

Claude Certified Architect - Foundations

CCA-F
Official exam code CCAR-F

Solution architects building production Claude apps (6+ months practical experience).

60	720	5	120 min	\$125	12 mo
SCORED ITEMS	TO PASS, OF 1000	WEIGHTED DOMAINS	TIME LIMIT	LIST PRICE	VALIDITY

This is a ledger, not a brochure

Most exam guides are read once. This one is meant to be marked up: you record where you actually are against each domain, work the two that cost you most, log two timed mocks, and end on a line that says book or do not book. Print it, or annotate it on screen.

Every *exam* figure — item count, cut score, domain weights, fee, validity — is transcribed from Anthropic's published blueprint, version 1.0-2026-07. The study-pace and mock-timing figures are this site's own, not Anthropic's. It does not reproduce exam content, and it is not the official guide.

Take the free CCA-F diagnostic

Read the full guide online



What the CCA-F weighs

Weights are the share of scored items drawn from each domain. The item counts are those weights applied to a 60-question exam, which is the number worth budgeting study hours against.

D1	Agentic Architecture & Orchestration		27%	~16 items
D2	Tool Design & MCP Integration		18%	~11 items
D3	Claude Code Configuration & Workflows		20%	~12 items
D4	Prompt Engineering & Structured Output		20%	~12 items
D5	Context Management & Reliability		15%	~9 items

Where the exam actually is

Agentic Architecture & Orchestration alone carries 27% — about 16 of the 60 items. The three heaviest domains together are 67% of the paper.

How the result comes back

Pass or fail with a scaled score from 100 to 1000, plus the percentage of items you answered correctly in each domain. Those per-domain percentages are reported to help you read your performance; they do not decide the pass.

Tick what you know. Cross what you do not.

That is the whole system. Two marks, and you add them up.

- ☒ **Tick.** You could explain this to someone else without looking it up.
- ☐ **Cross.** You could not. This is what you study.

Then add up your ticks

Each domain page has a score box. Count your ticks, write the number in, and compare it with how many questions that domain is worth. A domain worth 16 questions where you ticked 2 of 6 topics is where your exam is won or lost.

Do it in this order

1. **Take the free diagnostic first, before you revise.** Write your score on page 5. Revising first only tells you what looks familiar.
2. **Go through pages 6 to 10 and mark every topic** tick or cross. Be hard on yourself here. A generous tick costs you an exam fee later.
3. **Study the two domains with the most crosses and the highest question count.**
4. **Sit a full mock, then mark the same pages again** in the second column. Anything that went from tick to cross is what the real exam would have caught.

Take the free diagnostic

Open the 60-question mock

How the questions are written

Every question follows the same shape. Knowing the shape is worth a few marks on its own, because you stop misreading what is being asked.

1. THE SITUATION	A short description of a real setup and what is going wrong. Two or three sentences. Every number in it matters.
2. THE CATCH	A phrase that rules most answers out on its own — "as a first step", "without adding infrastructure", "easiest to maintain". Read this before you read the answers.
3. THE QUESTION	"Which is most effective" and "what would you check first" are different questions. Answer the one asked.
4. HOW MANY TO PICK	Some questions want one answer, some want two or more. The question always tells you. Check this first.
5. THE ANSWERS	Usually one right, one that does too much, one that fixes a different problem, and one that is simply wrong.



Wrong answers that look right

LOOKS RIGHT	IS ACTUALLY RIGHT
D1 Parsing natural language for loop termination	Check stop_reason: 'tool_use' vs 'end_turn' The API's stop_reason field is the authoritative loop-termination signal. 'tool_use' means continue (run the requested tool); 'end_turn' means the model has signaled completion. Read the structured field; don't infer from prose.
D1 Arbitrary iteration caps as the primary stopping rule	Let stop_reason drive the loop, use a cap as a safety net stop_reason is the primary terminator. The cap is a guard against pathological loops, not the design. If you're hitting the cap regularly, your prompt or tool surface is the problem, not the cap.
D2 18+ tools registered on a single agent	4-5 tools per agent; spawn subagents for additional surfaces Keep the tool surface narrow and use subagents (via the Task tool) for orthogonal capabilities. Each subagent has its own 4-5 tools scoped to its job.

Write down where you are before you revise

A number you wrote two weeks ago is the only proof the fortnight worked. Take the diagnostic without revising first, and fill the top row in.

Your diagnostic

DATE	QUESTIONS RIGHT	OUT OF	WEAKEST DOMAIN
		10	

How you did in each domain

The middle column is how many questions that domain is worth on the real exam. Write in how many you got right, and the difference is your gap — the number of marks currently sitting on the table.

DOMAIN	QUESTIONS ON THE EXAM	YOU GOT RIGHT	GAP
D1 Agentic Architecture & Orchestration	about 16		
D2 Tool Design & MCP Integration	about 11		
D3 Claude Code Configuration & Workflows	about 12		
D4 Prompt Engineering & Structured Output	about 12		
D5 Context Management & Reliability	about 9		
Total	60		

Your mocks

A mock score is a guide to how ready you are, not a prediction of your official result. It uses the same 720 pass mark so you can compare week to week, but only Anthropic scores the real exam.

DATE	QUESTIONS RIGHT	OUT OF	PASSED?	WHAT WENT WRONG
		60		
		60		

Sit a full 60-question mock

Agentic Architecture & Orchestration

D1

27% of the exam · about 16 of 60 items

YOUR SCORE FOR THIS DOMAIN

After the diagnostic _____ of 16

After the mock _____ of 16

TOPICS IN THIS DOMAIN	KNOW IT	AFTER MOCK
Agentic Loops An agentic loop turns one API response into a controlled, multi-turn system. The harness reads stop_reason, dispatches tools, appends observations, and decides whether to continue.	☐	☐
Subagents Subagents are specialized, isolated agents spawned by a coordinator to handle domain-specific tasks while preserving conversation isolation. They do NOT inherit memory, the coordinator passes context explicitly in the prompt.	☐	☐
Stop Reason Read stop_reason, not the prose: end_turn means Claude finished, tool_use means it wants a tool. How to end a conversation loop correctly.	☐	☐
Session State & Persistence Session state is the running message list (and any external persistence) that gives an agentic loop continuity.	☐	☐
Human-in-the-Loop Escalation Escalation is the deterministic handoff path when policy thresholds, low confidence, or repeated failure conditions are hit.	☐	☐
Subagent State Handoff Subagents hand off state three ways: structured messages (default - typed JSON envelopes through the orchestrator), tool results (tightest sync chain when receiver is a tool), or shared files (size or binary fallback only).	☐	☐

NOTES – WHAT YOU GOT WRONG, AND WHY

Tool Design & MCP Integration

D2

18% of the exam · about 11 of 60 items

YOUR SCORE FOR THIS DOMAIN

After the diagnostic _____ of 11

After the mock _____ of 11

TOPICS IN THIS DOMAIN	KNOW IT	AFTER MOCK
Tool Calling Tool calling is how Claude decides to invoke external functions and pass structured arguments. Tools are routing infrastructure, good descriptions reduce the need for classifiers or few-shot examples.	☐	☐
Tool Choice tool_choice is the parameter that controls whether Claude can decide to use a tool ("auto"), must call any tool ("any"), or must call a specific tool by name.	☐	☐
Model Context Protocol MCP is a communication standard that lets Claude access pre-built tools, resources, and prompts from specialized servers without you writing integration code.	☐	☐
SDK Hooks (Pre/PostToolUse) Hooks are deterministic code that runs before or after tool calls. They enforce policy that prompt-only patterns can't guarantee, refund limits, identity verification, audit logging.	☐	☐
Tool Evaluation & Testing	☐	☐
Built-In Tools vs Custom Tools vs Skills vs MCP Built-in tools, custom tools, Agent Skills, and MCP servers are complementary layers, not a power ranking — one agent often uses several at once.	☐	☐

NOTES — WHAT YOU GOT WRONG, AND WHY

Claude Code Configuration & Workflows

D3

20% of the exam · about 12 of 60 items

YOUR SCORE FOR THIS DOMAIN

After the diagnostic _____ of 12

After the mock _____ of 12

TOPICS IN THIS DOMAIN	KNOW IT	AFTER MOCK
CLAUDE.md Hierarchy CLAUDE.md is persistent project memory loaded automatically by Claude Code. A three-level hierarchy (user → project → directory) lets you set global defaults, team standards, and per-directory rules.	☐	☐
Plan Mode Plan Mode forces Claude Code to draft and approve a multi-step plan before executing. Use it for non-trivial changes (3+ steps); skip for trivial edits. The exam pattern: when do you require Plan Mode vs. direct execution.	☐	☐
Skills Skills are reusable, on-demand task workflows defined in markdown with YAML frontmatter. Unlike CLAUDE.md (always loaded), skills load only when Claude matches the description to the current task.	☐	☐
Checkpoints & Session Management	☐	☐

NOTES – WHAT YOU GOT WRONG, AND WHY

Prompt Engineering & Structured Output

D4

20% of the exam · about 12 of 60 items

YOUR SCORE FOR THIS DOMAIN

After the diagnostic _____ of 12

After the mock _____ of 12

TOPICS IN THIS DOMAIN	KNOW IT	AFTER MOCK
System Prompts & Instructions System prompts establish role, constraints, format, and tool guidance. Anatomy: role, capability boundaries, style/format rules, tool guidance, examples.	☐	☐
Prompt Caching Prompt caching reduces cost (~90%) on repeated context like long system prompts and tool definitions. Cache breakpoints, TTL, and cache_control field placement are exam patterns.	☐	☐
Batch API Message Batches API: 50% discount for async, non-time-sensitive workloads.	☐	☐
Structured Outputs Structured outputs guarantee Claude returns JSON matching a schema instead of natural language. Force a tool with a JSON schema; the API enforces structure, not your parser.	☐	☐
Vision & Multimodal	☐	☐
Streaming	☐	☐

NOTES – WHAT YOU GOT WRONG, AND WHY

Context Management & Reliability

D5

15% of the exam · about 9 of 60 items

YOUR SCORE FOR THIS DOMAIN

After the diagnostic _____ of 9

After the mock _____ of 9

TOPICS IN THIS DOMAIN	KNOW IT	AFTER MOCK
Context Window Management Context window management is how you keep long conversations within limits without dropping critical facts. Patterns: case-facts blocks, progressive summarization, retrieval.	☐	☐
Case Facts Block A case-facts block is an immutable set of transactional data (customer ID, order details, refund amount, policy limits) included at the top of every prompt during a multi-turn conversation.	☐	☐
	☐	☐
	☐	☐

NOTES – WHAT YOU GOT WRONG, AND WHY

Work every concept, free

Two weeks, in the order that pays

People already working with Claude day to day report 2-4 weeks at 60-90 minutes a day.
From a cold start, 6-8 weeks. The order matters more than the hours.

WHEN	WHAT	RECORDED ON
Day 1	Diagnostic, cold. No revision first. Write the per-domain percentages down.	Page 5
Day 1	Mark every domain in the first-pass column, honestly.	Pages 6–10
Days 2–6	Agentic Architecture & Orchestration first — 27% of the paper, and the single largest block of marks available.	Page 6
Day 7	Full 60-question mock , timed, no notes. The mock here runs tighter than the real 120 minutes on purpose — train faster than race pace.	Page 5
Days 8–12	Claude Code Configuration & Workflows , plus whatever the mock exposed. Re-mark in the second column.	Page 8
Day 13	Second timed mock. Different questions, same conditions.	Page 5
Day 14	The decision on page 12. Book, or take another week.	Page 12

Fourteen days, at a glance

1

MEASURE

☐

2

WORK

☐

3

WORK

☐

4

WORK

☐

5

WORK

☐

6

WORK

☐

7

MOCK

☐

8

WORK

☐

9

WORK

☐

10

WORK

☐

11

WORK

☐

12

WORK

☐

13

MOCK

☐

14

DECIDE

☐

Tick each square as the day is done. Two squares are timed mocks; the last one is the decision on page 12.

Day 1 starts here

Day 7 mock



Booking, sitting, and what happens if it goes wrong

Registering

Registration and scheduling run through the Anthropic Partner Academy and Pearson VUE: open your exam's page on the Partner Academy, read the terms and exam policy, register and check out (the fee shown reflects any partner-tier discount), create your Pearson VUE account from the confirmation, then choose a date and either online proctoring or a test centre. You can cancel or reschedule up to 24 hours ahead; inside 24 hours the fee is forfeit.

If you do not pass

Retakes wait 14 days after the first attempt, 30 after the second, 90 after the third, up to four attempts in a rolling twelve months. The fee applies to every attempt. Limits are per exam, so a fail on one does not block registering for another.

Pacing, on the day

120 minutes for 60 scored items is 2.0 minutes each. Scenario items run long, so bank time early. Check yourself at the quarter marks.



The decision

Second mock cleared 720 with room

☐ yes☐ no

No domain still carrying empty boxes above 15% weight

☐ yes☐ no

Sat both mocks under time, without notes

☐ yes☐ no

Three yes: book it. Any no: one more week. Booked for

Sit the mock that decides it

Registration walkthrough

On the day

Government-issued photo ID, and the name on it must match your registration exactly. The workspace has to be clear of notes, books, phones, secondary monitors and any recording device. You accept a confidentiality agreement before the exam begins; declining ends the session with no refund.

Afterwards

The credential is valid 12 months. On-time renewal is a free, non-proctored assessment on the Partner Academy; let it lapse and you retake the full exam at the full fee.

