

Databricks-Generative-AI-Engineer-Associate

Übungsfragen: Databricks Certified Generative AI Engineer Associate & Databricks-Generative-AI-Engineer-Associate Dateien Prüfungsunterlagen



2026 Die neuesten ZertPruefung Databricks-Generative-AI-Engineer-Associate PDF-Versionen Prüfungsfragen und Databricks-Generative-AI-Engineer-Associate Fragen und Antworten sind kostenlos verfügbar: https://drive.google.com/open?id=1uFIYKYhKLova_quffHxmzrUzHraBEuc

Tun Sie, was Sie gesagt haben, was Beginn des Erfolgs ist. Weil Sie die schwierige IT-Zertifizierungsprüfung ablegen wollen, sollen Sie sich bemühen, um das Zertifikat zu bekommen. Die Fragenkataloge zur Databricks Databricks-Generative-AI-Engineer-Associate Prüfung von ZertPruefung sind sehr gut. Mit Ihr können Sie Ihren Erfolg ganz leicht erzielen. Sie sind ganz zuverlässig. Ich glaube, Sie werden die Prüfung 100% bestehen.

Databricks Databricks-Generative-AI-Engineer-Associate Prüfungsplan:

Thema	Einzelheiten
Thema 1	Evaluation and Monitoring: This topic is all about selecting an LLM choice and key metrics. Moreover, Generative AI Engineers learn about evaluating model performance. Lastly, the topic includes sub-topics about inference logging and usage of Databricks features.
Thema 2	Data Preparation: Generative AI Engineers covers a chunking strategy for a given document structure and model constraints. The topic also focuses on filter extraneous content in source documents. Lastly, Generative AI Engineers also learn about extracting document content from provided source data and format.
Thema 3	- Governance: Generative AI Engineers who take the exam get knowledge about masking techniques, guardrail techniques, and legal - licensing requirements in this topic.

>> Databricks-Generative-AI-Engineer-Associate Quizfragen Und Antworten <<

Databricks Databricks-Generative-AI-Engineer-Associate Fragen und Antworten, Databricks Certified Generative AI Engineer Associate

Prüfungsfragen

Warum wählen viele ZertPrüfung? Weil er Bequemlichkeiten und Anwendbarkeit bringen. Das hat von der Praxis überprüft. Die Lernmaterialien zur Databricks Databricks-Generative-AI-Engineer-Associate Zertifizierungsprüfung von ZertPrüfung ist den allen bekannt. Viele Kandidaten sind nicht selbstsicher, die Databricks Databricks-Generative-AI-Engineer-Associate Zertifizierungsprüfung zu bestehen. Deshalb sollen Sie die Materialien zur Databricks Databricks-Generative-AI-Engineer-Associate Zertifizierungsprüfung haben. Mit ihm können Sie mehr Selbstbewusstsein haben und sich gut auf die Prüfung vorbereiten.

Databricks Certified Generative AI Engineer Associate Databricks-Generative-AI-Engineer-Associate Prüfungsfragen mit Lösungen (Q42-Q47):

42. Frage

A Generative AI Engineer is responsible for developing a chatbot to enable their company's internal HelpDesk Call Center team to more quickly find related tickets and provide resolution. While creating the GenAI application work breakdown tasks for this project, they realize they need to start planning which data sources (either Unity Catalog volume or Delta table) they could choose for this application. They have collected several candidate data sources for consideration:

call_rep_history: a Delta table with primary keys `representative_id`, `call_id`. This table is maintained to calculate representatives' call resolution from fields `call_duration` and `call_start_time`.

transcript Volume: a Unity Catalog Volume of all recordings as a *.wav files, but also a text transcript as *.txt files.

call_cust_history: a Delta table with primary keys `customer_id`, `call_id`. This table is maintained to calculate how much internal customers use the HelpDesk to make sure that the charge back model is consistent with actual service use.

call_detail: a Delta table that includes a snapshot of all call details updated hourly. It includes `root_cause` and `resolution` fields, but those fields may be empty for calls that are still active.

maintenance_schedule - a Delta table that includes a listing of both HelpDesk application outages as well as planned upcoming maintenance downtimes.

They need sources that could add context to best identify ticket root cause and resolution.

Which TWO sources do that? (Choose two.)

- A. **call_detail**
- B. `maintenance_schedule`
- C. **transcript Volume**
- D. `call_cust_history`
- E. `call_rep_history`

Antwort: A,C

Begründung:

In the context of developing a chatbot for a company's internal HelpDesk Call Center, the key is to select data sources that provide the most contextual and detailed information about the issues being addressed. This includes identifying the root cause and suggesting resolutions. The two most appropriate sources from the list are:

* Call Detail (Option D):

* Contents: This Delta table includes a snapshot of all call details updated hourly, featuring essential fields like `root_cause` and `resolution`.

* Relevance: The inclusion of `root_cause` and `resolution` fields makes this source particularly valuable, as it directly contains the information necessary to understand and resolve the issues discussed in the calls. Even if some records are incomplete, the data provided is crucial for a chatbot aimed at speeding up resolution identification.

* Transcript Volume (Option E):

* Contents: This Unity Catalog Volume contains recordings in .wav format and text transcripts in .txt files.

* Relevance: The text transcripts of call recordings can provide in-depth context that the chatbot can analyze to understand the nuances of each issue. The chatbot can use natural language processing techniques to extract themes, identify problems, and suggest resolutions based on previous similar interactions documented in the transcripts.

Why Other Options Are Less Suitable:

* A (Call Cust History): While it provides insights into customer interactions with the HelpDesk, it focuses more on the usage metrics rather than the content of the calls or the issues discussed.

* B (Maintenance Schedule): This data is useful for understanding when services may not be available but does not contribute directly to resolving user issues or identifying root causes.

* C (Call Rep History): Though it offers data on call durations and start times, which could help in assessing performance, it lacks direct information on the issues being resolved.

Therefore, Call Detail and Transcript Volume are the most relevant data sources for a chatbot designed to assist with identifying and resolving issues in a HelpDesk Call Center setting, as they provide direct and contextual information related to customer issues.

43. Frage

What is an effective method to preprocess prompts using custom code before sending them to an LLM?

- A. It is better not to introduce custom code to preprocess prompts as the LLM has not been trained with examples of the preprocessed prompts
- **B. Write a MLflow PyFunc model that has a separate function to process the prompts**
- C. Rather than preprocessing prompts, it's more effective to postprocess the LLM outputs to align the outputs to desired outcomes
- D. Directly modify the LLM's internal architecture to include preprocessing steps

Antwort: B

Begründung:

The most effective way to preprocess prompts using custom code is to write a custom model, such as an MLflow PyFunc model. Here's a breakdown of why this is the correct approach:

* **MLflow PyFunc Models:** MLflow is a widely used platform for managing the machine learning lifecycle, including experimentation, reproducibility, and deployment. A PyFunc model is a generic Python function model that can implement custom logic, which includes preprocessing prompts.

* **Preprocessing Prompts:** Preprocessing could include various tasks like cleaning up the user input, formatting it according to specific rules, or augmenting it with additional context before passing it to the LLM. Writing this preprocessing as part of a PyFunc model allows the custom code to be managed, tested, and deployed easily.

* **Modular and Reusable:** By separating the preprocessing logic into a PyFunc model, the system becomes modular, making it easier to maintain and update without needing to modify the core LLM or retrain it.

* **Why Other Options Are Less Suitable:**

* **A (Modify LLM's Internal Architecture):** Directly modifying the LLM's architecture is highly impractical and can disrupt the model's performance. LLMs are typically treated as black-box models for tasks like prompt processing.

* **B (Avoid Custom Code):** While it's true that LLMs haven't been explicitly trained with preprocessed prompts, preprocessing can still improve clarity and alignment with desired input formats without confusing the model.

* **C (Postprocessing Outputs):** While postprocessing the output can be useful, it doesn't address the need for clean and well-formatted inputs, which directly affect the quality of the model's responses.

Thus, using an MLflow PyFunc model allows for flexible and controlled preprocessing of prompts in a scalable way, making it the most effective method.

44. Frage

A Generative AI Engineer is deciding between using LSH (Locality Sensitive Hashing) and HNSW (Hierarchical Navigable Small World) for indexing their vector database. Their top priority is semantic accuracy. Which approach should the Generative AI Engineer use to evaluate these two techniques?

- A. Compare the Bilingual Evaluation Understudy (BLEU) scores of returned results for a representative sample of test inputs
- **B. Compare the cosine similarities of the embeddings of returned results against those of a representative sample of test inputs**
- C. Compare the Levenshtein distances of returned results against a representative sample of test inputs
- D. Compare the Recall-Oriented-Understudy for Gisting Evaluation (ROUGE) scores of returned results for a representative sample of test inputs

Antwort: B

Begründung:

The task is to choose between LSH and HNSW for a vector database index, prioritizing semantic accuracy.

The evaluation must assess how well each method retrieves semantically relevant results. Let's evaluate the options.

* **Option A:** Compare the cosine similarities of the embeddings of returned results against those of a representative sample of test inputs

* **Cosine similarity** measures semantic closeness between vectors, directly assessing retrieval accuracy in a vector database. Comparing returned results' embeddings to test inputs' embeddings evaluates how well LSH or HNSW preserves semantic relationships, aligning with the priority.

* **Databricks Reference:** "Cosine similarity is a standard metric for evaluating vector search accuracy" ("Databricks Vector Search Documentation," 2023).

* **Option B:** Compare the Bilingual Evaluation Understudy (BLEU) scores of returned results for a representative sample of test inputs

* **BLEU** evaluates text generation (e.g., translations), not vector retrieval accuracy. It's irrelevant for indexing performance.

- * Databricks Reference: "BLEU applies to generative tasks, not retrieval" ("Generative AI Cookbook").
- * Option C: Compare the Recall-Oriented-Understudy for Gisting Evaluation (ROUGE) scores of returned results for a representative sample of test inputs
- * ROUGE is for summarization evaluation, not vector search. It doesn't measure semantic accuracy in retrieval.
- * Databricks Reference: "ROUGE is unsuited for vector database evaluation" ("Building LLM Applications with Databricks").
- * Option D: Compare the Levenshtein distances of returned results against a representative sample of test inputs
- * Levenshtein distance measures string edit distance, not semantic similarity in embeddings. It's inappropriate for vector-based retrieval.
- * Databricks Reference: No specific support for Levenshtein in vector search contexts.

Conclusion: Option A (cosine similarity) is the correct approach, directly evaluating semantic accuracy in vector retrieval, as recommended by Databricks for Vector Search assessments.

45. Frage

A Generative AI Engineer has been asked to design an LLM-based application that accomplishes the following business objective: answer employee HR questions using HR PDF documentation.
Which set of high level tasks should the Generative AI Engineer's system perform?

- A. Calculate averaged embeddings for each HR document, compare embeddings to user query to find the best document. Pass the best document with the user query into an LLM with a large context window to generate a response to the employee.
- B. Create an interaction matrix of historical employee questions and HR documentation. Use ALS to factorize the matrix and create embeddings. Calculate the embeddings of new queries and use them to find the best HR documentation. Use an LLM to generate a response to the employee question based upon the documentation retrieved.
- C. Split HR documentation into chunks and embed into a vector store. Use the employee question to retrieve best matched chunks of documentation, and use the LLM to generate a response to the employee based upon the documentation retrieved.
- D. Use an LLM to summarize HR documentation. Provide summaries of documentation and user query into an LLM with a large context window to generate a response to the user.

Antwort: C

Begründung:

To design an LLM-based application that can answer employee HR questions using HR PDF documentation, the most effective approach is option D. Here's why:

* Chunking and Vector Store Embedding: HR documentation tends to be lengthy, so splitting it into smaller, manageable chunks helps optimize retrieval. These chunks are then embedded into a vector store (a database that stores vector representations of text). Each chunk of text is transformed into an embedding using a transformer-based model, which allows for efficient similarity-based retrieval.

* Using Vector Search for Retrieval: When an employee asks a question, the system converts their query into an embedding as well. This embedding is then compared with the embeddings of the document chunks in the vector store. The most semantically similar chunks are retrieved, which ensures that the answer is based on the most relevant parts of the documentation.

* LLM to Generate a Response: Once the relevant chunks are retrieved, these chunks are passed into the LLM, which uses them as context to generate a coherent and accurate response to the employee's question.

* Why Other Options Are Less Suitable:

* A (Calculate Averaged Embeddings): Averaging embeddings might dilute important information. It doesn't provide enough granularity to focus on specific sections of documents.

* B (Summarize HR Documentation): Summarization loses the detail necessary for HR-related queries, which are often specific. It would likely miss the mark for more detailed inquiries.

* C (Interaction Matrix and ALS): This approach is better suited for recommendation systems and not for HR queries, as it's focused on collaborative filtering rather than text-based retrieval.

Thus, option D is the most effective solution for providing precise and contextual answers based on HR documentation.

46. Frage

A Generative AI Engineer at an automotive company would like to build a question-answering chatbot for customers to inquire about their vehicles. They have a database containing various documents of different vehicle makes, their hardware parts, and common maintenance information.

Which of the following components will NOT be useful in building such a chatbot?

- A. Embedding model

- B. Response-generating LLM
- C. Vector database
- D. Invite users to submit long, rather than concise, questions

Antwort: D

Begründung:

The task involves building a question-answering chatbot for an automotive company using a database of vehicle-related documents. The chatbot must efficiently process customer inquiries and provide accurate responses. Let's evaluate each component to determine which is not useful, per Databricks Generative AI Engineer principles.

* Option A: Response-generating LLM

* An LLM is essential for generating natural language responses to customer queries based on retrieved information. This is a core component of any chatbot.

* Databricks Reference: "The response-generating LLM processes retrieved context to produce coherent answers" ("Building LLM Applications with Databricks," 2023).

* Option B: Invite users to submit long, rather than concise, questions

* Encouraging long questions is a user interaction design choice, not a technical component of the chatbot's architecture. Moreover, long, verbose questions can complicate intent detection and retrieval, reducing efficiency and accuracy - counter to best practices for chatbot design. Concise questions are typically preferred for clarity and performance.

* Databricks Reference: While not explicitly stated, Databricks' "Generative AI Cookbook" emphasizes efficient query processing, implying that simpler, focused inputs improve LLM performance. Inviting long questions doesn't align with this.

* Option C: Vector database

* A vector database stores embeddings of the vehicle documents, enabling fast retrieval of relevant information via semantic search. This is critical for a question-answering system with a large document corpus.

* Databricks Reference: "Vector databases enable scalable retrieval of context from large datasets" ("Databricks Generative AI Engineer Guide").

* Option D: Embedding model

* An embedding model converts text (documents and queries) into vector representations for similarity search. It's a foundational component for retrieval-augmented generation (RAG) in chatbots.

* Databricks Reference: "Embedding models transform text into vectors, facilitating efficient matching of queries to documents" ("Building LLM-Powered Applications").

Conclusion: Option B is not a useful component in building the chatbot. It's a user-facing suggestion rather than a technical building block, and it could even degrade performance by introducing unnecessary complexity. Options A, C, and D are all integral to a Databricks-aligned chatbot architecture.

47. Frage

.....

Die angemessene Bildung ist die Garantie des Erfolgs. Deshalb ist es sehr wichtig, diese Prüfung auszuwählen. Wir bieten Ihnen die neuesten und effektivsten Prüfungsfragen und Testantworten zur Databricks Databricks-Generative-AI-Engineer-Associate Zertifizierung. Und Wir ZertPruefung garantieren 100% Erfolg. Und der Preis ist auch günstig. Die Bezahlung bei Paypal kann Ihre Sicherheit garantieren. Wir ZertPruefung bieten die besten Lernhilfe für Databricks Databricks-Generative-AI-Engineer-Associate Prüfungen und auch aktualisieren immer die Prüfungsunterlagen.

Databricks-Generative-AI-Engineer-Associate Zertifizierungsantworten: https://www.zertpruefung.ch/Databricks-Generative-AI-Engineer-Associate_exam.html

- Databricks-Generative-AI-Engineer-Associate Testengine ☐ Databricks-Generative-AI-Engineer-Associate Testengine ☐ ☐ Databricks-Generative-AI-Engineer-Associate Echte Fragen ☐ Öffnen Sie die Webseite ☐ de.fast2test.com ☐ und suchen Sie nach kostenloser Download von (Databricks-Generative-AI-Engineer-Associate) ☐ Databricks-Generative-AI-Engineer-Associate Echte Fragen
- Kostenlose Databricks Certified Generative AI Engineer Associate vce dumps - neueste Databricks-Generative-AI-Engineer-Associate examcollection Dumps ☐ Sie müssen nur zu ☐ www.itzert.com ☐ gehen um nach kostenloser Download von > Databricks-Generative-AI-Engineer-Associate < zu suchen ☐ Databricks-Generative-AI-Engineer-Associate Tests
- Databricks-Generative-AI-Engineer-Associate Prüfungen ☐ Databricks-Generative-AI-Engineer-Associate Ausbildungsressourcen ➔ Databricks-Generative-AI-Engineer-Associate Schulungsangebot ☐ Geben Sie ☐ www.zertpruefung.ch ☐ ein und suchen Sie nach kostenloser Download von ➤ Databricks-Generative-AI-Engineer-Associate ☐ ☐ Databricks-Generative-AI-Engineer-Associate Prüfungsübungen
- Databricks-Generative-AI-Engineer-Associate Übungsmaterialien - Databricks-Generative-AI-Engineer-Associate

Databricks-Generative-AI-Engineer-Associate □ □Databricks-Generative-AI-Engineer-Associate Lerntipps

- Kostenlose und neue Databricks-Generative-AI-Engineer-Associate Prüfungsfragen sind auf Google Drive freigegeben von rufung verfügbar: https://drive.google.com/open?id=1uFIYKYhKLova_quffHxmzrUzHraBEuc

P.S. Kostenlose und neue Databricks-Generative-AI-Engineer-Associate Prüfungsfragen sind auf Google Drive freigegeben von ZertPruefung verfügbar: https://drive.google.com/open?id=1uFIYKYhKLova_quffHxmzrUzHraBEuc