

# AIP-C01 Practice Exam - AIP-C01 Official Practice Test



BONUS!!! Download part of CertkingdomPDF AIP-C01 dumps for free: <https://drive.google.com/open?id=1j-CpfWrJuWua5fkGbAzinYLEJMG6sKCI>

In today's rapidly changing Amazon industry, the importance of obtaining Amazon AIP-C01 certification has become increasingly evident. With the constant evolution of technology, staying competitive in the job market requires professionals to continuously upgrade their skills and knowledge. The CertkingdomPDF is committed to completely assisting you in exam preparation with AIP-C01 Questions. Success in the AWS Certified Generative AI Developer - Professional (AIP-C01) certification exam is crucial in the tech sector, where the stakes are high, and a single mistake can have significant consequences.

For offline practice, our AWS Certified Generative AI Developer - Professional (AIP-C01) desktop practice test software is ideal. This AWS Certified Generative AI Developer - Professional (AIP-C01) software runs on Windows computers. The AWS Certified Generative AI Developer - Professional (AIP-C01) web-based practice exam is compatible with all browsers and operating systems. No software installation is required to go through the web-based AWS Certified Generative AI Developer - Professional (AIP-C01) practice test.

>> AIP-C01 Practice Exam <<

## AIP-C01 Official Practice Test, AIP-C01 Testking

CertkingdomPDF also offers the AIP-C01 web-based practice exam with the same characteristics as desktop simulation software but with minor differences. It is online AIP-C01 Certification Exam which is accessible from any location with an active internet connection. This Amazon AIP-C01 Practice Exam not only works on Windows but also on Linux, Mac, Android, and iOS.

Additionally, you can attempt the Amazon AIP-C01 practice test through these browsers: Opera, Safari, Firefox, Chrome, MS Edge, and Internet Explorer.

## Amazon AIP-C01 Exam Syllabus Topics:

| Topic   | Details                                                                                                                                                                                                                                                                                                                                     |
|---------|---------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Topic 1 | <ul style="list-style-type: none"><li>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.</li></ul>                                                              |
| Topic 2 | <ul style="list-style-type: none"><li>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.</li></ul>                                              |
| Topic 3 | <ul style="list-style-type: none"><li>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.</li></ul> |
| Topic 4 | <ul style="list-style-type: none"><li>AI Safety, Security, and Governance: This domain addresses input</li><li>output safety controls, data security and privacy protections, compliance mechanisms, and responsible AI principles including transparency and fairness.</li></ul>                                                           |
| Topic 5 | <ul style="list-style-type: none"><li>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.</li></ul>                                               |

## Amazon AWS Certified Generative AI Developer - Professional Sample Questions (Q48-Q53):

### NEW QUESTION # 48

A company uses Amazon Bedrock to implement a Retrieval Augmented Generation (RAG)-based system to serve medical information to users. The company needs to compare multiple chunking strategies, evaluate the generation quality of two foundation models (FMs), and enforce quality thresholds for deployment.

Which Amazon Bedrock evaluation configuration will meet these requirements?

- A. Set up a pipeline that uses multiple retrieve-only evaluation jobs to assess retrieval quality. Create separate evaluation jobs for both FMs that use Amazon Nova Pro as the LLM-as-a-judge model. Evaluate based on faithfulness and citation precision metrics.
- B. Create a retrieve-and-generate evaluation job that uses custom precision-at-k metrics and an LLM-as-a-judge metric with a scale of 1-5. Include each chunking strategy in the evaluation dataset. Use a supported version of Anthropic Claude Sonnet to evaluate responses from both FMs.
- C. Create a separate evaluation job for each chunking strategy and FM combination. Use Amazon Bedrock built-in metrics for correctness and completeness. Manually review scores before deployment approval.
- D. Create a retrieve-only evaluation job that uses a supported version of Anthropic Claude Sonnet as the evaluator model. Configure metrics for context relevance and context coverage. Define deployment thresholds in a separate CI/CD pipeline.

### Answer: B

Explanation:

Option B is the correct evaluation configuration because it enables end-to-end assessment of both retrieval and generation quality while supporting direct comparison of chunking strategies and foundation models.

Amazon Bedrock evaluation jobs are designed to support RAG workflows by evaluating how well retrieved context supports accurate and high-quality model outputs.

A retrieve-and-generate evaluation job evaluates the complete RAG pipeline, not just retrieval. This is essential for medical information use cases, where both the relevance of retrieved content and the correctness of generated responses directly impact user safety and trust. Including multiple chunking strategies in the evaluation dataset allows side-by-side comparison under identical prompts and conditions.

Custom precision-at-k metrics measure how effectively the retrieval component surfaces relevant chunks, while an LLM-as-a-judge

metric provides qualitative scoring of generated responses. Using a numeric scale enables consistent, repeatable evaluation and supports automated quality gates. Amazon Bedrock supports LLM-based evaluators to score dimensions such as accuracy, completeness, and relevance.

Using the same evaluator model to assess outputs from both FMs ensures consistent scoring and eliminates evaluator bias. This configuration allows the company to define quantitative thresholds that must be met before deployment, enabling automated promotion through CI/CD pipelines.

Option A evaluates retrieval only and cannot assess generation quality. Option C introduces manual review, which does not scale and delays deployment. Option D separates retrieval and generation evaluation, making it harder to correlate chunking strategies with final output quality.

Therefore, Option B best meets the requirements for systematic evaluation, comparison, and quality enforcement in an Amazon Bedrock-based RAG system.

### NEW QUESTION # 49

A company is building a multicloud generative AI (GenAI)-powered secret resolution application that uses Amazon Bedrock and Agent Squad. The application resolves secrets from multiple sources, including key stores and hardware security modules (HSMs). The application uses AWS Lambda functions to retrieve secrets from the sources. The application uses AWS AppConfig to implement dynamic feature gating. The application supports secret chaining and detects secret drift. The application handles short-lived and expiring secrets. The application also supports prompt flows for templated instructions. The application uses AWS Step Functions to orchestrate agents to resolve the secrets and to manage secret validation and drift detection.

The company finds multiple issues during application testing. The application does not refresh expired secrets in time for agents to use. The application sends alerts for secret drift, but agents still use stale data. Prompt flows within the application reuse outdated templates, which cause cascading failures. The company must resolve the performance issues.

Which solution will meet this requirement?

- A. Use Amazon Bedrock Agents only. Configure Amazon Bedrock guardrails to restrict prompt variation. Use an inline JSON schema for a single agent's workflow definition to chain tool calls.
- **B. Use Step Functions Map states to run agent workflows in parallel. Pass updated secret metadata through Lambda function outputs. Use AWS AppConfig to version all prompt flows to gate and roll back faulty templates.**
- C. Use a centralized Amazon EventBridge pipeline to invoke each agent. Store intermediate prompts in Amazon DynamoDB. Resolve agent ordering by using TTL-based backoff and retries.
- D. Use Amazon EventBridge Pipes to invoke resolvers based on Amazon CloudWatch log patterns. Store response metadata in DynamoDB with TTL and versioned writes. Use Amazon Q Developer to dynamically generate fallback prompts.

**Answer: B**

Explanation:

Option A is the correct solution because it directly addresses all identified failure modes while preserving the existing Step Functions-based orchestration architecture with minimal redesign.

Using Step Functions Map states enables parallel execution of secret resolution workflows, which improves refresh latency for short-lived and expiring secrets. This ensures that secrets are refreshed in time before downstream agents require them. Passing updated secret metadata through Lambda outputs guarantees that subsequent steps always consume the latest resolved values, preventing agents from using stale data even after drift alerts are generated.

Versioning prompt flows in AWS AppConfig is critical to resolving cascading failures caused by outdated templates. AppConfig natively supports version control, validation, staged rollout, and rollback of configuration artifacts. By gating prompt flows through AppConfig, the company can immediately roll back faulty templates and prevent agents from reusing outdated instructions.

This solution maintains clear separation of concerns: Step Functions handle orchestration and parallelism, Lambda handles secret retrieval and metadata propagation, and AppConfig governs prompt lifecycle management. No additional event pipelines or custom retry coordination layers are required.

Option B oversimplifies the architecture and does not address secret lifecycle or drift. Option C introduces event-driven ordering complexity without solving prompt versioning. Option D introduces unnecessary tooling and dynamic prompt generation risk. Therefore, Option A best resolves performance, correctness, and stability issues while minimizing operational overhead.

### NEW QUESTION # 50

A company uses an AI assistant application to summarize the company's website content and provide information to customers. The company plans to use Amazon Bedrock to give the application access to a foundation model (FM).

The company needs to deploy the AI assistant application to a development environment and a production environment. The solution must integrate the environments with the FM. The company wants to test the effectiveness of various FMs in each

environment. The solution must provide product owners with the ability to easily switch between FMs for testing purposes in each environment.

Which solution will meet these requirements?

- A. Create a separate AWS CDK application for each environment. Configure the applications to invoke the Amazon Bedrock FMs by using the `aws_bedrock.FoundationModel.fromFoundationModelId()` method. Create a separate pipeline in AWS CodePipeline for each environment.
- **B. Create one AWS CDK application. Configure the application to invoke the Amazon Bedrock FMs by using the `aws_bedrock.FoundationModel.fromFoundationModelId()` method. Create a pipeline in AWS CodePipeline that has a deployment stage for each environment that uses AWS CodeBuild deploy actions.**
- C. Create one AWS CDK application. Create multiple pipelines in AWS CodePipeline. Configure each pipeline to have its own settings for each FM. Configure the application to invoke the Amazon Bedrock FMs by using the `aws_bedrock.ProvisionedModel.fromProvisionedModelArn()` method.
- D. Create one AWS CDK application for the production environment. Configure the application to invoke the Amazon Bedrock FMs by using the `aws_bedrock.ProvisionedModel.fromProvisionedModelArn()` method. Create a pipeline in AWS CodePipeline. Configure the pipeline to deploy to the production environment by using an AWS CodeBuild deploy action. For the development environment, manually recreate the resources by referring to the production application code.

**Answer: B**

Explanation:

Option C best satisfies the requirement for flexible FM testing across environments while minimizing operational complexity and aligning with AWS-recommended deployment practices. Amazon Bedrock supports invoking on-demand foundation models through the `FoundationModel` abstraction, which allows applications to dynamically reference different models without requiring dedicated provisioned capacity. This is ideal for experimentation and A/B testing in both development and production environments. Using a single AWS CDK application ensures infrastructure consistency and reduces duplication.

Environment-specific configuration, such as selecting different foundation model IDs, can be externalized through parameters, context variables, or environment-specific configuration files. This allows product owners to easily switch between FMs in each environment without modifying application logic.

A single AWS CodePipeline with distinct deployment stages for development and production is an AWS best practice for multi-environment deployments. It enforces consistent build and deployment steps while still allowing environment-level customization. AWS CodeBuild deploy actions enable automated, repeatable deployments, reducing manual errors and improving governance. Option A increases complexity by introducing multiple pipelines and relies on provisioned models, which are not necessary for FM evaluation and experimentation. Provisioned throughput is better suited for predictable, high-volume production workloads rather than frequent model switching.

Option B creates unnecessary operational overhead by duplicating CDK applications and pipelines, making long-term maintenance more difficult.

Option D directly conflicts with infrastructure-as-code best practices by manually recreating development resources, which increases configuration drift and reduces reliability.

Therefore, Option C provides the most flexible, scalable, and AWS-aligned solution for testing and switching foundation models across development and production environments.

## NEW QUESTION # 51

A company uses an application to process customer support tickets. The company wants to integrate AI- powered sentiment analysis and auto-response generation into the application by using Amazon Bedrock. The company wants to prioritize urgent issues and reduce initial response times by 40% compared to manual responses. The solution must process 100 concurrent webhook requests with response times under 500 ms.

The solution must maintain 99.9% availability across multiple AWS Regions and authenticate all incoming requests. The company must avoid any authentication failures. The company does not want to modify the existing application infrastructure, which includes several ticketing systems that use multiple webhook authentication methods. The solution must support scaling to handle occasional spikes up to 250,000 daily tickets during peak periods. Which solution will meet these requirements?

- **A. Use an Amazon API Gateway REST API with a Regional endpoint to receive webhook requests and invoke AWS Lambda functions. Configure Lambda authorizers to validate all the webhook authentication methods. Configure the Lambda functions to call Amazon Bedrock to perform sentiment analysis and generate responses. Store results in Amazon DynamoDB global tables to provide multi- Region availability.**
- B. Create AWS Lambda function URLs for each ticketing system. Configure the function URLs with the NONE authentication type. Configure separate Lambda functions to verify webhook signatures by using Hash-based Message Authentication Code (HMAC) validation in the function code. Deploy the functions to multiple Regions and use AWS Global Accelerator to route traffic. Use Amazon Bedrock to perform sentiment analysis and generate responses. Return responses

through webhook callbacks.

- C. Set up an Amazon SQS queue in each Region to receive webhook messages. Use the SQS queue to invoke AWS Lambda functions that call Amazon Comprehend to perform sentiment analysis and Amazon Lex to generate responses. Use Amazon EventBridge to retry message delivery to the application API.
- D. Deploy an AWS AppSync GraphQL API to multiple Regions. Configure API tokens to authenticate incoming requests. Create GraphQL mutation resolvers that publish events to Amazon EventBridge. Configure EventBridge rules to invoke AWS Lambda functions that use Amazon Bedrock to perform sentiment analysis and generate responses. Use Amazon CloudFront to reduce latency.

**Answer: A**

Explanation:

To handle high concurrency (100+ requests) with sub-500 ms response times and diverse authentication methods without infrastructure changes, Amazon API Gateway with Lambda authorizers is the optimal choice. The Lambda authorizers can evaluate multiple authentication tokens or signatures centrally before the request reaches the processing logic, preventing unauthorized traffic and potential authentication failures at scale. AWS Lambda integrated with Amazon Bedrock provides the scalability to handle ticket surges (up to 250,000 daily) without over-provisioning resources. For high availability (99.9%) and multi-region resilience, storing the resulting sentiment and responses in Amazon DynamoDB global tables ensures that data is accessible across regions with minimal latency. Option B is less secure due to the "NONE" auth type, and Option C introduces queuing latency that may exceed the 500 ms target.

## NEW QUESTION # 52

A company is using Amazon Bedrock and Anthropic Claude 3 Haiku to develop an AI assistant. The AI assistant normally processes 10,000 requests each hour but experiences surges of up to 30,000 requests each hour during peak usage periods. The AI assistant must respond within 2 seconds while operating across multiple AWS Regions.

The company observes that during peak usage periods, the AI assistant experiences throughput bottlenecks that cause increased latency and occasional request timeouts. The company must resolve the performance issues.

Which solution will meet this requirement?

- A. Purchase provisioned throughput and sufficient model units (MUs) in a single Region. Configure the application to retry failed requests with exponential backoff.
- B. Set up auto scaling AWS Lambda functions in each Region. Implement client-side round-robin request distribution. Purchase one model unit (MU) of provisioned throughput as a backup.
- C. Implement token batching to reduce API overhead. Use cross-Region inference profiles to automatically distribute traffic across available Regions.
- D. Implement batch inference for all requests by using Amazon S3 buckets across multiple Regions. Use Amazon SQS to set up an asynchronous retrieval process.

**Answer: C**

Explanation:

Option B is the correct solution because it directly addresses both throughput bottlenecks and latency requirements using native Amazon Bedrock performance optimization features that are designed for real-time, high-volume generative AI workloads.

Amazon Bedrock supports cross-Region inference profiles, which allow applications to transparently route inference requests across multiple AWS Regions. During peak usage periods, traffic is automatically distributed to Regions with available capacity, reducing throttling, request queuing, and timeout risks. This approach aligns with AWS guidance for building highly available, low-latency GenAI applications that must scale elastically across geographic boundaries.

Token batching further improves efficiency by combining multiple inference requests into a single model invocation where applicable. AWS Generative AI documentation highlights batching as a key optimization technique to reduce per-request overhead, improve throughput, and better utilize model capacity. This is especially effective for lightweight, low-latency models such as Claude 3 Haiku, which are designed for fast responses and high request volumes.

Option A does not meet the requirement because purchasing provisioned throughput in a single Region creates a regional bottleneck and does not address multi-Region availability or traffic spikes beyond reserved capacity. Retries increase load and latency rather than resolving the root cause.

Option C improves application-layer scaling but does not solve model-side throughput limits. Client-side round-robin routing lacks awareness of real-time model capacity and can still send traffic to saturated Regions.

Option D is unsuitable because batch inference with asynchronous retrieval is designed for offline or non-interactive workloads. It cannot meet a strict 2-second response time requirement for an interactive AI assistant.

Therefore, Option B provides the most effective and AWS-aligned solution to achieve low latency, global scalability, and high



throughput during peak usage periods.

## NEW QUESTION # 53

.....

Amazon AIP-C01 practice materials are highly popular in the market compared with other materials from competitors whether on the volume of sales or content as well. All precise information on the AWS Certified Generative AI Developer - Professional AIP-C01 Exam Questions and high accurate questions are helpful. To help you have a thorough understanding of our AIP-C01 training prep, free demos are provided for your reference.

**AIP-C01 Official Practice Test:** <https://www.certkingdompdf.com/AIP-C01-latest-certkingdom-dumps.html>

- AIP-C01 Free Test Questions ☐ AIP-C01 Flexible Testing Engine ☐ Reliable AIP-C01 Test Syllabus ☐ Simply search for ✓ AIP-C01 ☐ for free download on **【 www.prep4away.com 】** ☐ AIP-C01 Flexible Testing Engine
- AIP-C01 Reliable Exam Materials ☐ Exam AIP-C01 Guide Materials ☐ Valid AIP-C01 Exam Experience ☐ Immediately open ⇒ [www.pdfvce.com](http://www.pdfvce.com) ⇐ and search for ( AIP-C01 ) to obtain a free download ☐ Best AIP-C01 Study Material
- 100% Pass 2026 Amazon Latest AIP-C01: AWS Certified Generative AI Developer - Professional Practice Exam ☐ Enter ⇒ [www.prepawayexam.com](http://www.prepawayexam.com) ⇐ and search for ( AIP-C01 ) to download for free ☐ Exam AIP-C01 Guide Materials
- Free PDF Quiz 2026 Amazon AIP-C01: AWS Certified Generative AI Developer - Professional Updated Practice Exam ☐ Easily obtain ➤ AIP-C01 ☐ for free download through 「 [www.pdfvce.com](http://www.pdfvce.com) 」 ☐ Test AIP-C01 Engine Version
- Valid AIP-C01 Exam Experience ☐ AIP-C01 Reliable Exam Materials ☐ AIP-C01 Flexible Testing Engine ↗ Open website ➡ [www.vce4dumps.com](http://www.vce4dumps.com) ☐ and search for 《 AIP-C01 》 for free download ☐ Exam AIP-C01 Labs
- AIP-C01 Preparation Materials and AIP-C01 Study Guide: AWS Certified Generative AI Developer - Professional Real Dumps ☐ Go to website ☀ [www.pdfvce.com](http://www.pdfvce.com) ☀ ☐ open and search for ☀ AIP-C01 ☀ ☐ to download for free ☐ ☐ AIP-C01 Pass Rate
- 100% Pass 2026 Amazon Latest AIP-C01: AWS Certified Generative AI Developer - Professional Practice Exam ☀ Search for { AIP-C01 } and download exam materials for free through ➤ [www.prep4away.com](http://www.prep4away.com) ◀ ☐ Valid AIP-C01 Test Pattern
- 100% Pass Amazon Marvelous AIP-C01 Practice Exam ☐ Open ▷ [www.pdfvce.com](http://www.pdfvce.com) ◁ enter ➡ AIP-C01 ☐ and obtain a free download ☐ Valid AIP-C01 Exam Experience
- Free PDF Quiz 2026 Efficient Amazon AIP-C01: AWS Certified Generative AI Developer - Professional Practice Exam ☐ Search for 《 AIP-C01 》 on [ [www.verifeddumps.com](http://www.verifeddumps.com) ] immediately to obtain a free download ☐ 100% AIP-C01 Correct Answers
- Exam AIP-C01 Vce ☐ Valid AIP-C01 Test Pattern ☐ AIP-C01 Flexible Testing Engine ☐ Easily obtain ☐ AIP-C01 ☐ for free download through ☀ [www.pdfvce.com](http://www.pdfvce.com) ☀ ☐ ☐ AIP-C01 Updated Dumps
- 100% Pass Amazon Marvelous AIP-C01 Practice Exam ☐ Search on ( [www.examcollectionpass.com](http://www.examcollectionpass.com) ) for ➤ AIP-C01 ☐ to obtain exam materials for free download ☐ Latest AIP-C01 Exam Bootcamp
- [yoursocialpeople.com](http://yoursocialpeople.com), [mattieoouf315479.blogrelation.com](http://mattieoouf315479.blogrelation.com), [arunqvjg290458.bloggazzo.com](http://arunqvjg290458.bloggazzo.com), [ragingbookmarks.com](http://ragingbookmarks.com), [safadvm003518.yomoblog.com](http://safadvm003518.yomoblog.com), [tasneemaiqw347875.plpwiki.com](http://tasneemaiqw347875.plpwiki.com), [ledbookmark.com](http://ledbookmark.com), [bushraxnwn134909.vblogetin.com](http://bushraxnwn134909.vblogetin.com), [cyrusstm319603.buscawiki.com](http://cyrusstm319603.buscawiki.com), [antonnwpp184353.daneblogger.com](http://antonnwpp184353.daneblogger.com), Disposable vapes

2026 Latest CertkingdomPDF AIP-C01 PDF Dumps and AIP-C01 Exam Engine Free Share: <https://drive.google.com/open?id=1j-CpfWrJuWua5fkGbAzinYLEJMG6sKCI>