

AI-200 Originale Fragen & AI-200 Musterprüfungsfragen



Wir sind der Schnellste, der Prüfungsfragen und Antworten von Microsoft AI-200 Prüfung erhält. Unser ZertFragen bietet Ihnen die Testfragen und Antworten von Microsoft AI-200 Zertifizierungsprüfung, die von den IT-Experten durch Experimente und Praxis erhalten werden und über IT-Zertifizierungserfahrungen über 10 Jahre verfügt. ZertFragen verspricht, dass Sie das Microsoft AI-200 Zertifikat schneller und leichter erhalten, als Sie durch die anderen Webseiten.

Microsoft AI-200 Exam Syllabus Topics:

Section	Objectives
Implement and monitor AI workloads	- Deploy AI models and services - Monitor performance and troubleshoot issues
Plan and manage Azure AI solutions	- Select appropriate Azure AI services - Monitor and optimize AI solutions - Plan security and compliance requirements
Implement Azure AI solutions	- Implement natural language processing solutions - Implement knowledge mining with Azure AI Search - Implement computer vision solutions - Implement generative AI solutions using Azure OpenAI

>> AI-200 Originale Fragen <<

Valid AI-200 exam materials offer you accurate preparation dumps

Seit Jahren ist Microsoft AI-200 Prüfung eine sehr populäre Prüfung. Heutzutage wird Microsoft Zertifizierung immer wichtiger. Als von IT-Industrie international anerkannte Prüfung wird AI-200 eine der wichtigsten Prüfungen in Microsoft. Sie können viele Vorteile bekommen, wenn Sie das AI-200 Zertifikat bekommen. Microsoft AI-200 Dumps von ZertFragen gilt als das unentbehrliche Gerät, womit Sie die Microsoft AI-200 Prüfung vorbereiten, weil es den besten Nachschlag für Microsoft AI-200 Zertifizierungsprüfung ist.

Microsoft Developing AI Cloud Solutions on Azure AI-200 Prüfungsfragen mit Lösungen (Q128-Q133):

128. Frage

Case Study 1 - Fabrikam Inc.

Background

Fabrikam Inc. is a global retail analytics company that provides AI-driven demand forecasting and product recommendation services to online retailers. The company is modernizing its solution to run entirely on Microsoft Azure.

The platform ingests transaction data, generates embeddings for semantic retrieval, performs vector similarity search, and returns product recommendations through containerized microservices. Developers use Python and Azure SDKs. Operations teams manage container orchestration, scaling, monitoring, and security.

The solution must meet strict performance, scalability, and security requirements.

Current environment

Application architecture

The Recommendation engine is a customer-facing HTTP API running as a containerized Python application. The engine is deployed to Azure Container Apps (ACA).

Embeddings are stored in Azure Database for PostgreSQL by using pgvector.

Semantic retrieval uses metadata filtering combined with vector similarity search.

Azure Managed Redis is used as a caching layer.

Front-end and API workloads are deployed to Azure Container Apps (ACA).

Batch model retraining workloads run in Azure Kubernetes Service (AKS).

Container and CI/CD

Container images are stored in Azure Container Registry (ACR).

CI/CD uses ACR Tasks to build images on commit.

ACA environments support revision management.

AKS workloads are deployed by using Kubernetes manifest files stored in Git.

Monitoring

Logs are collected in Azure Monitor.

Teams inspect container logs and Kubernetes events when troubleshooting.

Developers write KQL queries to analyze latency spikes.

Business requirements

Customer experience: Maintain a seamless, low-latency recommendation experience for end- users, even during unpredictable seasonal traffic spikes.

Operational cost efficiency: Minimize compute expenditures by deallocating resources during periods of inactivity and by preventing runaway scaling costs.

Data integrity and freshness: Ensure that product recommendations always reflect the most current catalog metadata and pricing to prevent customer dissatisfaction.

Security and compliance: Adhere to a Zero Trust security model by eliminating long-lived credentials and centralizing the management of all sensitive secrets.

Global scalability: Support the rapid ingestion of millions of new product embeddings daily without degrading query performance for existing retailers.

Technical requirements

Performance: Semantic search latency must remain under 200 milliseconds at peak load.

Database optimization: Use pgvector for embeddings and implement metadata filtering to reduce compute overhead. Configure compute and memory appropriately for vector workloads to ensure high-dimensional index residency in RAM and efficient mathematical throughput. Vector similarity calculations must be performed only against products that satisfy mandatory metadata constraints.

Database performance: Database connections must support high concurrency with minimal latency through the implementation of connection optimization.

Data load strategy: To ensure maximum ingestion throughput, secondary indexes must be applied only after bulk loading of embeddings is complete.

Caching: Redis cache entries must expire automatically after 10 minutes. Implement a reactive mechanism to invalidate cache entries upon metadata updates.

Identity: Use managed identities for all service-to-service and service-to-database authentication.

Plain-text credentials in configuration files are strictly prohibited.

Secret management: All secrets must be stored centrally. Secrets must be rotated automatically by using a centralized lifecycle policy.

Scaling: Use Kubernetes event-driven autoscaling (KEDA) for event-driven scaling. The Recommendation API must scale based on HTTP traffic, while batch jobs must scale based on queue length and support scale-to-zero.

CI/CD: All images must be stored in Azure Container Registry. Use ACR Tasks to automate image builds triggered by source code commits.

Monitoring: Use KQL to analyze performance telemetry and troubleshoot microservice connectivity failures. Inspect logs and events when troubleshooting AKS and ACA.

Hotspot Question

You need to configure the database resources for the Azure Database for PostgreSQL instance.

How should you complete the configuration to meet the business and technical requirements? To answer, select the appropriate options in the answer area.

NOTE: Each correct selection is worth one point.

Configuration

- Increase compute vCores.
- Increase max_connections.
- Increase memory allocation.

- Enable read replica.
- Increase max_connections.
- Increase memory allocation.

- Enable storage autoscale.
- Increase compute vCores.
- Increase backup retention.



Configuration

- Increase compute vCores.
- Increase max_connections.
- Increase memory allocation.

- Enable read replica.
- Increase max_connections.
- Increase memory allocation.

- Enable storage autoscale.
- Increase compute vCores.
- Increase backup retention.



cache, ensuring the high-dimensional index remains resident in RAM for rapid semantic retrieval.

Box 3: Enable storage autoscale

Enable storage autoscale is the best action to support the continuous ingestion of transaction-based embeddings.

Continuous Ingestion Demands Dynamic Space: Continuous transaction processing causes vector databases (such as Azure Database for PostgreSQL with pgvector or Azure SQL Database) to expand constantly over time.

Preventing Ingestion Failures: If storage reaches capacity limits, the database switches into a read-only state. This immediately fails and halts all incoming real-time embedding write operations. Enabling storage autoscale allows the environment to dynamically provision storage on the fly without downtime.

Reference:

<https://dl.acm.org/doi/10.1145/3695053.3731013>

<https://learn.microsoft.com/en-us/training/paths/develop-ai-solutions-azure-database-postgresql/>

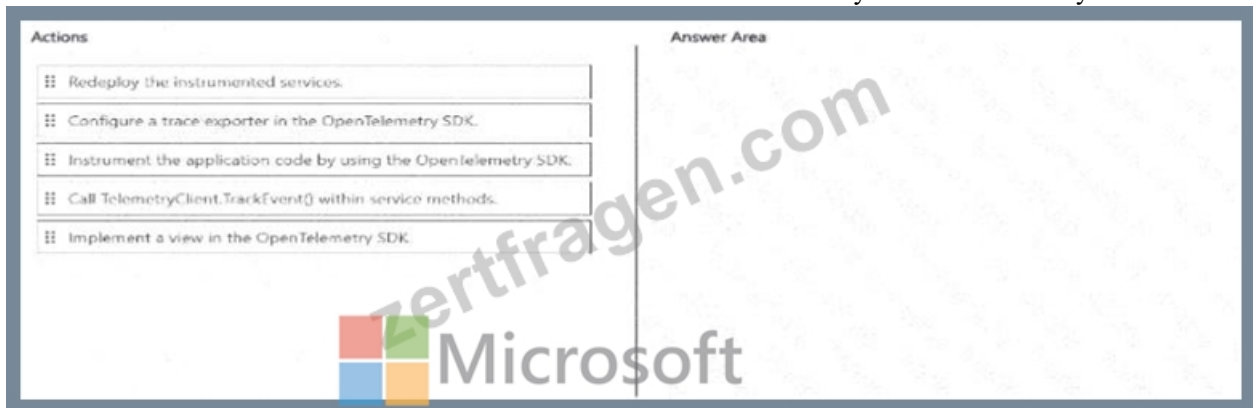
<https://learn.microsoft.com/en-us/azure/architecture/guide/technology-choices/vector-search>

129. Frage

You need to implement trace correlation according to the business requirements.

Which three actions should you perform in sequence? To answer, move the appropriate actions from the list of actions to the answer area and arrange them in the correct order.

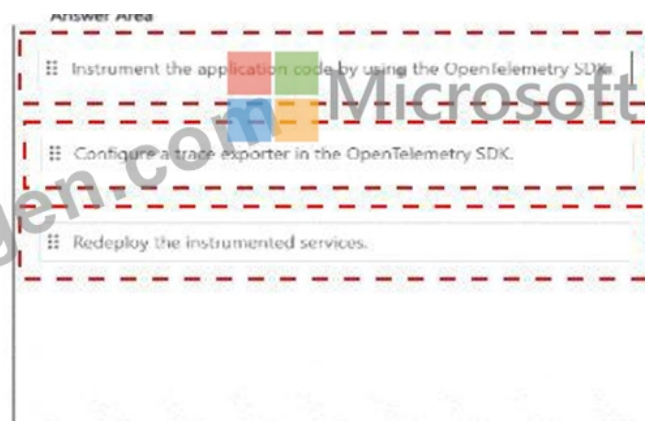
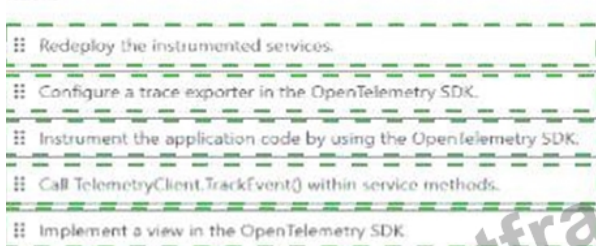
NOTE: More than one order of answer choices is correct. You will receive credit for any of the correct orders you select.



Antwort:

Begründung:

NOTES:



Explanation:



Verified answer: 1) Instrument the application code with the OpenTelemetry SDK; 2) configure an Azure Monitor trace exporter in

the OpenTelemetry SDK; 3) redeploy the instrumented services.

Detailed Explanation: Trace correlation requires the application to create OpenTelemetry spans and an exporter to send those spans to Azure Monitor. Instrumentation must therefore be present before the updated service is deployed. The previous Application Insights SDK instrumentation does not satisfy the case-study requirement that all tracing use OpenTelemetry. Configuring an OpenTelemetry view is unrelated to basic distributed-trace export, and calling TrackEvent is an Application Insights SDK pattern rather than the required OpenTelemetry approach.

Study Guide Alignment: Security and operations: Key Vault, App Configuration, managed identity, OpenTelemetry, Azure Monitor, and KQL-based troubleshooting.

Official Microsoft Learn References: AI-200 Study Guide | Enable Azure Monitor OpenTelemetry

130. Frage

You are developing a microservices-based application that uses Azure Container Apps. The application consists of several containerized services that handle tasks, such as processing orders, managing inventory, and generating reports.

You must secure the container apps. All apps must reside in the same virtual network, share the same Dapr configuration, and share the same logging location.

Apps must support the configuration of the amount of memory and compute resources available to containers.

You need to configure the Azure Container App

How should you complete the CLI command? To answer, select the appropriate options in the answer area.

NOTE: Each correct selection is worth one point.

-n MyContainerApp -g MyResourceGroup --location eastus2

env create add-on

env create add-on

-n MyContainerApp -g MyResourceGroup --location eastus2

--enable-mtls
--internal-only
--enable-workload-profiles

Microsoft

Antwort:

Begründung:

-n MyContainerApp -g MyResourceGroup --location eastus2

env create add-on

env create add-on

-n MyContainerApp -g MyResourceGroup --location eastus2

--enable-mtls
--internal-only
--enable-workload-profiles

Microsoft

Explanation:

Azure Container App CLI command

az containerapp env create

-n MyContainerApp -g MyResourceGroup --location eastus2

--enable-workload-profiles

Microsoft

Verified answer: `az containerapp env create ... --enable-workload-profiles`.

Detailed Explanation: The shared boundary for Container Apps is the managed environment. Apps in one environment can share the environment's virtual network integration, observability destination, and Dapr- related environment configuration. Enabling workload profiles provides selectable compute profiles and resource sizing options for workloads in that environment. The CLI must therefore create a Container Apps environment and enable workload profiles. Creating an individual container app alone would not establish the common environment-level network, logging, and resource-profile boundary required by the scenario.

Study Guide Alignment: Containerized Azure workloads: registry builds, App Service containers, Container Apps revision/scaling behavior, and AKS deployment choices.

131. Frage

Hotspot Question

You are reviewing the message-processing code in a backend worker service.

You review the following Python code that initializes and starts a Service Bus processor.

```
from azure.servicebus import ServiceBusClient, ServiceBusReceiveMode

01 import json
02 from azure.servicebus import ServiceBusClient, ServiceBusReceiveMode
03 from azure.identity import DefaultAzureCredential
04 credential = DefaultAzureCredential()
05 with client.get_queue_receiver(
06     queue_name="inference-requests",
07     receive_mode=ServiceBusReceiveMode.PEEK_LOCK,
08     max_wait_time=30
09 ) as receiver:
10     for msg in receiver:
11         try:
12             payload = json.loads(str(msg))
13             result = run_inference(payload)
14             save_result(result, payload.get("document_id"))
15             receiver.complete_message(msg)
16         except json.JSONDecodeError:
17             receiver.dead_letter_message(
18                 msg,
19                 reason="MalformedPayload",
20                 error_description="Message body is not valid JSON"
21             )
22         except ModelNotFoundError:
23             receiver.dead_letter_message(
24                 msg,
25                 reason="UnsupportedModel",
26                 error_description=f"Model {payload.get('model')} is not available"
27             )
28         except TransientServiceError:
29             receiver.abandon_message(msg)
```

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

Service Bus processor behavior and SDK defaults

Statement	Yes	No
Messages are removed from the queue as soon as they are received by the processor.	☐	☐
If message processing fails due to a transient service error, the message can be retried by another consumer.	☐	☐
Messages with invalid JSON payloads are retried until the maximum delivery count is reached.	☐	☐

Antwort:

Begründung:

Answer Area

Service Bus processor behavior and SDK defaults

Statement	Yes	No
Messages are removed from the queue as soon as they are received by the processor.	☐	☒
If message processing fails due to a transient service error, the message can be retried by another consumer.	☒	☐
Messages with invalid JSON payloads are retried until the maximum delivery count is reached.	☐	☒

Explanation:

Box 1: No

No, messages are not removed from the queue as soon as they are received when using PEEK_LOCK mode.

How PEEK_LOCK Works

The Lock: The Service Bus processor locks the message exclusively for your worker for a specific duration.

The Invisible State: Other workers cannot see or process the message while the lock is active.

The Removal: The message is only permanently removed from the queue after it is successfully processed and explicitly completed.

Box 2: Yes

Yes, the message can be retried by another consumer if it fails due to a TransientServiceError.

Explicit abandonment: The code calls `receiver.abandon_message(msg)` when catching a TransientServiceError.

Lock release: Abandoning a message immediately releases the PEEK_LOCK held by the current worker.

Immediate availability: The message returns to the active queue and becomes instantly visible to other competing consumers.

Box 3: No

No, messages with invalid JSON payloads will not be retried until the maximum delivery count is reached. They will be sent to the dead-letter queue (DLQ) immediately on the very first attempt.

Immediate Dead-Lettering: The code catches the `json.JSONDecodeError` and explicitly calls `receiver.dead_letter_message(msg)`.

Instant Termination: This API call tells Azure Service Bus to move the message to the DLQ right away, bypassing the normal delivery count logic.

No Re-delivery: Because the message is moved to the DLQ, the broker removes it from the main queue, stopping any further processing attempts.

Reference:

<https://learn.microsoft.com/en-us/python/api/azure-servicebus/azure.servicebus.aio.servicebusreceiver>

<https://learn.microsoft.com/en-us/azure/service-bus-messaging/service-bus-python-how-to-use-queues>

132. Frage

Hotspot Question

You are building a semantic search feature for a chatbot. You store document embeddings in Redis.

You review the following Python code that connects to Redis and stores an embedding value:

```

01 import numpy as np
02 from redis import Redis
03 embedding = np.array([0.12, 0.98, 0.44], dtype=np.float32).tobytes()
04 r = Redis(
05     host="myredis.redis.cache.windows.net",
06     port=6380,
07     password="accesskey",
08     ssl=True
09 )
10 r.hset("doc:1", mapping={"embedding": embedding})
11 r.expire("doc:1", 600)

```

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

NOTE: Each correct selection is worth one point.

Vector search fundamentals

Statement	Yes	No
The embedding is stored as a hash field.	☐	☐
The code enables similarity search on the embedding field.	☐	☐
The key will expire after 10 minutes.	☐	☐

Antwort:

Begründung:

Vector search fundamentals

Statement	Yes	No
The embedding is stored as a hash field.	☒	☐
The code enables similarity search on the embedding field.	☐	☒
The key will expire after 10 minutes.	☒	☐

Explanation:

Box 1: Yes

The line `r.hset("doc:1", mapping={"embedding": embedding})` will store the embedding as a field within a Redis Hash.

Box 2: No

Box 3: Yes

Reference:

<https://redis.io/docs/latest/commands/expire/>

133. Frage

.....

Sind Sie ein IT-Mann? Haben Sie sich an der populären IT-Zertifizierungsprüfung beteiligt? Wenn ja, würde ich Ihnen sagen, dass Sie wirklich glücklich sind. Unsere Schulungsunterlagen zur Microsoft AI-200 Zertifizierungsprüfung von ZertFragen werden Ihnen helfen, die Microsoft AI-200 Prüfung 100% zu bestehen. Das ist eine echte Nachricht. Wollen Sie Fortschritte in der IT-Branche machen, wählen Sie doch ZertFragen. Unsere Microsoft AI-200 Dumps können Ihnen zum Bestehen allen Zertifizierungsprüfungen verhelfen. Sie sind außerdem billig. Wenn Sie nicht glauben, gucken Sie mal und Sie werden das Wissen.

AI-200 Musterprüfungsfragen: https://www.zertfragen.com/AI-200_pruefung.html

- AI-200 Der beste Partner bei Ihrer Vorbereitung der Developing AI Cloud Solutions on Azure i Suchen Sie jetzt auf ☐ www.zertpruefung.de ☐ nach **【 AI-200 】** um den kostenlosen Download zu erhalten ☐ AI-200 Simulationsfragen
- Microsoft AI-200: Developing AI Cloud Solutions on Azure braindumps PDF - Testking echter Test ☐ Suchen Sie einfach auf www.itzert.com nach kostenloser Download von **➡ AI-200** ☐ AI-200 Testantworten
- AI-200 Examengine ☐ AI-200 Vorbereitungsfragen ☐ AI-200 Prüfungen ☐ Suchen Sie jetzt auf ☐ www.zertpruefung.ch ☐ nach **➡ AI-200** ☐ um den kostenlosen Download zu erhalten ☐ AI-200 German
- AI-200 Testengine ☐ AI-200 German ☐ AI-200 Vorbereitung ☐ Öffnen Sie die Website www.itzert.com ☐ Suchen Sie **➤ AI-200** ☐ Kostenloser Download ☐ AI-200 Prüfungsfragen
- Neuester und gültiger AI-200 Test VCE Motoren-Dumps und AI-200 neueste Testfragen für die IT-Prüfungen ☐ Öffnen Sie **➤** www.zertpruefung.ch ☐ geben Sie { AI-200 } ein und erhalten Sie den kostenlosen Download ☐ AI-200 Vorbereitung
- AI-200 Prüfungsfragen Prüfungsvorbereitungen, AI-200 Fragen und Antworten, Developing AI Cloud Solutions on Azure ☐ ☐ Suchen Sie jetzt auf **➡** www.itzert.com ☐ nach **➡ AI-200** ☐ um den kostenlosen Download zu erhalten ☐ AI-200 Demotesten
- AI-200 Prüfungs-Guide ☐ AI-200 Prüfungsvorbereitung ☐ AI-200 Fragenpool ☐ Geben Sie **➡** www.deutschpruefung.com ☐ ein und suchen Sie nach kostenloser Download von **☀ AI-200** ☐ ☀ ☐ AI-200 PDF
- Zertifizierung der AI-200 mit umfassenden Garantien zu bestehen ☐ URL kopieren www.itzert.com ☐ Öffnen und suchen Sie **➡ AI-200** ☐ Kostenloser Download ☐ AI-200 Prüfungs-Guide
- Kostenlos AI-200 dumps torrent - Microsoft AI-200 Prüfung prep - AI-200 examcollection braindumps ☐ Öffnen Sie die

Webseite (www.itzert.com) und suchen Sie nach kostenloser Download von ➡ AI-200 ☐☐☐ AI-200
Simulationsfragen

- AI-200 Deutsch Prüfungsfragen □ AI-200 Examengine □ AI-200 Testengine □ Suchen Sie jetzt auf ➡ www.itzert.com □ nach ▷ AI-200 ◁ und laden Sie es kostenlos herunter □AI-200 Fragenpool
- AI-200 Prüfungsfragen Prüfungsvorbereitungen, AI-200 Fragen und Antworten, Developing AI Cloud Solutions on Azure □
□ Suchen Sie auf der Webseite 【 www.zertsoft.com 】 nach ➡ AI-200 □□□ und laden Sie es kostenlos herunter □AI-200 Simulationsfragen
- www.stes.tyc.edu.tw, www.stes.tyc.edu.tw, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, app.parler.com, www.stes.tyc.edu.tw, www.stes.tyc.edu.tw, www.stes.tyc.edu.tw, myportal.utt.edu.tt, myportal.utt.edu.tt, myportal.utt.edu.tt, Disposable vapes