

Microsoft.AI-103-JPN.v2026-09-07.q66

試験コード：	AI-103-JPN
試験名称：	Developing AI Apps and Agents on Azure (AI-103日本語版)
認証ベンダー：	Microsoft
無料問題の数：	66
バージョン：	v2026-09-07
ページの閲覧量：	105
問題集の閲覧量：	674

<https://www.jpnsiken.com/shiken/Microsoft.AI-103-JPN.v2026-09-07.q66.html>

質問: 1

事例研究1 - コントソ社

概要

会社情報

Contoso, Ltdは、Microsoft Foundryを使用して生成型AIおよびエージェントベースのソリューションを構築、展開、管理する多国籍小売企業です。

既存の環境

アイデンティティ環境

Contosoは、エージェントが組織のリソースやサービスにアクセスできるようにするID管理、認証、認可機能にMicrosoft Entra IDを使用しています。

Contosoは最近、既存のAIソリューションを最適化および保守するために、Agent1Dev Teamという新しいAIエンジニアリングチームを結成しました。

チームは、ソリューションアーキテクト、DevOpsエンジニア、セキュリティエンジニアと協力して、AIアプリケーションの設計、実装、監視、およびセキュリティ対策を行います。

Contosoには、AIソリューションの導入前にその妥当性を検証する「Agent1Test Team」というチームも存在する。

生成環境

Contosoは、Project1とProject2という名前の2つのプロジェクトを含むMicrosoft Foundry環境を構築しています。

プロジェクト1

Project1には、顧客からの製品に関する問い合わせやトラブルシューティングの依頼に対応する、Agent1という名前のカスタマーサポートエージェントが含まれています。

Agent1には以下の設定があります。

- Agent1は基本モデルのデプロイメントを使用します。

安全性評価パイプラインは有効になっていません。

ツール呼び出し承認ワークフローは有効になっていません。

会話メモリの制約は設定されていません。

Agent1は、デジタルサポートチャネルを使用して顧客とやり取りし、Contoso製品に関する一般的な質問に回答します。

Project1は、欧州連合(EU)域内にあるAzureリージョンにデプロイされています。

Agent1開発チームは、Project1を使用してAgent1の最適化と保守を行います。

プロジェクト2

Project2には、既に展開済みの動画生成モデルが含まれています。Contoso社のマーケティング部門はProject2にアクセスでき、このモデルを用いて動画制作ソリューションを開発する予定です。

解決策の開発は未完了です。

データ環境

Contosoは、AIアプリケーションをサポートするAzureリソースに製品関連情報を保存します。

Azure環境には、Contoso製品の製品詳細シートを保存するstorage1という名前のAzure Blob Storageアカウントが含まれています。

製品仕様書には、仕様、機能説明、および製品サポート情報が記載されており、エージェント1はこれらを使用して顧客からの質問に回答できます。製品仕様書はPDF形式で保存されます。

問題提起

コントソ氏は以下の問題点を指摘している。

- Agent1はContoso製品に関する一般的な知識しか持っていません。
- Agent1との最近のチャットのやり取りについて、感情分析が行われました。分析結果はまだ処理されていません。
- Agent1は、顧客からの質問に回答する際に、storage1に保存されている製品シートの詳細な製品情報を使用しません。

コントソ社の財務部門によると、仕入先からの請求書は、契約書に定められた条件と一致していることを確認するため、手作業で精査する必要があるとのことです。請求書には表やロゴ、様々なレイアウトが含まれているため、一貫した処理が困難です。

要件

計画されている変更

コントソ氏は以下の変更を実施する予定です。

プロジェクト1向けに、仕入先請求書の視覚的なレイアウトとテキスト内容の両方を評価することで請求書を分析し、仕入先との契約条件と照合できるソリューションを実装する。

- Agent1で使用される基本モデルのデプロイメントを更新し、モデルのバージョンを標準化して、継続性と一貫した応答を確保します。
 - Agent1がstorage1に保存されている製品シートから詳細な製品情報を取得して使用できるようにします。
 - Agent1が顧客からの質問に回答するために使用できる製品シートのインデックス作成ソリューションを実装する。
- 動画制作ソリューションの開発を完了する。

技術要件

Contoso社は、以下の技術要件を特定しました。

- Agent1で使用されるモデル展開は、拡張性が高く、高スループットの生成型AIワークロードをサポートし、予約済みのスループット容量を必要とせずに、変動する顧客サポートトラフィックを処理するために動的に拡張する必要があります。
- 製品シートは、セマンティック検索とベクトル検索を可能にするインデックス作成パイプラインを使用して処理され、Agent1が関連する製品情報を取得できるようにする必要があります。

製品シートの情報に基づいて生成される回答は、関連性があり、完全かつ正確でなければなりません。

- エージェント1は、製品シートを使用して、製品の詳細に関する自然言語の質問に答えることができなければなりません。
- エージェント1が使用するモデルのバージョンは、安定した応答を確保するために一貫している必要があります。
- モデルによって処理されるデータは、EU域内に留まらなければならない。

セキュリティおよびコンプライアンス要件

Contosoは、以下のセキュリティおよびコンプライアンス要件を特定しています。

- APIキーを使用してFoundryにデプロイされたモデルにアクセスしてはなりません。

Azureリソースへのアクセスは、最小権限の原則に従う必要があります。

Contosoの開発者は、Microsoft Entra認証を使用してMicrosoft Foundryリソースへの認証を行う必要があります。

- Project1へのアクセス権は、SC_Agent1_Devという名前のセキュリティグループを使用して、Agent1Devチームのメンバーに割り当てる必要があります。
- Project1へのアクセス権は、SC_Agent1_Testという名前のセキュリティグループを使用して、Agent1Testチームのメンバーに割り当てる必要があります。
- Agent1は、顧客データを含む文書がストレージ1の製品シートリポジトリに誤って追加された場合でも、顧客情報を決して開示してはならない。

製品仕様書には、テキストが埋め込まれた画像が含まれている場合があります。Agent1は、画像内に隠されている可能性のある悪意のある命令から保護される必要があります。

ビジネス要件

Contosoは以下のビジネス要件を特定しました。

- Agent1とやり取りするユーザーは、今後のやり取りにおいてパーソナライズされた体験を得る必要があります、そのためにはAgent1が会話の文脈を保持し、過去のやり取りから関連情報を想起できる機能が必要です。

- エージェント1は、Contosoが販売する製品に関する質問のみに回答しなければならない。

Agent1は、Contoso製品に関する顧客からの質問にのみ回答するように設定する必要があります。

解決策はビジネス要件を満たす必要があります。あなたは何をすべきでしょうか？

A. システムメッセージの指示を修正します。

B. 少数の撮影例を追加する。

C. トップpサンプリングを適用する。

D. 温度パラメータの値を増加させます。

正解: ([正解を表示します](#))

Scenario: Contoso identifies the following business requirements:

Agent1 must answer questions only about the products sold by Contoso.

System Message Instructions: This dictates the agent's persona, boundaries, and rules. By explicitly configuring the system prompt (e.g., instructing the agent: "You are a product assistant.

Only answer questions using the provided product detail sheets. If you do not know the answer based on the provided documents, state that you do not know"), you prevent the model from answering with its general pre-trained knowledge.

Incorrect:

[Not B]

Few-shot examples: While useful for enforcing output formats or response tone, this does not explicitly stop a model from retrieving outside knowledge on unfamiliar topics.

[Not C, Not D]

Temperature & top-p: These parameters control the creativity and randomness of the model's responses. Decreasing them makes the model's outputs more deterministic and factual, but they do not actively restrict the model from generating information outside of the intended domain.

Reference:

<https://techcommunity.microsoft.com/blog/azure-ai-foundry-blog/announcing-safety-system-messages-in-azure-ai-studio-and-azure-openai-studio/4146991>

質問: 2

事例研究1 - Contoso, Ltd

概要

会社情報

Contoso, Ltdは、Microsoft Foundryを使用して生成型AIおよびエージェントベースのソリューションを構築、展開、管理する多国籍小売企業です。

既存の環境

アイデンティティ環境

Contosoは、エージェントが組織のリソースやサービスにアクセスできるようにするID管理、認証、認可機能にMicrosoft Entra IDを使用しています。

Contosoは最近、既存のAIソリューションを最適化および保守するために、Agent1Dev Teamという新しいAIエンジニアリングチームを結成しました。

チームは、ソリューションアーキテクト、DevOpsエンジニア、セキュリティエンジニアと協力して、AIアプリケーションの設計、実装、監視、およびセキュリティ確保を行います。

Contosoには、AIソリューションの導入前にその妥当性を検証する「Agent1Test Team」というチームも存在する。

生成環境

Contosoは、Project1とProject2という名前の2つのプロジェクトを含むMicrosoft Foundry環境を構築しています。

プロジェクト1

Project1には、顧客からの製品に関する問い合わせやトラブルシューティングの依頼に対応する、Agent1という名前のカスタマーサポートエージェントが含まれています。

Agent1には以下の設定があります。

- Agent1は基本モデルのデプロイメントを使用します。

安全性評価パイプラインは有効になっていません。

ツール呼び出し承認ワークフローは有効になっていません。

会話メモリの制約は設定されていません。

Agent1は、デジタルサポートチャネルを使用して顧客とやり取りし、Contoso製品に関する一般的な質問に回答します。

Project1は、欧州連合(EU)域内にあるAzure リージョンにデプロイされています。

Agent1開発チームは、Project1を使用してAgent1の最適化と保守を行います。

プロジェクト2

Project2には、既に展開済みの動画生成モデルが含まれています。Contoso社のマーケティング部門はProject2にアクセスでき、このモデルを用いて動画制作ソリューションを開発する予定です。

解決策の開発は未完了です。

データ環境

Contosoは、AIアプリケーションをサポートするAzure リソースに製品関連情報を保存します。

Azure環境には、Contoso製品の製品詳細シートを保存するstorage1という名前のAzure Blob Storageアカウントが含まれています。

製品仕様書には、仕様、機能説明、および製品サポート情報が記載されており、エージェント1はこれらを使用して顧客からの質問に回答できます。製品仕様書はPDF形式で保存されます。

問題提起

コントソ氏は以下の問題点を指摘している。

- Agent1はContoso製品に関する一般的な知識しか持っていません。

- Agent1との最近のチャットのやり取りについて、感情分析が行われました。分析結果はまだ処理されていません。

- Agent1は、顧客からの質問に回答する際に、storage1に保存されている製品シートの詳細な製品情報を使用しません。

コントソ社の財務部門によると、仕入先からの請求書は、契約書に定められた条件と一致していることを確認するため、手作業で精査する必要があるとのことです。請求書には表やロゴ、様々なレイアウトが含まれているため、一貫した処理が困難です。

要件

計画されている変更

コントソ氏は以下の変更を実施する予定です。

プロジェクト1向けに、仕入先請求書の視覚的なレイアウトとテキスト内容の両方を評価することで請求書を分析し、請求書の詳細が仕入先との契約条件と照合できるようにするソリューションを実装する。

- Agent1で使用される基本モデルのデプロイメントを更新し、モデルのバージョンを標準化して、継続性と一貫した応答を確保します。

- Agent1がstorage1に保存されている製品シートから詳細な製品情報を取得して使用できるようにします。

- Agent1が顧客からの質問に回答するために使用できる製品シートのインデックス作成ソリューションを実装する。

・動画制作ソリューションの開発を完了する。

技術要件

Contoso社は、以下の技術要件を特定しました。

- Agent1で使用するモデル展開は、拡張性が高く、高スループットの生成型AIワークロードをサポートし、予約済みのスループット容量を必要とせずに、変動する顧客サポートトラフィックを処理するために動的に拡張する必要があります。

- 製品シートは、セマンティック検索とベクトル検索を可能にするインデックス作成パイプラインを使用して処理され、Agent1が関連する製品情報を取得できるようにする必要があります。

製品シートの情報に基づいて生成される回答は、関連性があり、完全かつ正確でなければなりません。

- エージェント1は、製品シートを使用して、製品の詳細に関する自然言語の質問に答えることができません。

- エージェント1が使用するモデルのバージョンは、安定した応答を確保するために一貫している必要があります。

- モデルによって処理されるデータは、EU域内に留まらなければならない。

セキュリティおよびコンプライアンス要件

Contosoは、以下のセキュリティおよびコンプライアンス要件を特定しています。

- APIキーを使用してFoundryにデプロイされたモデルにアクセスしてはなりません。

Azureリソースへのアクセスは、最小権限の原則に従う必要があります。

Contosoの開発者は、Microsoft Entra認証を使用してMicrosoft Foundryリソースへの認証を行う必要があります。

- Project1へのアクセス権は、SC_Agent1_Devという名前のセキュリティグループを使用して、Agent1Devチームのメンバーに割り当てる必要があります。

- Project1へのアクセス権は、SC_Agent1_Testという名前のセキュリティグループを使用して、Agent1Testチームのメンバーに割り当てる必要があります。

- Agent1は、顧客データを含む文書がストレージ1の製品シートリポジトリに誤って追加された場合でも、顧客情報を決して開示してはならない。

製品仕様書には、テキストが埋め込まれた画像が含まれている場合があります。Agent1は、画像内に隠されている可能性のある悪意のある命令から保護される必要があります。

ビジネス要件

Contosoは以下のビジネス要件を特定しました。

- Agent1とやり取りするユーザーは、今後のやり取りにおいてパーソナライズされた体験を得る必要があります、そのためにはAgent1が会話の文脈を保持し、過去のやり取りから関連情報を想起できる機能が必要です。

- エージェント1は、Contosoが販売する製品に関する質問のみに回答しなければならない。

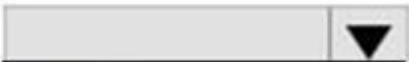
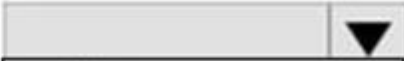
ホットスポットに関する質問

マーケティング部門がProject2に導入されたモデルを使用して動画を生成できることを確認する必要があります。

Pythonコードをどのように完成させるべきですか？回答するには、回答欄で適切なオプションを選択してください。

注：正解ごとに1ポイントが加算されます。

Answer Area

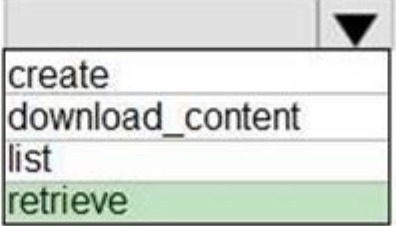
```
import time
from openai import OpenAI
client = OpenAI(
    base_url= "https://Contoso.openai.azure.com/openai/v1/",
    api_key=...,
)
video = client.videos. (
    model=deployment_name,
    prompt= "A video of our products",
)
while video.status not in ["completed", "failed", "cancelled"]:
    time.sleep(20)
    video = client.videos. (video.id)

    print(video.status)
```

The image shows a code editor with a Python script for interacting with the OpenAI API. The script imports the OpenAI module and creates a client. It then uses the client's videos endpoint to create a video, wait for it to complete, and print its status. Two dropdown menus are visible, one for the initial video creation and one for the subsequent status check. The dropdown menus are open, showing the following options: create, download_content, list, and retrieve. The Microsoft logo is visible in the background.

正解:

Answer Area Microsoft

```
import time
from openai import OpenAI
client = OpenAI(
    base_url="https://Contoso.openai.azure.com/openai/v1/",
    api_key=...,
)
video = client.videos. (
    model=deployment_name,
    prompt="A video of our products",
)
while video.status not in ["completed", "failed", "cancelled"]:
    time.sleep(20)
    video = client.videos. (video.id)

print(video.status)
```

Explanation:

Box 1: create

client.videos.create(...) initializes the generation job. This returns a video object containing a unique id and a status of "queued" or "in_progress".

Box 2: retrieve

client.videos.retrieve(video.id) refreshes the video object state from the server by fetching its updated completion progress and status until it finishes Reference:

<https://developers.openai.com/api/reference/python/resources/videos/methods/create>

質問: 3

お客様は、AI1 という名前の Azure AI Foundry インスタンスを含む Azure サブスクリプションをお持ちです。

システムログファイルに記録された問題を自動的にトリアージして解決するアプリをお持ちですね。

問題解決のために連携するインシデントマネージャーエージェントとDevOpsエージェントを作成します。

インシデントマネージャーエージェントがDevOpsエージェントに作業を割り当てられるようにする必要があります。ソリューションは開発作業を最小限に抑えるものでなければなりません。

あなたはどうすべきでしょうか？

- A. AI1のプロンプトフローを設定します。
- B. AI1 用の接続エージェントを設定します。
- C. インシデントマネージャーエージェント用のプラグインを作成します。
- D. DevOpsエージェント用のプラグインを作成します。

正解: **B** ([コメントを发表する](#))

Ensuring that an incident manager agent can assign work to a DevOps agent in Azure AI Foundry is primarily achieved through a Connected Agents configuration. This setup allows a "main" agent to delegate tasks to "specialized" agents via natural language or defined function calls.

Key Implementation Steps

1. Configure Connected Agents

Within the Azure AI Foundry portal, you must register the DevOps agent as a "Connected Agent" of the incident manager.

2. Define Tool-Based Delegation

Use the ConnectedAgentToolDefinition in the Azure Python SDK or C# SDK to programmatically link them. This exposes the DevOps agent to the incident manager as a callable "tool".

3. Implement Handoff Orchestration Patterns

Choose an orchestration pattern that fits the triage-to-resolve workflow.

4. Manage Context and State

5. Enable Permissions and Roles

Reference:

<https://learn.microsoft.com/en-us/azure/ai-foundry/agents/how-to/connected-agents>

質問: 4

Azure OpenAIリソースを含むAzureサブスクリプションをお持ちです。

Azure AI Agent Service を使用してエージェントを構築する予定です。このエージェントは、以下の動作を実行します。

ユーザーからの書面および口頭による質問を解釈する。

質問に対する回答を生成する。

回答を音声として出力します。

エージェント用のプロジェクトを作成する必要があります。

何をすべきでしょうか？

A. Azureポータル

B. ランゲージスタジオ

C. Azure AI Foundry

D. スピーチスタジオ

正解: ([正解を表示します](#))

Azure AI Foundry is a platform for designing, customizing, managing, and supporting AI applications and agents. It acts as an AI app factory, providing a unified environment with tools, models, and deployment pipelines for various AI tasks. It enables teams to build and operate AI solutions, including those powered by generative AI, while ensuring security, governance, and cost-efficiency.

Reference:

<https://learn.microsoft.com/en-us/azure/ai-foundry/what-is-azure-ai-foundry>

質問: 5

法務チームは、非構造化契約書から条項と推論フィールドを抽出する必要があります。必須フィールドの1つは契約終了日ですが、これは明示的に記載されていないため、開始日と契約期間から導き出す必要があります。チームにはラベル付きのトレーニングデータがありません。どのツールが最適でしょうか？

A. ラベル付きサンプルでトレーニングされた Azure AI Document Intelligence カスタムモデル

B. 事前構築済みコントラクトまたはゼロショットカスタムアナライザーを使用したAzure Content Understanding

C. Azure AI Searchのセマンティックランカー

D. Foundry Tools の Azure Translator

正解: **B** ([コメントを发表する](#))

Azure Content Understanding handles unstructured documents, works zero-shot without labelled data, and can infer fields that are not explicitly present, such as deriving a contract end date from the start date and term. That matches every requirement in the scenario.

質問: **6**

Azure AI Content Safety を構成して、生成アプリケーションでユーザーがアップロードした画像をモデレートするようにします。テスト中に、同僚が、ある画像が暴力の深刻度スコア 3 を返したと主張しました。

なぜこのような結果はあり得ないのか？

A. コンテンツセーフティは画像ではなくテキストのみを審査します。

B. 画像コンテンツモデルは、深刻度値0、2、4、6のみを返します。

C. 重症度3は自傷行為のカテゴリのみに適用されます

D. 画像モデレーションは常にブロックまたは許可のバイナリフラグを返し、深刻度は返しません。

正解: **B** ([コメントを发表する](#))

The Azure AI Content Safety image model uses a trimmed severity scale that returns only 0, 2, 4, and 6, so an odd value such as 3 cannot occur for image content. The text model uses the full 0 to 7 scale, which is where odd-numbered severities appear.

質問: **7**

事例研究1 - コントソ社

概要

会社情報

Contoso, Ltdは、Microsoft Foundryを使用して生成型AIおよびエージェントベースのソリューションを構築、展開、管理する多国籍小売企業です。

既存の環境

アイデンティティ環境

Contosoは、エージェントが組織のリソースやサービスにアクセスできるようにするID管理、認証、認可機能にMicrosoft Entra IDを使用しています。

Contosoは最近、既存のAIソリューションを最適化および保守するために、Agent1Dev Teamという新しいAIエンジニアリングチームを結成しました。

チームは、ソリューションアーキテクト、DevOpsエンジニア、セキュリティエンジニアと協力して、AIアプリケーションの設計、実装、監視、およびセキュリティ対策を行います。

Contosoには、AIソリューションの導入前にその妥当性を検証する「Agent1Test Team」というチームも存在する。

生成環境

Contosoは、Project1とProject2という名前の2つのプロジェクトを含むMicrosoft Foundry環境を構築しています。

プロジェクト1

Project1には、顧客からの製品に関する問い合わせやトラブルシューティングの依頼に対応する、Agent1という名前のカスタマーサポートエージェントが含まれています。

Agent1には以下の設定があります。

- Agent1は基本モデルのデプロイメントを使用します。

安全性評価パイプラインは有効になっていません。

ツール呼び出し承認ワークフローは有効になっていません。

会話メモリの制約は設定されていません。

Agent1は、デジタルサポートチャネルを使用して顧客とやり取りし、Contoso製品に関する一般的な質問に回答します。

Project1は、欧州連合(EU)域内にあるAzureリージョンにデプロイされています。

Agent1開発チームは、Project1を使用してAgent1の最適化と保守を行います。

プロジェクト2

Project2には、既に展開済みの動画生成モデルが含まれています。Contoso社のマーケティング部門はProject2にアクセスでき、このモデルを用いて動画制作ソリューションを開発する予定です。

解決策の開発は未完了です。

データ環境

Contosoは、AIアプリケーションをサポートするAzureリソースに製品関連情報を保存します。

Azure環境には、Contoso製品の製品詳細シートを保存するstorage1という名前のAzure Blob Storageアカウントが含まれています。

製品仕様書には、仕様、機能説明、および製品サポート情報が記載されており、エージェント1はこれらを使用して顧客からの質問に回答できます。製品仕様書はPDF形式で保存されます。

問題提起

コントソ氏は以下の問題点を指摘している。

- Agent1はContoso製品に関する一般的な知識しか持っていません。
- Agent1との最近のチャットのやり取りについて、感情分析が行われました。分析結果はまだ処理されていません。
- Agent1は、顧客からの質問に回答する際に、storage1に保存されている製品シートの詳細な製品情報を使用しません。

コントソ社の財務部門によると、仕入先からの請求書は、契約書に定められた条件と一致していることを確認するため、手作業で精査する必要があるとのことです。請求書には表やロゴ、様々なレイアウトが含まれているため、一貫した処理が困難です。

要件

計画されている変更

コントソ氏は以下の変更を実施する予定です。

プロジェクト1向けに、仕入先請求書の視覚的なレイアウトとテキスト内容の両方を評価することで請求書を分析し、仕入先との契約条件と照合できるソリューションを実装する。

- Agent1で使用される基本モデルのデプロイメントを更新し、モデルのバージョンを標準化して、継続性と一貫した応答を確保します。
- Agent1がstorage1に保存されている製品シートから詳細な製品情報を取得して使用できるようにします。
- Agent1が顧客からの質問に回答するために使用できる製品シートのインデックス作成ソリューションを実装する。
- 動画制作ソリューションの開発を完了する。

技術要件

Contoso社は、以下の技術要件を特定しました。

- Agent1で使用されるモデル展開は、拡張性が高く、高スループットの生成型AIワークロードをサポートし、予約済みのスループット容量を必要とせずに、変動する顧客サポートトラフィックを処理するために動的に拡張する必要があります。
- 製品シートは、セマンティック検索とベクトル検索を可能にするインデックス作成パイプラインを使用して処理され、Agent1が関連する製品情報を取得できるようにする必要があります。

製品シートの情報に基づいて生成される回答は、関連性があり、完全かつ正確でなければなりません。

- エージェント1は、製品シートを使用して、製品の詳細に関する自然言語の質問に答えることができなければなりません。
- エージェント1が使用するモデルのバージョンは、安定した応答を確保するために一貫している必要があります。
- モデルによって処理されるデータは、EU域内に留まらなければならない。

セキュリティおよびコンプライアンス要件

Contosoは、以下のセキュリティおよびコンプライアンス要件を特定しています。

- APIキーを使用してFoundryにデプロイされたモデルにアクセスしてはなりません。
- Azureリソースへのアクセスは、最小権限の原則に従う必要があります。
- Contosoの開発者は、Microsoft Entra認証を使用してMicrosoft Foundryリソースへの認証を行う必要があります。
- Project1へのアクセス権は、SC_Agent1_Devという名前のセキュリティグループを使用して、Agent1Devチームのメンバーに割り当てる必要があります。
 - Project1へのアクセス権は、SC_Agent1_Testという名前のセキュリティグループを使用して、Agent1Testチームのメンバーに割り当てる必要があります。
 - Agent1は、顧客データを含む文書がストレージ1の製品シートリポジトリに誤って追加された場合でも、顧客情報を決して開示してはならない。
- 製品仕様書には、テキストが埋め込まれた画像が含まれている場合があります。Agent1は、画像内に隠されている可能性のある悪意のある命令から保護される必要があります。
- ビジネス要件
- Contosoは以下のビジネス要件を特定しました。
- Agent1とやり取りするユーザーは、今後のやり取りにおいてパーソナライズされた体験を得る必要があります、そのためにはAgent1が会話の文脈を保持し、過去のやり取りから関連情報を想起できる機能が必要です。
 - エージェント1は、Contosoが販売する製品に関する質問のみに回答しなければならない。
- ホットスポットに関する質問
- Agent1のモデル展開を技術要件を満たすように構成する必要があります。
- 何を設定すればよいですか？回答するには、回答欄で適切なオプションを選択してください。
- 注：正解ごとに1ポイントが加算されます。

Deployment type:

Standard
Global Standard
Global Provisioned

Version update policy:

Once the current version expires
Opt out of automatic model version upgrades
Upgrade once a new default version becomes available

正解:

Answer Area

Deployment type:

Standard
Global Standard
Global Provisioned

Version update policy:

Once the current version expires
Opt out of automatic model version upgrades
Upgrade once a new default version becomes available

質問: 8

ホットスポットに関する質問

あなたは、エージェントを含み、KV1という名前のAzureキーコンテナを使用する、Project1という名前のMicrosoft Foundryプロジェクトを作成する予定です。Project1からKV1への接続を設定する必要があります。

上腕二頭筋コードをどのように完成させるべきですか？回答するには、回答欄で適切なオプションを選択してください。

注：正解ごとに1ポイントが加算されます。

Answer Area

```
resource existingKeyVault 'Microsoft.KeyVault/vaults@2024-11-01'
existing = {
  name: 'KV1'
  scope: resourceGroup()
}
resource connection
'Microsoft.CognitiveSevices/accounts/connections@2025-04-01-
preview' ={
  name: '${aiFoundryName}-keyvault'
  parent: aiFoundry
  properties: {
    category:
  target: existingKeyVault.id
  authType:
  isSharedToAll:
    metadata: {
      ApiType: 'Azure'
      ResourceId: existingKeyVault.id
      location: existingKeyVault.location
    }
  }
```

▼

AzureAIService

AzureKeyVault

AzureOpenAI

▼

AccountKey

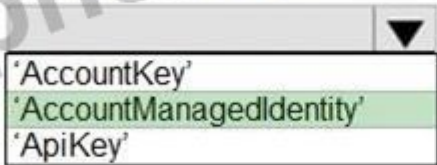
AccountManagedIdentity

ApiKey



正解:

Answer Area

```
resource existingKeyVault 'Microsoft.KeyVault/vaults@2024-11-01'
existing = {
  name: 'KV1'
  scope: resourceGroup()
}
resource connection
'Microsoft.CognitiveSevices/accounts/connections@2025-04-01-
preview' = {
  name: '${aiFoundryName}-keyvault'
  parent: aiFoundry
  properties: {
    category: 
    target: existingKeyVault.id
    authType: 
    isSharedToAll:
      metadata: {
        ApiType: 'Azure'
        ResourceId: existingKeyVault.id
        location: existingKeyVault.location
      }
  }
}
```



Explanation:

Box 1: AzureKeyVault

To configure a connection from a Microsoft Foundry project to an Azure Key Vault, you must use the AzureKeyVault connection category.

Box 2: AccountManagedIdentity

To connect a Microsoft Foundry project to an Azure Key Vault, you must use AccountManagedIdentity as the authType.

Microsoft Foundry secures its outbound connections to core Azure storage and security infrastructure using Microsoft Entra ID role-based access control (RBAC). Using AccountManagedIdentity completely eliminates the need to handle or rotate explicit credentials or api keys within your configuration.

Reference:

<https://learn.microsoft.com/en-us/azure/foundry/how-to/set-up-key-vault-connection>

質問: 9

あなたは、大量のトラフィックを処理するチャットアプリを提供するMicrosoft Foundryプロジェクトを持っています。

ほとんどの問い合わせは簡単なよくある質問(FAQ)ですが、中には高度な論理的思考を必要とするものもあります。

複雑な質問に対する回答の質を低下させることなく、一般的な問い合わせに対するコストと遅延を削減する必要があります。

あなたはすべきでしょうか？

A. すべてのリクエストをより小さなモデルにルーティングします。

B. リクエストを異なるモデルにルーティングするモデルカスケードを使用します。

- C. すべてのリクエストに対して max_tokens パラメータの値を増やします。
- D. すべてのリクエストを最も能力の高いモデルにルーティングします。

正解: ([正解を表示します](#))

One should absolutely use a model cascade to route requests to different models based on complexity. This architectural pattern is highly effective for high-volume chat applications because it directly addresses the trade-off between operational cost, API latency, and response quality.

Using a model cascade router ensures that your high-volume Microsoft Foundry application scales efficiently by reserving expensive computational power exclusively for queries that actually require advanced cognitive processing.

Note:

1. Analyze Request Complexity

Implement a lightweight intent classifier or routing layer at the entry point of your Microsoft Foundry project. This router quickly inspects incoming user prompts using basic heuristic keyword matching, semantic embeddings, or a highly optimized, fast model (like Phi-3 or GPT-4o-mini) to categorize the query as either a "Simple FAQ" or a "Complex Reasoning" request.

2. Route to the Optimal TierTier 1 (Fast & Cheap):

Route standard, predictable FAQ requests to a smaller, cost-effective model or a local cache/vector database lookup. This keeps latency in milliseconds and drastically lowers token costs.

Tier 2 (Advanced Reasoning): Route multi-step logic, coding, or highly contextual queries to a frontier model (like GPT-4o).

3. Implement Fallback LogicDesign the cascade to be dynamic. If the smaller Tier 1 model generates a response with low confidence, or if the user asks a follow-up question that invalidates the simple FAQ status, seamlessly upgrade the conversation loop to the Tier 2 model.

Reference:

<https://medium.com/@sujathamudadla1213/what-is-the-primary-purpose-of-a-model-cascade-in-machine-learning-0b145a7bc6e2>

質問: 10

Microsoft Foundry エージェントが、以下の内容を含む Azure Search インデックスからの応答を解析します。

- 商品名と商品コードを検索できるテキストフィールド
- 製品説明の埋め込みを格納するベクトル場

ユーザーが以下の方法でインデックスを照会できるようにする必要があります。

- 正確な製品名またはコード
- 製品の自然言語による説明

何を設定すればよいですか？

- A. ベクター検索のみ
- B. ハイブリッド検索
- C. キーワード検索のみ
- D. 意味検索のみ

正解: ([正解を表示します](#))

To meet your requirements, you need to configure a Hybrid Search with Semantic Ranking in Azure AI Search. This setup combines keyword matching for exact identifiers with vector search for natural language queries, delivering the most accurate grounding data to your Microsoft Foundry agent.

Reference:

<https://www.tredence.com/blog/searchsmart-enhancing-rag-with-azure-ai-search-service>

質問: 11

ホットスポットに関する質問

顧客サポートのトリージングプロセス用のワークフローを含むMicrosoft Foundryプロジェクトがあります。

質問ノードがあり、ユーザーの回答はVar01という名前のローカル変数に保存されます。

以下のPower Fx式を作成する必要があります。

- Var01 に以下の内容が含まれることを保証する if/else 条件式

価値

- 保存されたユーザー応答を返す送信メッセージ式

大文字

式はどのように設定すればよいですか？回答するには、回答欄で適切なオプションを選択してください。

注：正解ごとに1ポイントが加算されます。

Answer Area

If/else condition expression:

IsBlank(Local.Var01)
IsEmpty(Local.Var01)
Not(IsBlank(Local.Var01))

Send message expression:

{Local.Var01}
{Upper(Local.Var01)}
{Upper(Var01)}

正解:

Answer Area

If/else condition expression:

IsBlank(Local.Var01)
IsEmpty(Local.Var01)
Not(IsBlank(Local.Var01))

Send message expression:

{Local.Var01}
{Upper(Local.Var01)}
{Upper(Var01)}

質問: 12

あなたは、ユーザーがアップロードした複数のファイルを結合して処理するエージェントを構築する予定です。

エージェントの開発にAzure AI Agent Serviceを使用するかどうかを評価しています。

このサービスにアップロードできるファイルの最大サイズはどれくらいですか？

A. 1GB

B. 10GB

C. 100 GB

D. 1TB

正解: ([正解を表示します](#))

<https://learn.microsoft.com/en-us/azure/ai-services/agents/quotas-limits>

質問: 13

エージェントを含む Microsoft Foundry プロジェクトがあり、GitHub リポジトリを使用しています。リポジトリには、エージェントの評価設定を定義する File1 という名前の YAML ファイルが含まれています。プルリクエスト (PR) が開かれたときに File1 で定義された評価を実行する GitHub Actions ワークフローを作成する必要があります。ワークフローはどのように構成すればよいのでしょうか？

A. project-endpoint をプロジェクトのエンドポイントに設定します。

B. evaluation-config を YAML ファイルのパスに設定します。

C. model-ployment-name をデプロイされたモデルに設定します。

D. tenant-id を Microsoft Entra テナント ID に設定します

正解: A ([コメントを发表する](#))

The correct configuration choice is to set project-endpoint to the endpoint of the project.

azure-ai-project-endpoint: This parameter is required and must be set to the URL endpoint of your Microsoft Foundry project. This aligns with setting project-endpoint.

Reference:

<https://learn.microsoft.com/en-us/azure/foundry/how-to/evaluation-github-action>

質問: 14

Project1 という名前の Microsoft Foundry プロジェクトがあり、その中に以下のものが含まれています。

- 外部APIを呼び出すOpenAPIツール

- 外部 API の API キーを格納する Connection1 という名前のプロジェクト接続。エージェントが OpenAPI ツールを呼び出すと、API は 401 認証エラーを返し、トレースでは API キー ヘッダーが送信されていないことが示されます。

OpenAPIツールが、すべてのリクエストにConnection1のAPIキーを自動的に含めるようにする必要があります。

あなたはすべきでしょうか？

A. IDパススルーを有効にして、ツールが呼び出し元のMicrosoft Entraトークンを使用するようにします。

B. OpenAPI仕様にAPIキーヘッダーを手動で追加します。

C. ツールをProject1のデフォルト接続を使用するように設定します。

D. ツールを接続1に接続します。

正解: ([正解を表示します](#))

To make sure that the OpenAPI tool automatically includes the API key from the connection on all requests, you should explicitly define the security schemes and security requirements in your OpenAPI specification file.

Incorrect:

[Not D]

Even if you successfully link the tool to the custom connection in the portal UI, the Microsoft Foundry agent service will intentionally redact or withhold the API key if it cannot find a corresponding matching structure in the tool's code definition.

Reference:

<https://learn.microsoft.com/en-us/azure/foundry/agents/how-to/tools/openapi>

質問: 15

ドラッグアンドドロップ問題

デプロイ済みのチケットトリアージエージェントを含むMicrosoft Foundryプロジェクトがあります。
エージェントがツールを必要とする場合でも、ツールを呼び出さずに応答することがあるということに気づくでしょう。
エージェントが実行中にツールを呼び出すようにする必要があります。
Pythonコードを完成させるにはどうすればよいでしょうか？適切な値を正しいターゲットにドラッグして答えてください。各値は1回、複数回、または全く使用されない場合があります。コンテンツを表示するには、ペイン間の分割バーをドラッグするか、スクロールする必要がある場合があります。
注：正解ごとに1ポイントが加算されます。

Values

⋮ "auto"

⋮ "required"

⋮ "response_format"

⋮ "tool_choice"

⋮ "tools"

⋮ "type"

Answer Area

```
run_payload = {  
    "assistant_id": agent_id,  
    Value : Value ,  
    "metadata": {  
        "scenario": "ticket-triage"  
    }  
}
```

正解:

Values

⋮ "auto"
⋮ "required"
⋮ "response_format"
⋮ "tool_choice"
⋮ "tools"
⋮ "type"

Answer Area

```
run_payload = {  
    "assistant_id": agent_id,  
    "tool_choice": "required",  
    "metadata": {  
        "scenario": "ticket-triage"  
    }  
}
```



質問: 16

エージェントは、ツールセットとして公開されている組織の内部在庫システムを呼び出す必要があります。ツール定義は、各エージェントで再定義するのではなく、一元的に管理され、複数のエージェント間で再利用できるようにしたいと考えています。Foundryのどの機能がこれに適していますか？

- A. モデルコンテキストプロトコル(MCP)サーバーをエージェントツールとして追加する
- B. モデルのコンテキストウィンドウを拡大する
- C. 在庫管理システムのAPIスキーマに基づいてモデルを微調整する
- D. コードインタープリタツールを有効にする

正解: A ([コメントを发表する](#))

A Model Context Protocol (MCP) server exposes a reusable, centrally maintained set of tools that any agent can connect to, and Foundry Agent Service supports adding MCP servers from the tools catalogue. This is the standard pattern for connecting agents to external systems through shared tool definitions.

有効的な**AI-103-JPN**問題集はJPNTest.com提供され、**AI-103-JPN**試験に合格することに役に立ちます！JPNTest.comは今最新**AI-103-JPN**試験問題集を提供します。JPNTest.com AI-103-JPN試験問題集はもう更新されました。ここで**AI-103-JPN**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/AI-103-JPN-mondaishu> **159問、30%ディスカウント**、特別な割引コード: **JPNshiken**」

質問: 17

事例研究1 - コントソ社

概要

会社情報

Contoso, Ltdは、Microsoft Foundryを使用して生成型AIおよびエージェントベースのソリューションを構築、展開、管理する多国籍小売企業です。

既存の環境

アイデンティティ環境

Contosoは、エージェントが組織のリソースやサービスにアクセスできるようにするID管理、認証、認可機能にMicrosoft Entra IDを使用しています。

Contosoは最近、既存のAIソリューションを最適化および保守するために、Agent1Dev Teamという新しいAIエンジニアリングチームを結成しました。

チームは、ソリューションアーキテクト、DevOpsエンジニア、セキュリティエンジニアと協力して、AIアプリケーションの設計、実装、監視、およびセキュリティ対策を行います。

Contosoには、AIソリューションの導入前にその妥当性を検証する「Agent1Test Team」というチームも存在する。

生成環境

Contosoは、Project1とProject2という名前の2つのプロジェクトを含むMicrosoft Foundry環境を構築しています。

プロジェクト1

Project1には、顧客からの製品に関する問い合わせやトラブルシューティングの依頼に対応する、Agent1という名前のカスタマーサポートエージェントが含まれています。

Agent1には以下の設定があります。

- Agent1は基本モデルのデプロイメントを使用します。

安全性評価パイプラインは有効になっていません。

ツール呼び出し承認ワークフローは有効になっていません。

会話メモリの制約は設定されていません。

Agent1は、デジタルサポートチャネルを使用して顧客とやり取りし、Contoso製品に関する一般的な質問に回答します。

Project1は、欧州連合(EU)域内にあるAzure リージョンにデプロイされています。

Agent1開発チームは、Project1を使用してAgent1の最適化と保守を行います。

プロジェクト2

Project2には、既に展開済みの動画生成モデルが含まれています。Contoso社のマーケティング部門はProject2にアクセスでき、このモデルを用いて動画制作ソリューションを開発する予定です。

解決策の開発は未完了です。

データ環境

Contosoは、AIアプリケーションをサポートするAzure リソースに製品関連情報を保存します。

Azure環境には、Contoso製品の製品詳細シートを保存するstorage1という名前のAzure Blob Storageアカウントが含まれています。

製品仕様書には、仕様、機能説明、および製品サポート情報が記載されており、エージェント1はこれらを使用して顧客からの質問に回答できます。製品仕様書はPDF形式で保存されます。

問題提起

コントソ氏は以下の問題点を指摘している。

- Agent1はContoso製品に関する一般的な知識しか持っていません。

- Agent1との最近のチャットのやり取りについて、感情分析が行われました。分析結果はまだ処理されていません。

- Agent1は、顧客からの質問に回答する際に、storage1に保存されている製品シートの詳細な製品情報を使用しません。

コントソ社の財務部門によると、仕入先からの請求書は、契約書に定められた条件と一致していることを確認するため、手作業で精査する必要があるとのことです。請求書には表やロゴ、様々なレイアウトが含まれているため、一貫した処理が困難です。

要件

計画されている変更

コントソ氏は以下の変更を実施する予定です。

プロジェクト1向けに、仕入先請求書の視覚的なレイアウトとテキスト内容の両方を評価することで請求書を分析し、仕入先との契約条件と照合できるソリューションを実装する。

- Agent1で使用される基本モデルのデプロイメントを更新し、モデルのバージョンを標準化して、継続性と一貫した応答を確保します。
- Agent1がstorage1に保存されている製品シートから詳細な製品情報を取得して使用できるようにします。
- Agent1が顧客からの質問に回答するために使用できる製品シートのインデックス作成ソリューションを実装する。
- 動画制作ソリューションの開発を完了する。

技術要件

Contoso社は、以下の技術要件を特定しました。

- Agent1で使用されるモデル展開は、拡張性が高く、高スループットの生成型AIワークロードをサポートし、予約済みのスループット容量を必要とせずに、変動する顧客サポートトラフィックを処理するために動的に拡張する必要があります。
- 製品シートは、セマンティック検索とベクトル検索を可能にするインデックス作成パイプラインを使用して処理され、Agent1が関連する製品情報を取得できるようにする必要があります。

製品シートの情報に基づいて生成される回答は、関連性があり、完全かつ正確でなければなりません。

- エージェント1は、製品シートを使用して、製品の詳細に関する自然言語の質問に答えることができません。
- エージェント1が使用するモデルのバージョンは、安定した応答を確保するために一貫している必要があります。
- モデルによって処理されるデータは、EU域内に留まらなければならない。

セキュリティおよびコンプライアンス要件

Contosoは、以下のセキュリティおよびコンプライアンス要件を特定しています。

- APIキーを使用してFoundryにデプロイされたモデルにアクセスしてはなりません。

Azureリソースへのアクセスは、最小権限の原則に従う必要があります。

Contosoの開発者は、Microsoft Entra認証を使用してMicrosoft Foundryリソースへの認証を行う必要があります。

- Project1へのアクセス権は、SC_Agent1_Devという名前のセキュリティグループを使用して、Agent1Devチームのメンバーに割り当てる必要があります。
- Project1へのアクセス権は、SC_Agent1_Testという名前のセキュリティグループを使用して、Agent1Testチームのメンバーに割り当てる必要があります。
- Agent1は、顧客データを含む文書がストレージ1の製品シートリポジトリに誤って追加された場合でも、顧客情報を決して開示してはならない。

製品仕様書には、テキストが埋め込まれた画像が含まれている場合があります。Agent1は、画像内に隠されている可能性のある悪意のある命令から保護される必要があります。

ビジネス要件

Contosoは以下のビジネス要件を特定しました。

- Agent1とやり取りするユーザーは、今後のやり取りにおいてパーソナライズされた体験を得る必要があります、そのためにはAgent1が会話の文脈を保持し、過去のやり取りから関連情報を想起できる機能が必要です。
- エージェント1は、Contosoが販売する製品に関する質問のみに回答しなければならない。

財務部門から報告された問題を解決する請求書レビューソリューションを提案する必要があります。提案には何を含めるべきでしょうか？

- A.** チャット完了
- B.** Foundry Tools の Azure Document Intelligence
- C.** Foundry Tools における Azure コンテンツの理解
- D.** 画像解析

正解: ([正解を表示します](#))

Scenario:

Problem Statements: Contoso identifies the following issues:

The finance department at Contoso reports that vendor invoices must be reviewed manually to ensure that the invoices match the terms defined in the vendor contracts. The invoices contain tables, logos, and varied layouts that make the documents difficult to process consistently.

The correct tool to include in your solution is Azure Document Intelligence in Foundry Tools.

Structured Data Extraction: Azure Document Intelligence is a cloud-based AI service specifically optimized for extracting structured data, key-value pairs, and text from complex documents such as vendor invoices, receipts, and contracts.

Layout and Table Handling: Invoices frequently contain unpredictable tables, complex spacing, logos, and mixed image-and-text variations. Azure Document Intelligence uses dedicated, prebuilt invoice models and layout analysis capabilities engineered to map and preserve the semantic structure of tabular data across inconsistent layouts.

Confidence Scoring for Manual Review: The service generates confidence scores for extracted fields. This allows you to automatically flag low-confidence Extractions and seamlessly route complex anomalies to human operators for manual validation against vendor contracts.

Incorrect:

[Not C]

Azure Content Understanding: While useful for multi-modal analysis and schema-based image querying, it currently lacks the specialized, highly optimized parsing engine required for complex multi-page tables. Using it for varied invoice processing often results in inconsistent line-item mapping and tabular errors compared to the robust, specialized Document Intelligence invoice models.

Reference:

<https://medium.com/@meetapa/part-9-building-an-intelligent-invoice-processing-system-with-microsoft-foundry-and-ai-agents-0cd181df4b2e>

質問: 18

Azure OpenAIリソースを含むAzureサブスクリプションをお持ちです。

GPT-4モデルをリソースにデプロイします。

モデルの基礎データとして使用されるファイルをアップロードできることを確認する必要があります。

作成すべきリソースは、どの2種類でしょうか？それぞれの正解は、解決策の一部を示しています。

注：正解ごとに1ポイントが加算されます。

A. Azure AI ドキュメント インテリジェンス

B. Azure AI Search

C. Azure SQL

D. Azure Blob Storage

E. Azure AI Bot Service

正解: ([正解を表示します](#))

Azure OpenAI On Your Data enables you to run advanced AI models such as GPT-35-Turbo and GPT-4 on your own enterprise data without needing to train or fine-tune models.

For some data sources such as uploading files from your local machine (preview) or data contained in a blob storage account (preview), Azure AI Search is used.

Reference:

<https://learn.microsoft.com/en-us/azure/ai-services/openai/concepts/use-your-data>

質問: 19

事例研究1 - コントソ社

概要

会社情報

Contoso, Ltdは、Microsoft Foundryを使用して生成型AIおよびエージェントベースのソリューションを構築、展開、管理する多国籍小売企業です。

既存の環境

アイデンティティ環境

Contosoは、エージェントが組織のリソースやサービスにアクセスできるようにするID管理、認証、認可機能にMicrosoft Entra IDを使用しています。

Contosoは最近、既存のAIソリューションを最適化および保守するために、Agent1Dev Teamという新しいAIエンジニアリングチームを結成しました。

チームは、ソリューションアーキテクト、DevOpsエンジニア、セキュリティエンジニアと協力して、AIアプリケーションの設計、実装、監視、およびセキュリティ対策を行います。

Contosoには、AIソリューションの導入前にその妥当性を検証する「Agent1Test Team」というチームも存在する。

生成環境

Contosoは、Project1とProject2という名前の2つのプロジェクトを含むMicrosoft Foundry環境を構築しています。

プロジェクト1

Project1には、顧客からの製品に関する問い合わせやトラブルシューティングの依頼に対応する、Agent1という名前のカスタマーサポートエージェントが含まれています。

Agent1には以下の設定があります。

- Agent1は基本モデルのデプロイメントを使用します。

安全性評価パイプラインは有効になっていません。

ツール呼び出し承認ワークフローは有効になっていません。

会話メモリの制約は設定されていません。

Agent1は、デジタルサポートチャネルを使用して顧客とやり取りし、Contoso製品に関する一般的な質問に回答します。

Project1は、欧州連合(EU)域内にあるAzureリージョンにデプロイされています。

Agent1開発チームは、Project1を使用してAgent1の最適化と保守を行います。

プロジェクト2

Project2には、既に展開済みの動画生成モデルが含まれています。Contoso社のマーケティング部門はProject2にアクセスでき、このモデルを用いて動画制作ソリューションを開発する予定です。

解決策の開発は未完了です。

データ環境

Contosoは、AIアプリケーションをサポートするAzureリソースに製品関連情報を保存します。

Azure環境には、Contoso製品の製品詳細シートを保存するstorage1という名前のAzure Blob Storageアカウントが含まれています。

製品仕様書には、仕様、機能説明、および製品サポート情報が記載されており、エージェント1はこれらを使用して顧客からの質問に回答できます。製品仕様書はPDF形式で保存されます。

問題提起

コントソ氏は以下の問題点を指摘している。

- Agent1はContoso製品に関する一般的な知識しか持っていません。

- Agent1との最近のチャットのやり取りについて、感情分析が行われました。分析結果はまだ処理されていません。

- Agent1は、顧客からの質問に回答する際に、storage1に保存されている製品シートの詳細な製品情報を使用しません。

コントソ社の財務部門によると、仕入先からの請求書は、契約書に定められた条件と一致していることを確認するため、手作業で精査する必要があるとのことです。請求書には表やロゴ、様々なレイアウトが含まれているため、一貫した処理が困難です。

要件

計画されている変更

コントソ氏は以下の変更を実施する予定です。

プロジェクト1向けに、仕入先請求書の視覚的なレイアウトとテキスト内容の両方を評価することで請求書を分析し、仕入先との契約条件と照合できるソリューションを実装する。

- Agent1で使用する基本モデルのデプロイメントを更新し、モデルのバージョンを標準化して、継続性と一貫した応答を確保します。

- Agent1がstorage1に保存されている製品シートから詳細な製品情報を取得して使用できるようにします。

- Agent1が顧客からの質問に回答するために使用できる製品シートのインデックス作成ソリューションを実装する。

動画制作ソリューションの開発を完了する。

技術要件

Contoso社は、以下の技術要件を特定しました。

- Agent1で使用するモデル展開は、拡張性が高く、高スループットの生成型AIワークロードをサポートし、予約済みのスループット容量を必要とせずに、変動する顧客サポートトラフィックを処理するために動的に拡張する必要があります。

- 製品シートは、セマンティック検索とベクトル検索を可能にするインデックス作成パイプラインを使用して処理され、Agent1が関連する製品情報を取得できるようにする必要があります。

製品シートの情報に基づいて生成される回答は、関連性があり、完全かつ正確でなければなりません。

- エージェント1は、製品シートを使用して、製品の詳細に関する自然言語の質問に答えることができません。

- エージェント1が使用するモデルのバージョンは、安定した応答を確保するために一貫している必要があります。

- モデルによって処理されるデータは、EU域内に留まらなければならない。

セキュリティおよびコンプライアンス要件

Contosoは、以下のセキュリティおよびコンプライアンス要件を特定しています。

- APIキーを使用してFoundryにデプロイされたモデルにアクセスしてはなりません。

Azureリソースへのアクセスは、最小権限の原則に従う必要があります。

Contosoの開発者は、Microsoft Entra認証を使用してMicrosoft Foundryリソースへの認証を行う必要があります。

- Project1へのアクセス権は、SC_Agent1_Devという名前のセキュリティグループを使用して、Agent1Devチームのメンバーに割り当てる必要があります。

- Project1へのアクセス権は、SC_Agent1_Testという名前のセキュリティグループを使用して、Agent1Testチームのメンバーに割り当てる必要があります。

- Agent1は、顧客データを含む文書がストレージ1の製品シートリポジトリに誤って追加された場合でも、顧客情報を決して開示してはならない。

製品仕様書には、テキストが埋め込まれた画像が含まれている場合があります。Agent1は、画像内に隠されている可能性のある悪意のある命令から保護される必要があります。

ビジネス要件

Contosoは以下のビジネス要件を特定しました。

- Agent1とやり取りするユーザーは、今後のやり取りにおいてパーソナライズされた体験を得る必要があります、そのためにはAgent1が会話の文脈を保持し、過去のやり取りから関連情報を想起できる機能が必要です。

- エージェント1は、Contosoが販売する製品に関する質問のみに回答しなければならない。

Agent1がストレージ1から関連製品情報を取得できるように、インデックス作成パイプラインを構成する必要があります。このソリューションは技術要件を満たしている必要があります。

どの2つの組み込みスキルを使うべきでしょうか？それぞれの正解は、解決策の一部を示しています。

注：正解ごとに1ポイントが加算されます。

A. Azure OpenAI 埋め込み

B. 固有表現認識

C. テキスト分割

D. マージ

E. 言語検出

F. キーワード抽出

正解: ([正解を表示します](#))

The most essential skills for this scenario are Azure OpenAI Embedding and Text Split.

Azure OpenAI Embedding: This skill is critical for generating the vector representations (embeddings) of your text, which directly enables the required vector search capability.

Text Split: This skill is essential because LLMs and embedding models have strict token limits.

Breaking large product detail sheets into smaller chunks ensures the text fits into the embedding model and improves the accuracy of semantic search.

Incorrect:

[Not B]

Entity Recognition: This extracts specific entities like names, dates, or locations. While helpful for advanced filtering, it is not a foundational requirement to enable basic semantic or vector search.

[Not D]

Merge: This skill combines text from multiple fields into a single string. Since product sheets are already unified documents, splitting and chunking them is the priority rather than merging separate fields.

[Not E]

Language Detection: This identifies the language of the input text. Unless your product sheets are completely multilingual and require conditional routing to different language models, this skill is secondary.

[Not F]

Key Phrase Extraction: This pulls out main talking points or keywords. This is primary used for traditional keyword tagging or basic search indexing, whereas your requirement specifically dictates vector and semantic-based retrieval.

Scenario:

Technical Requirements;

*-> The product sheets must be processed by using an indexing pipeline that enables semantic and vector search, so that Agent1 can retrieve the relevant product information.

Data environment: The Azure environment contains an Azure Blob Storage account named storage1 that stores product detail sheets for all the Contoso products.

Planned changes: Enable Agent1 to retrieve and use the detailed product information from the product sheets stored in storage1.

Reference:

<https://www.rheininsights.com/blog/en/Retrieval+Augmented+Generation+with+Azure+AI+Search+and+Atlassian+Confluence.php>

質問: 20

ホットスポットに関する質問

感情分析と光学文字認識(OCR)を実行するために使用する新しいリソースを作成する必要があります。ソリューションは以下の要件を満たす必要があります。

単一のキーとエンドポイントを使用して、複数のサービスにアクセスできます。

今後利用する可能性のあるサービスの請求を一本化します。

- 将来的にFoundry ToolsでAzure Visionの使用をサポートする。

新しいリソースを作成するためのHTTPリクエストはどのように完了すればよいですか？回答するには、回答欄で適切なオプションを選択してください。

注：正解ごとに1ポイントが加算されます。

Answer Area

	▼	https://management.azure.com/subscriptions/ xxxxxxx-xxxx-xxxx-xxxx- xxxxxxxxxxxx/resourceGroups/RG1
PATCH		
POST		
PUT		

/providers/Microsoft.CognitiveService/accounts/CS1?api-version=2021-04-30

```
{  
  "location": "West US",  
  "kind": "
```

Need to create a Cognitive Services Multi-Service resource using the PUT method. This specific resource type provides a single endpoint and key for multiple AI services, consolidates billing, and supports Azure Vision.

Box 2: CognitiveServices

Using the CognitiveServices kind creates a multi-service resource. This fulfills all the requirements by providing a single key and endpoint for multiple services, consolidating billing, and enabling access to features like Azure AI Vision and Text Analytics under one roof.

Reference:

<https://learn.microsoft.com/en-us/azure/foundry/how-to/develop/sdk-overview>

質問: 21

あなたは、受信したサポートチケットをカテゴリ別に分類するMicrosoft Foundryプロジェクト用のプロンプトを開発しています。

モデルを保持したり、知識を永続的に保存したりすることなく、正しい分類がどのようなものかモデルに示すことで、精度を向上させる必要があります。

どの迅速エンジニアリング手法を用いるべきでしょうか？

A. 検索拡張生成(RAG)

B. ゼロショット学習

C. 思考の連鎖

D. 少数ショット学習

正解: D ([コメントを发表する](#))

Few-shot prompting is the best approach for this project. This technique improves classification accuracy by including a few high-quality examples directly inside the prompt, giving the model a clear pattern to follow without modifying its weights or storing data permanently.

Zero retraining: Works entirely through in-context learning during the API call.

No permanent storage: The knowledge disappears as soon as the inference request is completed.

Immediate accuracy boost: Demonstrates formatting, nuances, and edge cases directly to the model.

Reference:

<https://www.linkedin.com/pulse/top-interview-questions-answers-prompt-engineering-nitin-sharma-ka4fc>

質問: 22

ドラッグアンドドロップ問題

エージェントを含む Microsoft Foundry プロジェクトがあります。このエージェントはスレッドとファイルアップロードを使用し、Azure OpenAI モデルのデプロイメントを呼び出します。

負荷テスト中に、呼び出しが断続的に失敗し、HTTP 429 レート制限超過エラーが返されます。

一部のユーザーによるアップロードが失敗し、HTTP 400 ファイルサイズ超過エラーが発生します。

エラーを軽減し、呼び出しの失敗を減らす必要があります。解決策は、サービスおよびモデルの制限内に収まるものでなければなりません。

各エラーを解決するにはどうすればよいでしょうか？回答するには、適切な操作を正しいエラーにドラッグしてください。各操作は、1回、複数回、またはまったく使用しない場合があります。コンテンツを表示するには、ペイン間の分割バーをドラッグするか、スクロールする必要がある場合があります。

注：彗星の選択はそれぞれ1ポイントです。

Actions	Answer Area
Microsoft Increase tenant-wide quotas.	
Move large content to files and use file search.	HTTP 429:
Use additional agent tools to reduce the message size.	HTTP 400:
Implement exponential backoff and jitter in the retry logic.	
Split content into smaller files before uploading the files.	

正解:

Actions	Answer Area
Microsoft Increase tenant-wide quotas.	
Move large content to files and use file search.	HTTP 429: Implement exponential backoff and jitter in the retry logic.
Use additional agent tools to reduce the message size.	HTTP 400: Split content into smaller files before uploading the files.
Implement exponential backoff and jitter in the retry logic.	
Split content into smaller files before uploading the files.	

Explanation:

Box 1: Implement exponential backoff and jitter in the retry logic

To remedy intermittent HTTP 429 rate limit errors while staying within service limits, you must implement a retry policy that uses exponential backoff and jitter.

Because the agent uses threads and file uploads, load testing triggers short-window concurrency spikes (bursting). Azure OpenAI evaluates rate limits in small slices (like 1- to 10-second windows), meaning overlapping file processing or concurrent thread steps will trigger a 429 even if your total Tokens-Per-Minute (TPM) or Requests-Per-Minute (RPM) look safe overall.

Box 2: Split content into smaller files before uploading the files.

Instead of uploading a large document in a single standard files.create request, utilize the Chunked Uploads REST API path. This splits the file payload on your client application layer before streaming it to Azure.

Note: To resolve the HTTP 400 file size exceeded error in your Microsoft Foundry agent, you need to bypass the strict payload constraints of standard single-request uploads. In Azure OpenAI and Microsoft Foundry Agent architectures, a multipart standard upload has a rigid request body limit (typically 30 MB or 50 MB).

Reference:

<https://learn.microsoft.com/en-us/answers/questions/2265002/getting-error-after-deployed-a-model-in-azure-ai-f>

[https://learn.microsoft.com/en-us/answers/questions/5521436/getting-400-error\(request-body-too-large\)-for-batc](https://learn.microsoft.com/en-us/answers/questions/5521436/getting-400-error(request-body-too-large)-for-batc)

質問: 23

あるチームは、Foundryカタログからモデルを可能な限り迅速にセットアップし、消費したトークン分だけを支払い、仮想マシン(VM)の割り当て量を一切設定せずにデプロイしたいと考えています。どのデプロイオプションが最適でしょうか？

- A. マネージドコンピューティング(マネージドエンドポイント)の展開
- B. サーバーレスAPIデプロイメント、別名モデル・アズ・ア・サービス(MaaS)
- C. Azure Kubernetes Service 上でのセルフホスト型デプロイメント
- D. 開発者ワークステーション上で実行されるローカルコンテナ

正解: **B** ([コメントを发表する](#))

Serverless API deployment runs on Microsoft-managed infrastructure, bills per input and output token, and needs no VM quota, so it is the fastest path that matches token-based billing. It is the standard choice when you want pay-per-use economics without managing compute.

質問: 24

Azure OpenAIのGPT-4モデルを使用して歌詞を生成するチャットボットを構築します。

回答には、モデル学習中に取り込まれた可能性のある人気曲の歌詞が含まれていないことを確認する必要があります。

どのAzure AIコンテンツ安全APIを使用すべきでしょうか？

- A. カスタムカテゴリ
- B. テキスト分析
- C. 接地性検出
- D. 保護された素材のテキスト検出

正解: ([正解を表示します](#))

The Protected material text detection feature in Azure AI Content Safety is specifically designed to identify and block responses that may contain copyrighted or protected content, such as lyrics from popular songs. Since GPT-4 may have been trained on publicly available data, this feature helps ensure that the chatbot does not generate copyrighted lyrics by detecting and filtering out protected material.

質問: 25

注 :このセクションには、同じシナリオと問題に関する複数の質問セットが含まれています。各質問には、問題に対する固有の解決策が提示されています。提示された解決策が、提示された目標を満たしているかどうかを判断する必要があります。セット内の複数の解決策が問題を解決できる場合もあります。また、セット内のどの解決策も問題を解決できない場合もあります。

このセクションの質問に回答すると、前のセクションに戻ることはできません。そのため、これらの質問は復習画面には表示されません。

画像のアップロードを受け付け、抽出した画像テキストを使用して応答を生成する、マルチモーダルなAI生成モデルをお持ちです。

ユーザーが安全でない画像をアップロードしたり、モデルを操作するための隠された指示を画像に埋め込んだりできることが分かります。

リスクを軽減するための対策を実施する必要があります。

解決策 :ユーザープロンプト用のプロンプトシールドを設定します。

これは目標を達成していると言えるでしょうか？

- A. はい
- B. いいえ

正解: ([正解を表示します](#))

Correct:

* You configure a prompt shield for documents.

Prompt Shield for Documents: Highly Effective (Critical Defense)

How it helps: This shield specifically scans untrusted, third-party data inputs (like external documents or text extracted from uploaded images).

Mechanism: It evaluates the extracted image text before it is sent to the LLM to identify hidden jail

* You configure a prompt shield for user prompts.

Prompt Shield for User Prompts: Partially Effective (Defense in Depth)

How it helps: This shield targets direct jailbreak attempts written manually by the user in the text prompt field accompanying the upload.

Mechanism: It prevents the user from typing supporting instructions that prime the model to execute the hidden instructions found within the image.

* You configure image moderation to block unsafe content before processing the images.

Implementing rigorous image moderation is one of the most effective ways to secure multimodal AI systems against these threats. Moderation acts as a necessary gatekeeper, preventing malicious inputs from ever reaching the generative model.

Incorrect:

* You configure protected material detection.

Protected Material Detection: Ineffective for this Threat

Why it does not help: This feature is designed to scan model outputs to prevent the generation of copyrighted text, proprietary source code, or licensed imagery.

Limitation: It does not scan inputs for adversarial instructions and will not prevent a user from manipulating the model's logic.

Reference:

<https://www.upgrad.com/blog/what-is-multimodal-ai/>

<https://learn.microsoft.com/en-us/azure/ai-services/content-safety/concepts/jailbreak-detection>

質問: **26**

ホットスポットに関する質問

App1という名前のPythonアプリケーションがあり、それがProject1という名前のMicrosoft Foundryプロジェクトと統合されています。

App1が以下の要件を満たしていることを確認する必要があります。

- Microsoft Entra マネージド ID を使用して認証します

- Azure OpenAI Responses API を使用して、デプロイ済みのモデルにプロンプトを送信します。Python コードをどのように完成させるべきですか？回答するには、回答領域で適切なオプションを選択してください。

注：正解ごとに1ポイントが加算されます。

```

from azure.identity import DefaultAzureCredential
from azure.ai.projects import AIProjectClient

credential =  ()


|                        |
|------------------------|
| AzureKeyCredential     |
| ClientSecretCredential |
| DefaultAzureCredential |



project_client = AIProjectClient(
    endpoint="https://contosoai-services.ai.azure.com/api/projects/project1",
    credential=credential,
)

with project_client.get_openai_client() as openai_client:
    response = openai_client.responses. (


|          |
|----------|
| compact  |
| create   |
| retrieve |



        model="trail-guide-chat",
        input="Create a 3-day hiking itinerary near Seattle.",
    )

print(response.output_text)

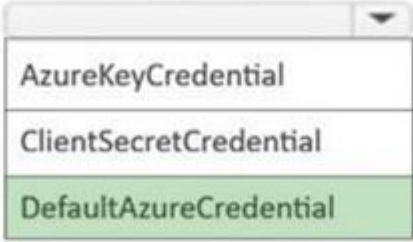
```



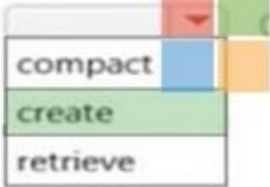
正解:

Answer Area

```
from azure.identity import DefaultAzureCredential
from azure.ai.projects import AIProjectClient

credential =  ()

project_client = AIProjectClient(
    endpoint="https://contosoai.services.ai.azure.com/api/projects/project1",
    credential=credential,
)

with project_client.get_openai_client() as openai_client:
    response = openai_client.responses.

    model="trail-guide-chat",
    input="Create a 3-day hiking itinerary near Seattle.",
)

print(response.output_text)
```

Microsoft

質問: 27

エージェントを含む Microsoft Foundry プロジェクトがあります。

エージェントは、データ取得ツールとして Azure AI Search を使用します。

エージェントがドキュメント内および埋め込み画像内のテキストに基づいて応答を生成できるように、PDF を Azure AI Search インデックスに取り込む予定です。

利用者は、出典ファイルへのリンクを含む引用を求めている。

インデックス作成の際には、画像が内蔵の光学文字認識(OCR)機能の入力として使用できる構造に抽出されるようにする必要があります。

どのインデックス作成方法を採用すべきでしょうか？

- A. 画像データを正規化された画像コレクションに抽出するインデクサー
- B. OCR入力を再構築するシェイパースキル
- C. インデックスのコンテンツフィールドに対してOCRスキルを直接実行するスキルセット
- D. 画像データを検索可能なフィールドに書き込むための outputFieldMappings パラメータ

正解: A ([コメントを发表する](#))

The best configuration step an indexer to extract image data into a normalized_images collection.

Cracking Foundation: In Azure AI Search, document cracking automatically separates textual data from visual content. To make embedded images available to image-processing skills like OCR, the indexer must be explicitly configured with an imageAction configuration property (such as generateNormalizedImages or generateNormalizedImagePerPage).

Input Structure for OCR: This indexer step extracts embedded images from the PDF and restructures them into an internal normalized_images array (accessed via the path

/document/normalized_images/*). The built-in OCR skill strictly requires this normalized image collection format as its input payload.

Incorrect:

[Not C]

Content Field Limitations: The /document/content field generated during document cracking holds only the raw text extracted from the file. It does not contain the binary data or structural properties of embedded images.

Targeting Errors: Pointing an OCR skill directly at the raw text content field would fail to analyze the actual images, leaving the agent unable to ground answers in visual elements like diagrams or embedded infographics.

Reference:

<https://docs.azure.cn/en-us/search/tutorial-skillset>

質問: 28

ドラッグアンドドロップ問題

Azure AI Search を使用する Web アプリケーションをお持ちです。

アクティビティをレビューしたところ、予想を上回る検索クエリ数が確認されました。クエリキーが侵害された疑いがあります。

検索エンドポイントへの不正アクセスを防止し、ユーザーがドキュメントコレクションに対して読み取り専用アクセスのみを許可されるようにする必要があります。また、アプリケーションのダウンタイムを最小限に抑えるソリューションが求められます。

どの3つの行動を順番に実行すべきでしょうか？回答するには、行動リストから適切な行動を回答欄に移動させ、正しい順序に並べ替えてください。

Actions

Regenerate the primary admin key

Regenerate the secondary admin key

Change the app to use the secondary admin key

Add a new query key

Change the app to use the new key

Delete the compromised key

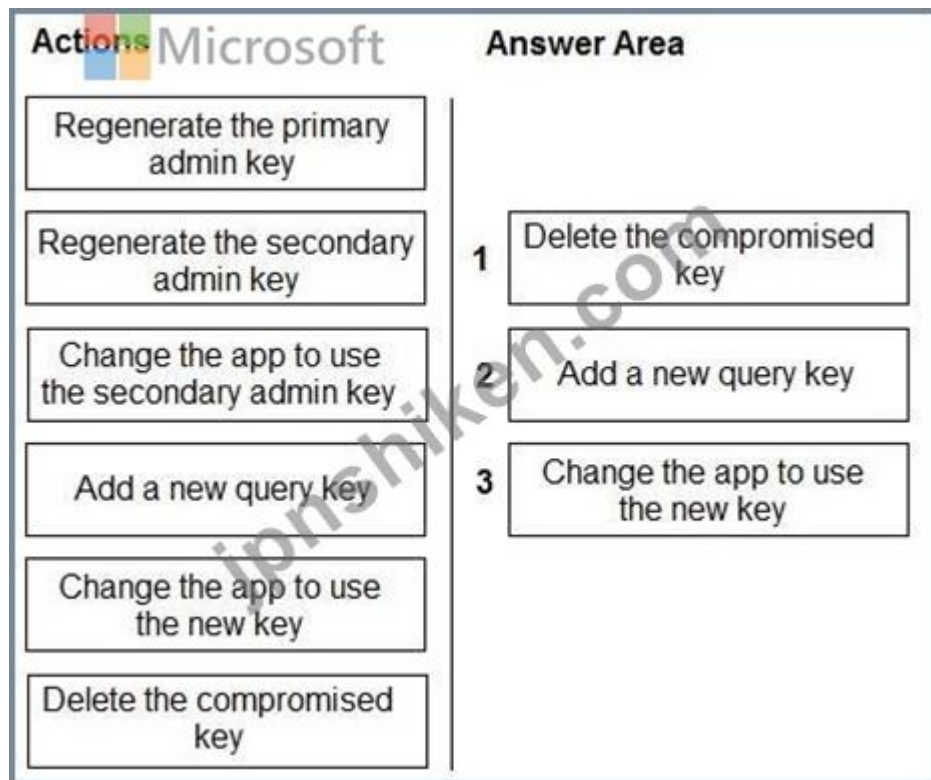
Answer Area

1

2

3

正解:



Explanation:

Enforces Read-Only Permissions: Query keys are specifically designed to provide read-only access to the documents collection of an index. Admin keys provide full read-write administrative privileges and should never be distributed to consumer-facing applications.

Zero Downtime: Azure AI Search lets you generate up to 50 individual query keys. Creating a new one allows the app to stay online throughout the entire key rotation process

Reference:

<https://learn.microsoft.com/en-us/azure/search/search-security-api-keys>

質問: 29

Foundry Tools の Azure Constant Understanding を使用する App1 という名前の請求書処理アプリケーションがあります。

あなたは、以下の要件を満たす必要がある、Pipeline1という名前の新しいコンテンツ理解パイプラインを構築しています。

請求書とそれに関連する発注書を比較する

- 音声を静的なベンダー連絡先文書と照合して検証する
- 不一致箇所を含む単一の構造化出力を返す

Pipeline1を設定し、パイプラインを単一のアナライザーエンドポイントとして公開する必要があります。何を設定すればよいでしょうか？

- A. 標準モードで提供されるベンダー契約書を分析中の追加文書として使用する単一ファイルタスク。
- B. 抽出されたフィールドに対して信頼度スコアが有効になっている標準モードの単一ファイルタスク。
- C. ベンダー契約ファイルを参考データとして使用するプロモードの複数ファイルタスク
- D. 請求書と発注書をアナライザーへの入力として使用する標準モードの複数ファイルタスク

正解: C ([コメントを发表する](#))

Multiple-file task: Required over a single-file task. Your pipeline needs to reconcile data across two separate active transactional documents (the invoice and the purchase order) within a single analyzer request.

Pro mode is required instead of standard mode. Pro mode is specifically designed for advanced scenarios requiring multi-step reasoning, cross-file analysis, and validation against a knowledge base. Note that Pro mode currently supports classify and generate fields but does not support confidence scores for specific extracted fields.

Vendor contract files as reference data: Required. Because the vendor contracts are static compliance documents, they should be uploaded and treated as the analyzer's background knowledge base (reference data) to guide the validation logic.

Reference:

<https://learn.microsoft.com/en-us/azure/ai-services/content-understanding/concepts/standard-pro-modes>

質問: 30

Agent1という名前のカスタムエージェントが作成されました。

Microsoft Foundryを使用して、Agent1へのアクセスを制御し、アクティビティを監視する必要があります。

まず最初に何をすべきでしょうか？

A. Application Insights リソースをプロビジョニングします。

B. Microsoft Foundry プレイグラウンドに Agent1 を追加します。

C. Microsoft Foundry プロジェクトに Agent1 を追加します。

D. Microsoft Foundry プロジェクトを作成します。

正解: ([正解を表示します](#))

To monitor and control access to a custom agent in Azure, you must first create a Microsoft Foundry project. Once the project is created, you register your custom agent within it to enable management capabilities such as access control and activity monitoring.

Reference:

<https://learn.microsoft.com/en-us/azure/foundry/control-plane/register-custom-agent>

質問: 31

注 :このセクションには、同じシナリオと問題に関する複数の質問セットが含まれています。各質問には、問題に対する固有の解決策が提示されています。提示された解決策が、提示された目標を満たしているかどうかを判断する必要があります。セット内の複数の解決策が問題を解決できる場合もあります。また、セット内のどの解決策も問題を解決できない場合もあります。

このセクションの質問に回答すると、前のセクションに戻ることはできません。そのため、これらの質問は復習画面には表示されません。

エージェントを含む Microsoft Foundry プロジェクトがあります。このエージェントは、取得したポリシー文書から概要を生成します。

ユーザーからは、取得したコンテンツに規定の規制条項が含まれている場合でも、一部の回答では必要な規制条項が省略されているとの報告がある。

回答の完全性を改善する必要があります。

解決策：回答の完全性を採点し、定義されたしきい値を下回る回答をブロックする評価フローを実行します。

これは目標を達成していると言えるでしょうか？

A. はい

B. いいえ

正解: ([正解を表示します](#))

Correct:

* You add a reflection pass that regenerates the response if the required clauses are missing.

This is Self-Correction Strategy: A reflection pass allows an agent to evaluate its own initial output against specified constraints (e.g., checking for the presence of mandatory regulatory clauses). If the required text is missing, the agent triggers a programmatic self-correction or regeneration loop to include them before final delivery.

Incorrect:

* You increase the value of the max_tokens parameter.

Increasing the max_tokens parameter prevents the response from being cut off mid-sentence due to length constraints. However, it does not force the model's logic to explicitly include missing information that it chose to leave out earlier in the text.

* You increase the value of the temperature parameter.

Raising the temperature parameter increases randomness and creativity. For rigid compliance tasks like summarizing regulatory documents, higher temperature actually increases the risk of hallucination and omission.

* You run an evaluation flow that scores responses for completeness and blocks responses that fall below a defined threshold.

Evaluation Flow Block: Running an evaluation flow to score completeness and blocking bad responses identifies and stops low-quality outputs, but it does not fix or actively improve the response completeness. It simply filters failures out of the system.

Reference:

<https://pub.towardsai.net/reflection-with-llm-how-to-make-ai-review-its-own-work-2db122fca1d8>

有効的な**AI-103-JPN**問題集はJPNTTest.com提供され、**AI-103-JPN**試験に合格することに役に立ちます！JPNTTest.comは今最新**AI-103-JPN**試験問題集を提供します。JPNTTest.com AI-103-JPN試験問題集はもう更新されました。ここで**AI-103-JPN**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/AI-103-JPN-mondaishu> **159問、30%ディスカウント**、特別な割引コード: **JPNshiken**」

質問: **32**

あなたは、複数のエージェントを含む Project1」という名前のMicrosoft Foundryプロジェクトを計画しています。

各エージェントは同じAzure AI Searchリソースにアクセスします。

Project1内でAzure AI Searchの認証情報を一元管理するためのソリューションを提案してください。このソリューションは、すべてのエージェントに実装する必要があります。何をおすすめしますか？

- A.** Azure AI Search リソースに対してロールベースのアクセス制御 (RBAC) を有効にします。
- B.** Azure AI Search リソースのキーベースのアクセス制御を無効にします。
- C.** Azure AI Search リソースへの接続を追加します。
- D.** Azure AI Search リソースに接続するマネージド プライベート エンドポイントを作成します。

正解: ([正解を表示します](#))

To best manage security and centrally handle credentials across multiple agents, you should add a connection to the Azure AI Search resource at the Azure AI Foundry project level.

Why This Works

Central Hub: The project acts as the single security perimeter for all your agents.

Credential Masking: Agents inherit access without hardcoding secrets, API keys, or connection strings in their code.

Identity Management: It allows you to leverage Microsoft Entra ID (formerly Azure AD) for role-based access control (RBAC).

How to Implement It

1. Navigate to your Azure AI Foundry portal.
2. Select your specific project from the dashboard.
3. Open the "Management Center" or "Project settings" tab.
4. Click on "Connected resources" or "Connections".
5. Add the Azure AI Search resource.
6. Choose Entra ID (managed identity) over API keys for maximum security.

Reference:

<https://partner.microsoft.com/en-us/blog/article/azure-updates-december-2025>

質問: 33

あなたはコンテンツ管理システムを設計しています。

読解力や学習能力に障害のあるユーザー(例えば、失読症など)にとって最適な読書体験を提供できるよう、最適化されたソリューションを構築する必要があります。また、開発の手間を最小限に抑えることも重要です。

ソリューションにはどのAzureサービスを含めるべきですか？

- A. Foundry Tools の Azure Document Intelligence
- B. Foundry ToolsにおけるAzure言語
- C. Azure AI Immersive Reader
- D. Foundry Tools の Azure Translator

正解: [\(正解を表示します\)](#)

Include Azure AI Immersive Reader in your solution. It is an applied AI service that provides built- in accessibility features like read-aloud (text-to-speech), real-time translation, line focusing, and customizable typography. It requires no machine learning expertise and integrates easily to minimize development effort.

Reference:

<https://azure.microsoft.com/en-us/products/ai-services/ai-immersive-reader>

質問: 34

ドラッグアンドドロップ問題

あなたは、Foundry Tools の Azure Content Understanding を使用してマーケティング動画を分析する Microsoft Foundry プロジェクトを持っています。

動画のセグメンテーションが有効になっています。

各ビデオセグメントの配色を記述するJSONフィールドを出力するようにアナライザーを設定する必要があります。

アナライザーの設定方法を教えてください。適切な値を正しいターゲットにドラッグして設定します。各値は、1回、複数回、または全く使用しない場合があります。コンテンツを表示するには、ペイン間の分割バーをドラッグするか、スクロールする必要がある場合があります。

注：正解ごとに1ポイントが加算されます。

Microsoft
Answer Area

Values

classify

generate

group

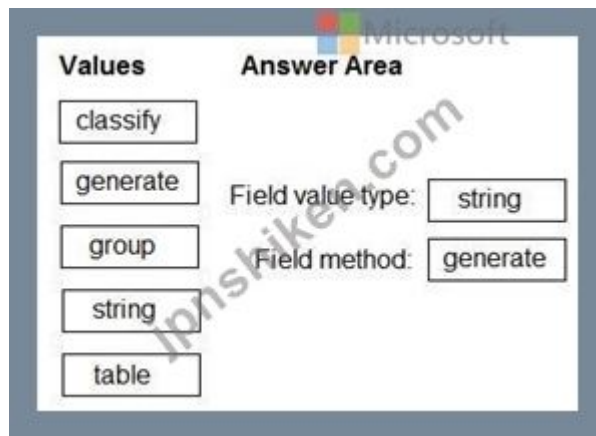
string

table

Field value type:

Field method:

正解:



Explanation:

Box 1: string

The Type (string): While you want the output schema to contain a description of the color palette or style, Azure Content Understanding schemas map the underlying data types using fundamental types like string, number, or array. The field will output the final result in the JSON response under a key like "valueString" Box 2: generate

The Method (generate): The extract method is only used for pulling literal text exactly as it appears in content (such as speech-to-text transcripts or OCR). Because analyzing visual aesthetics and synthesizing descriptive text about a "color scheme" requires multimodal AI reasoning, you must use the generate method. This instructs the underlying model to analyze the segment and synthesize a completely new descriptive insight.

References:

<https://learn.microsoft.com/en-us/azure/ai-services/content-understanding/video/elements>

質問: 35

ホットスポットに関する質問

Microsoft Foundry Agent Service を使用してカスタマーサポートエージェントを作成するための計画を提案する必要があります。エージェントは以下の要件を満たす必要があります。

- 複数の会話間でユーザーの設定を保持する。
 - ユーザーが直接アップロードすることで、文脈的な根拠を提供できるようにする
- チャット中にドキュメントを閲覧する。

それぞれの要件に対して、どのFoundry機能を推奨すべきでしょうか？回答するには、回答欄で適切なオプションを選択してください。

注：正解ごとに1ポイントが加算されます。

Answer Area



Microsoft

To retain user preferences across conversations, use:

Agent memory that uses persistent storage
Conversation history
Orchestration-managed session context

To enable users to provide contextual grounding during chats, use the:

Azure AI Search tool
Code interpreter tool
File search tool

正解:

Answer Area

Microsoft

To retain user preferences across conversations, use:

Agent memory that uses persistent storage
Conversation history
Orchestration-managed session context

To enable users to provide contextual grounding during chats, use the:

Azure AI Search tool
Code interpreter tool
File search tool

質問: 36

ホットスポットに関する質問

PaymentAgent という名前のエージェントを含む Microsoft Foundry プロジェクトがあります。

PaymentAgentには、外部APIを使用して顧客への払い戻しを行う機能ツールが含まれています。

YAML形式でワークフローを作成しています。

ワークフローが人間の承認のために一時停止し、承認が得られた後にのみ返金処理に進むようにする必要があります。

ワークフロー定義はどのように完成させるべきですか？回答するには、回答欄で適切なオプションを選択してください。

注：正解ごとに1ポイントが加算されます。

Answer Area



steps:

- id: propose_refund
type: agent
agent: PaymentAgent
- id: approval

type:

ask_question
basic_chat
data_transformation

id: execute_refund

type: agent

agent: PaymentAgent

condition:

approval == "approved"
propose_refund.output != null
true

正解:

Answer Area

Microsoft

steps:

- id: propose_refund

type: agent

agent: PaymentAgent

- id: approval

type:

ask_question

basic_chat

data_transformation

- id: execute_refund

type: agent

agent: PaymentAgent

condition:

approval == "approved"

propose_refund.output != null

true

質問: 37

Microsoft Foundry プロジェクトにチャット アプリがあり、Azure AI Search のベクトル化インデックスがあります。

以下の要件を満たすには、インデックスに接続する必要があります。

複雑な質問では、複数の情報チャンクから情報を取得する必要があります。

複数ターンにわたる会話は、情報検索計画に影響を与える必要がある。

遅延を削減するため、データ取得は並列で実行する必要があります。

どの検索方法を用いるべきでしょうか？

- A. 反復検索
- B. エージェント型検索拡張生成(RAG)
- C. 思考の連鎖
- D. 古典的な検索拡張生成(RAG)

正解: [\(正解を表示します\)](#)

An agentic Retrieval-Augmented Generation (RAG) architecture is the best fit for your requirements. Traditional, linear RAG pipelines struggle with complex, multi-turn dependencies, whereas an agentic approach natively solves them through iterative reasoning and tool orchestration.

Here is how an agentic RAG pattern specifically satisfies each the three technical requirements:

1. Complex Questions & Multi-Chunk Retrieval

The Challenge: Complex questions often require aggregating distinct pieces of information scattered across different documents or sections.

The Agentic Solution: An agent uses a Reasoning and Acting (ReAct) loop. It evaluates the user's query, breaks it down into sub-questions, and executes multiple distinct search queries. It then synthesizes the information from these diverse chunks before generating a final response.

2. Multi-Turn Conversations & Retrieval Planning

The Challenge: Standard RAG often just passes the latest user message or a simple chat history summary to the search index, which can lose context or misinterpret the user's true intent in long conversations.

The Agentic Solution: The agent acts as a query planner. It maintains the state of the conversation and dynamically decides if a retrieval is needed, what specific keywords or vectors to target based on past turns, and how to reformulate the query to bridge context gaps.

3. Parallel Retrievals for Low Latency

The Challenge: Sequential lookups drastically increase time-to-first-token (TTFT), harming the chat user experience.

The Agentic Solution: Modern AI agent frameworks allow the LLM to emit multiple tool calls simultaneously. The application layer can intercept these calls and execute parallel asynchronous queries against your Azure AI Search vectorized index, drastically minimizing latency.

Reference:

<https://arxiv.org/html/2501.09136v1>

質問: 38

エージェントを含む Microsoft Foundry プロジェクトがあります。

エージェントの知識源は、Azure Blob Storageに保存されているスキャン済みのPDFトラブルシューティングガイドのセットです。ガイドページは2列レイアウトと表で構成されています。

PDFを処理するには、Foundry ToolsのAzure Content Understandingを使用します。

処理済みのコンテンツをインデックスに取り込み、検索拡張生成(RAG)に利用し、抽出したフィールドを後続の自動化のために保存する予定です。

関係者は、抽出された各フィールド値が元のPDFのどこから来たのかを確認できなければならず、信頼性の低い抽出結果は手動レビューのために回送する必要がある。

コンテンツ理解ドキュメントアナライザーの出力には、フィールドごとの信頼度スコアと、ソースドキュメント内の位置を示すソース情報が含まれていることを確認する必要があります。

あなたはどうすべきでしょうか？

A. enableSegmentをtrueに設定します。

B. ラベル付きサンプルを提供する。

C. estimateFieldSourceAndConfidence を有効にします。

D. アナライザーを、すべてのフィールドに対して生成抽出を使用するように設定します。

正解: ([正解を表示します](#))

To fulfill all your requirements using Azure Content Understanding in Foundry Tools, you need to configure a custom document analyzer with specific flags, set up an index ingestion pipeline, and build a downstream human-in-the-loop validation rule.

*-> 1. Enable Confidence Scores and Source Grounding

To force the analyzer to provide per-field confidence metrics and precise layout/bounding box coordinates for verification, you must opt-in to the estimateFieldSourceAndConfidence parameter within your configuration.

Option A (Global): Set estimateFieldSourceAndConfidence = true in the main analyzer config to evaluate all fields.

Option B (Field-Level): Set estimateSourceAndConfidence = true under individual field schemas. This ensures the generated JSON response populates the bounding box coordinates, page numbers, and a confidence score 0.0 to 1.0 for every extracted entity.

2. Configure Document Extraction for Two-Column & Table Layouts

3. Build the Ingestion Pipeline (RAG vs. Automation Dual-Path)

4. Implement Threshold Routing and Source Verification

Reference:

<https://learn.microsoft.com/en-us/azure/ai-services/content-understanding/document/overview>

質問: 39

2026年半ばに、Azure上で全く新しいエージェントアプリケーションの開発を開始し、Foundryエージェントの一般提供されサポートされているエントリーポイントを利用したいと考えています。どのAPIをターゲットにすべきでしょうか？

- A. アシスタントAPI
- B. Foundry Agent Service を介したレスポンス API
- C. 従来のテキスト補完API
- D. LUISランタイムAPI

正解: ([正解を表示します](#))

The Responses API is the generally available single entry point for Foundry Agent Service and supports both prompt agents and hosted agents. It is the API Microsoft directs new agentic development towards.

質問: 40

あなたは、Agent1 という名前のエージェントを含む Microsoft Foundry プロジェクトを持っています。

エージェントは正常に実行されますが、Foundry Control Planeにはエラー率、実行回数、トークン使用量の値が表示されず、トレースタブも空です。

「検出されたコントロールプレーン」に、Agent1 の適切な値が表示されていることを確認する必要があります。

あなたはどうすべきでしょうか？

- A. Agent1を新しいバージョンにアップデートしてください。
- B. Foundry Control Planからエージェントを再起動する
- C. エージェント1にログ分析ワークスペースを割り当てます。
- D. Agent1 の Application Insights を有効にします。

正解: ([正解を表示します](#))

To resolve this issue, you must connect and configure an Azure Application Insights resource for your Microsoft Foundry project.

The Foundry Control Plane, its Agent Monitoring Dashboard, and the Traces tab rely directly on telemetry data stored within the connected Application Insights instance. If this resource is missing, unlinked, or improperly configured, the dashboard cannot display runs, error rates, token usage, or transaction spans.

Reference:

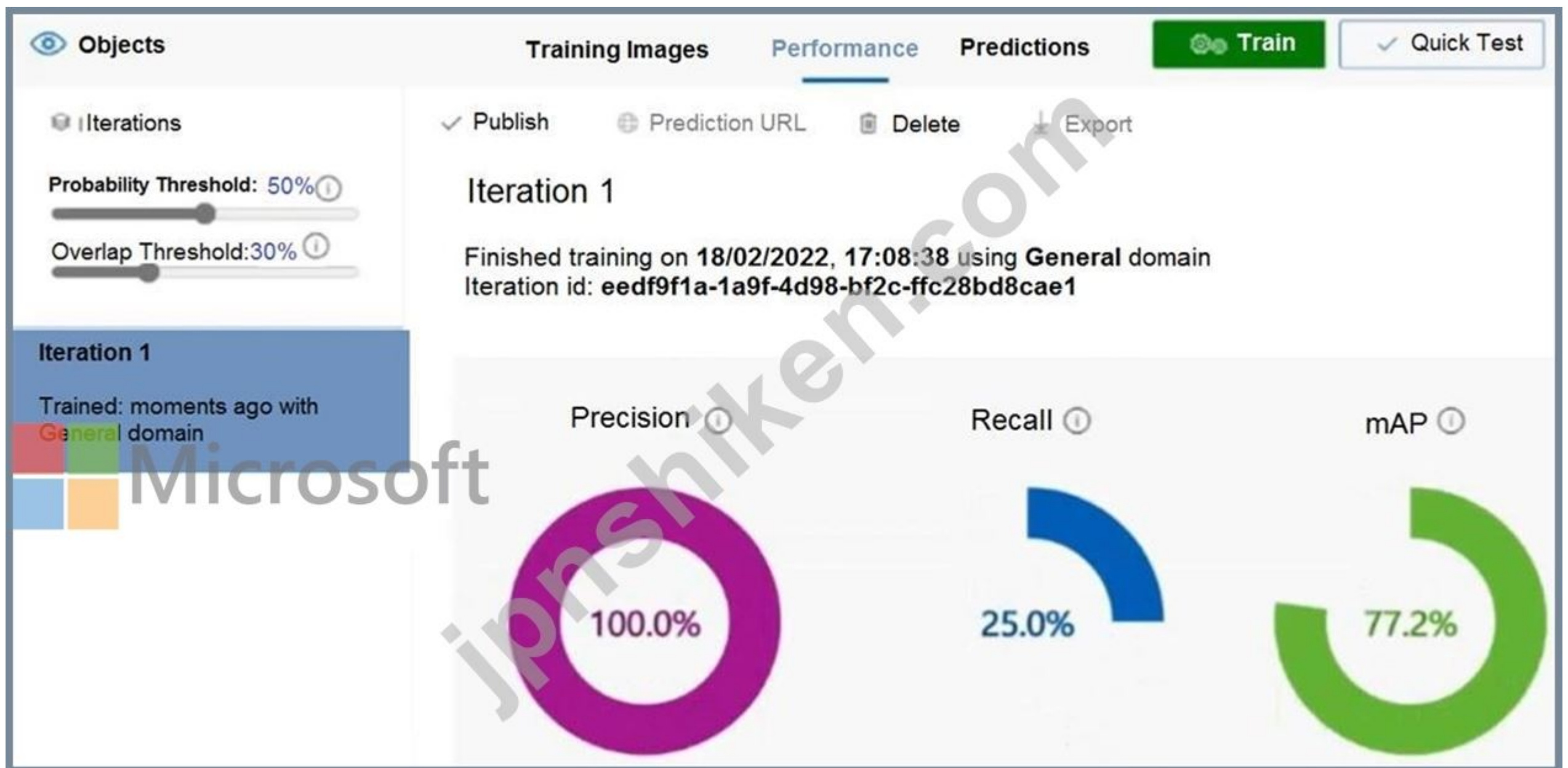
<https://techcommunity.microsoft.com/blog/azure-ai-foundry-blog/monitoring--observability-in-microsoft-foundry/4517250>

質問: 41

ホットスポットに関する質問

あなたは画像内の物体を検出するモデルを構築しています。

訓練データに基づくモデルの性能を以下の図に示します。



図に示された情報に基づいて、各記述を完成させる選択肢をドロップダウンメニューを使用して選択してください。

注：正解ごとに1ポイントが加算されます。

Answer Area



The percentage of false positives is:

▼
0
25
30
50
100

The value for the number of true positives divided by the total number of true positives and false negatives is:

▼
0
25
30
50
100

正解:

Answer Area



The percentage of false positives is:

▼
0
25
30
50
100

The value for the number of true positives divided by the total number of true positives and false negatives is:

▼
0
25
30
50
100

Explanation:

Box 1: 0

The percentage of false positives is 0%.

Because the model made zero incorrect positive predictions, the count and percentage of false positives must be exactly zero.

Box 2: 25

The value for the number of true positives divided by the total number of true positives and false negatives is 25% (or 0.25).

High Precision (100%): Every single object your model detects is correct; it generates zero false alarms.

Low Recall (25%): Your model misses 75% of the actual objects it was supposed to find.

Reference:

<https://medium.com/grabngoinfo/how-to-evaluate-the-performance-of-a-binary-classification-model-6e7193dcbbf9>

質問: 42

Azure AI Studioでは、GPT-35 Turboモデルを使用してCompletionsプレイグラウンドを利用します。
プロンプトに以下のコードが含まれています。

```
function F(n)
{
var f = [0, 1];
for (var i = 2; i < n; i++) f[i] = f[i-1] + f[i-2];
return f;
}
```

コードの説明を作成するにはモデルが必要です。解決策はコストを最小限に抑えるものでなければなりません。
あなたはすべきでしょうか？

- A. モデルを GPT-4-32k に変更します。
- B. プロンプトに // 関数 F は何をしますか? を追加します。
- C. プロンプトに関数 F(説明) を追加します。
- D. 温度パラメータを1に設定します。

正解: [\(正解を表示します\)](#)

By adding a comment like // what does function F do? to the prompt, you are instructing the GPT model to provide an explanation of the function's behavior. This approach directly asks the model to explain the code, which will help generate an explanation without requiring changes to the model or additional resource-intensive configurations.

質問: 43

ドラッグアンドドロップ問題

あなたは、工場生産ラインで製造された不良部品を検出するアプリケーションを開発しています。これらの部品は、あなたのビジネスに特有のものです。
一般的な障害を検出するには、Azure Custom Vision API を使用する必要があります。
どの3つの行動を順番に実行すべきでしょうか？回答するには、行動リストから適切な行動を回答欄に移動させ、正しい順序に並べ替えてください。

Actions

Initialize the training dataset.

Train the classifier model.

Create a project.

Upload and tag images.

Train the object detection model.

Answer Area

1

2

3

正解:

Actions

Initialize the training dataset.

Train the classifier model.

Create a project.

Upload and tag images.

Train the object detection model.

Answer Area

1 Create a project.

2 Upload and tag images.

3 Train the object detection model.

Explanation:

Step 1: Create a project

You must first set up a new project in the Custom Vision portal or via the SDK to store your data and configuration.

Step 2: Upload and tag images

Next, you need to upload photos of your specific factory components and draw bounding boxes around the faults to tag them.

Step 3: Train the object detection model

Because you need to identify and locate specific, localizable faults on a component, you must train an object detection model rather than a general classifier model.

Reference:

<https://thegroundtruth.blog/tag/computervision/>

質問: 44

Azure Blob Storageに保存されている請求書を取り込む、Pipeline1という名前のAzure AI Searchインデックス作成パイプラインを構築しています。請求書はスキャンされた画像として保存されています。

ユーザーが請求書の各項目を横断して請求書データを検索できるようにする必要があります。

Pipeline1のスキルセットに、どの組み込みスキルを追加すべきでしょうか？

A. テキスト分割

B. テキスト翻訳

C. 光学文字認識(OCR)

D. 画像解析

正解: C ([コメントを发表する](#))

You should add the Azure Machine Learning OCR skill (historically known as the Microsoft.Skills.Vision.OcrSkill) to your skillset pipeline.

Text Extraction: Scanned images contain unsearchable pixels rather than text.

Data Ingestion: The OCR skill extracts the printed text and layout information from the images.

Downstream Processing: It feeds raw text into subsequent skills (like Entity Recognition) or directly into your search index fields.

Reference:

<https://techcommunity.microsoft.com/blog/azurearchitectureblog/from-large-semi-structured-docs-to-actionable-data-reusable-pipelines-with-adi-a/4474054>

質問: 45

カスタマーサポート担当者は、毎週新しい製品ドキュメントで更新される社内ナレッジベースに基づいて回答する必要があり、その回答は常に最新の内容を反映していなければなりません。どちらのアプローチを採用すべきでしょうか？

- A. 毎週新しいドキュメントに基づいて基本モデルを微調整する
- B. 文書のインデックスに対して検索拡張生成(RAG)を実装する
- C. モデルの温度を上げて、最近の事実の想起を改善する
- D. 文書に直接選好最適化(DPO)を適用する

正解: **B** ([コメントを發表する](#))

RAG injects fresh, frequently changing knowledge at inference time by retrieving relevant content from an index and grounding the model's answer in it, which is exactly the "chat with your data" pattern this scenario describes. Microsoft's guidance is direct on the boundary between the two techniques.

Use fine-tuning when you need to change model behavior, style, or task performance, rather than add fresh knowledge.

質問: 46

マルチモーダルエージェントは、ユーザーがアップロードした画像を処理します。攻撃者は、エージェントがそれを読み取って実行することを期待して、画像内に隠しテキスト指示を埋め込みます。Azure AI Content Safety のどの機能が、このような攻撃を検出するように設計されていますか？

- A. 暴力カテゴリのコンテンツフィルタが「高」に設定されています
- B. プロンプトシールドは、文書および間接的なプロンプト挿入を検出します。
- C. アナライザーの画像解像度しきい値を上げる
- D. エージェントのファイル検索ツールを無効にする

正解: ([正解を表示します](#))

Prompt Shields detects prompt injection attacks, including indirect injection where malicious instructions are embedded in content the model processes, such as text hidden inside an image.

有効的な**AI-103-JPN**問題集はJPNTTest.com提供され、**AI-103-JPN**試験に合格することに役に立ちます！JPNTTest.comは今最新**AI-103-JPN**試験問題集を提供します。JPNTTest.com AI-103-JPN試験問題集はもう更新されました。ここで**AI-103-JPN**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/AI-103-JPN-mondaishu> **159問、30%ディスカウント**、特別な割引コード: **JPNshiken**」

質問: 47

App1 という名前のアプリケーションがあり、Foundry Tools の Azure Speech を使用してライブ通話を文字起こししています。

トランスクリプトには、英語とスペイン語の両方が含まれていることがよくあります。App1 は、各セグメントを Foundry Tools の Azure Translator に送信し、別の言語に翻訳します。

場合によっては、複数の言語が混在する部分では、翻訳が不完全または不正確になることがあります。

翻訳エラーを減らす必要があります。解決策は、文字起こし全体が確実に正しく翻訳されることを保証するものでなければなりません。

翻訳ツールにセグメントを送信する前に、何をすべきですか？

- A. 文書翻訳を使用して、トランスクリプト全体を単一の文書として翻訳します。
- B. 複数の言語が混在するセグメントを単一言語のセグメントに分割し、各セグメントを個別に翻訳します。
- C. 翻訳依頼の自動言語検出を有効にします。
- D. すべてのセグメントの翻訳依頼において、原文言語として英語を指定してください。

正解: ([正解を表示します](#))

To fix incomplete or incorrect translations from mixed-language transcripts, you must identify and isolate the languages at the sentence or phrase level before sending the text to Azure Translator.

Azure Translator performs best when a single request contains only one source language. When it receives a mixed-language segment under a single source language code, it often fails to parse the secondary language correctly.

Here is the step-by-step pipeline you should implement inside your application before hitting the Translator API.

1. Split Segments into Sentences

Live call transcript segments can contain multiple sentences. Do not send the raw, multi-sentence segment directly to the translator if it contains mixed languages. Break the segment into individual sentences. Use regex punctuation rules or a lightweight sentence-splitting library.

2. Run Language Detection Per Sentence

Azure Translator has a built-in Detect API, but for mixed-language live calls, executing a dedicated detection step per sentence provides better control.

Reference:

<https://medium.com/neural-engineer/azure-ai-speech-to-text-real-time-transcription-58bfd5fd1a28>

質問: 48

エージェントを含む Microsoft Foundry プロジェクトがあります。このエージェントは、取得したポリシー文書から概要を生成します。

応答の完全性を改善する必要があります。この解決策は、応答が返される前に、アプリケーションコードのロジックに実装されなければなりません。

あなたはすべきでしょうか？

A. レスポンスを返す前に再試行評価を追加します。

B. max_tokens パラメータの値を減らします。

C. Retrieval Augmented Generation (RAG) に切り替えます。

D. より小規模な展開でモデルを置き換えます。

正解: ([正解を表示します](#))

To enhance response completeness in your Microsoft Foundry agent, you must intercept the retrieved documents and the generated summary within your backend application logic before returning the payload to the user.

1. Implement Completeness Verification Logic

Add a verification step in your orchestration code (e.g., in your Python/Semantic Kernel or LangChain pipeline) that compares the generated summary against the retrieved chunks.' Map Key Assertions: Extract main policy rules from retrieved text. Cross-Reference Entities: Verify all key entities are in the summary.

Check Scope Coverage: Ensure every retrieved document is represented.

Scan for Gaps: Identify critical missing constraints or exceptions.

2. Apply Application-Level Mitigation Strategies

If the verification step detects that the summary is incomplete, use your code to correct it before the final response leaves your system.

Reference:

<https://dev.to/moonrunnerkc/how-i-built-a-verification-layer-for-copilot-clis-multi-agent-output-4b7h>

質問: 49

あなたは、社内文書の検索に Azure AI Search を使用するアプリケーションを開発しています。

Azure AI Search では、ドキュメントレベルのフィルタリングを実装する必要があります。

解決策に含めるべき3つの行動はどれですか？それぞれの正解は、解決策の一部を示しています。

注：正解ごとに1ポイントが加算されます。

A. 各インデックスエントリに許可されたグループを追加する

- B. グループごとに1つのインデックスを作成します。
- C. 検索リクエストとともに、Microsoft Entra ID からアクセストークンを送信します。
- D. すべてのグループを取得します。
- E. ユーザーのグループメンバーシップを取得します
- F. 検索リクエストのフィルタとしてグループを提供する

正解: ([正解を表示します](#))

[A]

Add allowed groups to each index entry: Each document in your search index must include a field (e.g., group_ids) listing the security groups or principals that have permission to view that document.

[E]

Retrieve the group memberships of the user: When a user submits a query, your application must query Microsoft Entra ID to fetch the list of all security groups the user belongs to.

[F]

Supply the groups as a filter for the search requests: Your application will then translate the user's group memberships into an OData filter (using the search.in() function) and append it to the search API query, restricting the results only to documents the user has access to.

Reference:

<https://learn.microsoft.com/en-us/azure/search/search-security-trimming-for-azure-search>

質問: 50

あなたは、GPTリアルタイムモデルを使用するApp1という名前の顧客サポートWebアプリをMicrosoft Foundryで構築しています。
アプリ1は以下をサポートする必要があります：

- Azure OpenAI を使用した、リアルタイムで低遅延の音声会話
- ユーザーからのストリーミング音声入力と再生音声応答

クライアントアプリケーションでリアルタイムオーディオストリーミングをサポートし、レイテンシが約100ミリ秒を目標とする接続方法を設定する必要があります。
どの接続方法を使用すべきですか？

- A. RTMP
- B. WebRTC
- C. SIP
- D. WebSocket

正解: D ([コメントを发表する](#))

WebSockets is the best connection method for this application because it enables full-duplex, bi- directional streaming over a single TCP connection, meeting the strict ~100 ms latency requirement for real-time audio.

Reference:

<https://www.chat-data.com/blog/implement-openai-realtime-api-for-chatgpt-voice>

質問: 51

お客様は、AI1 という名前の Azure OpenAI リソースを含む Azure サブスクリプションをお持ちです。
AI1を使用して特定の質問に対して生成的な回答を提供するチャットボットを構築します。
組み込みの安全機能を回避することを目的とした質問は、確実にブロックされるようにする必要があります。
Azure AIのコンテンツ安全機能のうち、どれを実装すべきでしょうか？

- A. 中程度のテキストコンテンツ
- B. 保護対象テキストの検出
- C. プロンプトシールド
- D. オンライン活動を監視する

正解: **C** ([コメントを发表する](#))

Prompt Shields in Azure AI Content Safety are specifically designed to detect and block prompts or questions that attempt to bypass safety features in generative AI models, such as those used in chatbots. This feature ensures that the chatbot adheres to safety and ethical guidelines by preventing inappropriate or harmful content generation.

質問: **52**

あなたは薬局向けの音声エージェントを開発しています。音声による質問を受け付け、合成音声で応答し、一般的な音声モデルでは聞き間違いやすい専門的な医薬品名を正確に認識する必要があります。どの機能を構成すべきでしょうか？

- A. エージェント内の音声テキスト変換のためのカスタム音声モデル
- B. デフォルトのテキスト読み上げ音声のみ
- C. Azure Translatorを使用して音声を変換します
- D. 音声ストリームに対するAzure AI Document Intelligence

正解: ([正解を表示します](#))

A custom speech model improves recognition of specialised vocabulary, such as medicine names, by adapting speech-to-text (STT) to your domain, and AI-103 covers integrating speech, including custom speech models, as an agent modality. This directly addresses the accuracy problem in the scenario.

質問: **53**

副操縦士は、会話の以前のやり取りに依存する複数の質問に答える必要があります。システムが複雑な質問をそれぞれ焦点を絞ったサブクエリに分割し、それらを並列実行し、回答を作成する前に結果を意味的に再ランク付けするようにしたい場合、Azure AI Search のどの機能がこれを提供しますか？

- A. スコアリングプロファイルを使用したクラシックなキーワード検索
- B. エージェントによる検索
- C. 高い類似度閾値を持つ単一ベクトルクエリ
- D. 単一キーワードクエリにおけるトップk値の増加

正解: **B** ([コメントを发表する](#))

Agentic retrieval in Azure AI Search uses a large language model to decompose a complex query into focused subqueries, runs them in parallel, semantically reranks each set of results, and merges them into a unified response, while taking conversation history into account. That is precisely the behaviour the scenario describes.

質問: **54**

注 :このセクションには、同じシナリオと問題に関する複数の質問セットが含まれています。各質問には、問題に対する固有の解決策が提示されています。提示された解決策が、提示された目標を満たしているかどうかを判断する必要があります。セット内の複数の解決策が問題を解決できる場合もあります。また、セット内のどの解決策も問題を解決できない場合もあります。

このセクションの質問に回答すると、前のセクションに戻ることはできません。そのため、これらの質問は復習画面には表示されません。

エージェントを含む Microsoft Foundry プロジェクトがあります。このエージェントは、取得したポリシー文書から概要を生成します。

ユーザーからは、取得したコンテンツに規定の規制条項が含まれている場合でも、一部の回答では必要な規制条項が省略されているとの報告がある。

回答の完全性を改善する必要があります。

解決策：必要な句が欠落している場合にレスポンスを再生成するリフレクションパスを追加します。

これは目標を達成していると言えるでしょうか？

A. はい

B. いいえ

正解: (正解を表示します)

Correct:

* You add a reflection pass that regenerates the response if the required clauses are missing.

This is Self-Correction Strategy: A reflection pass allows an agent to evaluate its own initial output against specified constraints (e.g., checking for the presence of mandatory regulatory clauses). If the required text is missing, the agent triggers a programmatic self-correction or regeneration loop to include them before final delivery.

Incorrect:

* You increase the value of the max_tokens parameter.

Increasing the max_tokens parameter prevents the response from being cut off mid-sentence due to length constraints. However, it does not force the model's logic to explicitly include missing information that it chose to leave out earlier in the text.

* You increase the value of the temperature parameter.

Raising the temperature parameter increases randomness and creativity. For rigid compliance tasks like summarizing regulatory documents, higher temperature actually increases the risk of hallucination and omission.

* You run an evaluation flow that scores responses for completeness and blocks responses that fall below a defined threshold.

Evaluation Flow Block: Running an evaluation flow to score completeness and blocking bad responses identifies and stops low-quality outputs, but it does not fix or actively improve the response completeness. It simply filters failures out of the system.

Reference:

<https://pub.towardsai.net/reflection-with-llm-how-to-make-ai-review-its-own-work-2db122fca1d8>

質問: 55

エージェントを含む Microsoft Foundry プロジェクトがあります。

エージェントは、表と埋め込まれたQRコードを含む、スキャンされたPDF形式の仕入先請求書を取り込みます。

エージェントは、抽出された出力においてPDFのレイアウトを保持し、後続の処理でセクションや表を参照できるようにする必要があります。

Foundry Tools で Azure Content Understanding を呼び出す予定です。

言語モデルの展開を必要とせずに、コンテンツとレイアウト要素を抽出し、QRコードを検出する必要があります。

どの内蔵アナライザーを使用すべきですか？

A. 事前構築済みドキュメントフィールドスキーマ

B. プリビルド読み取り

C. 事前構築済みドキュメント検索

D. 事前構築済みレイアウト

正解: D (コメントを发表する)

To extract content, preserve tables and document layout, and detect embedded QR codes without deploying a large language model (LLM), you should use the built-in prebuilt-layout analyzer.

Note:

Unlike schema-driven extraction models in Content Understanding that utilize generative AI orchestration, the Layout analyzer is a highly efficient machine-learning-based model. It natively outputs structural geometry and decodes barcodes without requiring an active LLM deployment or provisioned throughput.

Structural Preservation: It extracts headers, paragraphs, and nested sections, returning precise spatial bounding boxes for every single element. Downstream agents can utilize this geometric metadata to anchor or cross-reference sections accurately.

Advanced Table Mapping: It maps intricate, multi-page invoice tables, capturing text alongside row and column indices. You can configure the output structure format natively into Markdown or HTML tables to maintain formatting cleanliness.

Built-in QR and Barcode Decoding: By default, the configuration parameter enableBarcode is set to true. The analyzer scans the scanned PDF image, isolates 2D code regions, and appends the decoded string payload into the output JSON alongside text blocks.

Zero LLM Dependency: It does not route text to foundational models like GPT-4o for its extraction, keeping processing latency low and lowering operational costs significantly.

Reference:

<https://learn.microsoft.com/en-us/azure/ai-services/content-understanding/quickstart/content-understanding-studio>

質問: 56

製品に関する質問に回答するSemantic Kernelエージェントがあります。

あなたは、注文状況に関する質問に答えられるように、エージェントの機能を拡張することを計画しています。

エージェントが注文状況に関する質問にできるだけ迅速に回答できるようにする必要があります。また、開発の手間を最小限に抑えるソリューションが求められます。

解決策には何を含めるべきでしょうか？

A. 種子

B. プラグイン

C. 関数

D. エージェント

正解: B ([コメントを发表する](#))

To enable the agent to answer questions about order status, you need a plugin.

In Semantic Kernel, a plugin acts as a container for the specific capabilities (functions) that allow an agent to interact with external data or systems, such as an order database or API.

To help me give you a more specific example or guide, could you tell me:

The data source where order info is kept (e.g., a SQL database, a REST API, or Shopify) The programming language you're using for the Kernel (C# or Python) Reference:

<https://learn.microsoft.com/en-us/semantic-kernel/frameworks/agent/agent-functions>

質問: 57

あなたは、Microsoft Foundry プロジェクト内のマルチモーダルチャットモデルのデプロイメントにリクエストを送信する App1 という名前の Web アプリケーションを持っています。

ユーザーメッセージには、テキストと画像の両方を含めることができます。

現在、App1では画像URLがメッセージコンテンツ内にプレーンテキストとして含まれているため、モデルはそれを画像として認識できません。

トレース結果によると、リクエストにはマルチモーダルコンテンツ配列ではなく、単一のテキストメッセージが含まれていることがわかった。

モデルが画像を正しく処理できるようにするには、テキスト部分と画像参照の両方を含む構造化配列としてメッセージを送信する必要があります。

あなたはどうすべきでしょうか？

A. ユーザーメッセージコンテンツ配列に、タイプ: text とタイプ: image_url のアイテムを含めるように設定します。

B. 画像をbase64でエンコードし、エンコードされたデータをユーザーメッセージのコンテンツ文字列内に含めます。

C. リクエストメタデータセクションに画像URLを追加すると、モデルが処理の問題を自動的に解決できます。

D. システムメッセージ内に画像URLを配置し、タイプをimage_urlに設定して、モデルが初期化時に画像をロードするようにします。

正解: A ([コメントを发表する](#))

To fix this, you must change your request payload structure from a single string to a structured content array. Multimodal models in Azure AI Foundry expect an array of distinct objects for text and images.

Reference:

<https://towardsdatascience.com/multimodal-ai-search-for-business-applications-65356d011009/>

質問: 58

大量の標準的な製品説明文を、迅速かつ一貫性を保ち、かつ低コストで20言語に翻訳するアプリケーションが必要です。テキストは簡潔で、業界特有のニュアンスは含まれていません。どの手法が最適でしょうか？

A. カスタムシステムプロンプトを使用した大規模言語モデル翻訳フロー

B. Foundry Tools の Azure Translator

C. 20の対象言語それぞれに対して個別に微調整されたモデル

D. Azure AI ドキュメントインテリジェンス

正解: ([正解を表示します](#))

Azure Translator in Foundry Tools provides deterministic, broad-language translation at scale and at lower cost, which suits high-volume, straightforward content.

質問: 59

エージェントを含む Microsoft Foundry プロジェクトがあります。

スキャンされたテキスト、表、複数列レイアウトを含む、様々な形式のドキュメントを処理する必要があります。抽出されたコンテンツはドキュメントの構造を保持し、後続の処理のためにMarkdown形式に変換する必要があります。

最初に何を設定すべきですか？

A. Azure Language in Foundry Tools のテキスト分析モデルのデプロイ

B. 生成型チャット完了リクエスト

C. マルチモーダルモデルを使用するAzure OpenAI Responses API呼び出し

D. Foundry Tools の Azure コンテンツ理解アナライザー

正解: D ([コメントを发表する](#))

The first step should be an Azure Content Understanding analyzer within Microsoft Foundry Tools.

This specific tool is designed to handle the complexities in this scenario--scanned text, complex tables, and multi-column layouts--by converting unstructured documents into a structured Markdown format. This preserves the document's original hierarchy and structural relationships, which is critical for accurate downstream reasoning by an agent.

Incorrect:

[Not A]

Azure Language is designed for text-based natural language processing (NLP) like sentiment analysis or PII redaction on raw text strings. It cannot parse scanned images, optical character recognition (OCR), tables, or complex multi-column visual layouts.

[Not C]

The first step in this specific Microsoft Foundry agent configuration should not be an Azure OpenAI Responses API call.While multimodal models can read images, they frequently hallucinate data, fail to capture complex multi-column reading orders, and degrade table structures when processing dense, mixed-format documents.

Reference:

<https://learn.microsoft.com/en-us/azure/ai-services/content-understanding/document/elements>

質問: 60

ホットスポットに関する質問

内部Q&Aエージェントを含むMicrosoft Foundryプロジェクトがあります。

ユーザーがエージェントに質問した際に、以下のような問題が報告されています。

- 次の回答が増加しました：関連情報なし」
見つかった"

- ピーク時間帯にHTTP 429レート制限超過エラーが定期的発生

各問題の原因が、モデルの利用不可、リソース制限、または推論の失敗のいずれであるかを特定する必要があります。

どうすればよいですか？回答するには、回答欄で適切な選択肢を選んでください。

注：正解ごとに1ポイントが加算されます。

Answer Area

Metrics to enable:

Model Availability Rate and Provisioned Utilization

Only Tokens Cache Match Rate

Only Total Requests filtered to status code 200

Time To Response and Total Tokens

Diagnostic log to collect

AllMetrics

audit

RequestResponse

trace



正解:

Metrics to enable:

Model Availability Rate and Provisioned Utilization

Only Tokens Cache Match Rate

Only Total Requests filtered to status code 200

Time To Response and Total Tokens

Diagnostic log to collect

AllMetrics

audit

RequestResponse

trace



質問: 61

エージェントを含む Microsoft Foundry プロジェクトがあります。このエージェントは、取得したポリシー文書から概要を生成します。

応答の完全性を改善する必要があります。この解決策は、応答が返される前に、アプリケーションコードのロジックに実装されなければなりません。

あなたはすべきでしょうか？

- A. レスポンスを返す前に再試行評価を追加します。
- B. 温度パラメータの値を下げる。
- C. presence_penaltyパラメータの値を増やす
- D. より小規模な展開でモデルを置き換えます。

正解: **A** ([コメントを发表する](#))

You should implement an evaluation and retry loop in your application code.

You must wrap the agent call inside a conditional code loop that evaluates the output against required criteria (such as a checklist of regulatory clauses), and programmatically forces a retry if information is omitted.

Reference:

<https://learn.microsoft.com/en-us/azure/foundry/agents/how-to/tools/ai-search>

有効的な**AI-103-JPN**問題集はJPNTest.com提供され、**AI-103-JPN**試験に合格することに役に立ちます！JPNTest.comは今最新**AI-103-JPN**試験問題集を提供します。JPNTest.com AI-103-JPN試験問題集はもう更新されました。ここで**AI-103-JPN**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/AI-103-JPN-mondaishu> **159**問、**30%ディスカウント**、特別な割引コード: **JPNshiken**」

質問: **62**

ホットスポットに関する質問

デプロイ済みのチャットモデルを含むMicrosoft Foundryプロジェクトがあります。

Pythonサービスがあり、そこからモデルにAPIリクエストが送信されます。このサービスは、生成された出力を承認済みの応答パターンと比較する自動検証システムと統合されています。

関係者によると、わずかな文言の違いが検証結果の不一致を引き起こしているとのことだ。

出力の安定性を向上させるには、リクエストパラメータを更新する必要があります。解決策は推論の質を最大限に高めるものでなければなりません。

Pythonコードをどのように完成させるべきですか？回答するには、回答欄で適切なオプションを選択してください。

注：正解ごとに1ポイントが加算されます。

Answer Area

```
message = client.messages.create(  
    model="deployment-name",  
    messages=[  
        {"role": "user", "content": "Summarize the release notes in 3 bullet points."}  
    ],  
    max_tokens=800,  
    temperature=

|   |   |
|---|---|
|   | ▼ |
| 0 |   |
| 1 |   |
| 2 |   |

,  
    thinking={"type": "enabled"},  
    output_config={"effort":

|          |   |
|----------|---|
|          | ▼ |
| "high"   |   |
| "low"    |   |
| "medium" |   |

}  
)
```



正解:

Answer Area

```
message = client.messages.create(
    model="deployment-name",
    messages=[
        {"role": "user", "content": "Summarize the release notes in 3 bullet points."},
    ],
    max_tokens=800,
    temperature=
    thinking={"type": "enabled"},
    output_config={"effort":
    },
```

0

1

2

"high"

"low"

"medium"

質問: 63

Foundry Tools の Azure Speech リソースがあり、そこにカスタムの音声認識モデルがカスタムエンドポイントにデプロイされています。エージェントはこのエンドポイントを使用して、リアルタイムの音声認識を実行します。

お客様のカスタム音声認識モデルの有効期限が近づいています。

モデルの有効期限が切れた場合、どのような動作が想定されますか？

A. 新しいカスタムモデルがデプロイされるまで、音声認識リクエストは4xxエラーを返します。

B. 音声認識リクエストは、モデルが手動で削除されるまで、期限切れのカスタムモデルを引き続き使用します。

C. 音声認識要求は、同じロケールの最新の基本モデルにフォールバックします。

D. カスタムモデルは、モデルの有効期限が切れると自動的に削除されます。

正解: (正解を表示します)

When the custom speech-to-text model expires, the real-time endpoint will automatically fall back to using the most recent base model for that locale.

Because of this automated fallback design, your agent's real-time speech recognition streams will not fail or throw a connection error. However, you will likely experience a drop in transcription accuracy, as the fallback base model lacks the domain-specific vocabulary, acronyms, or unique audio adaptations built into your custom model.

Immediate Impact Summary

Real-Time Endpoints: Continue to process requests. The endpoint swaps the expired custom model for the newest standard base model behind the scenes.

Batch Transcriptions (If used): Any batch transcription jobs explicitly targeting the expired custom model ID will fail with a 4xx error code.

Customization Loss: Specific jargon, formatting rules, or accents trained into your model will temporarily stop applying to incoming agent audio.

Reference:

<https://learn.microsoft.com/en-us/azure/ai-services/speech-service/how-to-custom-speech-model-and-endpoint-lifecycle>

質問: 64

PDF形式のポリシーマニュアルを取り込むAzure AI Searchインデクサーをお持ちです。

クライアントアプリケーションは、テキストと画像の両方について境界ポリゴンを持つページレベルの引用を表示する必要があります。

Azure AI Search スキルセットに、組み込みのマルチモーダルコンテンツ抽出スキルを 1 つ追加する必要があります。ソリューションは、次の要件を満たす必要があります。

テキストと画像の位置メタデータを提供する。

複数ページにまたがる表を抽出する。

何を追加すればいいですか？

A. ドキュメントレイアウト

B. 文書抽出

C. Azureコンテンツ理解

D. GenAIプロンプト

正解: ([正解を表示します](#))

To meet the requirements, you need to use the Azure Content Understanding skill in your Azure AI Search skillset.

Reference:

<https://learn.microsoft.com/en-us/azure/search/cognitive-search-skill-content-understanding>

質問: 65

あなたは、顧客サポートプラットフォーム向けに、Microsoft Foundry 上で音声処理ソリューションを構築しています。

このプラットフォームは通話内容をリアルタイムで文字起こしするため、社内の管理者は通話中に文字起こしを確認し、問題点を検知することができます。通話音声は電話システムから連続ストリームとして配信されます。

通話の文字起こしが音声ストリームの再生開始からわずか数秒以内に表示されるようにする必要があります。

あなたはどうすべきでしょうか？

A. カスタムニューラルボイスを使用してテキスト読み上げ機能を使用します。

B. 音声翻訳を使用して、複数の言語の文字起こしを生成します。

C. 録音された音声ファイルに対してバッチ文字起こしジョブを実行します。

D. リアルタイム音声認識を使用して、ストリーミング音声入力を処理します。

正解: D ([コメントを发表する](#))

You should use Real-Time Speech-to-Text (Streaming STT) for this project.

Why Real-Time STT is Required

Low Latency: It transcribes audio within seconds.

Continuous Input: It processes live, ongoing audio streams.

Live Monitoring: Supervisors need immediate text to flag issues.

Reference:

<https://learn.microsoft.com/en-us/azure/ai-services/speech-service/get-started-speech-to-text>

質問: 66

製品サポートマニュアルをお持ちですね。

マニュアルに基づいて製品サポートチャットボットを構築する必要があります。開発の手間とコストを最小限に抑えるソリューションが求められます。

何を使うべきでしょうか？

A. Azure AI Search を使用するグラウンディングデータ付き Azure OpenAI GPT-4

B. Azure AI ドキュメントインテリジェンス

C. Azure AI 言語カスタム質問応答

D. Azure AI Phi-3-medium(微調整済み)

正解: ([正解を表示します](#))

Azure OpenAI On Your Data makes it easier for developers to connect, ingest and ground their enterprise data to create personalized copilots (preview) rapidly. It enhances user comprehension, expedites task completion, improves operational efficiency, and aids decision- making.

Azure OpenAI On Your Data enables you to run advanced AI models such as GPT-35-Turbo and GPT-4 on your own enterprise data without needing to train or fine-tune models. You can chat on top of and analyze your data with greater accuracy. You can specify sources to support the responses based on the latest information available in your designated data sources.

Reference:

<https://learn.microsoft.com/en-us/azure/ai-services/openai/concepts/use-your-data>

有効的な**AI-103-JPN**問題集はJPNTest.com提供され、**AI-103-JPN**試験に合格することに役に立ちます！JPNTest.comは今最新**AI-103-JPN**試験問題集を提供します。JPNTest.com AI-103-JPN試験問題集はもう更新されました。ここで**AI-103-JPN**問題集のテストエンジンを手に入れます。最新版のアクセス、<https://www.jpntest.com/shiken/AI-103-JPN-mondaishu> 159問、30%**ディスカウント**、特別な割引コード: **JPNshiken**」