

The AI Bootcamp

4

Keep the conversation going

In this chapter you'll learn

4 follow-up lines that fix most bad replies, and the point where a thread stops helping you

The AI Bootcamp

THE PROBLEM

Most AI education is a pile of disconnected tips. A LinkedIn carousel here, a YouTube tutorial there, a \$2,000 cohort course that goes out of date before the refund window closes. None of it is built to be read in order, so nothing compounds. You finish twelve videos and can recite twelve facts, not perform one new skill.

THE FIX

The AI Bootcamp is built the other way. 90 skills, packed 2 to 3 per chapter, 36 chapters, 3 tracks. Every chapter states what it needs from you and what it hands to the chapter after it, so skill 47 only works because skill 12 already landed. Read it in order and by the end you are not collecting facts about AI, you are running work through it.

LETTER FROM THE EDITOR

We built this because we kept meeting smart people who felt behind on AI for no reason other than nobody had organized the material for them.

This book does not assume you are technical, and it does not water anything down either. Every chapter earns its place: a real skill, a real exercise, a one-page companion you can flip back to in month three when you have forgotten the details but still need the answer.

If a chapter loses you, or a chapter changes how you work, I want to hear about it either way.

Kris

Editor, AI Central

Write to me at editor@thecentral.ai

THE 3 TRACKS

AI 101

Get running

Pick a tool, set it up, write a clear instruction, use it safely at work.

AI 202

Prompt properly

The framework behind every good instruction, and catching AI when it is wrong.

AI 303

Build without watching it

Connect AI to your tools, and automations that run while you are not looking.

IN THIS LESSON YOU'LL LEARN

By the end of this chapter you can take a mediocre first reply and steer it into a usable one in 2 or 3 lines, you will know the difference between repairing a thread and restarting it, and you will know when a conversation has stopped being worth continuing.

This is the exact point where most people quit. They type one thing, get back something flat, decide the tool is overrated and never open it again. The reply they judged was a first draft, and the second attempt is where nearly all the value in this course lives.

CONTENTS	4.1	The first reply is a draft, not a verdict
	4.2	4 lines fix most of what comes back
	4.3	Name the one thing that is wrong
	4.4	A thread has a working life

TOPICS REQUIRED

Chapter 3. You need the task you decided was worth handing over, and the reply it gave you. Ideally the disappointing one: this chapter is more useful when you have something imperfect to work on.

AFTER THIS CHAPTER

Chapter 5 assumes you can steer a reply after it arrives. That is what lets it concentrate on the opening instruction instead of teaching rescue moves.

4.1 The first reply is a draft, not a verdict

You typed something real, you got back something fine but not good, and a voice said: this is overrated. That moment is the single biggest reason people stop. It is also the easiest one to fix.

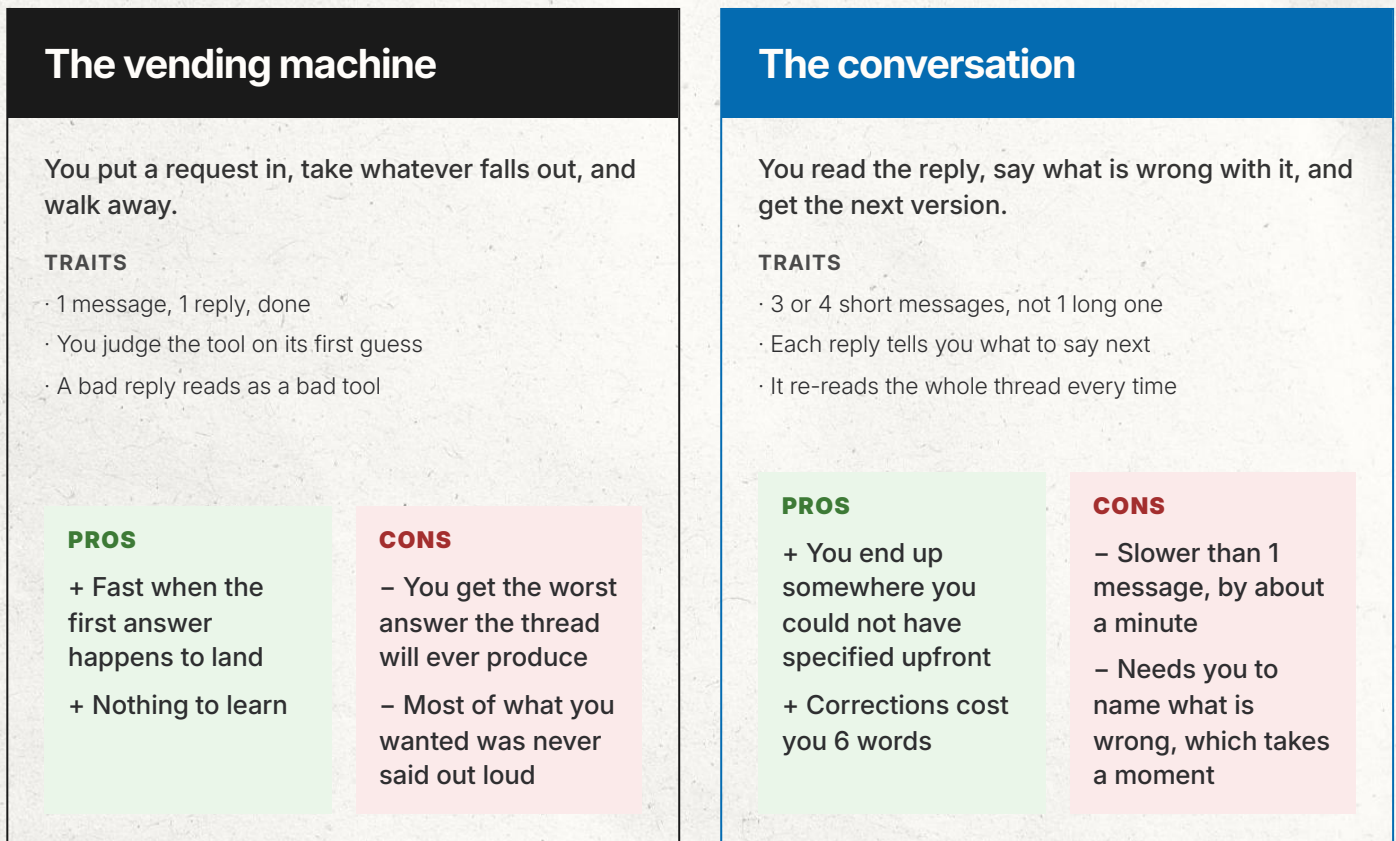


Fig 4.1 Two ways of using the identical box. The second one is not more advanced, it is just the way almost nobody uses it in week 1.

4.1

The moment. One flat reply, and the conclusion that the tool cannot do your job. Almost everybody has it.

What it actually was. A first guess made with the least information it will ever have about you.

It re-reads everything. Each new message is answered against the whole thread, not on its own. Follow-ups compound.

Here is the moment this chapter exists to fix. You handed over a real task in Chapter 3, you got back something competent and bland, and somewhere in the back of your head a verdict formed: this is overrated. Most people reach that verdict in week 1, and most of them stop there.

The reply you judged was a first draft written by someone who has read one paragraph about your work and has never met your audience. It is not a verdict on the tool. It is an opening position, and the conversation it opens is the part nobody starts.

The mechanic underneath is worth knowing, because it changes how the next message feels to write. When you send a second message, the assistant does not read that message alone. It re-reads the entire thread from the top, your first ask, its reply, your correction, and answers again with more information than it had the first time. Nothing is wasted, and you are not nagging. This is the intended way to use it.

Think about what you would do if a colleague sent you that same draft. You would not conclude they cannot write. You would say: too long, and it is aimed at the wrong reader, it goes to the board not the team. Then you would get a better version back in 10 minutes and think nothing of it.

The reason that feels awkward here is that the box looks like a search bar, and nobody argues with a search bar. You type, you get results, you either use them or try different words. That habit is doing real damage in the first week, because it stops you at the one place where continuing is cheap and obvious.

So do this before you read any further. Open the reply you got in Chapter 3, resist starting a new chat, and type one more line into the same thread.

4.2 4 lines fix most of what comes back

Bad replies fail in a small number of ways, and each way has a short line that repairs it. None of these need special wording. You can use all 4 today.

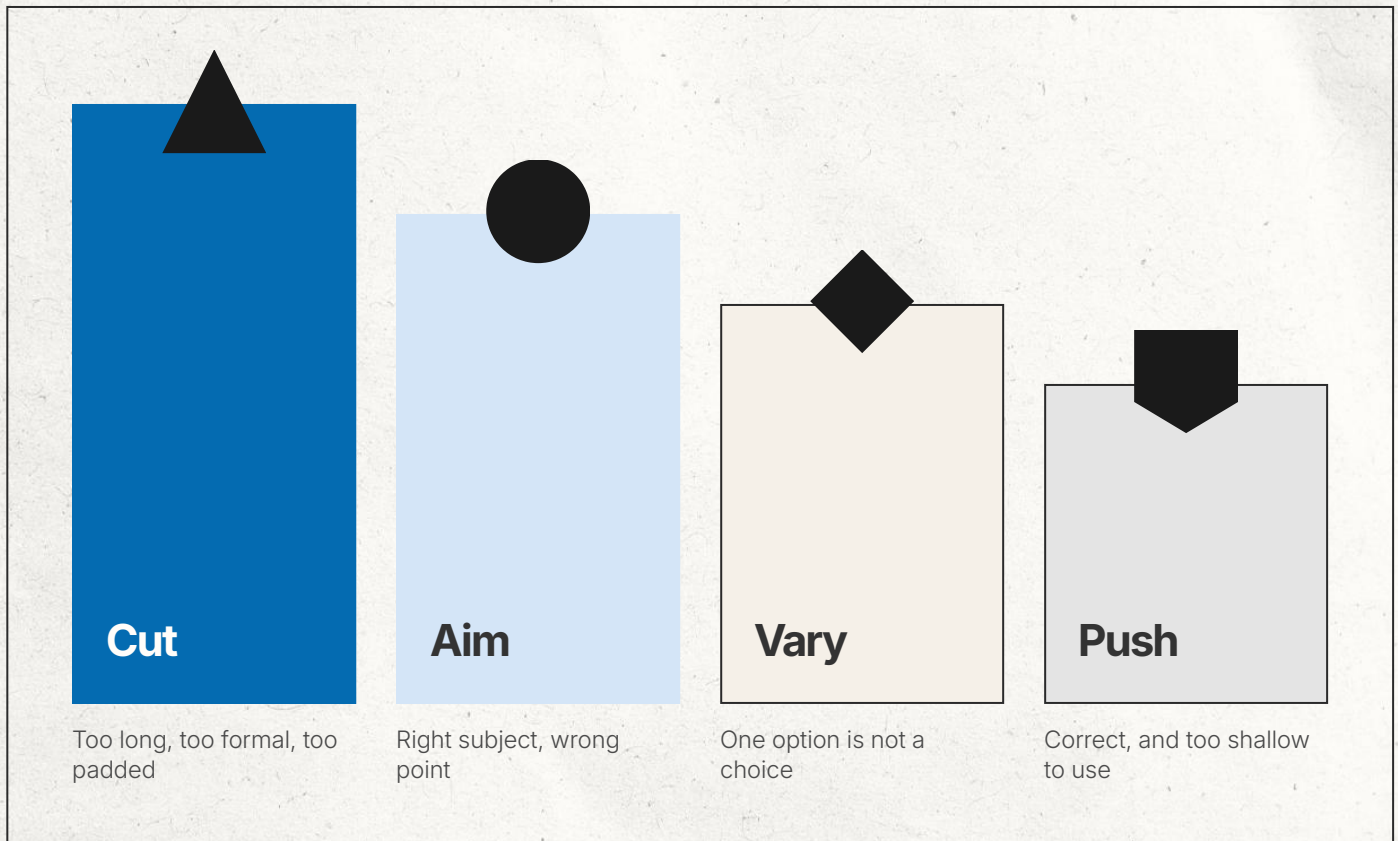


Fig 4.2 The 4 repairs, and the line that performs each one. The bar heights are visual rhythm only, not a ranking: which one you need depends entirely on what came back.

4.2

Name the defect. "Make it better" gets you a different draft. "Half the length" gets you a better one.

One repair per message.
Send the biggest problem first. Fixing 4 things at once tends to fix none of them well.

Point at the sentence. Paste back the line you dislike and say what is wrong with it. That one changes, the rest stays.

Most disappointing replies fail in 1 of 4 ways, and each has a line short enough to type without thinking about it.

Cut. The most common failure is bulk: an opening paragraph restating your question, a closing paragraph summarising itself, and hedging in between. The lines that fix it are blunt. "Half that length." "Cut the introduction and the summary, keep the middle." "Less formal, write it the way I would say it out loud." Length and register are the 2 cheapest things to change and the 2 that most improve a draft. Making it sound specifically like you, rather than merely less stiff, is Chapter 10.

Aim. Sometimes it is the right subject and the wrong point, which is the failure that most often gets read as stupidity. It is not: you knew what mattered and never said it. The repair is to say it now. "You missed the point. The point is that the deadline already slipped, not that the budget is tight. Try again." Naming what it missed is the entire move, and it is why vague dissatisfaction gets you vague changes.

Vary. One option is not a choice, and you cannot tell whether a draft is good until you have seen a worse one. "Give me 3 versions, different angles, one line each, no explanation." Then you get to be an editor instead of an author: "version 2, but end it the way version 3 ends." That sentence is often the whole of the work.

Push. The last failure is a reply that is correct and too shallow to use, which happens because it answered the question you asked rather than the one you meant. "What did you leave out that I would regret not knowing?" is the most useful 12 words in this chapter. Getting it to genuinely argue with you, rather than just add more, is Chapter 21.

Pick whichever of the 4 matches what is actually in front of you and send it now, in your own words. None of these are magic phrases, and Chapter 5 covers how to open better so you need them less often.

4.3 Name the one thing that is wrong

There is a wrong way to keep going, and it is the way most people do it: adding one more requirement every turn until the whole thing collapses. The research on this is unusually clear.

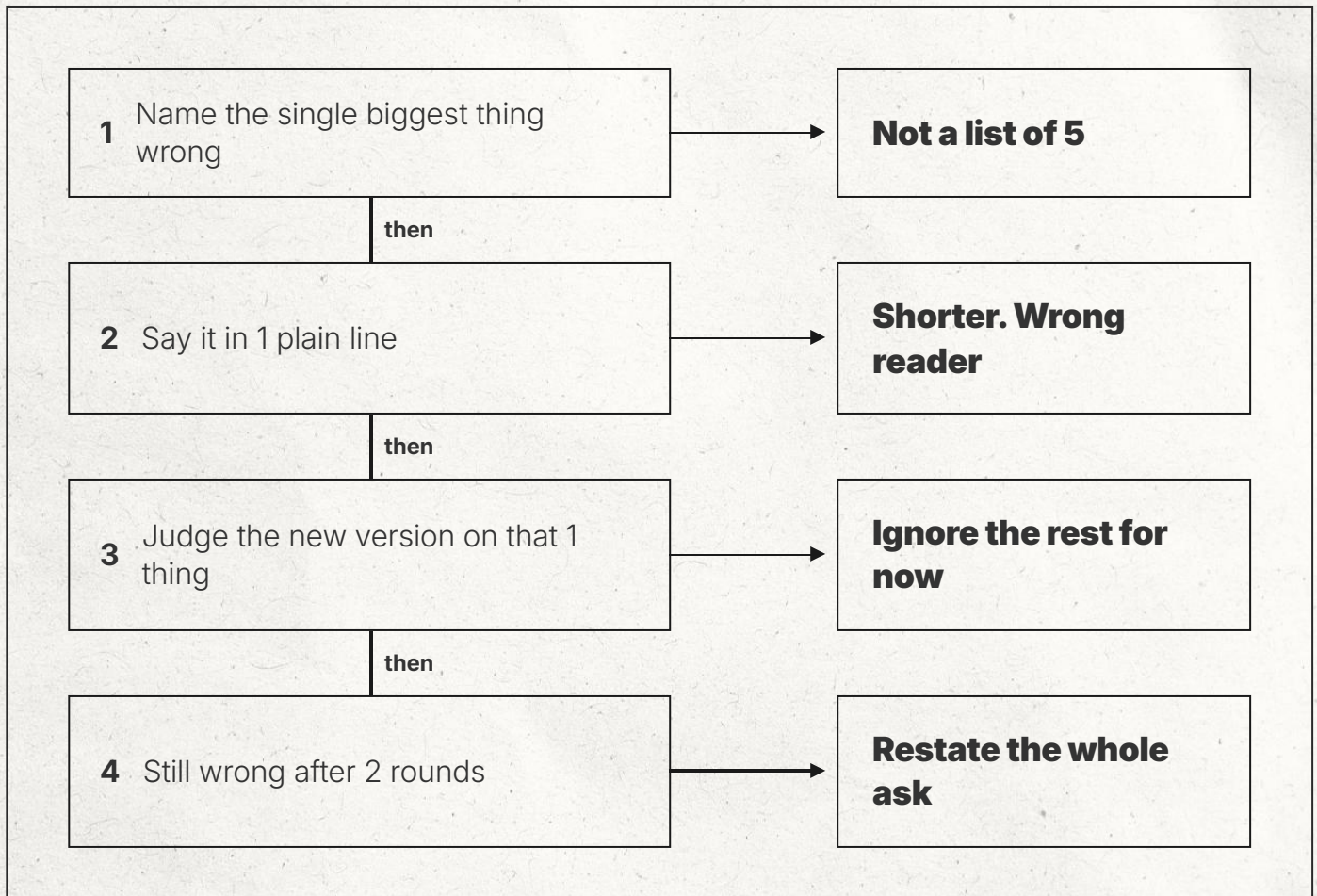


Fig 4.3 The loop, and the exit from it. Step 4 is the one people skip, which is why threads die slowly instead of getting restarted quickly.

4.3

ChatGPT Edit any of your own messages to rewind, or use Branch in new chat from a message menu to explore a second route without losing the first.

Claude The pencil icon on your message forks the thread from that point. Small arrows underneath switch between the versions.

Gemini Edit text next to your prompt, then Update, and it answers again. Regenerate works on the most recent reply only.

Copilot No rewind to speak of. Starting a new chat session clears the context it was carrying, which is the reset.

Keeping the conversation going does not mean feeding it your requirements a spoonful at a time. That specific habit is measurably worse than saying everything at once.

A study presented at ICLR in 2026 ran more than 200,000 simulated conversations across 15 models from 8 providers, including the ones you are using. When the same task was delivered a piece at a time across several turns instead of in a single well-specified message, performance dropped by an average of 39%. The interesting part is why. The models did not get less capable, they got less reliable: they committed to an assumption in the first turn or two, and once a conversation had taken a wrong turn, it did not find its way back.

So draw a line between 2 things that look identical from the outside. Iterating on an output is the move: it has produced a draft, you tell it what is wrong with the draft, it produces a better one. Drip-feeding the input is the trap: you withhold half of what you needed and release it turn by turn, and every turn it guesses at what you have not said yet and builds on the guess.

That gives you a rule you can actually apply. Two rounds of correction with no real progress means the problem is at the root, not in the draft, and no further patching will reach it. Stop. Write the whole ask again in a single message, this time including everything the last 4 turns taught you that you needed to say.

Better still, do not retype it into the bottom of a polluted thread. All 4 assistants let you go back and change what you said, to different degrees, and the notes beside this page give the specifics. Editing your earlier message rewinds the conversation to that point, so the wrong turns stop feeding into what comes next. Correcting an assistant mid-thread leaves the mistake sitting in view, and it keeps building on what it can still see.

Try the loop once on the thread you already have open. Name 1 defect, send 1 line, and count the rounds.

4.4 A thread has a working life

Chapter 2 told you to start a new chat when you change topics. Here is the reason, because it is not tidiness and it is not about keeping your sidebar neat.

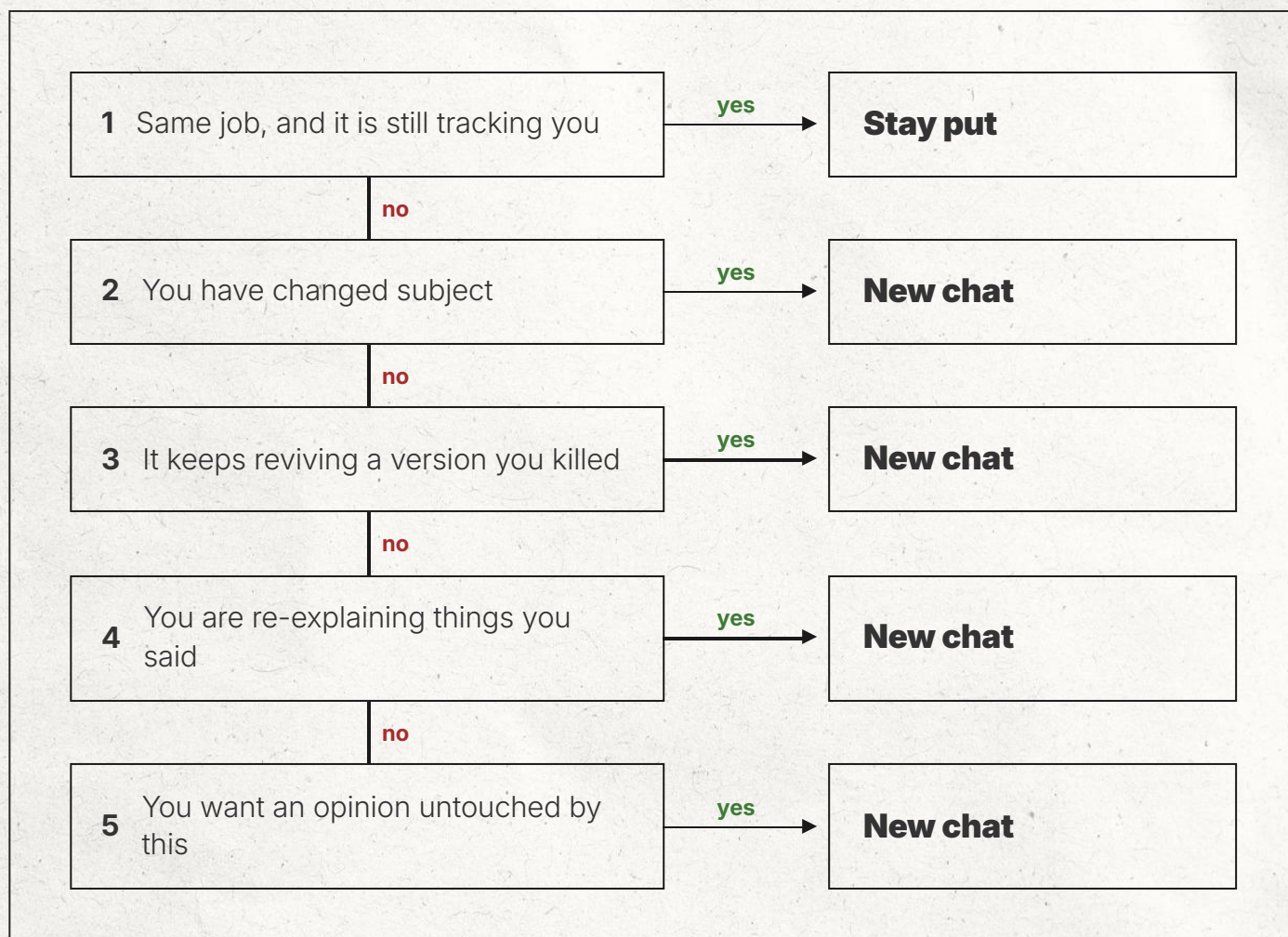


Fig 4.4 Read down until a line is true. Almost every reason to start over is a version of the same thing: something in the window no longer applies but is still being read.

4.4

It thins before it fills. Attention gets less reliable long before the thread hits any published limit.

Old material competes. The draft you abandoned an hour ago is still in view, and still pulling answers toward it.

Carry it in 3 lines. What you are making, who it is for, what you have already ruled out. That is the whole handover.

A thread has a working life. It starts rough, gets sharper for a while as the assistant learns what you actually want, and then at some point quietly starts getting worse.

Part of that is capacity. Everything in a conversation sits inside a fixed budget, and long threads eventually exceed it. None of the 4 assistants simply cut you off any more: Claude summarises the earlier part of a long conversation to keep going, Copilot compacts the older turns, and the others manage it in comparable ways. That keeps the thread alive, and it also means what you carefully specified at message 6 is no longer being read in the words you wrote.

The larger part is that quality falls off well before capacity does. Independent testing of 18 leading models, including the ones behind all 4 assistants, found that performance grows steadily less reliable as the input gets longer, even on tasks that are trivially easy when asked in isolation. Separate work has shown position matters too: material at the start and the end of a long input gets attended to far more reliably than material buried in the middle, which is exactly where your constraint from 20 messages ago is now sitting.

Then there is the finding that turns this from a filing preference into a real reason. Irrelevant material does not sit there neutrally, it actively drags the answer toward itself, and a single stray passage is enough to measurably hurt accuracy. So the version you rejected, the tangent you took, the half-answered question about a different client: all of it is still in the room, and all of it is competing with the thing you actually want.

That is the case for 1 thread per topic. When the subject changes, the previous subject is not context, it is interference. When it keeps offering you back a version you killed 10 messages ago, it is not being stubborn, it is reading a document in which that version still exists. Start a new chat and give it 3 lines: what you are making, who it is for, and what you have already ruled out. That takes 20 seconds and it beats 40 messages of accumulated debris.

One thing not to lean on here. All 4 assistants now remember something about you between conversations, and it is on by default in most of them, but what that memory holds is preferences and scattered facts, not the reasoning of a particular thread. It will not carry your half-finished argument into the new chat for you. Making context carry properly between conversations is Chapter 11. For now, write the 3 lines yourself.

READING ENDS HERE

Everything before was reading. Everything after is doing.

3 exercises, About 10 minutes and one real decision that opens Chapter 5.

Do these before Chapter 5

3 exercises. About 10 minutes

Open the reply you got in Chapter 3. Write the single biggest thing wrong with it in one sentence, then send exactly that sentence as a follow-up. Note what changed.

Run 2 more follow-ups on the same thread using the 4 moves. Write down the line you sent and whether it worked.

THE FOLLOW-UP THAT WORKED

THE FOLLOW-UP THAT DID NOTHING

Write the 3-line handover for that thread: what you are making, who it is for, what you have already ruled out. You will paste it into Chapter 5.

Chapter 4 on one page

THE 4 MOVES

Cut (too long, too formal) · Aim (right subject, wrong point) · Vary (ask for 3 versions) · Push (what did you leave out)

LINES THAT WORK

"Half that length." · "You missed the point. The point is X. Try again." · "Give me 3 versions, one line each." · "What did you leave out that I would regret not knowing?"

THE 2-ROUND RULE

Two corrections with no real progress means the problem is the ask, not the draft. Restate the whole thing, or edit your original message and rewind

ONE THREAD PER TOPIC

Old material in a thread is not neutral. It competes with the answer you want, and attention thins before the thread is ever full

NEEDED BEFORE CHAPTER 5

One thread you actually steered, and the 3-line handover you wrote for it

BRING TO CHAPTER 5

The thread you steered and the 3-line handover you wrote. Chapter 5 uses that handover as the raw material for a proper opening instruction.

BOTTOM LINE

You can now take a flat first reply and steer it into something usable, and you know when to stop steering and start again.

NEXT CHAPTER

Write an instruction it can actually follow

Chapter 5 works on the opening message instead of the repair. Context, task, format: the 3 things a clear ask contains, in plain English, so you spend fewer rounds fixing what you could have said first.

BRING WITH YOU

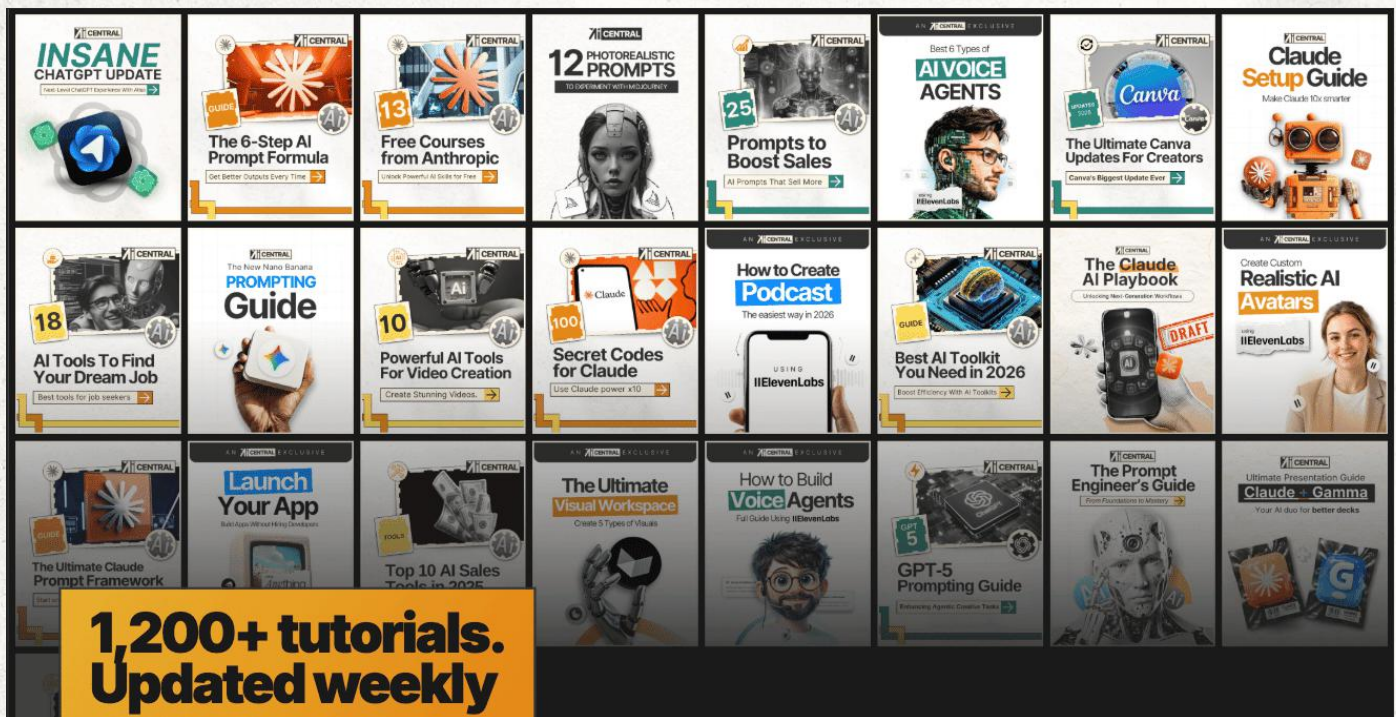
- ☐ The 3-line handover from exercise 3
- ☐ The follow-up line that worked, and the one that did nothing

GO FURTHER

The Ultimate AI Library

Every guide, prompt pack and playbook AI Central has published, in one place. The reference shelf this course keeps pointing you to.

explore the library



WHO MAKES THIS



AI Central Media is a London-based editorial team. Our flagship publication covers practical AI for senior professionals, and 100+ AI companies, SaaS platforms and growth teams have advertised with us since 2023.

100+

AI companies, SaaS platforms, education brands and growth teams have advertised with us since 2023

613K

Accounts reached every month, across the 7 channels we operate

151

Countries our readers are in, plus all 50 US states

PROUD PARTNER OF

GAMMA

IIElevenLabs

guidde.

Outskill

Luma AI

HubSpot

attio

S Synthflow

Notion

Replit

We turn attention into pipeline.

visit thecentral.ai

