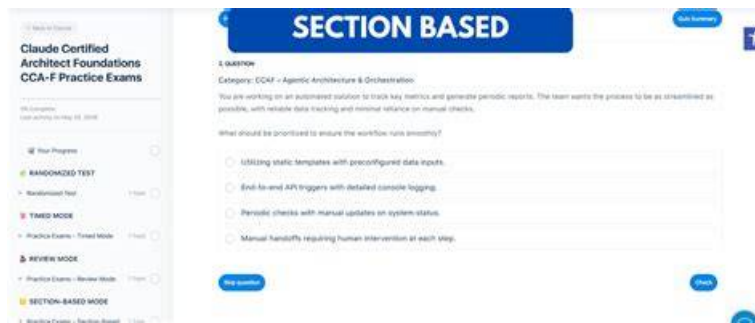


CCAR-F Trusted Exam Resource - CCAR-F Reliable Practice Materials



The practice materials of the exam with low quality may complicate matters of the real practice exam. So, you must know about our CCAR-F question torrent. Our study material is not same as other dumps or study tools, it not only has good quality but also has cheap price. We have most professional team to compiled and revise CCAR-F Exam Question, in order to try our best to help you pass the exam and get a better condition of your life and your work.

Anthropic CCAR-F Exam Syllabus Topics:

Section	Weight	Objectives
Claude Code Configuration & Workflows	20%	- Developer productivity workflows - Integrating Claude Code into development processes - Claude Code usage and configuration
Tool Design & MCP Integration	18%	- Model Context Protocol (MCP) concepts and integration - Tool safety, reliability, and usability - Designing effective tools for Claude applications
Agentic Architecture & Orchestration	27%	- Agent coordination and orchestration patterns - Designing agentic systems and workflows - Selecting appropriate Claude architectures
Prompt Engineering & Structured Output	20%	- Structured output generation and validation - Improving Claude response quality and consistency - Prompt design strategies
Context Management & Reliability	15%	- Managing context windows and information flow - Production deployment considerations - Evaluation and reliability strategies

>> CCAR-F Trusted Exam Resource <<

CCAR-F Trusted Exam Resource 100% Pass | Efficient CCAR-F: Claude Certified Architect - Foundations 100% Pass

Using our products does not take you too much time but you can get a very high rate of return. Our CCAR-F quiz guide is of high quality, which mainly reflected in the passing rate. We can promise higher qualification rates for our CCAR-F exam question than materials of other institutions. Because our products are compiled by experts from various industries and they are based on the true problems of the past years and the development trend of the industry. What's more, according to the development of the time, we will send the updated materials of CCAR-F Test Prep to the customers soon if we update the products. Under the guidance of our study materials, you can gain unexpected knowledge. Finally, you will pass the exam and get a Anthropic certification.

Anthropic Claude Certified Architect - Foundations Sample Questions (Q123-Q128):

NEW QUESTION # 123

You are using Claude Code to accelerate software development. Your team uses it for code generation, refactoring, debugging, and

documentation. You need to integrate it into your development workflow with custom slash commands, CLAUDE.md configurations, and understand when to use plan mode vs direct execution.

Your team is configuring MCP servers in Claude Code. You want to add a shared venue lookup server that all team members should have access to, and you personally want to add an experimental music playlist server that only you are testing. Which configuration approach correctly applies MCP server scopes?

- A. Add both servers to the project-level .mcp.json file
- **B. Add venue server to .mcp.json and playlist server to ~/.claude.json.**
- C. Add both servers to your local ~/.claude.json
- D. Add venue server to ~/.claude.json and playlist server to .mcp.json

Answer: B

Explanation:

Project-scoped MCP servers are stored in .mcp.json and can be shared through version control with the whole team. Personal local or user-scoped servers are stored in ~/.claude.json and remain private to the individual user.

NEW QUESTION # 124

Your automated review calls the Claude API for each pull request, using tool_use with a report_findings tool that returns a JSON array of finding objects. Each object contains file_path, line_number, severity, category, and description. During testing on a large pull request touching more than 30 files, the response reaches the max_tokens limit and is truncated in the middle of the JSON, causing your pipeline's parser to fail. What is the most effective way to handle this?

- A. Add retry logic that detects truncated JSON and resends the request with instructions to report only critical- and high-severity findings.
- B. Increase max_tokens to the model's maximum and instruct Claude to keep each finding description under 50 words.
- **C. Split the review into multiple API calls that each analyze a subset of the changed files, and then merge the resulting findings arrays.**
- D. Switch from tool_use to prompting Claude to return findings as a Markdown list.

Answer: C

Explanation:

Option A reduces the maximum output required from any single response while preserving the structured schema and complete severity range. The pipeline can partition files into coherent groups, execute bounded reviews, validate each returned array, and merge and deduplicate findings using stable fields such as file path, line number, category, and description.

Anthropic's stop-reason documentation confirms that max_tokens means generation reached the configured output limit and the response must be treated as incomplete. Structured output constraints can guarantee schema-valid generation when completion succeeds, but they cannot create unlimited output capacity. A large findings array can still exceed the available token budget.

Option B may postpone the failure but provides no durable guarantee for still-larger pull requests, and aggressively shortening descriptions may eliminate necessary evidence. Option C abandons machine-validated structure without reducing the amount of generated content. Option D deliberately suppresses medium- or low-severity findings and repeats an oversized request rather than addressing its scope.

Partitioning establishes predictable output bounds, supports targeted retries, retains every required finding category, and prevents a single truncated response from invalidating the complete review.

NEW QUESTION # 125

You are using Claude Code to accelerate software development. Your team uses it for code generation, refactoring, debugging, and documentation. You need to integrate it into your development workflow with custom slash commands, CLAUDE.md configurations, and understand when to use plan mode vs direct execution.

Your team has connected a custom MCP server that provides DevOps workflow templates. The server exposes several MCP prompts (such as deploy_checklist and incident_response) in addition to tools.

How do these MCP prompts become accessible within Claude Code?

- A. They are added to Claude Code's tool registry alongside the server's tools, invoked automatically by the model when relevant to the task.
- B. They are surfaced as @-mentionable resources alongside files, fetched and attached to your message when referenced.
- C. They are automatically prepended to every conversation as additional system-level context, influencing Claude's behavior throughout the session.

- D. They appear as slash commands (e.g., `/mcp__servername__deploy_checklist`) that you can invoke, with arguments passed after the command name.

Answer: D

Explanation:

MCP prompts are exposed as user-invoked commands rather than autonomous tools or permanently loaded system instructions. Claude Code dynamically discovers prompts from connected MCP servers and displays them in the command list using the naming convention `/mcp__servername__promptname` .

Arguments are supplied as space-separated values after the command. When executed, the MCP server resolves the prompt and its returned content is injected into the active conversation. Anthropic's official documentation provides examples such as `/mcp__github__list_prs` and `/mcp__jira__create_issue "Bug in login flow" high` . (<https://code.claude.com/docs/en/mcp>) Option A would consume context continuously and incorrectly treat optional workflow templates as mandatory system instructions. Option B confuses MCP prompts with MCP tools: tools are model-callable operations, while prompts are reusable prompt templates invoked as commands. Option C describes MCP resources, which can be referenced and attached but are a distinct MCP capability. For the stated server, the team could invoke commands such as `/mcp__devops__deploy_checklist` or `/mcp__devops__incident_response service-name` . The exact server segment is derived from the configured server name, with normalization applied where necessary.

Official references/topics: MCP Prompts; Dynamic Prompt Discovery; MCP Slash-Command Naming; Prompt Arguments.

NEW QUESTION # 126

Production monitoring shows that follow-up queries like "summarize what we learned about market trends" consistently take 40+ seconds. Investigation reveals the coordinator spawns the synthesis subagent for each summarization request, passing 80K+ tokens of accumulated findings. The coordinator already has these findings in its context from orchestrating the research.

What's the most effective way to improve response time for these follow-up summaries?

- A. Enable prompt caching on the synthesis subagent to reduce the overhead of repeatedly transferring the same research findings.
- B. Have the coordinator handle straightforward summarization requests directly using its existing context, reserving subagent spawning for complex analysis.
- C. Pre-generate and cache summaries at multiple granularities whenever new findings accumulate.
- D. Spawn the synthesis subagent with reduced context and have it request specific findings from the coordinator on-demand.

Answer: B

Explanation:

The coordinator already holds the accumulated findings and can perform simple summarizations directly. Avoiding repeated spawning of the synthesis subagent for straightforward queries reduces token processing and latency, improving response speed while maintaining accuracy.

NEW QUESTION # 127

You are building a multi-agent research system using the Claude Agent SDK. A coordinator agent delegates to specialized subagents: one searches the web, one analyzes documents, one synthesizes findings, and one generates reports. The system researches topics and produces comprehensive, cited reports.

The document-analysis agent has a single `analyze_document` tool that takes a document and a free-text instruction parameter. During evaluation, requests such as "extract the key financial metrics" often return narrative summaries, while "summarize the methodology" sometimes returns raw data tables. The synthesis agent reports that 35% of analysis results require new requests with clarified instructions.

What is the most effective way to improve reliability?

- A. Have the coordinator pre-classify each analysis request before passing instructions to the document- analysis agent.
- B. Keep the single tool but add an `analysis_type` enum parameter requiring explicit selection between extraction, summarization, and verification modes.
- C. Split the generic tool into purpose-specific tools-`extract_data_points`, `summarize_content`, and `verify_claim_against_source`-each with defined input and output contracts.
- D. Enhance the tool description with detailed examples showing how different instruction phrasings should map to different output formats.

Answer: C

Option B removes ambiguity at the tool-interface level. Extraction, summarization, and claim verification are different operations with different success criteria and output structures. Separate tools allow each operation to have a precise name, purpose, parameter schema, response contract, and error behavior. Anthropic’s guidance on writing effective tools for agents recommends a small set of thoughtful tools targeted at specific, high-impact workflows rather than generic wrappers that force the model to infer operational meaning from loosely structured instructions. Tool descriptions and examples, option A, could improve behavior but leave the overloaded free-text contract intact. An `analysis_type` enum, option C, identifies the requested mode but still requires one tool to return substantially different result structures, increasing validation and downstream branching. Coordinator pre-classification, option D, adds another probabilistic decision without correcting the interface used by the analysis agent. Purpose-specific tools let the synthesis layer know exactly what output it will receive. Their schemas should require relevant fields—for example, data-point names and source locations for extraction, or claim, verdict, and evidence for verification—and should be tested against the observed failure cases.

• • • • •

CCAR-F Reliable Practice Materials: <https://www.vceprep.com/CCAR-F-latest-vce-prep.html>

- [illegible]