

TEAM400

Welcome to the

Claude for Work Training

Modules 1–3

Foundations

Understand what AI is, how Claude works,
and how to prompt it for business results

bit.ly/Claude400Slides

<https://claude400.ai/>

[Password: team400](#)

Michael Ridland | Founder @ Team 400

2026

Module 1

Welcome

What Is Claude & Why It Matters

Course overview

Today's agenda

Session 01

Foundations

What is Claude & why it matters
How Claude works
Prompting essentials

Modules 1–3

Session 02

Applied Skills

Writing, email & brainstorming
Projects & knowledge bases
Data analysis & Office integration.
Research, connectors

Modules 4–6

Session 03

Advanced

Skills & automation & MCP
Memory, Cowork & AI-first org

Modules 7–9

Morning session · 9:00 – 12:30

Morning agenda

ACCESS DETAILS

Login IT Training Event

Password itTraining2026

- 8:40 Registration
- 9:00 Kick-off & Welcome
- 9:20 – 10:15 Slides & Demonstration
- 10:30 – 10:45 Morning break
- 10:45 – 11:00 Slides & Demonstration
- 11:00 – 12:00 Exercises
- 12:00 – 12:30 Q&A
- 12:30 Lunch · full-day attendees continue

[bit.ly/Claude400Slides](https://claude400slides.com)

<https://claudeday.team400.ai/>

Password: team400

Trainers & Mentors



Michael



Libin



Darcy



Thai

TEAM 400

What is Claude?

Claude is built by **Anthropic**, a San Francisco-based AI safety company founded in 2021 by former OpenAI researchers.

Their core belief is that AI should be **helpful**, **harmless**, and **honest**.

- Fast Moving Products & Quality AI
- Less likely to make things up than other AI tools
- More likely to say “I don’t know” when uncertain
- Designed to follow nuanced instructions precisely

What Claude can do

Think & Reason

Analyse documents, compare options, build strategy

Write & Edit

Emails, reports, presentations in your voice

Research

Search the web, synthesise sources, cite everything

Analyse Data

Upload CSV/Excel, run calculations, produce charts

Create Files

Word docs, PowerPoint decks, spreadsheets, PDFs

Build Tools

Interactive dashboards, calculators, mini-apps

Connect

Google Workspace, Microsoft 365, Slack, Jira & more

Assistants **(New)**

When you combine all functionality inside Claude you can build your own AI employees.

Cowork **(New)**

Work with Claude as a coworker on long and complex knowledge based work.

The mental model

“Claude is a brilliant new hire who knows everything about your industry but nothing about your company.”

Your job is to give it **context**. The better the briefing, the better the output. That’s what we’ll practise today.

Module 1.5-2

How Claude Works

Actually Works (and Why It Matters)

Four disciplines

1

Delegation

Knowing what to give to Claude and what to keep for yourself. Where does AI give leverage - and where is human judgment non-negotiable?

2

Description

Communicating the task clearly enough that Claude can do it. Context, role, format, examples - the prompting skill.

3

Discernment

Evaluating what comes back. Is it correct? Plausible-but-wrong? Surface-level? This separates leverage from liability.

*snake example

4

Diligence

Being responsible for the work. Disclosing AI use, taking accountability for output, keeping human judgment in the loop.

*Deloitte example

How Claude/AI generates text

Claude is fundamentally a **next-token predictor**. It looks at all the text in the conversation and predicts the most likely next word. Then the next. Then the next. No database lookup, no rule engine.

Claude has a **training cutoff**. Anything after that date, Claude doesn't know natively. This is why Research mode and connectors matter.

Anthropic uses **Constitutional AI** - principles the model is trained to follow. This is why Claude tends to say "I don't know" rather than guess, and why it pushes back on harmful requests.

Three properties to know

1

Knowledge

Claude knows a lot, but not everything, and not what's recent. Depth varies, mainstream topics deeply, niche topics shallowly.

Implication: For current information, use Research or connectors. For your company's data, use Projects.

2

Working Memory

Each conversation has a context window, roughly 200K tokens (~500 pages). Fills up fast with long documents and multi-turn chats.

Implication: Summarise in chunks. Start fresh conversations when one becomes unwieldy. Use Projects for persistent context.

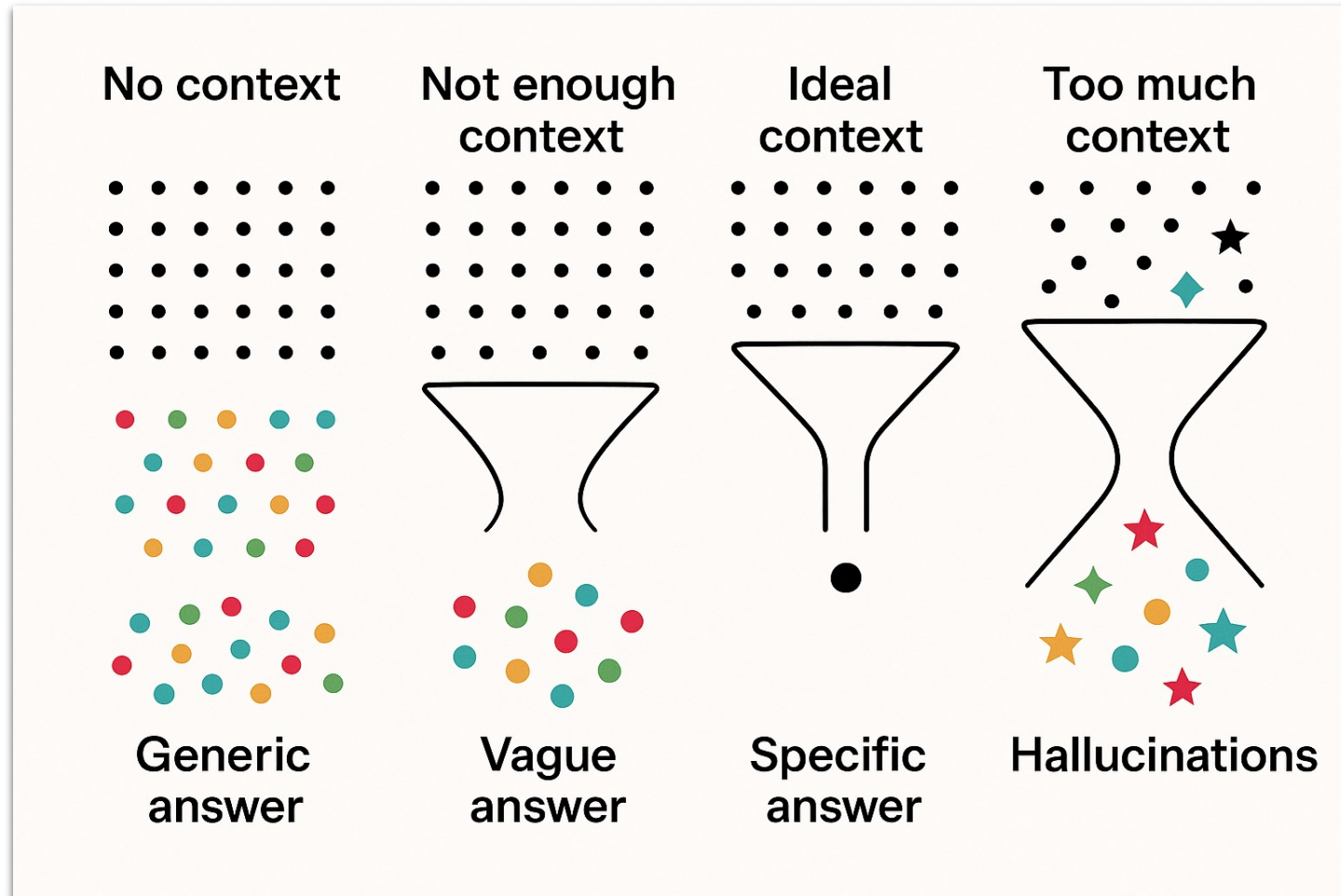
3

Steerability

Claude follows instructions remarkably well. You can change tone, format, depth, and behaviour through prompts and project instructions.

Implication: If Claude isn't behaving the way you want, give clearer instructions or set up a Project.

Context Engineering



Remember this

- Claude doesn't "look things up", it generates text that fits the pattern of the conversation
- Confidence is not correctness, your job is to be the calibration layer
- Context, examples, and clear instructions shape the output, that's what we'll practise next

Data privacy

No Training on Your Data

Anthropic does NOT train models on your data on Team, Enterprise, or API plans. Period.

*Also need to check your settings

Controlled Retention

Conversations stored for 30 days (Team) or configurable windows (Enterprise). Audit logs available.

Enterprise Certifications

SOC 2 Type II, ISO 27001, ISO 42001 (AI management). HIPAA-ready with BAA on Enterprise.

Module 3

Prompting

Essentials

Before & after

✗ VAGUE PROMPT

"Write me a marketing email."

Result: generic, off-brand, requires multiple revisions

✓ STRUCTURED PROMPT

"Write a product launch email for our CRM tool targeting CFOs. Tone: professional but warm. Include a 30-day trial CTA. 150 words max."

Result: on-brand, targeted, ready to send

6 techniques for effective prompting

1. Provide context

Before

"Tell me about climate change."

After

"Explain three major impacts of climate change on agriculture in tropical regions, with examples from the past decade."

2. Show examples of "good"

Before

"Please convert this technical statement to plain language: 'The platform implements end-to-end encryption protocols to safeguard data integrity.'"

After

"Here are two examples of how to convert technical jargon into plain language: [examples]. Now, please convert this statement..."

3. Specify output constraints

Before

"Design me a personal art portfolio website."

After

"Create a clean, modern single-page portfolio with these sections: Hero, About Me, Skills, Portfolio, Projects, Experience, and Contact."

6 techniques for effective prompting

4. Break complex tasks into steps

Before

"Analyze this quarterly sales data."

After

"I'd like to analyze this quarterly sales data. Please approach this by: 1. Identifying top-performing products, 2. Comparing current quarter to previous, 3. Highlighting unusual patterns or trends."

5. Ask it to think first

"Before answering, please think through this problem carefully. Consider the different factors involved, potential constraints, and various approaches before recommending the best solution."

Giving the AI space to think before responding encourages more thoughtful, comprehensive answers.

6. Define the AI's role

"Please explain how rainbows form from the perspective of an experienced science teacher speaking to a bright 10-year-old who's interested in science."

Defining the AI's role, tone, or style helps shape its approach to fit your specific needs and audience.

Secret Weapon: Ask the AI for help with prompting

"I'm trying to get you, Claude, to help me with [goal]. I'm not sure how to phrase my request to get the best results. Can you help me craft an effective prompt for this?"

When you're not sure how to ask for something, the AI can help improve your prompt. This is perhaps the most powerful technique of all!

The CRAFT framework



Context

Who you are, the situation,
background info



Role

What perspective Claude
should take



Action

What you want Claude to
actually do



Format

How to structure the output



Target

Who the output is for

CRAFT in action

[Context] I'm a marketing manager at a B2B SaaS company launching a new CRM feature next Tuesday.

[Role] Act as a direct-response copywriter who specialises in product announcements.

[Action] Write a launch announcement email with a compelling subject line and 3 key benefits.

[Format] 150 words max, short paragraphs, one clear CTA button at the end.

[Target] CFOs and finance directors at mid-market companies (200–500 employees).

Ideation & creativity tricks

The "List 20" Technique

Ask for 20+ ideas to push past the first 5 generic answers. The AI's best, most creative suggestions emerge in items 10–20.

"Give me 20 different angles for a pitch deck about sustainable packaging."

"Bait" Disagreement

State a deliberately one-sided opinion to provoke the AI into correcting you with nuanced reasoning.

"Obviously, Approach A is better than B for this, right?"

Deep Persona Role-Play

Go beyond "You are a marketer." Give a specific job title, years of experience, and industry context.

"You are a senior IT architect with 15 years in cloud migration at a Fortune 500 company."

Reverse Prompting

Ask the AI to write the optimal prompt for your goal. It often knows its own triggers better than you do.

"What would be the best prompt to give you to achieve [Goal]?"

Structuring & scaffolding

XML Tags for Prompt Architecture

Separate instructions, context, and examples with tags like <context>, <task>, <format>. Prevents confusion in long prompts.

*"<context>Q1 financials attached</context>
<task>Summarise the 3 biggest risks</task>"*

Few-Shot Conditioning

Provide 2–3 examples of the exact format and tone you want before the real task. The single most effective way to lock in a style.

"Here are 2 examples of the summary format I want: [Ex1] [Ex2]. Now summarise this report."

Chain-of-Thought Scaffolding

Force the AI to "show its work" before the final answer. Dramatically reduces logical errors in math, code, and strategy.

"Think step-by-step. Explain your reasoning before providing the final answer."

Anchor the Output Start

End your prompt with the first words of the desired output. The model is forced to complete the pattern you started.

*"The three main reasons are:
1."*

Iteration & self-correction

Recursive Self-Critique

After receiving an answer, ask the AI to identify 3 weaknesses, then rewrite to fix them. Dramatically reduces filler and hallucinations.

"Now critique your own response. What's weak? Rewrite fixing those specific flaws."

"Ask Me Questions First"

Before a complex task, let the AI interview you for missing context. Produces far better first drafts.

"Before you start, ask me any questions you need to give the best possible answer."

Explicit Negative Constraints

Strip away AI verbosity with direct "Do Not" instructions. Forces purely functional, direct output.

"Do not apologise. Do not use markdown. Do not include an intro or conclusion paragraph."

Context Handoff Documents

When a long chat degrades, ask the AI to summarise progress into a handoff doc for a fresh session.

"Summarise our progress, what worked, and what didn't in a HANDOFF.md so a new chat can continue."

Common mistakes

✗ Being too vague

✗ Packing everything into one prompt

✗ Not providing context

✗ Accepting first output

✗ Treating Claude as a search engine

✓ Add specifics: audience, length, tone, format

✓ Break complex tasks into multiple steps

✓ Include background, examples, constraints upfront

✓ Iterate: “Make it shorter”, “More formal”, “Add data”

✓ It’s a thinking partner, ask it to analyse, draft, compare

Exercise - Prompting

A

Review previous prompts

Have a look at some of your previous prompts if you do have any and see where there might be room for improvement based on these different techniques.

B

Iteration & Self Correction

Pick one of the iteration or self-correction tips and see if you can apply that anywhere and just test it out to get a feeling for it in Claude.

C

Ideation & Creativity

Pick one of the ideation and creative tricks and play around with it. See if you can use one of the tricks and what results you get in Claude.



Theo Hourmouzis  • 2nd
General Manager/Head of Australia & New Zealand
1d • 

+ Follow

...

Five days in as GM, ANZ at Anthropic. Some reflections.

The AI shift in this region is not coming. It's here. I sat with enterprise leaders this week who aren't running pilots, they're rewiring core operations on Claude. The conversations have well and truly moved from "what is this" to "how fast can we scale."

The founder energy is extraordinary. Hundreds of founders building on Anthropic, from pre-seed teams to companies scaling to millions of users. They picked Claude because it works, and they stay because we obsess over making it better.

The clarity of thinking, the speed, the willingness to bet on hard problems. Australian and New Zealand companies aren't watching this moment happen, they're driving it.

The founding ANZ team, incredibly sharp, hit the ground running, and we had colleagues from the US, EMEA and APAC in town to help us launch well. Cross-market support like that isn't standard, it made a difference.

Week one didn't change my view of why I joined, it sharpened it. Customer focus, safety rigour, and a high bar for the work. That's the company I signed up for.

To every customer, founder, and colleague who made this week what it was, thank you! We're just getting started and you.

If you're an ANZ leader thinking about how to build on Claude — let's talk.

 Craig Scroggie and 1,506 others



Like



Comment

Text tells

- "Buzzword salad" - watch for delve, tapestry, it's important to note, certainly!
- Too-perfect structure - uniform tone, identical patterns, unnaturally clean grammar
- No specificity - broad coverage, no concrete details, nothing you didn't already know
- Factual hallucinations - fabricated stats, sources, or quotes that don't exist

— [Rear Window](#)

Anthropic's new Australian boss lubed by AI slop

Theo Hourmouzis was named the local general manager of Anthropic last week. He wasted no time showing the limitations of the models.

Mark Di Stefano *Columnist*

May 4, 2026 – 1.42pm

 Save

 Share

 Gift this article

 Add us as a preferred source on Google

It took less than a week for Anthropic's new local boss **Theo Hourmouzis** to introduce himself and show us who he is. Hourmouzis [was poached from software giant Snowflake](#) to run the Australian office of the hottest AI company in the world.

Avoid AI Slop

Write like a human

- Edit AI output - don't just copy-paste. Add your voice, cut the filler, rewrite flat sentences
- Inject specifics - real names, concrete numbers, actual examples from your experience
- Break the pattern - vary sentence length, use fragments, add personality and opinion
- Remove the buzzwords - delete delve, tapestry, it's important to note, leverage

Prompt better, get better output

- Give Claude your actual context - audience, purpose, constraints, tone
- Show examples of what "good" looks like before asking it to write
- Ask for a specific structure - "3 paragraphs, no intro, lead with the data"
- Tell it what NOT to do - "No bullet lists, no filler phrases, no generic intros"

Verify everything

- Check every fact, stat, and quote - AI confidently fabricates all three
- Google any source it cites - if you can't find it, it probably doesn't exist
- Run the "would I stake my reputation on this?" test before publishing
- Cross-reference with primary sources, not just what the AI tells you

Use AI as a thinking partner

- Start with your own outline or draft - don't ask AI to generate from scratch
- Use AI to challenge, improve, and refine YOUR ideas, not replace them
- Ask it to poke holes in your argument rather than write one for you
- Treat output as a first draft that needs your expertise, not a finished product

Module 3(b)

Custom Instructions

Essentials

Custom Instructions

Set your standing rules once - Claude applies them to every new chat.

What they are

Standing instructions that tell Claude who you are and how you want it to respond. They load automatically at the start of every conversation.

Why they matter

Stop re-explaining yourself. Set your preferences once and every answer follows them - tone, format and rules included.

How to set them up

- 1 Click your initials (lower-left), then open Settings
- 2 Find “What preferences should Claude consider in responses?”
- 3 Paste your instructions into the text box
- 4 Save. They now apply to every new chat

Available on every plan. Keep it under ~500 words.

Rules worth setting (1–5)

Direct, sceptical and commercially useful - behaviours worth hard-coding.

1 No sycophancy

Don't agree by default. Challenge weak assumptions, bad logic and risky thinking, and name the failure mode.

2 Think like an operator

Prioritise execution over explanation: what to do, what to avoid, what's commercial, what creates leverage.

3 Australian context

Australian English, pricing, business norms and regulation unless I say otherwise.

4 Be blunt but useful

Direct language. No motivational tone, filler, hype or excessive caveats - the clearest useful answer, not the most agreeable.

5 Structure for decisions

When comparing options, use tables: upside, downside, risk, cost, complexity, speed and a recommendation.

Rules worth setting

Keep it tight, set it once, then refine it as you go.

6 Separate facts from judgement

Label what's fact, assumption, inference and opinion. Flag uncertainty instead of presenting it as fact.

7 Default to systems thinking

Analyse problems as systems: inputs, constraints, bottlenecks, incentives, feedback loops, failure points and second-order effects.

8 Tell who you are

About me: [your role, what you build, who you serve, and that you operate in Australia].

9 Avoid AI slop

No generic, buzzword-laden output. Be specific, grounded and usable, with no vague strategy language.

Your starting point

Copy this into Settings → “What preferences should Claude consider in responses?”

PASTE THIS — START WITH WHO YOU ARE

About me: *[your role, what you build, who you serve, and that you operate in Australia]*

Be direct, sceptical, and commercially practical. Do not agree with me by default. Challenge weak assumptions and point out risks clearly. Use Australian English and Australian business context. Prioritise execution, revenue, margin, delivery capacity, speed to cash, and operational leverage. Separate facts, assumptions, inference, and opinion. Analyse problems as systems with constraints, bottlenecks, incentives, failure points, and second-order effects. Avoid generic AI-sounding content, hype, filler, praise, and vague strategy language. Give the answer or recommendation first, then reasoning. Use tables when comparing options.

Exercise - Edit your own system prompt

Let Claude interview you, then draft instructions tailored to how you work.

How to use it

- 1 Exercises-2-Features-and-Projects.docx
- 2 Step 2 Custom Instructions

bit.ly/Claude400Slides

What's next

You learned

- What Claude is and why Anthropic built it
- Data privacy
- Prompting framework

UP NEXT

Modules 4–6

Applied Use Cases

- Writing, research & analysis
- Data, spreadsheets & file creation
- Projects & connectors