

# Databricks-Generative-AI-Engineer-Associate Echte Fragen & Databricks-Generative-AI-Engineer-Associate Demotesten



Das Databricks Databricks-Generative-AI-Engineer-Associate Zertifikat kann nicht nur Ihre Fähigkeiten, sondern auch Ihre Fachkenntnisse und Erfahrungen beweisen. Der Boss hat Sie doch nicht umsonst eingestellt. Zur Zeit braucht IT-Branche eine zuverlässige Ressourcen zur Databricks Databricks-Generative-AI-Engineer-Associate Zertifizierungsprüfung. PrüfungFrage ist eine gute Wahl. Sie können die Databricks Databricks-Generative-AI-Engineer-Associate Prüfung in kurzer Zeit bestehen, ohne viel Zeit und Energie zu verwenden, und eine glänzende Zukunft haben.

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

| Thema   | Einzelheiten                                                                                                                                                                                                                                                                                                                                                                                                                                                                |
|---------|-----------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Thema 1 | <ul style="list-style-type: none"><li>• Application Development: In this topic, Generative AI Engineers learn about tools needed to extract data, Langchain</li><li>• similar tools, and assessing responses to identify common issues. Moreover, the topic includes questions about adjusting an LLM's response, LLM guardrails, and the best LLM based on the attributes of the application.</li></ul>                                                                    |
| Thema 2 | <ul style="list-style-type: none"><li>• Governance: Generative AI Engineers who take the exam get knowledge about masking techniques, guardrail techniques, and legal</li><li>• licensing requirements in this topic.</li></ul>                                                                                                                                                                                                                                             |
| Thema 3 | <ul style="list-style-type: none"><li>• Assembling and Deploying Applications: In this topic, Generative AI Engineers get knowledge about coding a chain using a pyfunc mode, coding a simple chain using langchain, and coding a simple chain according to requirements. Additionally, the topic focuses on basic elements needed to create a RAG application. Lastly, the topic addresses sub-topics about registering the model to Unity Catalog using MLflow.</li></ul> |

>> Databricks-Generative-AI-Engineer-Associate Echte Fragen <<

## Databricks Databricks-Generative-AI-Engineer-Associate Demotesten - Databricks-Generative-AI-Engineer-Associate Zertifikatsfragen

Die Fragenpool zur Databricks Databricks-Generative-AI-Engineer-Associate Zertifizierungsprüfung von PrüfungFrage werden

nach dem gleichen Lernplan bearbeitet. Wir aktualisieren auch ständig unsere Fragenpool, die Prüfungsfragen und Antworten enthalten. Weil unsere Prüfungen den echten Prüfungen sehr ähnlich sind, ist unsere Erfolgsquote auch sehr hoch. Diese Tatsache ist nicht zu leugnen, Unsere Fragenpool zur Databricks Databricks-Generative-AI-Engineer-Associate Zertifizierung können den Kandidaten sehr helfen. Und unser Preis ist ganz rational, was jedem IT-Kandidaten passt.

## **Databricks Certified Generative AI Engineer Associate Databricks-Generative-AI-Engineer-Associate Prüfungsfragen mit Lösungen (Q15-Q20):**

### **15. Frage**

A Generative AI Engineer needs to design an LLM pipeline to conduct multi-stage reasoning that leverages external tools. To be effective at this, the LLM will need to plan and adapt actions while performing complex reasoning tasks. Which approach will do this?

- **A. Implement a framework like ReAct which allows the LLM to generate reasoning traces and perform task-specific actions that leverage external tools if necessary.**
- B. Use a Chain-of-Thought (CoT) prompting technique to guide the LLM through a series of reasoning steps, then manually input the results from external tools for the final answer.
- C. Train the LLM to generate a single, comprehensive response without interacting with any external tools, relying solely on its pre-trained knowledge.
- D. Encourage the LLM to make multiple API calls in sequence without planning or structuring the calls, allowing the LLM to decide when and how to use external tools spontaneously.

**Antwort: A**

Begründung:

The task requires an LLM pipeline for multi-stage reasoning with external tools, necessitating planning, adaptability, and complex reasoning. Let's evaluate the options based on Databricks' recommendations for advanced LLM workflows.

\* Option A: Train the LLM to generate a single, comprehensive response without interacting with any external tools, relying solely on its pre-trained knowledge

\* This approach limits the LLM to its static knowledge base, excluding external tools and multi-stage reasoning. It can't adapt or plan actions dynamically, failing the requirements.

\* Databricks Reference: "External tools enhance LLM capabilities beyond pre-trained knowledge" ("Building LLM Applications with Databricks," 2023).

\* Option B: Implement a framework like ReAct which allows the LLM to generate reasoning traces and perform task-specific actions that leverage external tools if necessary

\* ReAct (Reasoning + Acting) combines reasoning traces (step-by-step logic) with actions (e.g., tool calls), enabling the LLM to plan, adapt, and execute complex tasks iteratively. This meets all requirements: multi-stage reasoning, tool use, and adaptability.

\* Databricks Reference: "Frameworks like ReAct enable LLMs to interleave reasoning and external tool interactions for complex problem-solving" ("Generative AI Cookbook," 2023).

\* Option C: Encourage the LLM to make multiple API calls in sequence without planning or structuring the calls, allowing the LLM to decide when and how to use external tools spontaneously

\* Unstructured, spontaneous API calls lack planning and may lead to inefficient or incorrect tool usage. This doesn't ensure effective multi-stage reasoning or adaptability.

\* Databricks Reference: Structured frameworks are preferred: "Ad-hoc tool calls can reduce reliability in complex tasks" ("Building LLM-Powered Applications").

\* Option D: Use a Chain-of-Thought (CoT) prompting technique to guide the LLM through a series of reasoning steps, then manually input the results from external tools for the final answer

\* CoT improves reasoning but relies on manual tool interaction, breaking automation and adaptability. It's not a scalable pipeline solution.

\* Databricks Reference: "Manual intervention is impractical for production LLM pipelines" ("Databricks Generative AI Engineer Guide").

Conclusion: Option B (ReAct) is the best approach, as it integrates reasoning and tool use in a structured, adaptive framework, aligning with Databricks' guidance for complex LLM workflows.

### **16. Frage**

A Generative AI Engineer is creating an LLM-based application. The documents for its retriever have been chunked to a maximum of 512 tokens each. The Generative AI Engineer knows that cost and latency are more important than quality for this application. They have several context length levels to choose from.

Which will fulfill their need?

- A. context length 32768; smallest model is 14GB and embedding dimension 4096
- B. context length 2048; smallest model is 11GB and embedding dimension 2560
- C. context length 514; smallest model is 0.44GB and embedding dimension 768
- D. context length 512; smallest model is 0.13GB and embedding dimension 384

**Antwort: D**

Begründung:

When prioritizing cost and latency over quality in a Large Language Model (LLM)-based application, it is crucial to select a configuration that minimizes both computational resources and latency while still providing reasonable performance. Here's why D is the best choice:

\* Context length: The context length of 512 tokens aligns with the chunk size used for the documents (maximum of 512 tokens per chunk). This is sufficient for capturing the needed information and generating responses without unnecessary overhead.

\* Smallest model size: The model with a size of 0.13GB is significantly smaller than the other options.

This small footprint ensures faster inference times and lower memory usage, which directly reduces both latency and cost.

\* Embedding dimension: While the embedding dimension of 384 is smaller than the other options, it is still adequate for tasks where cost and speed are more important than precision and depth of understanding.

This setup achieves the desired balance between cost-efficiency and reasonable performance in a latency-sensitive, cost-conscious application.

### 17. Frage

A Generative AI Engineer is building a Generative AI system that suggests the best matched employee team member to newly scoped projects. The team member is selected from a very large team. The match should be based upon project date availability and how well their employee profile matches the project scope. Both the employee profile and project scope are unstructured text. How should the Generative AI Engineer architect their system?

- A. Create a tool for finding available team members given project dates. Embed team profiles into a vector store and use the project scope and filtering to perform retrieval to find the available best matched team members.
- B. Create a tool to find available team members given project dates. Create a second tool that can calculate a similarity score for a combination of team member profile and the project scope. Iterate through the team members and rank by best score to select a team member.
- C. Create a tool for finding team member availability given project dates, and another tool that uses an LLM to extract keywords from project scopes. Iterate through available team members' profiles and perform keyword matching to find the best available team member.
- D. Create a tool for finding available team members given project dates. Embed all project scopes into a vector store, perform a retrieval using team member profiles to find the best team member.

**Antwort: A**

Begründung:

\* Problem Context: The problem involves matching team members to new projects based on two main factors:

\* Availability: Ensure the team members are available during the project dates.

\* Profile-Project Match: Use the employee profiles (unstructured text) to find the best match for a project's scope (also unstructured text).

The two main inputs are the employee profiles and project scopes, both of which are unstructured. This means traditional rule-based systems (e.g., simple keyword matching) would be inefficient, especially when working with large datasets.

\* Explanation of Options: Let's break down the provided options to understand why D is the most optimal answer.

\* Option A suggests embedding project scopes into a vector store and then performing retrieval using team member profiles. While embedding project scopes into a vector store is a valid technique, it skips an important detail: the focus should primarily be on embedding employee profiles because we're matching the profiles to a new project, not the other way around.

\* Option B involves using a large language model (LLM) to extract keywords from the project scope and perform keyword matching on employee profiles. While LLMs can help with keyword extraction, this approach is too simplistic and doesn't leverage advanced retrieval techniques like vector embeddings, which can handle the nuanced and rich semantics of unstructured data. This approach may miss out on subtle but important similarities.

\* Option C suggests calculating a similarity score between each team member's profile and project scope. While this is a good idea, it doesn't specify how to handle the unstructured nature of data efficiently. Iterating through each member's profile individually could be computationally expensive in large teams. It also lacks the mention of using a vector store or an efficient retrieval mechanism.

\* Option D is the correct approach. Here's why:

\* Embedding team profiles into a vector store: Using a vector store allows for efficient similarity searches on unstructured data.

Embedding the team member profiles into vectors captures their semantics in a way that is far more flexible than keyword-based

matching.

\* Using project scope for retrieval: Instead of matching keywords, this approach suggests using vector embeddings and similarity search algorithms (e.g., cosine similarity) to find the team members whose profiles most closely align with the project scope.

\* Filtering based on availability: Once the best-matched candidates are retrieved based on profile similarity, filtering them by availability ensures that the system provides a practically useful result.

This method efficiently handles large-scale datasets by leveraging vector embeddings and similarity search techniques, both of which are fundamental tools in Generative AI engineering for handling unstructured text.

\* Technical References:

\* Vector embeddings: In this approach, the unstructured text (employee profiles and project scopes) is converted into high-dimensional vectors using pretrained models (e.g., BERT, Sentence-BERT, or custom embeddings). These embeddings capture the semantic meaning of the text, making it easier to perform similarity-based retrieval.

\* Vector stores: Solutions like FAISS or Milvus allow storing and retrieving large numbers of vector embeddings quickly. This is critical when working with large teams where querying through individual profiles sequentially would be inefficient.

\* LLM Integration: Large language models can assist in generating embeddings for both employee profiles and project scopes. They can also assist in fine-tuning similarity measures, ensuring that the retrieval system captures the nuances of the text data.

\* Filtering: After retrieving the most similar profiles based on the project scope, filtering based on availability ensures that only team members who are free for the project are considered.

This system is scalable, efficient, and makes use of the latest techniques in Generative AI, such as vector embeddings and semantic search.

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

**Antwort: C**

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.

### 19. Frage

A Generative AI Engineer is tasked with developing an application that is based on an open source large language model (LLM). They need a foundation LLM with a large context window.

Which model fits this need?

- A. DistilBERT
- B. DBRX
- C. Llama2-70B
- D. MPT-30B

**Antwort: C**

Begründung:

\* Problem Context: The engineer needs an open-source LLM with a large context window to develop an application.

\* Explanation of Options:

\* Option A: DistilBERT: While an efficient and smaller version of BERT, DistilBERT does not provide a particularly large context window.

\* Option B: MPT-30B: This model, while large, is not specified as being particularly notable for its context window capabilities.

\* Option C: Llama2-70B: Known for its large model size and extensive capabilities, including a large context window. It is also available as an open-source model, making it ideal for applications requiring extensive contextual understanding.

\* Option D: DBRX: This is not a recognized standard model in the context of large language models with extensive context windows.

Thus, Option C (Llama2-70B) is the best fit as it meets the criteria of having a large context window and being available for open-source use, suitable for developing robust language understanding applications.

### 20. Frage

.....

PrüfungFrage ist der beste Katalysator für den Erfolg der IT-Fachleute, Viele Kandidaten, die Databricks Databricks-Generative-AI-Engineer-Associate IT-Zertifizierungsprüfungen bestanden haben, haben Schulungsunterlagen von PrüfungFrage benutzt. Unser Expertenteam von PrüfungFrage hat die neuesten und effizientesten Prüfungsfragen und Antworten zur Databricks Databricks-Generative-AI-Engineer-Associate Zertifizierungsteste.

**Databricks-Generative-AI-Engineer-Associate Demotesten:** <https://www.pruefungfrage.de/Databricks-Generative-AI-Engineer-Associate-dumps-deutsch.html>

- Databricks-Generative-AI-Engineer-Associate Prüfungsunterlagen ☐ Databricks-Generative-AI-Engineer-Associate Musterprüfungsfragen ☐ Databricks-Generative-AI-Engineer-Associate Demotesten ☐ Öffnen Sie die Website **【** [www.pass4test.de](http://www.pass4test.de) **】** Suchen Sie ☒ Databricks-Generative-AI-Engineer-Associate ☐ ☒ Kostenloser Download ☐ ☐ Databricks-Generative-AI-Engineer-Associate Prüfungen
- Databricks-Generative-AI-Engineer-Associate Deutsch Prüfung ☐ Databricks-Generative-AI-Engineer-Associate Deutsch Prüfung ☐ Databricks-Generative-AI-Engineer-Associate Zertifizierungsantworten ☐ Öffnen Sie die Website  $\Rightarrow$  [www.itzert.com](http://www.itzert.com)  $\Leftarrow$  Suchen Sie ☒ Databricks-Generative-AI-Engineer-Associate ☐ ☒ Kostenloser Download ☐ ☐ Databricks-Generative-AI-Engineer-Associate Prüfungen
- Kostenlos Databricks-Generative-AI-Engineer-Associate Dumps Torrent - Databricks-Generative-AI-Engineer-Associate exams4sure pdf - Databricks Databricks-Generative-AI-Engineer-Associate pdf vce ☐ **【** [www.zertpruefung.ch](http://www.zertpruefung.ch) **】** ist die beste Webseite um den kostenlosen Download von  $\ll$  Databricks-Generative-AI-Engineer-Associate  $\gg$  zu erhalten ☐ ☐ Databricks-Generative-AI-Engineer-Associate Prüfungsübungen
- Databricks-Generative-AI-Engineer-Associate Ausbildungsressourcen ☐ Databricks-Generative-AI-Engineer-Associate Demotesten ☐ Databricks-Generative-AI-Engineer-Associate Online Prüfung ☐ Suchen Sie auf **【** [www.itzert.com](http://www.itzert.com) **】** nach **【** Databricks-Generative-AI-Engineer-Associate **】** und erhalten Sie den kostenlosen Download mühelos ☐ ☐ Databricks-Generative-AI-Engineer-Associate Prüfungsübungen
- Databricks-Generative-AI-Engineer-Associate Prüfungsfragen, Databricks-Generative-AI-Engineer-Associate Fragen und Antworten, Databricks Certified Generative AI Engineer Associate ☐ Öffnen Sie die Webseite  $\gg$  [www.zertfragen.com](http://www.zertfragen.com) ☐ und suchen Sie nach kostenloser Download von { Databricks-Generative-AI-Engineer-Associate } ☐ ☐ Databricks-Generative-AI-Engineer-Associate Prüfungsunterlagen
- Databricks-Generative-AI-Engineer-Associate Prüfungsfragen, Databricks-Generative-AI-Engineer-Associate Fragen und

[illegible]