

Microsoft Azure AI Fundamentals (Updated Version)

Microsoft AI-901

Version Demo

Total Demo Questions: 11

Total Premium Questions: 117

Buy Premium PDF

<https://passqueen.com>

support@passqueen.com

Topic Break Down

Topic	No. of Questions
Topic 1, Describe Artificial Intelligence workloads and considerations	19
Topic 2, Describe fundamental principles of machine learning on Azure	7
Topic 3, Describe features of computer vision workloads on Azure	11
Topic 4, Describe features of Natural Language Processing (NLP) workloads on Azure	20
Topic 5, Describe features of generative AI workloads on Azure	59
Topic 6, Mix Questions	1
Total	117

QUESTION NO: 1

When prompting a generative AI model, which two purposes can instructions serve? Each correct answer represents part of the solution.

NOTE: Each correct selection is worth one point.

- A. defines constraints on the model's responses
- B. defines the agent's role and behavior
- C. defines the Azure region where inference occurs
- D. selects which model to use
- E. defines the tokens per minute (TPM) allocation for the model

ANSWER: A B

Explanation:

Instructions are high-level directions that guide how a generative AI model should respond across an interaction. “defines constraints on the model's responses” is correct because instructions can establish requirements such as the expected response format, tone, length, permitted scope, safety requirements, and rules for handling information. These directions help produce outputs that are consistent with the application’s needs rather than treating each user message as an isolated request.

“defines the agent's role and behavior” is also correct because instructions can establish the context in which the model operates. For example, they can direct the model to act as a tutor, customer-support assistant, or technical writer, while specifying communication style, priorities, and how to respond when information is uncertain. Role and behavior guidance helps the model interpret later user prompts in a way that supports the intended experience. Microsoft documents that system messages can define an assistant's behavior, personality, and constraints, and that agent instructions establish an agent’s purpose and expected behavior. See [Azure OpenAI chat completions guidance](#) and [Azure AI Foundry Agent Service guidance](#).

QUESTION NO: 2

You have a Microsoft Foundry project that includes an agent named Agent1.

You must ensure that Agent1 invokes an Azure Function whenever it responds to user input.

Which `tool_choice` value should you configure for Agent1?

- A. auto
- B. none
- C. required

ANSWER: C

Explanation:

`required` is the correct `tool_choice` setting because it requires the model to call at least one configured tool as part of processing a user request. When Agent1 has an Azure Function configured as its available function tool, this setting ensures that the agent performs a tool invocation rather than deciding on its own whether a tool is needed before it produces a response. This is appropriate when function execution is a mandatory step in the agent workflow, such as retrieving current information, enforcing business rules, validating a request, or recording an action.

Tool choice is configured in the request or run settings used with an agent. The `required` value provides the mandatory tool-use behavior needed for this scenario. If several tools are available and a specific function must always be invoked, the application design should also ensure that the intended Azure Function is the applicable configured tool or use a named tool-selection mechanism where supported. Microsoft’s function-calling guidance describes how agents invoke customer-defined

functions and how the application executes the function call and submits its result. See [Function calling with Azure AI Foundry Agent Service](#) and [Azure AI Foundry Agent Service tools overview](#).

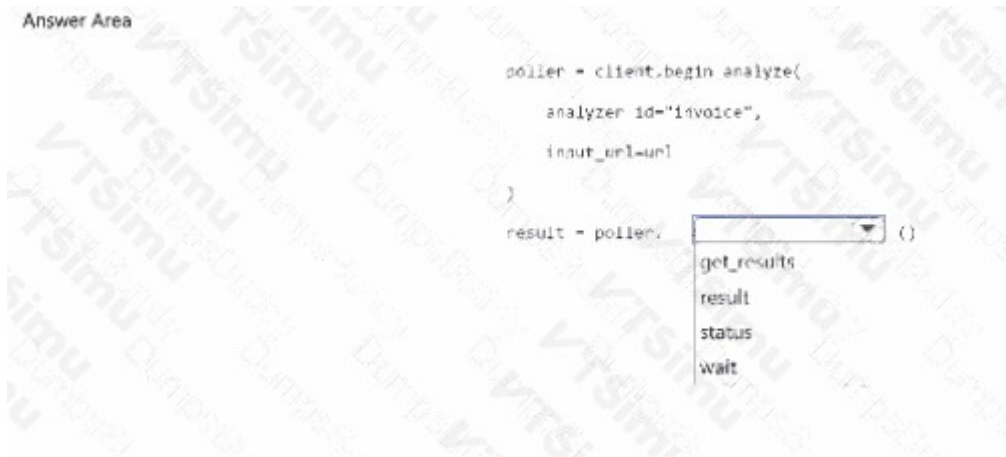
QUESTION NO: 3 - (SIMULATION)

You are developing an application that analyze invoices by using Azure Content Understanding in Foundry Tools.

You need to ensure that the application retrieves the analysis results after processing completes.

How should you complete the Python code? To answer, select the appropriate option in the answer area.

NOTE: Each correct selection is worth one point.



ANSWER: See the explanation for the answer

Explanation:

Select result.

```
poller = client.begin_analyze(  
    analyzer_id="invoice",  
    input_url=url  
)
```

```
result = poller.result()
```

`begin_analyze()` returns a long-running-operation poller. Calling `poller.result()` waits until the invoice analysis has completed and then retrieves the completed analysis result.

QUESTION NO: 4

A manufacturing company records equipment temperature readings every minute.

The company needs to detect readings that deviate significantly from the normal operating pattern so that technicians can investigate potential equipment failures.

Which machine learning capability should be used?

- A. anomaly detection
- B. regression
- C. binary classification
- D. clustering

ANSWER: A

Explanation:

Anomaly detection is the appropriate machine learning capability because it is designed to identify observations that differ substantially from the expected or typical pattern in a dataset. In this scenario, normal equipment behavior is represented by the usual range, trend, and variation of temperature readings collected over time. A reading, or sequence of readings, that departs materially from that learned baseline can be flagged as an anomaly for technicians to investigate.

This is particularly useful for predictive maintenance and operational monitoring, where unexpected temperature changes may be an early signal of equipment degradation, sensor issues, or developing failures. The model can continuously evaluate incoming telemetry and produce an anomaly score or alert when behavior is unusual, enabling a timely inspection workflow. Microsoft describes anomaly detection as identifying irregularities in time-series data, including scenarios such as monitoring equipment and detecting unexpected changes. For more information, see [Azure AI Anomaly Detector overview](#) and [Anomaly detection in Azure Machine Learning](#).

QUESTION NO: 5

Your company processes thousands of customer reviews daily.

You must determine whether each review conveys a positive or negative opinion and identify references to product names, companies, and cities.

Which two Azure AI Language capabilities should you use? Each correct answer provides part of the solution.

NOTE: Each correct selection is worth one point.

- A. sentiment analysis
- B. Named Entity Recognition (NER)
- C. summarization
- D. language detection
- E. key phrase extraction

ANSWER: A B

Explanation:

sentiment analysis is appropriate because Azure AI Language evaluates the emotional tone expressed in text. For customer reviews, it can classify the overall sentiment as positive, negative, neutral, or mixed, and it can also provide confidence scores. This enables the company to process a large volume of reviews consistently and identify reviews that reflect favorable or unfavorable customer experiences. Where needed, the service can additionally assess sentiment at the sentence level, which is useful when a single review discusses both positive and negative aspects of a product or service.

Named Entity Recognition (NER) is appropriate for finding and categorizing real-world entities in unstructured review text. Azure AI Language can recognize entity categories including organization and location, supporting the identification of company names and city names. It can also recognize product-related entities and extract the text spans that represent the entities found in each review. Together, these capabilities provide both the review's expressed opinion and structured information about the names and places it mentions. See [Microsoft Learn: sentiment analysis and opinion mining](#) and [Microsoft Learn: named entity recognition](#).

QUESTION NO: 6

You are developing a chatbot for a hotel. A user enters the message, "Book a room in Seattle for two people next Friday."

You need the chatbot to recognize the user's intent to make a reservation and extract the city, number of guests, and date from the message.

Which Azure AI capability should you use?

- A. conversational language understanding
- B. sentiment analysis
- C. Named Entity Recognition (NER)
- D. custom question answering

ANSWER: A

Explanation:

conversational language understanding is the appropriate Azure AI capability because it is designed to interpret natural-language messages in a conversational application by identifying both an intent and the entities associated with that intent. In this hotel scenario, a model can be trained with a reservation-related intent, such as booking a room, and configured to extract domain-specific entities such as the destination city, guest count, and check-in date. For the message provided, the service can recognize the booking intent and return values corresponding to Seattle, two people, and next Friday. This structured result enables the chatbot to continue the reservation workflow, for example by checking availability and asking for any missing booking details. Conversational language understanding supports custom schemas that define the intents and entities relevant to an application's business domain, which makes it well suited to task-oriented chatbots. Microsoft describes conversational language understanding as a cloud-based API for building natural-language understanding into conversational applications, including intent classification and entity extraction. See [Conversational language understanding overview](#) and [Microsoft Learn: intents in conversational language understanding](#).

QUESTION NO: 7 - (HOTSPOT)

For each of the following statements, select Yes if the statement is true. Otherwise, select No.

NOTE: Each correct selection is worth one point.

Answer Area

Statements

Voice Live returns only transcribed text.

Yes

☐

No

☐

Voice Live requires you to separately implement speech to text and text to speech services.

☐☐

Voice Live combines speech to text, reasoning, and text to speech into a single conversational experience.

☐☐

ANSWER:

Answer:

Answer Area

Statements

Voice Live returns only transcribed text.

Yes

☐

No

☒

Voice Live requires you to separately implement speech to text and text to speech services.

☐☒

Voice Live combines speech to text, reasoning, and text to speech into a single conversational experience.

☒☐

Answer:

Explanation:

Answer Area

Statements

Voice Live returns only transcribed text.

Yes

No

☐☒

Voice Live requires you to separately implement speech to text and text to speech services.

☐☒

Voice Live combines speech to text, reasoning, and text to speech into a single conversational experience.

☒☐

Statement 1:

Voice Live returns only transcribed text. = No Voice Live is not limited to transcription. Microsoft documentation states that the Voice Live API supports real-time bidirectional voice applications, including speech recognition, text-to-speech synthesis, avatar streaming, animation data, and audio processing. **Statement 2: Voice Live requires you to separately implement speech to text and text to speech services. = No** Voice Live provides a single real-time voice API experience rather than requiring separate STT and TTS implementations for the conversational loop. Microsoft describes live AI voice conversations as combining speech capabilities for real-time interaction, and the Voice Live API includes speech recognition and text-to-speech synthesis features. **Statement 3: Voice Live combines speech to text, reasoning, and text to speech into a single conversational experience. = Yes** This is correct. Microsoft's guidance explains that Azure OpenAI audio/realtime capabilities are for scenarios that combine audio with language understanding, reasoning, or generation in a single model call, and Voice Live supports real-time voice-enabled applications over WebSocket connections.

Explanation:

Azure AI Speech Voice Live is intended for building natural, real-time spoken conversations. In a Voice Live session, an application can stream a user's audio to the service, receive recognition of the user's spoken input, use an AI model to understand the conversation and generate an appropriate response, and deliver that response as synthesized speech. This end-to-end flow supports interactive voice experiences such as conversational assistants, where the user can speak naturally and receive a spoken response in the same ongoing session.

The correct selections are **No** for "Voice Live returns only transcribed text," **No** for "Voice Live requires you to separately implement speech to text and text to speech services," and **Yes** for "Voice Live combines speech to text, reasoning, and text to speech into a single conversational experience." The central value of Voice Live is its unified, low-latency conversational architecture. Rather than treating spoken input and spoken output as isolated operations that an application must independently coordinate for the core conversational loop, Voice Live brings speech input, model interaction, and speech output together through its real-time API experience.

Applications communicate with Voice Live through a WebSocket-based session and exchange real-time events, enabling continuous conversational context and responsive audio interaction. This design is especially useful where a natural turn-taking experience matters, including voice agents and other interactive applications. Microsoft's [Azure AI Speech Voice Live documentation](#) describes Voice Live as a unified voice-to-voice capability that integrates speech recognition, generative AI, and speech synthesis. The broader [Azure AI Speech documentation](#) provides context on the speech capabilities used in these conversational scenarios.

QUESTION NO: 8 - (HOTSPOT)

For each of the following statements, select Yes if the statement is true. Otherwise, select No. NOTE: Each correct selection is worth one point.

Answer Area

Statements

In Microsoft Foundry Agent Service, setting tool_choice to auto for an agent enables the agent to decide whether to call a tool.

Yes

No

☐☐

In Microsoft Foundry Agent Service, setting tool_choice to none for an agent means that the model decides whether to call a tool.

☐☐

In Microsoft Foundry Agent Service, setting tool_choice to required for an agent ensures that the agent must call one or more tools during each run.

☐☐

ANSWER:

Answer Area

Statements

In Microsoft Foundry Agent Service, setting `tool_choice` to `auto` for an agent enables the agent to decide whether to call a tool.

Yes

No

In Microsoft Foundry Agent Service, setting `tool_choice` to `none` for an agent means that the model decides whether to call a tool.

In Microsoft Foundry Agent Service, setting `tool_choice` to `required` for an agent ensures that the agent must call one or more tools during each run.

Explanation:

In Microsoft Foundry Agent Service, `tool_choice` controls the tool-calling behavior available to the model during an agent run. Setting it to `auto` gives the model discretion to determine whether a tool call is useful for the current request. Therefore, the statement describing `auto` as enabling the agent to decide whether to call a tool is true.

The value `none` explicitly specifies that no tools are to be called for the run. A statement that says this setting lets the model make the decision is not true, because the behavior has already been constrained to prevent tool use. Thus, the correct selection for the statement about `none` is No.

Using `required` directs the model to make one or more tool calls. This is appropriate when the agent must use a configured tool as part of fulfilling the request, rather than responding without invoking one. Consequently, the statement that describes `required` as ensuring that the agent calls one or more tools during each run is true.

These values are part of the documented tool-choice controls for Azure AI Foundry agent tool and function calling. See the Microsoft Learn guidance on [function calling with Azure AI Foundry agents](#) for the documented tool-calling behavior.

QUESTION NO: 9 - (SIMULATION)

You have a Microsoft Foundry project that contains a model deployment.

You are developing an application that sends an image and a user question to the model.

You need to send both text and image content in the same request so the model can return an answer.

How should you complete the Python code? To answer, drag the appropriate values to the correct targets. Each value may be used once, more than once, or not at all. You may need to drag the split bar between panes or scroll to view content.

NOTE: Each correct selection is worth one point.

Values

input_image

input_text

input_url

output_image

output_text

Answer Area

```
input=[
    {
        'role': 'user',
        'content': [
            {'type': 'text', 'text': 'What is in this image? Provide 3 bullet points.'},
            {'type': 'image_url', 'image_url': image_url}
        ]
    }
]
```

```
)
print(response.output_text)
```

ANSWER: See the explanation for the answer

Explanation:

Complete the multimodal request by setting the content types to:

`output_text` is the generated model response property and is not an input content type.

QUESTION NO: 10 - (DRAG DROP)

You are developing an application that extracts structured information from different types of content by using Azure Content Understanding in Foundry Tools.

You need to extract scanned invoices in the PDF format and voicemail recordings in the WAV format. Which type of analyzer should you use for each content type? To answer, drag the appropriate analyzer types to the correct content types. Each analyzer type may be used once, more than once,

or not at all. You may need to drag the split bar between panes or scroll to view content. NOTE: Each correct selection is worth one point.

The screenshot shows the 'Analyzer Types' pane on the left with four items: 'audio analyzer', 'document analyzer', 'image analyzer', and 'video analyzer'. The 'Answer Area' on the right has two categories: 'Scanned invoices' and 'Voicemail recordings'. Both categories have empty rectangular boxes next to them, indicating that no analyzers have been assigned yet.

ANSWER:

The screenshot shows the same interface as before, but with the correct analyzers assigned. In the 'Scanned invoices' category, the 'document analyzer' has been dragged. In the 'Voicemail recordings' category, the 'audio analyzer' has been dragged.

Explanation:

Azure Content Understanding uses analyzers that are aligned to the modality of the input content. For scanned invoices in PDF format, use a document analyzer. An invoice PDF is a document, even when it is scanned rather than digitally generated. The document analyzer can interpret the document's text and visual layout and return structured information from the invoice, such as invoice-related fields. This is the appropriate analyzer category for extracting information from document-based business forms and PDFs.

For voicemail recordings in WAV format, use an audio analyzer. A WAV file contains audio, and a voicemail's useful information is contained in the spoken language. The audio analyzer is intended for processing audio content and can derive information from the recording, including speech-based output that can be used as structured data. Selecting the analyzer based on the source modality ensures the service applies the correct processing pipeline to each file type.

Microsoft's Content Understanding documentation describes analyzers as the resources that define how content is processed and extracted, including document and audio scenarios. The analyzer selection should therefore follow the content type: document processing for the scanned invoice PDF and audio processing for the WAV voicemail recording. See the [Azure Content Understanding overview](#) and the [Content Understanding analyzer concepts](#) for details on analyzer capabilities and supported content modalities.

QUESTION NO: 11

You are building a generative AI assistant that enables employees to search internal technical documentation.

You must reduce unsupported answers and ensure harmful user input is evaluated before it is sent to the assistant.

Which two actions should you take? Each correct answer represents part of the solution.

NOTE: Each correct selection is worth one point.

A. Ground the model request with retrieved approved documentation.

B. Evaluate user input by using Azure AI Content Safety.

- C. Increase the temperature setting.
- D. Increase the maximum completion token limit.
- E. Remove system instructions from the assistant.

ANSWER: A B

Explanation:

Ground the model request with retrieved approved documentation. is appropriate because retrieval-augmented generation supplies the model with relevant, organization-approved material at query time. Rather than relying only on the model's pretrained knowledge, the assistant can use retrieved passages from the internal technical documentation as grounding context when it generates its response. This makes responses more traceable to authoritative internal sources and helps reduce answers that are unsupported by the available documentation. Microsoft describes this pattern for Azure OpenAI as using your data to ground responses in enterprise content: [Azure OpenAI On Your Data](#).

Evaluate user input by using Azure AI Content Safety. is also required because the input can be screened before it reaches the generative AI assistant. Azure AI Content Safety evaluates text for harmful-content categories and returns severity information that an application can use to block, route for review, or otherwise handle unsafe prompts according to its policy. Applying this evaluation at the input boundary establishes a safety control before model processing occurs. The supported harm categories and their use in text analysis are documented in [Azure AI Content Safety harm categories](#).