

KillTest

質量更高 服務更好



學習資料

<http://www.killtest.net>

一年免費更新服務

Exam : AI-901

Title : Microsoft Azure AI
Fundamentals (Updated
Version)

Version : DEMO

1.You have a Microsoft Foundry project that contains an agent named Agent1.
You need to ensure that Agent1 always calls an Azure function when the agent responds to user input.
To what should you set tool_choice for Agent1?

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

Answer: C

Explanation:

Microsoft's Foundry Agent Service documentation states that tool_choice provides deterministic control over tool calling:

auto means the model decides whether to call tools.

required means the model must call one or more tools.

none means the model does not call tools.

Therefore:

A. auto = Incorrect, because the model may or may not call the Azure function.

B. none = Incorrect, because this prevents tool/function calls.

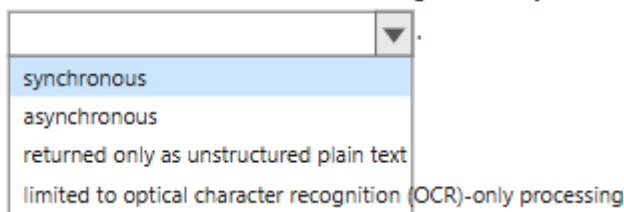
C. required = Correct, because it forces the agent to call a tool.

The Azure OpenAI function-calling documentation also confirms that tool_choice="auto" lets the model decide whether to call a function, while tool_choice="none" forces a user-facing response without a tool call.

2.HOTSPOT

Select the answer that correctly completes the sentence.

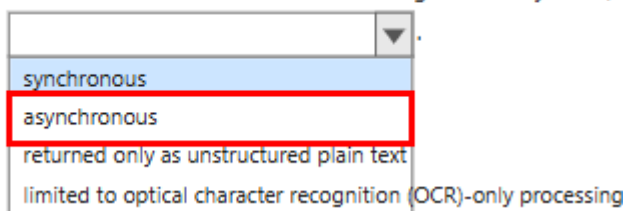
When content is submitted to Azure Content Understanding in Foundry Tools,
the analysis is



A screenshot of a dropdown menu with four options: 'synchronous', 'asynchronous', 'returned only as unstructured plain text', and 'limited to optical character recognition (OCR)-only processing'. The 'synchronous' option is currently selected and highlighted in blue.

Answer:

When content is submitted to Azure Content Understanding in Foundry Tools,
the analysis is



A screenshot of a dropdown menu with the same four options as above. In this image, the 'asynchronous' option is highlighted with a red rectangular box, indicating it is the correct answer.

Explanation:

When content is submitted to Azure Content Understanding in Foundry Tools, the analysis is asynchronous. This means the service does not return results immediately within the same HTTP request. Instead, it uses the standard Azure long-running operation (LRO) pattern — you call begin_analyze() to submit the content, which immediately returns a poller object, and then call poller.result() to wait for processing to complete and retrieve the structured extraction results.

Why the other options are wrong:

Synchronous is incorrect — the analysis pipeline involves multiple AI steps (OCR, speech transcription,

schema mapping) that take time; a blocking synchronous call is not supported.

Returned only as unstructured plain text is incorrect — Azure Content Understanding returns richly structured JSON output with named fields mapped to your defined schema, not plain unstructured text. Limited to OCR-only processing is incorrect — Content Understanding goes far beyond OCR; it supports document, audio, image, and video analyzers, and performs semantic field extraction using AI, not just character recognition.

This asynchronous design is consistent across all Azure AI services that perform complex, multi-step content processing.

3.DRAG DROP

You have a Microsoft Foundry project named project1 that contains an Azure OpenAI resource named Resource1.

To Resource1, you deploy a gpt-4.1-mini model by using a model deployment named my-mini-gpt.

You need to connect to my-mini-gpt from an application.

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. NOTE: Each correct selection is worth one point.

Values	Answer Area
gpt-4.1-mini	<pre> client = OpenAI(api_key="...", base_url="https://[project1].openai.azure.com/openai/v1/",) response = client.responses.create(model="[my-mini-gpt]", ...) </pre>
my-mini-gpt	
project1	
resource1	

Answer:

Values	Answer Area
gpt-4.1-mini	<pre> client = OpenAI(api_key="...", base_url="https://[resource1].openai.azure.com/openai/v1/",) response = client.responses.create(model="[my-mini-gpt]", ...) </pre>
my-mini-gpt	
project1	
resource1	

Explanation:

```

client = OpenAI(
api_key="...",
base_url="https://resource1.openai.azure.com/openai/v1/",
)
response = client.responses.create(
model="my-mini-gpt",
...
)

```

For Azure OpenAI in Microsoft Foundry, the base_url uses the Azure OpenAI resource name in the endpoint format:

https://<resource-name>.openai.azure.com/openai/v1/

In the question, the Azure OpenAI resource is named Resource1, so the first blank must be resource1. Microsoft documentation for Azure OpenAI v1 endpoints confirms that the endpoint must use the ...openai.azure.com/openai/v1/ path.

For the model parameter, Azure OpenAI requires the deployment name, not the underlying model name. Microsoft states that Azure OpenAI always requires the deployment name when calling APIs, even when the parameter is named model.

The deployed model is gpt-4.1-mini, but the deployment name is my-mini-gpt. Therefore, the second blank must be:

model="my-mini-gpt"

So the correct selections are:

base_url blank = resource1

model blank = my-mini-gpt

4.What are two purposes of instructions when prompting a generative AI model? Each correct answer presents 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:

Microsoft Foundry Agent Service documentation states that instructions define goals, constraints, and behavior for an agent. Therefore, instructions are used to guide how the generative AI model or agent should respond and behave.

Option A is correct because instructions can define constraints the model must follow.

Option B is correct because instructions can define the agent's role and behavior.

Options C, D, and E are incorrect because Azure region, model selection, and TPM allocation are configuration or deployment/resource settings, not purposes of prompt instructions.

5.You are developing an application that analyzes voicemail recordings by using Azure Content Understanding in Foundry Tools.

You need to extract a transcript and structured information from the recordings.

Which type of analyzer should you use?

- A. document analyzer
- B. video analyzer
- C. audio analyzer
- D. image analyzer

Answer: C

Explanation:

Voicemail recordings are audio content. Azure Content Understanding analyzers define what type of content to process, including documents, images, audio, or video, and what elements to extract, including transcripts and structured fields.

Microsoft's custom analyzer documentation also shows an audio example based on prebuilt-audio for

processing customer support call recordings, which is the same content type as voicemail recordings. Therefore, to extract a transcript and structured information from voicemail recordings, you should use an audio analyzer.