

FREE PREVIEW • ~30% OF THE FULL GUIDE

Unofficial Study Guide for the CCAO-F Exam

Claude Certified Associate — Foundations · sample edition

CCAO-F

60 Q • 120 MIN

PASS 720/1000

An independent study resource by Zehon.com · August 2026

This free preview contains complete, unabridged sections of the full guide: Day 1 of the 7-day plan, all of Domain 1, and 15 practice questions with full answer explanations.

Independence. This is an independent, unofficial study resource. It is not affiliated with, endorsed by, sponsored by, or approved by Anthropic, PBC or Pearson Education, Inc. “Claude,” the certification names, and the exam codes are trademarks of Anthropic, PBC, used here solely to identify the exam this guide helps you prepare for. “Pearson VUE” is a trademark of Pearson Education, Inc.

Practice questions. Every practice question is original and written to the publicly published exam objectives. Nothing here is an actual exam question, and nothing is derived from the content of any exam sitting.

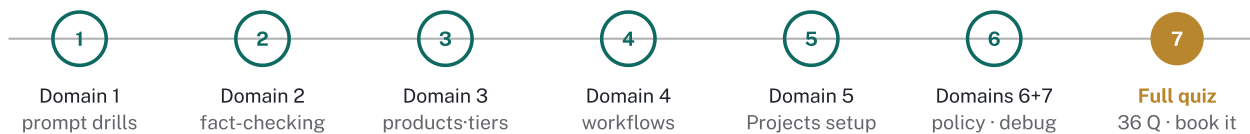
Accuracy. Exam logistics reflect the official Anthropic and Pearson VUE program pages as of August 2026 and can change — verify before booking.

Production note. Produced with AI assistance. This preview may be shared freely; the full edition may not be redistributed.

The 7-day plan

Read a domain, get your hands dirty, checkpoint with the quiz — every day.

~1.5–2 hours a day for 7 days. Every day pairs reading from the study guide with hands-on drills inside Claude.ai — the certification is about judgement in *using* the product, and the fastest way to build that judgement is to use it. You'll need a Claude.ai account: a paid plan lets you practice Research mode and RAG-expanded Projects; everything else works on Free, including up to 5 Projects.



Materials: the study guide · the 36-question quiz bank · a Claude.ai account (Free works; paid adds Research mode + RAG Projects)

One domain a day, drills in the product itself; day 7's quiz score tells you whether to book.

Day 1 — Prompting & Task Execution (Domain 1)

Read (40 min): guide intro (the "three judgement threads") + Domain 1, objectives 1–4.

Hands-on in Claude.ai (50 min):

- **Prompt skeleton drill.** Pick a real work task (an email, a summary, a plan). Prompt it twice: once as a bare one-liner, once with full Role → Context → Task → Format → Examples. Compare outputs side by side; note exactly which added element caused which improvement.
- **Decomposition drill.** Take a big deliverable (e.g. "a Q4 marketing plan" or "a project charter"). First ask for the whole thing in one prompt. Then redo it in stages: outline → one section at a time → assemble. Compare depth.
- **Iteration drill.** Take a mediocre output and improve it in 3 turns using *specific* feedback only (never "try again"). Then deliberately derail a chat with contradictory corrections and practice the recovery: summarize what you actually want, open a new chat, restart clean.
- **Task-type drill.** Run the same topic four ways — as analysis, research, drafting, brainstorming — and observe how your prompt wording must change (structured criteria vs. citations vs. tone/audience vs. "give me 20 varied ideas").

✓ **CHECKPOINT (10 MIN)**

Without looking, write the five prompt elements and one example of each. State when to iterate in-chat vs. restart. If shaky, reread objective 3.

 IN THE FULL EDITION — THE REST OF THE PLAN

- Day 2 — Evaluating & Validating Outputs (Domain 2)
- Day 3 — Products, Models & Context (Domain 3)
- Day 4 — Workflow Integration & Solution Design (Domain 4)
- Day 5 — Configuration & Knowledge Management (Domain 5)
- Day 6 — Governance & Responsible Use + Troubleshooting (Domains 6–7)
- Day 7 — Quick Reference, Full Quiz, Exam Logistics

The study guide

Seven domains, twenty-eight objectives — workplace judgement, not code.

Everything on this page maps to the 28 objectives of the official exam blueprint. This is a **non-technical, judgement-based exam**: scenarios put you in the shoes of a persona — a marketing associate, an operations analyst, a project manager, a teacher — and ask *what that person should do*. The right answer is rarely "which button do I press"; it is "which behavior fits the business context, the policy, and the workflow."

EXAM SNAPSHOT

60 Q

CLOSED BOOK

7 domains

28 OBJECTIVES

120 min

2 MIN / QUESTION

720

PASS, 100–1000 SCALE

\$99

LIST PRICE —
DISCOUNTS MAY
APPLY

0 code

NON-TECHNICAL ·
PERSONA-BASED

Pearson OnVUE, online proctored, registered via the Anthropic Partner Academy (partner-network employees). Credly badge valid 12 months. No guessing penalty — answer everything. Retakes: 14 days after attempt 1, 30 after attempt 2, 90 after attempt 3; max 4 attempts per rolling 12 months. Logistics reflect the official program pages as of August 2026 — verify before booking.

DOMAIN SIZE · OBJECTIVES OF 28

1 · Prompting & Task Execution		4 obj
2 · Evaluating & Validating Outputs		5 obj
3 · Choosing Product & Model		5 obj
4 · Workflow Integration & Design		4 obj
5 · Configuration & Knowledge Mgmt		3 obj
6 · Governance & Responsible Use		4 obj
7 · Troubleshooting		3 obj

Anthropic hasn't published official domain weights — objective count is the best available proxy for coverage. Study all seven; none is skippable.



THREE JUDGEMENT THREADS RUN THROUGH THE WHOLE EXAM

Trust but verify — Claude output containing facts, numbers, citations, names, or legal/medical/financial claims must be validated before it leaves your hands; any answer that ships unverified AI content externally is wrong. **Sensitive data gets handled, not pasted** — customer PII, health data, financials, or anything covered by regulation or company policy means checking policy, redacting/anonymizing, or using an approved enterprise environment, never "just paste it in." **Know your lane** — an associate configures Projects, writes prompts, and redesigns her own workflow; the moment a scenario needs an API, custom integration, automation at scale, or production deployment, the answer is **escalate to a developer or architect** — not build it herself, and not give up.

Domain 1 · Prompting & Task Execution **OBJECTIVES 1-4**

Getting good work out of Claude on the first pass and improving it deliberately.

1 · Create effective prompts for business and technical tasks

A good business prompt gives Claude the same briefing you'd give a capable new colleague. The reliable skeleton is **Role** → **Context** → **Task** → **Format** → **Examples**, with constraints riding along.

One briefing, five slots — like briefing a capable new colleague



Constraints ride along anywhere: length, tone, banned phrases, "say 'unknown' rather than guessing."

The skeleton in action: "Act as a senior B2B copywriter... our company sells payroll software to 50–200 person firms; this email goes to churned trial users... draft three subject lines and a 120-word re-engagement email... return a table with columns Subject / Preview text / Rationale."

Format instructions can also be limits and conventions: "Keep it under 150 words," "Use British spelling." Examples means one or two samples of what "good" looks like (few-shot) when style matters.

Marketing associate

Instead of "write a product blurb," provide the product one-pager, the persona, the channel, the word limit, and a past blurb the team loved.

Operations analyst

"You are a process analyst. Here is our current returns procedure [pasted]. Identify the three steps with the most handoffs and propose fixes. Output: numbered list, each with problem / proposal / risk."

Project manager

Paste the meeting transcript and ask for "decisions, owners, deadlines, open questions — as a four-column table."

Teacher

"Act as a curriculum designer for 9th-grade biology. Create a 45-minute lesson plan on osmosis: objectives, warm-up, activity, assessment."

DO

- State audience and purpose
- Give constraints — length, tone, banned phrases
- Paste or attach source material rather than assuming Claude knows your internal facts
- Tell Claude what to do when information is missing ("say 'unknown' rather than guessing")

DON'T

- Stack five unrelated asks in one sentence
- Leave format to chance when you need a specific deliverable
- Assume Claude knows internal jargon, current internal data, or events after its training cutoff

🔔 DECISION CUES

A persona gets a vague, generic output → the fix is **add context / role / format to the prompt** — not "switch models," not "Claude can't do this." The best answer usually includes providing **source material + explicit output format**, not just a longer instruction.

2 · Apply task decomposition techniques to structure complex requests

Big, multi-part deliverables come out better when broken into stages: outline → draft sections → assemble → polish. Decomposition also applies *within* a prompt ("First list assumptions, then analyze, then recommend") and *across* chats (research in one conversation, drafting in another).

Project manager

Project charter — step 1: extract goals and constraints from the kickoff notes; step 2: draft scope and out-of-scope; step 3: risk register; step 4: assemble and edit for the exec audience.

Marketing associate

Campaign plan → first audience analysis, then messaging pillars, then channel plan, then calendar — reviewing each stage before moving on.

Operations analyst

Vendor comparison → extract criteria from procurement policy, then score each vendor's proposal separately, then produce the comparison matrix.

DO

- Sequence steps so each output feeds the next
- Review intermediate outputs — catching a wrong assumption at the outline stage is cheap
- Ask Claude itself to propose a plan of attack for a fuzzy task before executing

DON'T

- Dump a 40-page requirement into one prompt and expect a finished, polished, correct deliverable in one shot
- Skip review between stages on high-stakes work

🔔 DECISION CUES

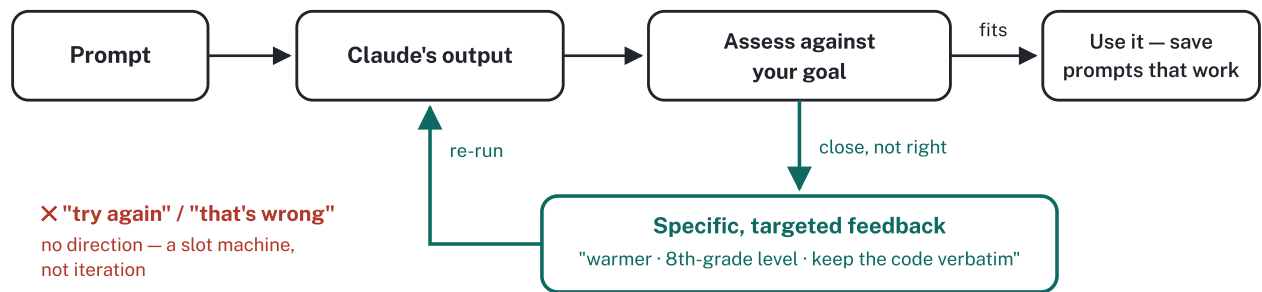
"The report Claude produced is shallow / misses sections" → **break the task into steps**, not "regenerate until it works."

✗ THE PROMPT-INFLATION DISTRACTOR

"Ask Claude to do everything at once with a longer prompt" — decomposition beats prompt inflation for genuinely complex tasks.

3 · Iterate prompts to improve output quality

Prompting is a conversation, not a slot machine. The professional loop: run → inspect the output against your goal → give *specific* feedback ("shorter, more direct, drop the second paragraph, add a CTA") → re-run. When a prompt will be reused (a weekly report template, a Project instruction), refine it until it works reliably, then save it.



Off the rails after many contradictory corrections? Summarize what you actually want, then restart in a NEW chat with a rewritten prompt.

The teal loop is deliberate iteration; the red text is the anti-pattern; the amber exit is the recovery move for a contaminated conversation.

Marketing associate

Gets an email that's too formal: replies "Make it warmer and more conversational; write at an 8th-grade reading level; keep the discount code sentence verbatim."

Operations analyst

Refining a recurring "summarize the weekly incident log" prompt: after three weeks of tweaks, moves the polished version into a Project's instructions so everyone benefits.

DO

- Change one thing at a time so you know what helped
- Tell Claude what was *right* as well as wrong
- Keep a personal library of prompts that worked
- Start a fresh chat when the conversation has accumulated contradictory corrections

DON'T

- Repeat the identical prompt hoping for luck
- Say only "that's wrong, try again" with no direction
- Iterate forever on a chat polluted by earlier misunderstandings — restart with a better first prompt instead

🔔 DECISION CUES

"Output is close but not right" → **iterate with specific feedback in the same chat.** "Output is fundamentally off-track after many corrections" → **start a new chat with a rewritten prompt.**

Reusable/recurring task → the mature answer involves saving the refined prompt (often as Project instructions).

4 · Adapt prompting strategies based on task type

Different work wants different prompting:

Analysis

Provide the data/document, ask for structured reasoning ("walk through step by step"), request assumptions and confidence levels, define the output (table, ranked list).

Research

Ask for sources and citations; use Research mode / web search for current facts; expect and verify; ask "what don't we know?"

Drafting

Heavy on audience, tone, format, examples; iterate on style.

Brainstorming

Ask for *quantity and variety* ("20 ideas, ranging from safe to wild"), defer judgement, then ask Claude to cluster and shortlist. Constraints kill brainstorming — add them in the filtering pass.

PM analyzing survey results

"Here are 200 free-text responses. Categorize into themes, count each, quote two representative responses per theme" — analysis framing, structured output.

Teacher brainstorming

"Give me 15 different ways to open a lesson on fractions for a low-energy Friday class" — volume first, pick later.

Comms drafting

Precise tone spec + a model paragraph from a previous well-received announcement.

DO

- Match the temperature of the ask to the task — precision language for analysis, generative language for ideation
- For research, explicitly request sources

DON'T

- Use a rigid single-answer prompt for brainstorming
- Accept research answers without citations
- Treat a creative draft as a factual document



DECISION CUES

Ideation scenario → the best answer **maximizes divergent options before filtering**. Research on current events / market data → **Research mode or web-enabled search with citations, then verification** — not Claude's memory alone.

 IN THE FULL EDITION — SIX MORE SECTIONS OF THIS DEPTH

- Domain 2 · Evaluating & Validating Outputs
- Domain 3 · Choosing the Right Claude Product & Model
- Domain 4 · Workflow Integration & Solution Design
- Domain 5 · Configuration & Knowledge Management
- Domain 6 · Governance & Responsible Use
- Domain 7 · Troubleshooting
- Quick reference — cram sheet

Practice bank – 36 questions

36 original, exam-style practice questions with full explanations — written to the published objectives, not taken from any exam.

Work them by domain as checkpoints (the 7-day plan says when), then re-run your misses on day 7. The answer key starts on its own page after the last question — answers follow on the next page.

Prompting & Task Execution

Q1 EASY

A marketing associate asks Claude to 'write a product blurb' and receives generic copy that could describe any product. What should she do first?

- A. Switch to a more capable model tier and resubmit the same prompt
- B. Report the poor output to her IT team as a product defect
- C. Rewrite the prompt to include the product details, target audience, channel, tone, and a word limit
- D. Regenerate the response several times until one version sounds better

Q2 MEDIUM

A project manager needs a complete project charter: goals, scope, out-of-scope, risks, and stakeholder plan, based on 30 pages of kickoff notes. Which approach is most likely to produce a high-quality result?

- A. Paste all 30 pages with the instruction 'write the full charter' and polish whatever comes back
- B. Work in stages: extract goals and constraints first, then draft each charter section, reviewing each output before moving on
- C. Ask Claude to write the charter without the notes to avoid overloading the context window
- D. Split the notes across five separate unrelated chats and merge the five charters manually

Q3 MEDIUM

Claude drafts an email that is accurate and well-structured, but too formal for the team's culture. What is the most effective next step?

- A. Reply in the same chat with specific direction, such as 'make it warmer and more conversational, keep the deadline sentence verbatim'
- B. Start a brand-new chat and rewrite the entire original prompt from scratch
- C. Tell Claude 'that's wrong, try again' until the tone improves
- D. Accept the draft, since tone is subjective and editing wastes the time savings

Q4 EASY

An operations lead wants fresh ideas for improving warehouse morale. Which prompting approach best fits a brainstorming task?

- A. Ask for the single best evidence-based recommendation with citations
- B. Provide a strict template and ask Claude to fill in one idea per required field
- C. Ask a series of yes/no questions about existing morale programs
- D. Ask for 20 varied ideas ranging from conventional to unconventional, deferring evaluation until afterward

Q5 HARD

After 40 turns of corrections and reversals, a knowledge worker finds Claude mixing up her earlier and current instructions, and each fix introduces new problems. What is the best move?

- A. Continue correcting in the same chat, since Claude has the most context there
- B. Escalate the conversation to a developer to debug the model's behavior
- C. Summarize the final requirements, then start a new chat with one clean, complete prompt
- D. Ask Claude within the same chat to ignore everything said previously

Evaluating & Validating Outputs

Q6 EASY

A project manager receives a Claude-drafted report containing the claim 'industry churn averages 12.4% (Gartner, 2025)' — but she cannot locate any such Gartner figure. What should she do?

- A. Keep the statistic, since Claude included a specific source and year
- B. Remove or replace the claim with a figure she can verify from a real source before the report goes out
- C. Discard the entire report, since one bad statistic means nothing else can be trusted
- D. Soften the wording to 'roughly 12%' so the missing source matters less

Q7 MEDIUM

In a Claude-generated market overview, which element should raise the MOST suspicion of hallucination?

- A. A general statement that competition in the sector has been increasing
- B. A recommendation section clearly labeled as Claude's own analysis
- C. A caveat noting that recent developments may not be reflected
- D. A direct quotation attributed to a named executive, with a specific date, that perfectly supports the report's argument

Q8 MEDIUM

Before publishing a statistics-heavy blog post drafted with Claude, what is the most reliable validation workflow for a marketing associate?

- A. Ask Claude whether it is confident in the draft and publish if it says yes
- B. Run the draft through Claude a second time and publish if the two versions agree
- C. Have Claude list every factual claim in a table, then confirm each against a primary source, fixing or deleting what fails
- D. Publish with a general disclaimer that figures are AI-generated estimates

Q9 HARD

A communications associate used Claude to draft two items: an internal recap of a team meeting, and a customer-facing notification about a service incident. How should review be handled?

- A. Both must be approved by the legal team, since both were AI-generated
- B. The recap needs a quick self-review; the customer notification requires review by the accountable experts (legal/communications) before sending
- C. Neither needs extra review if the drafts read accurately, since the associate knows the context
- D. The customer notification can go out first because timeliness outweighs review for incidents

Q10 EASY

An analyst has an accurate, detailed technical summary from Claude but must present the findings to the executive committee. What is the best use of Claude here?

- A. Ask Claude to rewrite the summary for executives — five bullets focused on business impact — then edit for organizational nuance
- B. Send the technical summary as-is, since accuracy is what matters most
- C. Regenerate the entire analysis from scratch with a simpler prompt
- D. Shorten it by deleting the second half of the document

Choosing the Right Claude Product & Model

Q11 EASY

Every week, a marketing associate re-uploads the brand guidelines and re-types the same tone instructions before asking Claude to draft the newsletter. What should she do instead?

- A. Save the instructions in a personal document and paste them faster each week
- B. Ask her admin to enable a longer context window
- C. Use Research mode so Claude can find the brand guidelines itself
- D. Create a Project with the tone rules as project instructions and the brand guidelines as project knowledge

Q12 EASY

An operations team needs to classify 10,000 short customer feedback comments by sentiment, quickly and at the lowest cost. Which model tier fits best?

- A. Opus-tier, because higher capability always yields better classifications
- B. Haiku-tier, because the task is simple, well-defined, and high-volume
- C. Whichever tier is newest, since recency determines quality
- D. Sonnet-tier, because a mid-tier model is the safe default for everything

Q13 MEDIUM

A strategy lead must synthesize twelve dense market reports into a nuanced investment recommendation for the board, where quality matters far more than speed or cost. Which choice best aligns model to task?

- A. Haiku-tier, to keep the cost of reading twelve documents down
- B. Sonnet-tier for the synthesis, because balance is always optimal
- C. The most capable (Opus-tier) model, accepting slower and costlier responses for the highest-quality reasoning
- D. Run the task on all tiers and submit whichever finishes first

Q14 MEDIUM

After weeks in one continuous chat covering an entire project, a PM notices Claude giving slower, vaguer answers and misremembering early decisions. What is the best response?

- A. Ask Claude to summarize the key decisions, save that summary (e.g., into a Project), and continue in a fresh chat
- B. Upgrade to a higher-priced plan so the conversation can continue indefinitely
- C. Repeat the early decisions verbatim every few messages for the rest of the project
- D. Conclude the model has degraded and file a support ticket

Q15 MEDIUM

An analyst on a Pro plan must produce a cited overview of how competitors reacted to a regulation passed last month, drawing on many web sources. Which feature should she reach for?

- A. A plain chat with no web access, relying on Claude's training knowledge
- B. An artifact, since the deliverable is a document
- C. Research mode (with web search enabled), which runs multiple searches and returns a report with citations she can verify
- D. Project instructions directing Claude to 'always be up to date'

 IN THE FULL EDITION — 21 MORE QUESTIONS

- 36 questions total, each with a full explanation in the answer key
- coverage across every domain, mapped to the published objectives

Answer key & explanations

Read every explanation, even for questions you got right — the distractor rationales are training too.

Q1 → C Generic output is the classic symptom of an underspecified prompt. The fix is to supply role, context, task, format, and constraints — Claude cannot know product specifics or audience unless told. Switching models, regenerating blindly, or escalating to IT do not address the missing information.

Q2 → B Task decomposition — sequencing steps so each reviewed output feeds the next — produces deeper, more accurate results for complex deliverables and catches wrong assumptions early. One giant prompt yields shallow coverage; omitting the source material invites fabrication; five disconnected charters create inconsistency rather than structure.

Q3 → A When output is close but not right, iterate in the same conversation with specific, directive feedback — naming what to change and what to preserve. Restarting is for derailed conversations, not near-misses; 'try again' gives Claude nothing to act on; and shipping a tonally wrong email to a known audience defeats the purpose.

Q4 → D Brainstorming benefits from divergent prompting: request quantity and variety first, then filter and cluster in a second pass. Asking for one 'best' answer, rigid templates, or yes/no questions all constrain idea generation — the opposite of what ideation needs. Analysis-style prompting with citations suits research tasks, not brainstorming.

Q5 → C A long conversation accumulated contradictory instructions — a contaminated context. The professional recovery is to distill what is actually wanted and restart clean, carrying only the summary. Continuing adds more contradictions; 'ignore everything' still leaves the conflicting history in the context window; and this is a prompting issue, not something a developer should debug.

Q6 → B A precise, citation-shaped statistic that cannot be sourced is a classic hallucination and must be verified or removed before the document ships. Citation-shaped text is not a real citation; softening the number still publishes a fabricated fact; and scrapping the whole report over-corrects — validation targets the unverifiable claims, not the entire draft.

Q7 → D Hallucinations tend to be highly specific, confident, and conveniently supportive: exact quotes, precise statistics, named studies. A dated, named, perfectly on-message quotation is exactly the kind of checkable specific to verify first. General trends, labeled opinions, and honest caveats are lower-risk content.

Q8 → C Effective fact-checking pairs an AI-assisted step (extracting checkable claims) with human verification against primary sources. Claude re-asserting confidence is not validation, agreement between two AI passes can simply repeat the same error, and a disclaimer does not make unverified external claims acceptable to publish.

Q9 → B Review depth scales with stakes and reach. External, legally significant communications require domain-authoritative human review regardless of who or what drafted them; a low-stakes internal recap needs only self-review. Sending everything to legal over-escalates and clogs the process, while skipping review on external incident comms is the serious failure to avoid.

Q10 → A Adapting a correct output for a new audience is a core refinement skill: state the audience explicitly and let Claude transform the content, then apply human judgement for nuance. Sending technical detail to execs fails the audience; regenerating discards good work; arbitrary truncation loses substance rather than adapting it.

Q11 → D Repeatedly re-supplying the same context is the textbook signal for a Claude Project: instructions hold the standing rules and knowledge holds the reference files, applied automatically to every chat in the project. Faster pasting keeps the waste; context window size is not the issue; Research mode is for multi-source current-information questions, not internal brand files.

Q12 → B Haiku-tier models are the fastest and cheapest and are designed for exactly this: high-volume, simple, well-defined tasks like classification. Opus would add cost and latency without meaningfully better sentiment labels; defaulting to mid-tier ignores the cost/speed requirement; model recency is not a selection principle.

Q13 → C Model selection aligns capability with what the task genuinely requires. Complex multi-document synthesis with board-level stakes is precisely where the most capable tier earns its cost. Under-buying with a fast tier risks shallow analysis that costs more in rework; picking by completion speed optimizes the wrong variable.

Q14 → A Long conversations eventually strain the context window, degrading recall and quality. The standard remedy is summarize-and-restart, persisting durable context in project knowledge. The model has not 'degraded'; constant repetition bloats the context further; and no plan removes the fundamental limit of an ever-growing single conversation.

Q15 → C Research mode is built for deep, current, multi-source questions: it conducts multiple searches and compiles cited findings, and it requires web search to be on. Training knowledge cannot cover last month's events; an artifact is an output container, not an information source; and no instruction can grant Claude knowledge it does not have access to.

Ready for the rest?

You've just read about 30% of the full guide — at full depth, nothing dumbed down.

The full Unofficial CCAO-F Study Guide includes

- Every domain at the same depth as Domain 1 you just read — diagrams, decision tables, trap pairs, and cue boxes throughout
- The complete 7-day plan: 7 days of guided study, hands-on drills, and checkpoints
- All 36 original practice questions with full answer-key explanations — including the distractor rationales
- The quick-reference sheets and night-before cram material

Get the full edition at zehon.com