

AIP-C01 aktueller Test, Test VCE-Dumps für AWS Certified Generative AI Developer - Professional



Mit ZertSoft können Sie ganz leicht die Amazon AIP-C01 Prüfung bestehen. Wenn Sie die Amazon AIP-C01 Schulungsunterlagen im ZertSoft wählen und Amazon AIP-C01 die Prüfungsfragen und Antworten zur Zertifizierungsprüfung herunterladen, werden Sie sicher selbstbewusster sein, dass Sie die Prüfung ganz leicht bestehen können. Obwohl es auch andere Prüfungsunterlagen zur Amazon AIP-C01 Zertifizierungsprüfung auf andere Websites gibt, versprechen wir Ihnen, dass unsere Produkte am besten sind. Unsere Übungsfragen-und antworten sind sehr präzise. Sie umfassen viele Wissensgebiete. Sie sind immer erneuert und ergänzt. Deshalb steht unser ZertSoft Ihnen eine genaue Prüfungsvorbereitung zur Verfügung. Wenn Sie ZertSoft wählen, können Sie viel Zeit ersparen, ganz leicht und schnell die Amazon AIP-C01 Zertifizierungsprüfung bestehen und so schnell wie möglich ein IT-Fachmann in der Amazon IT-Branche werden.

Amazon AIP-C01 Prüfungsplan:

Thema	Einzelheiten
Thema 1	• Testing, Validation, and Troubleshooting: This domain covers evaluating foundation model outputs, implementing quality assurance processes, and troubleshooting GenAI-specific issues including prompts, integrations, and retrieval systems.
Thema 2	• AI Safety, Security, and Governance: This domain addresses input • output safety controls, data security and privacy protections, compliance mechanisms, and responsible AI principles including transparency and fairness.
Thema 3	• Implementation and Integration: This domain focuses on building agentic AI systems, deploying foundation models, integrating GenAI with enterprise systems, implementing FM APIs, and developing applications using AWS tools.
Thema 4	• Operational Efficiency and Optimization for GenAI Applications: This domain encompasses cost optimization strategies, performance tuning for latency and throughput, and implementing comprehensive monitoring systems for GenAI applications.
Thema 5	• Foundation Model Integration, Data Management, and Compliance: This domain covers designing GenAI architectures, selecting and configuring foundation models, building data pipelines and vector stores, implementing retrieval mechanisms, and establishing prompt engineering governance.

>> AIP-C01 Lernressourcen <<

bestehen Sie AIP-C01 Ihre Prüfung mit unserem Prep AIP-C01 Ausbildung Material & kostenloser Dowload Torrent

ZertSoft hat einen guten Online-Service. Wenn Sie die Produkte von ZertSoft kaufen, wird ZertSoft Ihnen einen einjährigen kostenlos Update-Service rund um die Uhr bieten. Wir benachrichtigen Ihnen rechtzeitig die neuesten Prüfungsinformationen zur Amazon AIP-C01 Zertifizierung, so dass Sie sich gut auf die Amazon AIP-C01 Zertifizierungsprüfung vorbereiten können. Mit wenig Zeit und Geld können Sie die IT-Prüfung bestehen. Es ist sehr preisgünstig, ZertSoft zu wählen und somit die Amazon AIP-C01 Zertifizierungsprüfung nur einmal zu bestehen.

Amazon AWS Certified Generative AI Developer - Professional AIP-C01 Prüfungsfragen mit Lösungen (Q25-Q30):

25. Frage

A bank is developing a generative AI (GenAI)-powered AI assistant that uses Amazon Bedrock to assist the bank's website users with account inquiries and financial guidance. The bank must ensure that the AI assistant does not reveal any personally identifiable information (PII) in customer interactions.

The AI assistant must not send PII in prompts to the GenAI model. The AI assistant must not respond to customer requests to provide investment advice. The bank must collect audit logs of all customer interactions, including any images or documents that are transmitted during customer interactions.

Which solution will meet these requirements with the LEAST operational effort?

- **A. Configure Amazon Bedrock guardrails to apply a sensitive information policy to detect and filter PII. Set up a topic policy to ensure that the AI assistant avoids investment advice topics. Use the Converse API to log model invocations. Enable delivery and image logging to Amazon S3.**
- B. Use Amazon Macie to detect and redact PII in user inputs and in the model responses. Apply prompt engineering techniques to force the model to avoid investment advice topics. Use AWS CloudTrail to capture conversation logs.
- C. Use regex controls to match patterns for PII. Apply prompt engineering techniques to avoid returning PII or investment advice topics to customers. Enable model invocation logging, delivery logging, and image logging to Amazon S3.
- D. Use an AWS Lambda function and Amazon Comprehend to detect and redact PII. Use Amazon Comprehend topic modeling to prevent the AI assistant from discussing investment advice topics. Set up custom metrics in Amazon CloudWatch to capture customer conversations.

Antwort: A

Begründung:

Option C is the correct solution because Amazon Bedrock guardrails are purpose-built to enforce defense-in-depth safety controls for GenAI applications with minimal operational overhead. Guardrails provide managed, policy-based enforcement that operates before prompts are sent to the foundation model and after responses are generated, which directly satisfies the requirement that PII must not be sent to the model and must not appear in outputs.

By configuring a sensitive information policy, the application can automatically detect and redact PII in user inputs and model responses without building custom preprocessing pipelines. This approach is more reliable and scalable than regex or prompt engineering techniques, which are brittle and error-prone for sensitive data handling.

The topic policy capability in Amazon Bedrock guardrails allows the bank to explicitly block investment advice topics, ensuring regulatory compliance. This policy-based approach is safer and more auditable than attempting to steer the model only through prompt instructions.

Using the Converse API enables structured, standardized interactions with the model and supports consistent logging of requests and responses. Enabling delivery logging and image logging to Amazon S3 ensures that all customer interactions, including documents and images, are captured in a durable, auditable storage layer.

This directly supports compliance, regulatory audits, and forensic analysis.

Option A incorrectly relies on Amazon Macie, which is designed for data-at-rest discovery rather than real-time conversational filtering. Option B introduces custom Lambda pipelines and topic modeling, increasing operational complexity. Option D relies on regex and prompt engineering, which do not meet financial-grade compliance standards.

Therefore, Option C delivers the strongest security, governance, and auditability with the least operational effort.

26. Frage

An ecommerce company is developing a generative AI (GenAI) solution that uses Amazon Bedrock with Anthropic Claude to recommend products to customers. Customers report that some recommended products are not available for sale or are not relevant. Customers also report long response times for some recommendations.

The company confirms that most customer interactions are unique and that the solution recommends products not present in the product catalog.

Which solution will meet this requirement?

- A. Use prompt engineering to restrict model responses to relevant products. Use streaming inference to reduce perceived latency.
- **B. Create an Amazon Bedrock Knowledge Bases and implement Retrieval Augmented Generation (RAG). Set the PerformanceConfigLatency parameter to optimized.**
- C. Increase grounding within Amazon Bedrock Guardrails. Enable automated reasoning checks. Set up provisioned throughput.
- D. Store product catalog data in Amazon OpenSearch Service. Validate model recommendations against the catalog. Use Amazon DynamoDB for response caching.

Antwort: B

Begründung:

Option C is the correct solution because it directly addresses both correctness and performance issues by grounding the model's responses in authoritative product data using Retrieval Augmented Generation.

Amazon Bedrock Knowledge Bases are designed to connect foundation models to trusted enterprise data sources, ensuring that generated responses are constrained to known, validated content.

By ingesting the product catalog into a knowledge base, the GenAI application retrieves only products that actually exist in the catalog. This prevents hallucinated or unavailable recommendations, which is a common issue when models rely solely on prompt instructions without retrieval grounding. RAG ensures that the model's output is based on retrieved facts rather than learned generalizations.

Setting the PerformanceConfigLatency parameter to optimized enables Bedrock to prioritize lower-latency retrieval and inference paths, improving responsiveness for real-time recommendation scenarios. This directly addresses the reported performance issues without requiring provisioned throughput or caching strategies that are ineffective for mostly unique interactions.

Option A improves safety and latency predictability but does not ensure recommendations are limited to valid products. Option B relies on prompt constraints, which are not sufficient to prevent hallucinations. Option D introduces additional validation and caching layers but increases complexity and does not improve generation relevance.

Therefore, Option C best resolves both relevance and latency challenges using AWS-native, low-maintenance GenAI integration patterns.

27. Frage

A company has deployed an AI assistant as a React application that uses AWS Amplify, an AWS AppSync GraphQL API, and Amazon Bedrock Knowledge Bases. The application uses the GraphQL API to call the Amazon Bedrock RetrieveAndGenerate API for knowledge base interactions. The company configures an AWS Lambda resolver to use the RequestResponse invocation type.

Application users report frequent timeouts and slow response times. Users report these problems more frequently for complex questions that require longer processing.

The company needs a solution to fix these performance issues and enhance the user experience.

Which solution will meet these requirements?

- A. Change the RetrieveAndGenerate API to the InvokeModelWithResponseStream API. Update the application to use an Amazon API Gateway WebSocket API to support the streaming response.
- B. Update the application to send an API request to an Amazon SQS queue. Update the AWS AppSync resolver to poll and process the queue.
- C. Increase the timeout value of the Lambda resolver. Implement retry logic with exponential backoff.
- **D. Use AWS Amplify AI Kit to implement streaming responses from the GraphQL API and to optimize client-side rendering.**

Antwort: D

Begründung:

Option A is the best solution because it directly addresses both observed problems: user-perceived latency and resolver timeouts that occur more frequently for complex prompts. In the current design, an AWS AppSync Lambda resolver is configured with synchronous RequestResponse behavior. That means the client receives nothing until the entire retrieval and generation workflow completes. For longer-running knowledge base queries, this increases the likelihood of hitting request time limits in the synchronous path and creates a poor user experience because the UI appears stalled.

Using AWS Amplify AI Kit to implement streaming responses allows the application to return partial output incrementally as the model produces tokens. This improves perceived responsiveness because users can see the answer forming immediately, even when the full response takes longer. Streaming also reduces the impact of variable model latency and retrieval time because the client no longer waits for a single final payload before rendering content. From a troubleshooting perspective, streaming makes it easier to distinguish "slow generation" from "no response," and it provides faster feedback during testing of complex questions.

Option B is not sufficient because increasing timeouts and adding retries can worsen load and cost while still producing a stalled UI

experience. Retries also risk duplicating requests to the knowledge base and can amplify token usage. Option C introduces an awkward polling model for GraphQL interactions and adds significant operational complexity, while not inherently improving interactivity. Option D adds major architectural changes by replacing the knowledge base RetrieveAndGenerate call path with a different streaming invocation API and introducing a WebSocket layer, which is unnecessary when the goal is primarily to fix timeouts and improve UX within the existing AppSync and Amplify design.

Therefore, streaming through Amplify AI Kit is the most effective and lowest-friction improvement.
Thought for 24s

28. Frage

A medical company is building a generative AI (GenAI) application that uses Retrieval Augmented Generation (RAG) to provide evidence-based medical information. The application uses Amazon OpenSearch Service to retrieve vector embeddings. Users report that searches frequently miss results that contain exact medical terms and acronyms and return too many semantically similar but irrelevant documents. The company needs to improve retrieval quality and maintain low end-user latency, even as the document collection grows to millions of documents.

Which solution will meet these requirements with the LEAST operational overhead?

- A. Increase the dimensions of the vector embeddings from 384 to 1536. Use a post-processing AWS Lambda function to filter out irrelevant results after retrieval.
- B. Replace OpenSearch Service with Amazon Kendra. Use query expansion to handle medical acronyms and terminology variants during pre-processing.
- C. Configure hybrid search by combining vector similarity with keyword matching to improve semantic understanding and exact term and acronym matching.
- D. Implement a two-stage retrieval architecture in which initial vector search results are re-ranked by an ML model hosted on Amazon SageMaker.

Antwort: C

Begründung:

Option A is the correct solution because hybrid search directly addresses the core retrieval failure modes while maintaining low latency and minimal operational overhead. In medical and scientific domains, exact terminology, abbreviations, and acronyms (for example, drug names, procedures, or conditions) are critical.

Pure vector similarity search often underweights these exact matches, leading to missed results and excessive semantically related but irrelevant documents.

Amazon OpenSearch Service natively supports hybrid search, which combines keyword-based retrieval (such as BM25) with vector similarity search. Keyword search ensures precise matching for exact terms and acronyms, while vector search captures semantic meaning and contextual similarity. By blending these approaches, the retrieval system improves both precision and recall without introducing additional infrastructure.

Hybrid search operates within the same OpenSearch index and query path, which preserves low end-user latency even at large scale. This is especially important as the document collection grows to millions of documents. Because OpenSearch handles scoring and ranking internally, no additional orchestration layers or post-processing steps are required.

Option B increases computational cost and latency while failing to address exact-term recall. Option C introduces a new service and ingestion pipeline, increasing operational overhead and latency. Option D adds model hosting, re-ranking infrastructure, and complexity that is unnecessary when OpenSearch provides native hybrid retrieval.

Therefore, Option A delivers the best balance of retrieval quality, scalability, latency, and operational simplicity for medical RAG workloads.

29. Frage

A medical company uses Amazon Bedrock to power a clinical documentation summarization system. The system produces inconsistent summaries when handling complex clinical documents. The system performed well on simple clinical documents. The company needs a solution that diagnoses inconsistencies, compares prompt performance against established metrics, and maintains historical records of prompt versions.

Which solution will meet these requirements?

- A. Implement version control for prompts in a code repository with a test suite that contains complex clinical documents and quantifiable evaluation metrics. Use an automated testing framework to compare prompt versions and document performance patterns.
- B. Create multiple prompt variants by using Prompt management in Amazon Bedrock. Manually test the prompts with simple clinical documents. Deploy the highest performing version by using the Amazon Bedrock console.

- Antwort: A**

Option A lacks objective metrics and does not address complex documents. Option C focuses on live traffic experimentation but does not inherently diagnose prompt inconsistencies or preserve detailed historical evaluations. Option D adds medical entity analysis but introduces unnecessary service coupling and does not provide robust prompt version history or automated comparative benchmarking. Therefore, Option B is the most complete and disciplined solution.

• • • • •

[illegible]

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, 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, akademi.jadipns.com, Disposable vapes