AIの利用・開発に関する契約チェックリスト」	GenUでは、Amazon Bedrockを通じて、以下のデータ管理機能を提供しています： 1. データの入力制御と保護 Amazon Bedrockは、ガードレール機能を通じてコンテンツフィルタリングやプロンプトインジェクション対策などのリスク低減機能を提供しています。GenUではガードレール機能を有効化することにより、不適切な内容や悪意のある入力を検出し、制御することが可能です。 2. 外部データ連携時の制御 Knowledge Baseを使用した外部データとの連携時には、メタデータフィルタリング機能により、適切なデータ管理を実現します。また、Bedrock Agentsを使用した外部データ連携においては、IAMサービスロールによるアクセス制御やAWS Lambdaによるプログラムベースの権限制御が可能です。 3. データ品質の評価 Bedrockのモデル評価ツールにより、パフォーマンスやバイアスの評価が可能です。これにより、データ品質の継続的な監視と改善を支援します。
						●	●	●	●	●	プロジェクトの目的に照らし、必要な範囲でインプット（プロンプト、学習用の生データ等）の管理・セキュリティ体制を定めており、問題が生じないよう対策を講じる	インプット（プロンプト、学習用の生データ等）の管理・セキュリティ体制方針がわかる資料等				
						●	●	●	●	●	プロジェクトの目的に照らし、必要な範囲でインプット（プロンプト、学習用の生データ等）の保持期間及び消去の管理をしており、問題が生じないよう対策を講じる	インプット（プロンプト、学習用の生データ等）の保持期間及び消去の管理方針がわかる資料等				
						●	●	●	●	●	プロジェクトの目的に照らし、必要な範囲でインプット（プロンプト、学習用の生データ等）の外部提供管理として、生成AIシステムのユーザーに対して提供する義務を定めた上での管理をする	インプット（プロンプト、学習用の生データ等）をユーザーに対して提供する義務を定めていることがわかる資料等				
						●	●	●	●	●	プロジェクトの目的に照らし、必要な範囲でインプット（プロンプト、学習用の生データ等）の外部提供管理として、第三者提供に關しての管理をする	インプット（プロンプト、学習用の生データ等）の第三者提供に関する管理方針がわかる資料等				
						●	●	●	●	●	プロジェクトの目的に照らし、必要な範囲でインプット・処理結果（学習用データ、中間生成物、派生的知的財産等）の使用・利用に關しての管理をする	インプット・処理結果（学習用データ、中間生成物、派生的知的財産等）の使用・利用に関する管理方針がわかる資料等				
						●	●	●	●	●	プロジェクトの目的に照らし、必要な範囲でインプット・処理結果（学習用データ、中間生成物、派生的知的財産等）の外部提供に關しての管理をする	インプット・処理結果（学習用データ、中間生成物、派生的知的財産等）の外部提供に関する管理方針がわかる資料等				
						●	●	●	●	●	プロジェクトの目的に照らし、必要な範囲で学習前・学習全体を通じて、データのアクセスを管理するデータ管理・制限機能の導入検討を行う等、適切な保護措置を実施する	学習前・学習全体を通じて、データの保護措置の方針がわかる資料等				
	7		アウトプットの品質保証	基本項目	7	生成AIシステムの期待品質を満たすための取組を行っていること	●	●	●	●	●	生成AIをチャットで用いる場合、チャットでの利用の期待品質を定め、その基準を満たしているかの測定及び評価をする	ベンチマーク※等を使用した評価基準とその評価内容がわかる資料等 ※一定の信頼性のある評価基準に基づき独自の評価を否定するものではない	「要件定義書標準テンプレート」 3.非機能要件定義 3.12.テストに関する事項 3.17.保守に関する事項	デジタル「テキスト生成AI活用におけるリスクへの対策ガイドブック（alpha）」	本項目における具体的な品質評価の実施は、主として事業者が実施すべき内容となりますが、AWSは、以下の品質評価機能を提供しています： 1. モデル評価基盤 Amazon Bedrockのモデル評価機能では、事前に定義された評価基準に基づいてモデルの性能を評価できます。例えば、回答の正確性、有害なコンテンツの生成有無、プロンプトインジェクションへの耐性などを評価することが可能です。また、RAG評価機能を使用することで、検索結果の関連性、生成された回答の品質、および回答と検索結果の一貫性を自動的に評価できます。さらに、ユーザーが独自の評価データセットを使用して、特定のユースケースに対するモデルの適合性を評価することもできます。 なお、以下の事項については、事業者が実施する必要があります： - 具体的な品質要件の設定 - ユースケースに応じたテストケースの設計と評価の実施 - 期待品質の具体的な評価基準の設定 - 定期的な品質評価の実施 (参考) ・Amazon Bedrock の評価：https://aws.amazon.com/jp/bedrock/evaluations/ ・Amazon Bedrock で RAG 評価のサポートを開始（一般提供）：https://aws.amazon.com/jp/about-aws/whats-new/2025/03/amazon-bedrock-rag-evaluation-generally-available/

調達チェックシート						(参考) 要求事項の参考										AWSサンプル回答
分類	評価 観点 #	評価観点	評価・選定時の 項目の分類	要求 事項 #	要求事項	対策例の対象AI					対策例	対策例詳細	裏付けとなる情報の例	(参考) 「DS-120 デジタル・ガバメント推進標準ガイドライン実 践ガイドブック」の「各種テンプレート」との対応関係	(参考) 対策例・対策例詳細の参考文献	
						入力		出力								
						テキス トを含 む	音声を含 む	テキス トを含 む	画像を含 む	音声を含 む						
						●	●	●	●	●	生成AIシステムの期待品質を満たしているか、テストケースを用いて評価する。その際テストケースは複数提案する（バッチ処理、情報検索、自然言語からのプログラム作成等）	プロンプトを入力した生成AIシステムの出力を検査することで、生成AIシステムの明示された要求仕様が規定する機能を生成AIモデルが実現可能かを調べている。プロンプトと生成AIモデルの組み合わせを対象とした試行錯誤的な繰り返し（一種のプロトタイプ開発スタイル）を通して、機能要求満足性を評価する	バッチ処理で生成AIを用いる場合、テストケース例とテスト方針がわかる資料等		対策例：デジタル庁「テキスト生成AI活用におけるリスクへの対策ガイドブック（α版）」 対策例詳細：産総研「生成AI品質マネジメントガイドライン 第1版」	
						●	●	●	●	●	マーケティング等で実施されるアンケートによる定性的な満足度評価手法を生成AIシステムの機能要求満足性の評価に用いて、機能要求満足性を満たしているか検証する		マーケティング等で実施されるアンケートによる定性的な満足度評価手法を生成AIシステムの機能要求満足性の評価に用いて、機能要求満足性を満たしているかの検証方針がわかる資料等		産総研「生成AI品質マネジメントガイドライン 第1版」	
						●	●	●	●	●	生成AIシステム出力のサンプリング手法やその「パラメータ」を指定できる場合、システムプロンプトの設計のための試行錯誤の中で合わせて検討する		生成AIシステム出力のサンプリング手法やその「パラメータ」を指定できる場合、システムプロンプトの設計と合わせて指定する開発方針がわかる資料等		産総研「生成AI品質マネジメントガイドライン 第1版」	
						●	●	●	●	●	生成AIシステムの機能要求水準を満たすのに十分な検索精度を確保するよう対策を講じる（多くの場合、想定される問題領域のなるべく多くの範囲で、一定以上の精度が出ることがよいとされる） ※RAGを利用した生成AIシステムのみ本項目の対象		生成AIシステムの機能要求水準を満たすのに十分な検索精度を確保するための対策方針がわかる資料等		産総研「生成AI品質マネジメントガイドライン 第1版」	

調達チェックシート						(参考) 要求事項の参考										AWSサンプル回答
分類	評価 観点 #	評価観点	評価・選定時の 項目の分類	要求 事項 #	要求事項	対策例の対象AI					対策例	対策例詳細	裏付けとなる情報の例	(参考) 「DS-120 デジタル・ガバメント推進標準ガイドライン実 践ガイドブック」の「各種テンプレート」との対応関係	(参考) 対策例・対策例詳細の参考文献	
						入力		出力								
						テキス トを含 む	音声を含 む	テキス トを含 む	画像を含 む	音声を含 む						
	10	文化的・言語的考慮	任意追加項目 (加点項目)	14	生成AIシステムが日本の言語環境や文化 環境に即した出力が可能であること	●	●	●	●	●	日本の言語環境や文化環境を踏まえたアウトプットを出力可能な生成AIシステムを提供する		日本の言語環境や文化環境に応じた出力結果となっているかの確認方針がわかる資料、開発時に学習データ等で日本の言語環境や文化環境に即したものととなるよう、取組がなされていることがわかる資料、日本語能力を測るベンチマークを活用した評価結果がわかる資料等 ※広く認知された一定の信頼性のある評価基準の使用を否定するものではない	「要件定義書標準テンプレート」 3.非機能要件定義 3.1.ユーザビリティ及びアクセシビリティに関する事項		GenUIでは、Amazon Bedrockを通じて、以下の機能を提供しています： 1. 日本語対応モデル Amazon NovaやAnthropic Claude等、多言語対応モデルを提供しており、日本語での入出力が可能です。 2. 日本語出力の評価 Amazon Bedrock Model Evaluationを使用することで、日本語での出力品質を評価することが可能です。これにより、日本の言語環境に適したアウトプットが得られているかを確認することができます。 なお、特定の業界や組織固有の表現への対応、具体的な日本語評価基準の設定については、事業者側での対応が必要となります。 (参考) ・Amazon Bedrock：https://aws.amazon.com/jp/bedrock/ ・Amazon Bedrock の評価：https://aws.amazon.com/jp/bedrock/evaluations/
	11	環境への配慮	任意追加項目 (加点項目)	15	環境に配慮した生成AIシステムを開発・提供すること	●	●	●	●	●	電力効率が高い半導体を使用したハードを基に学習されたモデルを利用する等の環境への配慮をした生成AIシステムを提供する又は利用する	システムテストに相当するソフトウェアテストで、システムが所与の機能仕様を適切に実現していることを検査し時間効率性及び資源効率性を評価する	消費電力等環境性能が優れていることがわかる資料等	「調達仕様書標準テンプレート」 4.作業の実施内容に関する事項 4.17.その他	対策例詳細：産総研「生成AI品質マネジメントガイドライン 第1版」	GenUIは、AWSのクラウドインフラストラクチャ上で運用されています。AWSでは環境への配慮に関して、以下の取り組みを行っています： 1. エネルギー効率の高いデータセンター AWSのデータセンターは、一般的なオンプレミスの最大 4.1 倍の効率を実現し、ワークロードの二酸化炭素排出量を最大 99 % 削減できると推定されています。 2. 再生可能エネルギーの利用 2023 年に、業務全体で消費される電力の 100% を再生可能エネルギーで賄うという目標を達成しています。 3. 深層学習 (DL) および生成 AI 推論アプリケーション向けに自社設計された 第2世代 AWS Inferentia チップを採用した Inf2 インスタンス、他の同等の Amazon EC2 インスタンスに比べて、ワットあたりのパフォーマンスが最大 50% 向上します。Inferentia2 チップは、高度なシリコンプロセスとハードウェアとソフトウェアの最適化を使用して、DL モデルを大規模に実行する際に高いエネルギー効率を実現します。 4. 16 億の AWS Trainium2 チップを搭載した Amazon EC2 Trn2 インスタンスは、生成 AI 専用に構築されており、数千億から数兆以上のパラメータを持つモデルのトレーニングとデプロイのためにハイパフォーマンス EC2 インスタンスを提供します。Trn2 インスタンスは Trn1 インスタンスよりも 3 倍高いエネルギー効率があります。 5. Bedrockを利用するにあたって、ファインチューニングをすることで用途に応じてより小型のモデルや効率的なモデルを選択・適用しやすくなるため、推論時の計算量や消費電力の削減につながります。また、基置モデルの利用により、パラメータ数の少ない小規模モデルでも高い性能を実現でき、計算リソースと消費電力を削減できます。 (参考) ・エネルギー転換：https://aws.amazon.com/jp/energy-utilities/sustainability/ ・Inf2インスタンス：https://aws.amazon.com/jp/ec2/instance-types/inf2/ ・Trn2インスタンス：https://aws.amazon.com/jp/ec2/instance-types/trn2/ ・Amazon Bedrock でファインチューニングまたは継続的な事前トレーニングを使用してモデルをカスタマイズする： https://docs.aws.amazon.com/ja_jp/bedrock/latest/userguide/custom-model-fine-tuning.html ・Amazon Bedrock Model Distillation：https://aws.amazon.com/jp/bedrock/model-distillation/
生成AIシステムの基本機能要件	12	有害情報の出力制御	基本項目	16	テロや犯罪に関する情報や攻撃的な表現等、有害情報の出力制限等もしていること	●	●	●	●	●	有害情報を入力あるいは想定出力を含むテストデータを入力した際、生成AIシステムの出力に当該情報が含まれない、若しくは出力を拒否できるような仕組みを具備する	サイバー攻撃やテロ等の犯罪、CBRN (Chemical, Biological, Radiological, Nuclear) に利用され得る情報を入力あるいは想定出力を含むテストデータを入力した際、生成AIシステムの出力に有害情報が含まれない、若しくは出力を拒否することができるような仕組みを具備する	サイバー攻撃やテロ等の犯罪、CBRNに利用され得る情報を入力あるいは想定出力を含むテストデータを入力した際、生成AIシステムの出力に有害情報が含まれない、若しくは出力を拒否することができるような仕組みの実現方法がわかる資料等	「要件定義書標準テンプレート」 3.非機能要件定義 3.1.ユーザビリティ及びアクセシビリティに関する事項	AISIJ AIセーフティに関する評価観点ガイド (第1.10版) J	GenUI では、Amazon Bedrock Guardrails により、有害情報・差別的表現・攻撃的表現等の有害コンテンツをフィルタリングできます。サイバー攻撃・テロ・CBRN に利用され得る情報を含む入力に対しても、出力への流入防止や出力拒否が可能です。フィルタリングルールはカスタマイズ可能で、リアルタイムでの出力検証・制御を行います： なお、具体的なフィルタリングルールの設定や業務特性に応じた追加制御は事業者側での対応が必要です。 (参考) ・Amazon Bedrock ガードレールを使用して有害なコンテンツを検出してフィルタリングする：https://docs.aws.amazon.com/ja_jp/bedrock/latest/userguide/guardrails.html ・プロンプトインジェクションのセキュリティ：https://docs.aws.amazon.com/ja_jp/bedrock/latest/userguide/prompt-injection.html
											差別表現等ユーザーが精神的な被害を受け得る情報を入力あるいは想定出力を含むテストデータを入力した際、生成AIシステムの出力に有害情報が含まれない、若しくは出力を拒否することができるような仕組みを具備する	差別表現等ユーザーが精神的な被害を受け得る情報を入力あるいは想定出力を含むテストデータを入力した際、生成AIシステムの出力に有害情報が含まれない、若しくは出力を拒否することができるような仕組みの実現方法がわかる資料等			AISIJ AIセーフティに関する評価観点ガイド (第1.10版) J	
						●	●	●	●	●	生成AIシステムの出力の有害性スコア (攻撃的であるかどうか等の有害さを数値で表したもの) の測定及び評価をする	生成AIシステムの出力の有害性スコア及びスコアの測定方法がわかる資料等			AISIJ AIセーフティに関する評価観点ガイド (第1.10版) J	
						●		●	●	●	テキストの入力に、虚偽の情報や意味不明な内容、複雑な内容が含まれる場合、生成AIシステムの挙動が不安定になることがある。そのような入力に対しても、有害な出力を抑制できるような仕組みを具備する	画像やテキストの入力に、虚偽の情報や意味不明な内容、複雑な内容が含まれる場合、生成AIシステムの挙動が不安定になることがある。そのような入力に対しても、有害な出力を抑制できるような仕組みの実現方法がわかる資料等			AISIJ AIセーフティに関する評価観点ガイド (第1.10版) J	
	13	偽談情報の出力・誘導の防止	基本項目	17	偽談情報の出力の防止措置を取っていること	●	●	●	●	●	調達対象の生成AIシステムにおいて利用予定の生成AIモデルと異なる生成AIモデルに対し、事業を問う内容の同一のテストデータを入力した際、生成AIシステムの出力と意味的に同様の内容が出力されるか検証を行う		調達対象の生成AIシステムにおいて利用予定の生成AIモデルと異なる生成AIモデルに対し、事業を問う内容の同一のテストデータを入力した結果が記載された資料等	「要件定義書標準テンプレート」 3.非機能要件定義 3.1.ユーザビリティ及びアクセシビリティに関する事項	AISIJ AIセーフティに関する評価観点ガイド (第1.10版) J	GenUIでは、Amazon Bedrockを通じて、以下の偽談情報出力防止機能を提供しています： 1. Bedrock Guardrailsによる制御機能 事業を問う内容に対する出力の一貫性を確認することができます。また、レビュー版の機能として、数値的に厳密な自動推論チェックにより、生成AIモデルの出力におけるハルシネーションの事実誤認を防止することも可能です。 2. Model Evaluationによる検証 Amazon Bedrock Model Evaluationを使用することで、異なるモデル間での回答の整合性チェックや、Ground Truth (正解データ) との照合により、出力内容の正確性を評価することが可能です。 3. RAG (Retrieval-Augmented Generation) 機能 Amazon Bedrock Knowledge Baseを利用することで、外部データソースとの連携による事実確認が可能です。検索結果と生成内容の関連性スコアの提供、一貫性チェック、さらに正解データと検索結果の意味的な一貫性の評価を行うことができます。 なお、以下の事項については、事業者側での対応が必要となります： ・業務特性に応じた信頼できる情報源の選定 ・RAGで使用する正解データの準備と検証 ・具体的な評価基準の設定 (参考) ・Amazon Bedrock のガードレール：https://aws.amazon.com/jp/bedrock/guardrails/ ・数式的に正しい自動推論チェックにより、生成AIモデルのハルシネーションによる事実誤差を防ぐ (レビュー)：https://aws.amazon.com/jp/blogs/news/prevent-factual-errors-from-llm-hallucinations-with-mathematically-sound-automated-reasoning-checks-preview/ ・Amazon Bedrock の新しい RAG 評価機能と LLM-as-a-Judge 機能：https://aws.amazon.com/jp/blogs/news/new-rag-evaluation-and-llm-as-a-judge-capabilities-in-amazon-bedrock/
						●	●	●	●	●	偽談情報の出力防止に関する評価項目のスコアを測定した結果、スコアに問題がないことを確認する	生成AIシステムに対するプロンプトと出力データの関連性に係るスコアの測定及び評価をする	生成AIシステムに対するプロンプトと出力データの関連性に係るスコア及びスコアの測定方法がわかる資料等			AISIJ AIセーフティに関する評価観点ガイド (第1.10版) J
											生成AIシステムの出力データとRAGによる検索結果の関連性に係るスコアの測定及び評価をする ※RAGを利用した生成AIシステムのみ本項目の対象	生成AIシステムの出力データとRAGによる検索結果の関連性に係るスコア及びスコアの測定方法がわかる資料等			AISIJ AIセーフティに関する評価観点ガイド (第1.10版) J	
											生成AIシステムの出力データとRAGによる検索結果の一貫性に係るスコアの測定及び評価をする ※RAGを利用した生成AIシステムのみ本項目の対象	生成AIシステムの出力データとRAGによる検索結果の一貫性に係るスコア及びスコアの測定方法がわかる資料等			AISIJ AIセーフティに関する評価観点ガイド (第1.10版) J	
											正解データとRAGによる検索結果の意味的な一貫性に係るスコアの測定及び評価をする ※RAGを利用した生成AIシステムのみ本項目の対象	正解データとRAGによる検索結果の意味的な一貫性に係るスコア及びスコアの測定方法がわかる資料等			AISIJ AIセーフティに関する評価観点ガイド (第1.10版) J	
											出力にコンテンツ認証がなされる生成AIシステムにおいて、様々なプロンプトを入力した場合でも適切にコンテンツ認証がなされる仕組みを具備する	生成AIシステムへ入力した様々なプロンプトに対し、適切なコンテンツ認証がなされることができるよう仕組みの実現方法がわかる資料等			AISIJ AIセーフティに関する評価観点ガイド (第1.10版) J	
											RAGを利用した生成AIシステムで検索対象となるドキュメントや、評価のために評価者が用意する正解データが、事実に基づいたデータであることを確認する ※RAGを利用した生成AIシステムのみ本項目の対象	RAGを利用した生成AIシステムで検索対象となるドキュメントや、評価のために評価者が用意する正解データが、事実に基づいたデータであることを確認した記録等			AISIJ AIセーフティに関する評価観点ガイド (第1.10版) J	

調達チェックシート					(参考) 要求事項の参考										AWSサンプル回答	
分類	評価 観点 #	評価観点	評価・選定時の 項目の分類	要求 事項 #	要求事項	対策例の対象AI					対策例	対策例詳細	裏付けとなる情報の例	(参考) 「DS-120 デジタル・ガバメント推進標準ガイドライン実 践ガイドブック」の「各種テンプレート」との対応関係		(参考) 対策例・対策例詳細の参考文献
						入力		出力								
						テキス トを含 む	音声を含 む	テキス トを含 む	画像を含 む	音声を含 む						
			不特定外部者 (一般国民 等) による府省 庁外利用の場 合は基本項目と して適用	18	ユーザーに対し、生成AIシステムの出力を 人間から発せられた情報と区別できるよ うにすること	●	●	●	●	●	ユーザーに対し、生成AIシステムの出力を人間から発せ られた情報と区別できるような仕組みを具備する		ユーザーに対し、生成AIシステムの出力を人間から 発せられた情報と区別できるような仕組みの実現方 法がわかる資料等	「要件定義書標準テンプレート」 3.非機能要件定義 3.1.ユーザビリティ及びアクセシビリティに関する事項	AISIF AIセーフティに関する評価観点ガイド（第1.10 版）」	GenUでは、以下の機能を通じて生成AIシステムの出力を人間から発せられた情報と区別可能にしています： 1. 識別表示の実装 AWS LambdaとAPI Gatewayを活用することで、出力の冒頭または末尾に「This is an AI-generated response」などの文言を自動的に追加したり、カスタマイズ可能な免責事項や注意書きを挿 入したりすることができます。 なお、以下の事項については、事業者側での対応が必要となります： - ユーザーインターフェースでのAI生成表示の具体的なデザイン - 組織のポリシーに基づいたAI生成コンテンツの説明文や免責事項の作成 (参考) ・API Gateway で REST API のステージ変数をセットアップする： https://docs.aws.amazon.com/ja_jp/apigateway/latest/developerguide/how-to-set-stage-variables-aws-console.html
			不特定外部者 (一般国民 等) による府省 庁外利用の場 合は基本項目と して適用	19	生成AIシステムによる、ユーザーの意思決 定の不当な誘導を防止するよう制御等をし ていること	●	●	●	●	●	生成AIシステムの推奨や指示により、ユーザーが主観的 又は客観的に不利益となる行動に誘導されることがない よう対策を講じる		生成AIシステムの推奨や指示により、ユーザーが主 観的又は客観的に不利益となる行動に誘導される 出力結果になっていないかの確認方法がわかる資 料等	「要件定義書標準テンプレート」 3.非機能要件定義 3.1.ユーザビリティ及びアクセシビリティに関する事項	AISIF AIセーフティに関する評価観点ガイド（第1.10 版）」	GenUでは、Amazon Bedrockを通じて、エンドユーザーへの不当な誘導を防止するための以下の機能を提供しています： 1. Bedrock Guardrailsによる新御 Guardrailを使用することで、誘導的な表現の検出と制御、過度な推奨や指示の制限、商業的な誘導の制御を行うことができます。また、URLやリンクの出力についても制御することが可能です。 2. Lambda/API Gatewayによる実装 AWS LambdaとAPI Gatewayを活用することで、生成AI出力であることの明示や、推奨・提案を含む場合の警告表示、オプトイン/アウト機能の実装が可能です。 なお、以下の事項については、事業者側での対応が必要となります： - 具体的な誘導防止ルールの設定 - ユーザーインターフェースでのオプトイン/アウト機能の実装 - 業務特性に応じた推奨・提案の制限ルールの設定 (参考) ・Amazon Bedrock のガードレール： https://aws.amazon.com/jp/bedrock/guardrails/ ・API Gateway で REST API のステージ変数をセットアップする： https://docs.aws.amazon.com/ja_jp/apigateway/latest/developerguide/how-to-set-stage-variables-aws-console.html
						●	●	●	●	●	生成AIシステムの出力において、生成AIシステムによる誘 導（URLリンク等を含むメッセージ生成等）のオプトイン 又はオプトアウト手段が出力に含まれる仕組みを具備す る		生成AIシステムの出力において、生成AIシステムに よる誘導のオプトイン又はオプトアウト手段が出力に 含まれる仕組みの実現方法がわかる資料等			

評価チェックシート					(参考) 要求事項の参考										AWSサンプル回答	
分類	評価 観点 #	評価観点	評価・選定時の 項目の分類	要求 事項 #	要求事項	対策例の対象AI					対策例	対策例詳細	裏付けとなる情報の例	(参考) 「DS-120 デジタル・ガバメント推進標準ガイドライン実 践ガイドブック」の「各種テンプレート」との対応関係		(参考) 対策例・対策例詳細の参考文献
						入力		出力								
						テキス トを含 む	音声 を含む	テキス トを含 む	画像 を含む	音声 を含 む						
	14	公平性と包摂性	不特定外部者 （一般国民 等）による府省 庁外利用の場 合は基本項目 として適用	20	有害なバイアスや、不当な差別の出力制 限等をしていること	●	●	●	●	●	特定のグループ（人種、性別、国籍、年齢、政治的 信念、宗教等の多様な背景に関する）に対する有害なバ イアスの入力あるいは想定出力に含むテストデータを入 力した際、生成AIシステムが回答を拒否できる仕組みを 具備する	出力コンテンツに対するフィルター機能を導入し、安全 性を損なう情報、想定外のパーソナルデータ漏洩、著 作権者の権利侵害に相当する出力、あるいは想定外のバ イアスの出力を防ぐようにデザインする	特定のグループ（人種、性別、国籍、年齢、政治的 信念、宗教等の多様な背景に関する）に対する有害な バイアスの入力あるいは想定出力に含むテストデータを入 力した際、生成AIシステムが回答を拒否することができ るような仕組みの実現方法がわかる資料等	「要件定義書標準テンプレート」 3.非機能要件定義 3.1.ユーザビリティ及びアクセシビリティに関する事項	対策例： AISIif AIセーフティに関する評価観点ガイド（第 1.10版）」 対策例詳細： 産総研「生成AI品質マネジメントガイドラ イン 第1版」	GenUIでは、以下の機能を通じて、有害なバイアスや不当な差別を防止することが可能です： 1. Bedrock Guardrailsによるバイアス制御 Guardrailsを使用することで、人種、性別、国籍、年齢、政治的・宗教等の属性に基づく差別的表現の検出と制御が可能で す。また、ヘイトスピーチの検出や、ステレオタイプ的な表現の制御を行うことも、これらのポリシー設定は組織固有の 配慮事項や地域特性に応じてカスタマイズすることが可能です。 2. Model Evaluationによる評価機能 Amazon Bedrock Model Evaluationにより、特定グループに対する偏りの評価や出力の公平性評価、属性による影響度 の測定が可能です。また、定期的なバイアスチェックと評価結果の可視化を行うことができます。 3. 追加的な制御とモニタリング AWS Lambdaを活用することで、AIモデルへの入力前後でフィルタリング処理や事前定義された禁止語句、差別的パター ンを検出するカスタムロジックを実装できます。また、CloudWatchを用いて生成AIシステムの出力をリアルタイムでモ ニタリングし、潜在的に有害な推奨を特定、アラート機能による迅速な対応が可能です。 なお、以下の事項については、事業者側での対応が必要となります： - 業務特性に応じたバイアス評価基準の設定 - 具体的な差別防止ルールの実装 - 定期的なバイアス評価の実施 - 複数のモデルを使用する場合の比較評価と選択 (参考) ・Amazon Bedrock のガイドレー： https://aws.amazon.com/jp/bedrock/guardrails/ ・Amazon Bedrock の評価： https://aws.amazon.com/jp/bedrock/evaluations/ ・Amazon Bedrock エージェントがユーザーから引き出す情報を送信するように Lambda 関数を設定する： https://docs.aws.amazon.com/ja_jp/bedrock/latest/userguide/agents-lambda.html ・Amazon BedrockとAmazon CloudWatchの統合による生成系AIPラケーションのモニタリング： https://aws.amazon.com/jp/blogs/news/amazon-bedrock_and_amazon-cloudwatch_integration_for_genai/
●	●	●	●	●	出力が人の属性に影響されないと想定されるテストデータを入力した際、出力結果が属性による影響を受けないこ との対策を講じる	出力コンテンツに対するフィルター機能を導入し、安全性を損なう情報、想定外のパーソナルデータ漏洩、著作権者の権利侵害に相当する出力、あるいは想定外のバイアスの出力を防ぐようにデザインする	出力が人の属性に影響されないと想定されるテストデータを入力した際、出力結果が属性による影響を受けないことを確認しているテスト方針がわかる資料等									
●	●	●	●	●	特定のモデルの利用が主たる目的ではない場合、生成AIモデルごとに異なるテクノロジーやバイアスを持つことがあることを考慮し、複数の生成AIモデルを利用する	-	複数の生成AIモデルを利用するシステム設計であることがわかる資料等									
●	●	●	●	●	学習データ、モデルの学習過程において、生成AIシステムに不公正なバイアスがないかを検証する また、生成AIモデルが、法令等に基づき特定の分野における出力制限をかけられている可能性について、整理して提供する	利用目的に応じ、バイアスや出力制限が支障となるかを踏まえ、適切な生成AIを選択する	生成AIシステムの検証プロセス・検証方法が分かる資料、生成AIモデルの提供企業の所在国におけるAIに対する規律や当該生成AIモデルの生成結果のバイアスに対する評価結果等									
●	●	●	●	●	出力結果について有害なバイアスが含まれていないか定期的にモニタリングする	-	有害なバイアスが含まれていないか定期的にモニタリングする際の確認観点がわかる資料等									
●	●	-	●	-	特定のカテゴリ（例えば職業）に関わる人物画像の生成を求めた場合に、個人の属性（例えば性別、年齢、人種）に関する偏りのない出力を生成できる仕組みを具備する	-	特定のカテゴリ（例えば職業）に関わる人物画像の生成を求めた場合に、個人の属性（例えば性別、年齢、人種）に関する偏りのない出力を生成することができるような仕組みの実現方法がわかる資料等									
基本項目	21	出力が全てのユーザーにとって理解しやすいものとなるよう対策を講じていること	●	●	●	●	●	生成AIシステムの出力の流暢性スコア（出力が文法的に適切かを数値で示したもの）の測定及び評価をする	生成AIシステムの出力の流暢性スコア及びスコアの測定方針がわかる資料等	生成AIシステムの出力の流暢性スコア及びスコアの測定方針がわかる資料等	「要件定義書標準テンプレート」 3.非機能要件定義 3.1.ユーザビリティ及びアクセシビリティに関する事項	AISIif AIセーフティに関する評価観点ガイド（第1.10版）」 AISIif AIセーフティに関する評価観点ガイド（第1.10版）」 AISIif AIセーフティに関する評価観点ガイド（第1.10版）」	GenUIでは、Amazon Bedrockを通じて、出力の可読性を確保するための以下の機能を提供しています： 1. Bedrock Guardrailsによる品質制御 Guardrailsを使用することで、文法的な適切性の確保や文章構造の一貫性確保、専門用語や難解な表現の適切な使用制御が可能です。また、文法的に不適切な入力に対しても、適切な処理を行うことができます。 2. Model Evaluationによる評価 Amazon Bedrock Model Evaluationにより、精度、毒性、頑健性の観点からモデルの評価が可能です。ただし、流暢性スコアや読みやすさ、文法的な適切性に関する直接的な評価機能は含まれていません。 3. 出力パラメータの制御 Temperature（創造性）やTop P（出力の多様性）、最大トークン数の設定、Response Formatによる出力形式の指定など、出力の品質に影響を与えるパラメータを調整することができます。これにより、読みやすさや一貫性のある応答の生成が可能です。 なお、以下の事項については、事業者側での対応が必要となります： - 具体的な可読性基準の設定 - 業務要件に基づく品質評価基準の設定 - 定期的な出力品質の評価実施 - 流暢性や読みやすさを評価するための追加的な仕組みの構築 (参考) ・Amazon Bedrock のガイドレー： https://aws.amazon.com/jp/bedrock/guardrails/ ・Amazon Bedrock の評価： https://aws.amazon.com/jp/bedrock/evaluations/ ・Amazon Bedrock のよくある疑問（Amazon Bedrock でのモデル評価とは何ですか？） https://aws.amazon.com/jp/bedrock/faqs/ ・推論パラメータでレスポンスの生成に影響を与える： https://docs.aws.amazon.com/ja_jp/bedrock/latest/userguide/inference-parameters.html			
			●	●	●	●	●	生成AIシステムの出力の読み易さに関するスコアの測定及び評価をする	生成AIシステムの出力の読み易さに関するスコア及びスコアの測定方針がわかる資料等	生成AIシステムの出力の読み易さに関するスコア及びスコアの測定方針がわかる資料等						
			●	●	●	●	●	生成AIシステムに文法的に不適切なテストデータを入力した場合でも、出力は文法的に適切であり、人間が理解しやすいものとなるよう対策を講じる	生成AIシステムに文法的に不適切なテストデータを入力した場合でも、出力は文法的に適切であり、人間が理解しやすいものとなっているかの確認方法がわかる資料等	生成AIシステムに文法的に不適切なテストデータを入力した場合でも、出力は文法的に適切であり、人間が理解しやすいものとなっているかの確認方法がわかる資料等						
			●	●	●	●	●	生成AIシステムに想定ユースケース以外の出力を想定したテストデータが入力された場合でも、ユーザーの生命・身体・財産等や各種の権利に危害を生じさせる出力がされない、あるいは、出力を拒否することができるような仕組みの実現方法がわかる資料等	生成AIシステムに想定ユースケース以外の出力を想定したテストデータが入力された場合でも、ユーザーの生命・身体・財産等や各種の権利に危害を生じさせる出力がされない、あるいは、出力を拒否することができるような仕組みの実現方法がわかる資料等	生成AIシステムに想定ユースケース以外の出力を想定したテストデータが入力された場合でも、ユーザーの生命・身体・財産等や各種の権利に危害を生じさせる出力がされない、あるいは、出力を拒否することができるような仕組みの実現方法がわかる資料等						
			●	●	●	●	●	生成AIシステムに想定ユースケース以外の出力を想定したテストデータが入力された場合でも、ユーザーの生命・身体・財産等や各種の権利に危害を生じさせる出力がされない、あるいは、出力を拒否することができるような仕組みの実現方法がわかる資料等	生成AIシステムに想定ユースケース以外の出力を想定したテストデータが入力された場合でも、ユーザーの生命・身体・財産等や各種の権利に危害を生じさせる出力がされない、あるいは、出力を拒否することができるような仕組みの実現方法がわかる資料等	生成AIシステムに想定ユースケース以外の出力を想定したテストデータが入力された場合でも、ユーザーの生命・身体・財産等や各種の権利に危害を生じさせる出力がされない、あるいは、出力を拒否することができるような仕組みの実現方法がわかる資料等						
			●	●	●	●	●	生成AIシステムに想定ユースケース以外の出力を想定したテストデータが入力された場合でも、ユーザーの生命・身体・財産等や各種の権利に危害を生じさせる出力がされない、あるいは、出力を拒否することができるような仕組みの実現方法がわかる資料等	生成AIシステムに想定ユースケース以外の出力を想定したテストデータが入力された場合でも、ユーザーの生命・身体・財産等や各種の権利に危害を生じさせる出力がされない、あるいは、出力を拒否することができるような仕組みの実現方法がわかる資料等	生成AIシステムに想定ユースケース以外の出力を想定したテストデータが入力された場合でも、ユーザーの生命・身体・財産等や各種の権利に危害を生じさせる出力がされない、あるいは、出力を拒否することができるような仕組みの実現方法がわかる資料等						
15	目的外利用への対応	基本項目	目的外利用の防止を行い、仮に目的外利用された場合にも危害・不利益を発生しにくくなるよう出力制限等をしていること	22	目的外利用の防止を行い、仮に目的外利用された場合にも危害・不利益を発生しにくくなるよう出力制限等をしていること	●	●	●	●	●	生成AIシステムに想定ユースケース以外の出力を想定したテストデータが入力された場合でも、ユーザーの生命・身体・財産等や各種の権利に危害を生じさせる出力がされない、あるいは、出力を拒否することができるような仕組みの実現方法がわかる資料等	生成AIシステムに想定ユースケース以外の出力を想定したテストデータが入力された場合でも、ユーザーの生命・身体・財産等や各種の権利に危害を生じさせる出力がされない、あるいは、出力を拒否することができるような仕組みの実現方法がわかる資料等	生成AIシステムに想定ユースケース以外の出力を想定したテストデータが入力された場合でも、ユーザーの生命・身体・財産等や各種の権利に危害を生じさせる出力がされない、あるいは、出力を拒否することができるような仕組みの実現方法がわかる資料等	「要件定義書標準テンプレート」 3.非機能要件定義 3.10.情報セキュリティに関する事項	AISIif AIセーフティに関する評価観点ガイド（第1.10版）」 AISIif AIセーフティに関する評価観点ガイド（第1.10版）」	GenUIは、Amazon Bedrockを通じて、目的外利用の防止と、仮に目的外利用された場合の危害・不利益の抑制のため、以下の機能を提供しています： 1. 入出力の制御による目的外利用の抑制 Amazon Bedrock Guardrailsにより、想定ユースケース以外の入力や、ユーザーの生命・身体・財産等や各種の権利に危害を生じさせる出力を検出・ブロックし、もしくは応答を拒否することができます。利用範囲の限定 IAMによるアクセス制御や、システムプロンプトによる役割・利用範囲の制約により、想定された目的の範囲内での利用に限定する設計が可能です。 3. 利活用状況のモニタリング モデル呼び出しのログ記録やAmazon CloudWatchにより、目的外と疑われる入出力を継続的に監視し、必要な対応につなげることができます。 なお、以下の事項については、事業者側での対応が必要となります： - 生成AIシステムの利用方法や利用上の留意点、想定利用環境、ガイドライン等の整備と、定期的な更新 - 業務特性に応じた想定ユースケースの定義と、拒否トピックや入出力制限ルールの設定 (参考) ・Amazon Bedrock のガイドレー： https://aws.amazon.com/jp/bedrock/guardrails/
						●	●	●	●	●	生成AIシステムに想定ユースケース以外の出力を想定したテストデータが入力された場合でも、ユーザーの生命・身体・財産等や各種の権利に危害を生じさせる出力がされない、あるいは、出力を拒否することができるような仕組みの実現方法がわかる資料等	生成AIシステムに想定ユースケース以外の出力を想定したテストデータが入力された場合でも、ユーザーの生命・身体・財産等や各種の権利に危害を生じさせる出力がされない、あるいは、出力を拒否することができるような仕組みの実現方法がわかる資料等	生成AIシステムに想定ユースケース以外の出力を想定したテストデータが入力された場合でも、ユーザーの生命・身体・財産等や各種の権利に危害を生じさせる出力がされない、あるいは、出力を拒否することができるような仕組みの実現方法がわかる資料等			

調達チェックシート						(参考) 要求事項の参考										AWSサンプル回答
分類	評価 観点 #	評価観点	評価・選定時の 項目の分類	要求 事項 #	要求事項	対策例の対象AI					対策例	対策例詳細	裏付けとなる情報の例	(参考) 「DS-120 デジタル・ガバメント推進標準ガイドライン実 践ガイドブック」の「各種テンプレート」との対応関係	(参考) 対策例・対策例詳細の参考文献	
						入力		出力								
						テキス トを含 む	音声を含 む	テキス トを含 む	画像を含 む	音声を含 む						
			企画・開発にお いて生成AIEデ ルに学習を行わ せる場合は基本 項目として適用	25	学習段階において、知的財産権等の保護 が図られるよう対策を講じていること	●	●	●	●	●	知的財産権等の侵害等防止のため、適切 なデータの学習を行う対策を講じる	外部から学習データの提供を受ける場合、 当該データ提供者との契約・利用規約にお いて、提供されるデータを生成AIへの学習 用データとして利用することが制限されて いる場合がある。許諾なく学習データとし て用いると契約・利用規約違反となる可 能性があるため、契約・利用規約を確認し 、データの利用制限の有無や内容をチェッ クする	データ提供者との契約・利用規約を満た すための学習を行う際の対応方針がわか る資料等	「要件定義書標準テンプレート」 2.機能要件定義 2.4.データに関する事項 3.非機能要件定義 3.12.データマネジメントに関する事項	経産省「コンテンツ制作のための生成AI利 活用ガイドブ ック」	GenUIの標準的な構成は、Amazon Bed rock上の基盤モデルをそのまま利用する ものであり、独自の追加学習（ファイン チューニング）を前提としていません。追 加学習やカスタムモデルを行う場合に備 え、Amazon Bedrockを通じて以下を提 供しています： 1. 学習データのリージョン内処理と非共有 Amazon Bedrockのカスタムモデル/ファ インチューニング機能では、学習データ はお客様のリージョン内で扱われ、ベース モデルの改変や他者への共有は行われ ません。 2. アクセス制御 IAMにより、学習データやカスタムモ デルへのアクセスを必要最小限の権限に 制御することができます。
												学習データをそのまま出力させよう学習 方法をとらない	学習データをそのまま出力させない学習 方法の実現方法がわかる資料等		文化庁「AIと著作権に関するチェックリス ト & ガイダンス」	なお、以下の事項については、事業者側 での対応が必要となります： - 外部から提供を受けるデータの契約・利 用規約上の、学習利用の可否の確認 - 著作権法第30条の4の「非享受目的」要 件、及び「著作権者の利益を不当に害す こととなる場合」への該当性の確認 - AI学習データの収集を制限する技術的措 置（robots.txt等）の尊重、及び学習デ ータをそのまま出力させない学習方法の 採用
												自らが権利を有しておらず、また、著作 権で保護された著作物等を学習データと して収集・利用する場合、① 契約の締結 等の権利者の許諾を得て行われているか 、又は② 権利者の許諾なく著作権法第30 条の4の適用を受けて行われているか確 認する。 ②の場合、学習データとしての収集・利 用が同条の「非享受目的」要件を満たす か、「著作権者の利益を不当に害すること となる場合」に該当しないかを確認する。 また、「非享受目的」のみの利用に留め、 享受目的あるいは享受目的が併存すると 評価される利用はせず、「著作権者の利益 を不当に害することとなる場合」に該当 するような利用はしない。 ③に関して、例えば以下のような取組が 考えられる ・学習データには、著作権者から許諾を 得る等権利処理されたデータ（ライセン スの契約等）を利用する ②に関して、「著作権者の利益を不当に害 することとなる場合」に該当しないよう、 例えば以下のような取組が考えられる ・AI学習データの収集を制限する技術的 措置（ID/PWによるアクセス制限や “ robots.txt”による制限等）を尊重する	学習データの収集・利用等における権利 者の許諾の状況の確認と、許諾を得る際 の対応方針がわかる資料等 又は、行おうとする学習データの収集・ 利用等で著作権法30条の4の適用を受け ようとするための確認方針と、その上で 学習データの利用が「非享受目的」のみに 利用する仕組みの実現方法がわかる資料 等		文化庁「AIと著作権に関するチェックリス ト & ガイダンス」 経産省「コンテンツ制作のための生成AI利 活用ガイドブ ック」	(参考) ・Amazon Bedrock でファインチューニ ングまたは継続的な事前トレーニングを 使用してモデルをカスタマイズする： https://docs.aws.amazon.com/ja_jp/bed rock/latest/userguide/custom-model-fine- tuning.html ・Amazon Bedrock のためのアイデンティ ティアクセス管理：https://docs.aws.am azon.com/ja_jp/bedrock/latest/userguide/ security-iam.html
17	セキュリティ確保	基本項目	26	生成AIシステム全体の脆弱性対策とし、 不正操作による影響の防止措置を取って いること	●	●	●	●	●	●	生成AIシステム全体に関するセキュリ ティ確保の対策を講 じる	生成AIシステムに対する攻撃手法は日々 新たなものが生まれており、これらのリス クに対応するための仕組みを検討・開発し ている、また必要に応じて導入することが 可能である	最新の攻撃手法に対する対策についての 資料等	「要件定義書標準テンプレート」 3.非機能要件定義 3.4.性能に関する事項 3.5.信頼性に関する事項 3.10.情報セキュリティに関する事項	AISIF AIセーフティに関する評価観点 ガイド（第1.10版）」	GenUIでは、以下のセキュリティ対策機 能を提供しています： 1. システム負荷対策 Amazon Bedrockでは、APIリクエストに 対するサービスクォータやモデルごとの 入出力トークン数制限が設定されていま す。また、プロビジョンドスループットを 使用することで専用のキャパシティを確 保し、安定した処理能力を維持すること が可能でです。 2. プロンプトインジェクション対策 Bedrock Guardrailsにより、悪意のある 入力の検出と制御、システムプロンプト の保護、制御回避の試みの防止が可能で す。また、禁止リストによるプロンプト フィルタリングやパターンマッチングによ る不正入力の検出を、リアルタイムで実 施できます。 3. 最新の脅威への対応 Guardrailsの継続的な更新により、プロ ンプトインジェクションなどの既知の攻 撃手法への対策を強化しています。また 、AWS Security Bulletinを通じて最新 の脆弱性情報を提供し、必要なセキュ リティアップデートを通知します。 4. RAG利用時のセキュリティ対策 Knowledge Base利用時には、IAMロー ールによる細かなアクセス権限の設定や Vector Store接続時の認証管理により、 要機密情報の不適切な出力を防止でき ます。また、Bedrock Agents利用時 には、IAMサービスロールによるアクセ ス制御やLambda関数による追加的な権 限制御が可能です。 なお、セキュリティポリシーの策定、業 務要件に応じた入力制限の設定、インシ デント対応手順の整備は事業者側での 対応が必要です（総務省「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」も参照）。
												生成AIEデルが意図しない出力を行わな いよう、安全基準を事後学習させる	生成AIEデルが意図しない出力を行わな いよう、安全基準を事後学習させるため の開発方針とその方法が分かる資料等		対策例詳細：総務省「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」（LLM及びLLMを構成要素に 含むAIシステムが対象）	(参考) ・Amazon Bedrock のクォータ：https ://docs.aws.amazon.com/ja_jp/bedrock/ latest/userguide/quotas.html ・Amazon Bedrock のガイドルール：ht tps://aws.amazon.com/jp/bedrock/guar drails/ ・Amazon Bedrock のプロビジョンドス ループットでモデル呼び出し容量を増や す：https://docs.aws.amazon.com/ja_j p/bedrock/latest/userguide/prov-through put.html ・Amazon Bedrock ナレッジベースのサー ビスロールを作成する：https://docs. aws.amazon.com/ja_jp/bedrock/latest/ userguide/kb-permissions.html
												生成AIEデルが従うべき指示の優先度を 定義し、優先度の高い指示（例：システ ムプロンプト）を常に優先的に処理する よう、生成AIEデルを事後学習させる	生成AIEデルが従うべき指示の優先度を 定義し、優先度の高い指示を常に優先的 に処理するようにするための開発方針と その方法がわかる資料等		対策例詳細：総務省「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」（LLM及びLLMを構成要素に 含むAIシステムが対象）	
												生成AIシステムの意図しない動作の防止 のために、オーストラレータの権限等に 不備がないことの対策を講じる。生成A IEデルや連携システムを操作するオース トラレータに係る権限も必要最小限とす ることで、生成AIEデルが攻撃を受けた 場合の被害拡大を抑制する	生成AIEデルや連携システムを操作する オーストラレータに係る権限設定等に不 備がないかの確認方法がわかる資料等		対策例詳細：総務省「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」（LLM及びLLMを構成要素に 含むAIシステムが対象）	
												入力された要機密情報を参照又は学習 する場合、権限を有する者以外の回答の 生成に、当該入力された要機密情報が 用いられないような仕組みとする方法 がわかる資料等	入力された要機密情報を参照又は学習 する場合、権限を有する者以外の回答の 生成に、当該入力された要機密情報が 用いられないような仕組みとする方法 がわかる資料等		-	
						●	●	●	●	●	生成AIEデルへのプロンプト入力に関 するセキュリティ確保の対策を講じる（ 入力プロンプト全般）	プロンプト内容の有害性を生成AIEデル への入力前に検知する等、生成AIEデル への入力前段階での防御策を講じる 例えば、以下のような対策を講じる ・禁止リストを活用したプロンプト検知 ・有害な結果をもたらす可能性のある コードの実行や生成を促す入力、又は 生成AIシステムの利用目的に照らして意 図しない出力を生成させる指示がプロ ンプトに含まれていないか検証し、そ のような入力を検知した場合にはプロ ンプトの一部削除による無害化や、処 理の拒否等の措置を講じる ・プロンプトインジェクション等によ り、生成AIシステムのバックエンドシ ステムに意図していない操作が行われ ることがないよう対策する ・実装した生成AIシステムの防御策に 関して、プロンプトインジェクション 等により防御策の回避が可能とならな いよう対策する	プロンプト内容の有害性を生成AIEデル への入力前に検知する等、生成AIEデル への入力前段階での防御策の実装方針 がわかる資料等		対策例詳細：AISIF AIセーフティに 関する評価観点ガイド（第1.10版）」 総務省「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」（LL M及びLLMを構成要素を含むAIシステ ムが対象）	
												生成AIシステムが処理困難でコストのか かる複数のテストデータを繰り返し入力 した場合でも、生成AIシステムの許容で きないリソース低下やシステムの停止が 発生しないよう対策を講じる	生成AIシステムが処理困難でコストのか かる複数のテストデータを繰り返し入力 した場合でも、生成AIシステムの許容で きないリソース低下やシステムの停止が 発生しないかの確認方法がわかる資料 等		対策例詳細：AISIF AIセーフティに 関する評価観点ガイド（第1.10版）」	
												生成AIシステムへの膨大なアクセスによ る攻撃を抑制するためのレートリミット を導入する	生成AIシステムのレートリミットの仕組 みの実装方法がわかる資料等		対策例詳細：総務省「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」（LLM及びLLMを構成要素に 含むAIシステムが対象）	
						●	●	●	●	●	生成AIEデルへのプロンプト入力に関 するセキュリティ確保の対策を講じる（ 特にユーザーが入力するプロンプト）	生成AIEデルがシステムプロンプトを無 視してユーザー入力のみに従うよう生 成AIEデルに指示する（いわゆるjail breakを試みる）ユーザー入力を無 効化するようシステムプロンプトをデ ザインする	jail breakを試みるユーザー入力を無 効化する仕組みの実現方法がわかる資 料等		対策例詳細：産総研「生成AI品質マ ネジメントガイドライン 第1版」	
												生成AIシステムの目的と用途に応じて、 ユーザーを確認して、認証を通らない ユーザーの利用を拒否することによっ て、ユーザー入力プロンプトの実装が 明らかになるよう制御する仕組みの実 現方法がわかる資料等	ユーザーを確認して、認証を通らない ユーザーの利用を拒否することによっ て、ユーザー入力プロンプトの実装が 明らかになるよう制御する仕組みの実 現方法がわかる資料等		対策例詳細：産総研「生成AI品質マ ネジメントガイドライン 第1版」	
						●	●	●	●	●	生成AIEデルへのプロンプト入力に関 するセキュリティ確保の対策を講じる（ システムプロンプト）	システムプロンプトに制約事項やセキュ リティ上の注意事項等を設定すること で、生成AIEデルが意図しない出力を 行わないようにする	システムプロンプトに制約事項やセキュ リティ上の注意事項等を設定すること で、生成AIEデルが意図しない出力を 行わないようにする仕組みを提供する 技術の実装方法がわかる資料等		対策例詳細：総務省「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」（LLM及びLLMを構成要素に 含むAIシステムが対象）	
												システムプロンプトに出力を意図しない 機密情報（例：APIキー）等を直接記 述することを避け、生成AIEデルが必要 に応じて参照できるよう別個に管理す る	システムプロンプトには、出力を意図し ない機密情報等を直接記述することと 避け、生成AIEデルが必要に応じて参 照できるよう別個に管理する対策の対 応方針がわかる資料等		対策例詳細：総務省「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」（LLM及びLLMを構成要素に 含むAIシステムが対象）	

調達チェックシート						(参考) 要求事項の参考										AWSサンプル回答
分類	評価観点 #	評価観点	評価・選定時の項目の分類	要求事項 #	要求事項	対策例の対象AI					対策例	対策例詳細	裏付けとなる情報の例	(参考) 「DS-120 デジタル・ガバナメント推進標準ガイドライン実践ガイドブック」の「各種テンプレート」との対応関係	(参考) 対策例・対策例詳細の参考文献	
						入力		出力								
						テキストを含む	音声を含む	テキストを含む	画像を含む	音声を含む						
						●	●	●	●	●	生成AIモデルの出力に関するセキュリティ確保の対策を講じる	生成AIシステムの意図しない動作の防止のための出力の検証機能を具備し、出力を意図しない情報が出力に含まれていないか検証し、検知した場合には応答を拒否する仕組みの実現方法がわかる資料等	出力を意図しない情報が出力に含まれていないか検証し、検知した場合には応答を拒否する仕組みの実現方法がわかる資料等		対策例詳細：総務省「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」（LLM及びLLMを構成要素を含むAIシステムが対象）	
						●	●	●	●	●	生成AIモデルが生成AIモデルの外のデータを参照する場合のセキュリティ確保の対策を講じる	RAGにて要機密情報のように参照権限が設定されている情報を利用する場合、RAG用のデータ及びデータストアへの参照権限をユーザーや役割に応じて適切に設定する（主にテキスト出力を念頭）	RAG用のデータ及びデータストアへの参照権限をユーザーや役割に応じて適切に設定し、不備がないかの確認方法がわかる資料等		対策例詳細：総務省「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」（LLM及びLLMを構成要素を含むAIシステムが対象）	
											外部サービスの呼び出しを行うためのプラグインツールを活用する場合、プラグインツールがアクセスする外部サービスやデータベース等の制御を実施する	プラグインツールがアクセスする外部サービスやデータベース等の制御の実現方法がわかる資料等	プラグインツールがアクセスする外部サービスやデータベース等の制御の実現方法がわかる資料等		産総研「生成AI品質マネジメントガイドライン 第1版」	
											外部サービスの呼び出しを行うためのプラグインツールを活用する場合、機能の導入前に当該プラグインツールの情報（SBOM 等）を確認し、外部連携としての利用可否を判断する	プラグインツールを導入する場合の、プラグインツールの情報の確認方針がわかる資料等	プラグインツールを導入する場合の、プラグインツールの情報の確認方針がわかる資料等		産総研「生成AI品質マネジメントガイドライン 第1版」	
											Webサイトや外部のデータベース等、外部データを参照する場合には、これらに意図しない出力を行わせる指示が含まれていないか検証し、そのような指示を検知した場合には、処理の拒否等の措置を講じる	Webサイトや外部のデータベース等、外部データを参照する場合には、これらに意図しない出力を行わせる指示が含まれていないか検証し、そのような指示を検知した場合には、処理の拒否等の措置を講じる等の仕組みの実現方法がわかる資料等	Webサイトや外部のデータベース等、外部データを参照する場合には、これらに意図しない出力を行わせる指示が含まれていないか検証し、そのような指示を検知した場合には、処理の拒否等の措置を講じる等の仕組みの実現方法がわかる資料等		総務省「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」（LLM及びLLMを構成要素を含むAIシステムが対象）	
											生成AIシステムに、入力フロントと外部参照データを明確に区分させ、外部参照データに高い注意を払わせる	生成AIシステムに、入力フロントと外部参照データを明確に区分させ、外部参照データに高い注意を払わせる仕組みの実現方法がわかる資料等	生成AIシステムに、入力フロントと外部参照データを明確に区分させ、外部参照データに高い注意を払わせる仕組みの実現方法がわかる資料等		総務省「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」（LLM及びLLMを構成要素を含むAIシステムが対象）	
			生成AIシステムに対して高い信頼性が求められる場合は基本項目として適用	27	情報セキュリティインシデント・生成AIシステム特有のリスク発生を検知した後、その影響を最小限に抑えつつ、迅速な封じ込めと復旧能力を確保するための対策を講じていること	●	●	●	●	●	情報セキュリティインシデント・生成AIシステム特有のリスク発生時の対応のために、生成AIシステムの制御性を強化するための仕組みを具備する	改ざんされたモデルの別バージョンや代替モデルへの即時切り替えを実施する	情報セキュリティインシデント・生成AIシステム特有のリスク発生時の対応方法として、改ざんされたモデルの別バージョンや代替モデルへ即時切替する仕組みの実現方法がわかる資料等	「要件定義書標準テンプレート」 3.非機能要件定義 3.4.性能に関する事項 3.5.信頼性に関する事項 3.10.情報セキュリティに関する事項 3.16.運用に関する事項 3.17.保守に関する事項	AISIF AIインシデントレスポンス・アプローチブック	GenUIは、Amazon Bedrockを通じて、情報セキュリティインシデント・生成AIシステム特有のリスクへの検知後における、迅速な封じ込めと復旧（制御性の強化）を支えるため、以下の機能を提供しています： 1. 代替モデル・別バージョンへの切替 APIレベルでのモデル指定により、改ざんや劣化が疑われるモデルから、別バージョンや代替モデルへ迅速に切り替える構成が可能です。また、クロスバージョン推論やプロビジョンスルーブットにより可用性を維持することができます。 2. 検知のためのログ・監視 AWS CloudTrailによるAPI呼び出しの記録、Amazon CloudWatchによる異常検知、モデル呼び出しのログ記録により、インシデントやリスクの兆候を検知することができます。 3. 復旧の容易性 GenUIはAWS CDKで構成されるため、設定のロールバックや再デプロイにより既知の正常な状態へ復旧する運用を設計しやすく、RAG利用時はKnowledge Baseのデータ再同期により汚染データの影響を限定することができます。 なお、以下の事項については、事業者側での対応が必要となります： - 検知後の対応手順（即時切替・隔離・通知）の整備と実行演習 - 異常時に影響を受けない上位権限を持つ、隔離された制御インターフェースの設置 - 混入した悪意ある学習データの局所的除去や追加学習等のモデル汚染への対応、及び情報セキュリティインシデント対応体制との連携
											混入された悪意のある学習データの局所的な除去や追加学習を実施する	情報セキュリティインシデント・生成AIシステム特有のリスク発生時の対応方法として、混入された悪意のある学習データの局所的な除去や追加学習等の、生成AIモデルの汚染対策の方針がわかる資料等	情報セキュリティインシデント・生成AIシステム特有のリスク発生時の対応を想定して、異常時に影響を受けない上位権限を持ち、隔離された制御インターフェースを設置する		AISIF AIインシデントレスポンス・アプローチブック	
											異常時に影響を受けない上位権限を持ち、隔離された制御インターフェースを設置する	情報セキュリティインシデント・生成AIシステム特有のリスク発生時の対応を想定して、異常時に影響を受けない上位権限を持ち、隔離された制御インターフェースの設置等の対策の実装方針がわかる資料等	情報セキュリティインシデント・生成AIシステム特有のリスク発生時の対応を想定して、異常時に影響を受けない上位権限を持ち、隔離された制御インターフェースの設置等の対策の実装方針がわかる資料等		AISIF AIインシデントレスポンス・アプローチブック	
			基本項目	28	生成AIシステムの開発の過程を通じて、適切にセキュリティ対策を講じていること	●	●	●	●	●	生成AIシステムの開発の過程を通じて、採用する技術の特性に照らして適切なセキュリティ対策を講じる	セキュリティ対策についての資料等	セキュリティ対策についての資料等	「要件定義書標準テンプレート」 3.非機能要件定義 3.4.性能に関する事項 3.5.信頼性に関する事項 3.10.情報セキュリティに関する事項	総務省、経産省「AI事業者ガイドライン 第1.2版」	AWSでは、開発プロセスでのセキュリティ対策として以下を提供しています： 1. セキュリティ実装のガイダンス OWASP Top 10 for LLM Applicationsに基づく実装ガイドを提供しています： - プロンプトインジェクション対策の実装方法 - データ漏洩対策の実装方法 - 認証とセッション管理の実装方法 - その他のLLMアプリケーション特有のセキュリティ考慮事項 なお、以下の事項については、事業者側での対応が必要となります： - セキュアな開発プロセスの確立 - OWASP Top 10 for LLMに基づくセキュリティ評価の実施 - 脆弱性評価とペネトレーションテストの実施 - インシデント対応プランの策定 （参考） ・AWS で実現する安全な生成 AI アプリケーション - OWASP Top 10 for LLM Applications 2025 の活用例：https://aws.amazon.com/jp/blogs/news/secure-gen-ai-applications-on-aws-refer-to-owasp-top-10-for-llm-applications/

