

45 Documents, 5 AI Models, One Context Block

A prompt engineering breakdown of the CDD 7 Tools run on a fictional NQ civil contractor — covering model routing, context chaining, the local/cloud split for sensitive data, and the prompt patterns that drove a 4.87/5 quality score.

expertaiprompts.com · consultancydd.com · June 2026

CS-EAP-RC-2026-01

CASE SUBJECT Company: Redstone Civil Pty Ltd (fictional) Sector: NQ Civil Earthmoving & Construction Goal: Build document library — 0 to Tier 1 ready Prompts: 52 prompts across 7 tools Duration: ~7 hours total AI-assisted work Output: 45+ documents · 112,799+ words		AI TOOLS USED Claude: Compliance map, CVs, WHS docs, showcases ChatGPT: Tender drafts, grant applications, workflows Gemini: Audience, persona, brand, tender V2 narrative Qwen: Financial data — LOCAL only via LM Studio Copilot: WHS overview, psychosocial register		
52	4.87/5	112,799+	5	~7 hrs
prompts run	avg quality score	words produced	AI models used	total run time

PROMPT STRATEGY

The DNA Block: One Context, All Downstream Prompts

The most important prompt engineering decision in this case study was not about any individual prompt. It was the construction of a single master context block in Tool 1 — the Business DNA — that every subsequent prompt drew from.

Rather than re-explaining the business in each session, the DNA block was assembled once: company registration, owner profile, team, plant, services, past projects, compliance status, target clients, and business goals. This block was then pasted into every prompt across all seven tools, giving every AI model the same ground truth to work from.

PROMPT PATTERN — DNA BLOCK USAGE

```
# Tool 2-7 prompt structure

BUSINESS DNA: [paste full context block here]
CUSTOMER PERSONA: [paste relevant ICP block]
TASK: [specific deliverable for this session]

# Result: AI has full context without re-interviewing the business owner
```

Every improvement to the DNA block compounded. When WHS status moved from Code A to Code C mid-session, that update cascaded into every subsequent prompt automatically — the AI stopped flagging WHS as a gap and started referencing it as a strength.

MODEL ROUTING

The Local/Cloud Split: Why Five Models Were Used

Using five different AI models was not arbitrary. Each model was routed to the task category it handled best — and one hard rule governed everything: **sensitive financial data stayed on a local model only.**

Model	Deployment	Tasks Routed To It	Why
Qwen 3.5	LOCAL LM Studio	Financial sections: schedule of rates, grant budget, capacity statement, financial QA, accountant letter	Revenue figures, rate structures, and financial capacity data never leave the device. Required 25,000 token context + 3.5GB RAM.
Claude	Cloud	CVs, compliance map, WHS deep research, showcase documents, Gap Closure Plan, Before/After	Strong structured document output and precise instruction-following for complex multi-section documents.
ChatGPT 5.5	Cloud	Tender response drafts, grant applications, WHS policy, project case studies, LinkedIn post sequence	High-quality long-form narrative generation with consistent professional tone across multi-section documents.
Gemini Pro	Cloud	Audience definition, brand identity, tender V2 narrative drafting and QA, customer journey, Google Business Profile	Strong at persona-driven content and structured audience analysis. Used for V2 tender narrative after daily rates were confirmed.

Model	Deployment	Tasks Routed To It	Why
Copilot	Cloud	WHS Policy Overview, Psychosocial Risk Register	Produced refined WHS document outputs after initial layout issues were resolved in iteration 2.

CHAINING

How Outputs Became Inputs: The V1→V2 Tender Improvement

The most concrete example of prompt chaining in this case study was the tender application for Townsville City Council (TCW00639). The first run (V1) produced a complete 4-phase tender response — but the commercial section used North Queensland market benchmark rates instead of Redstone Civil's actual figures, because the schedule of rates hadn't been built yet.

Step 1	Tool 7 — Phase 3 (V1)	Qwen generated NQ market benchmark rates for the tender schedule of rates. Output: ranges (\$580–\$650/hr) flagged by financial QA as "unsubmittable — evaluators cannot process ranges."
Step 2	Tool 7 — Daily Rates Schedule	A separate Qwen session built the confirmed schedule (RC-RATES-2026-01) with single confirmed figures for all 8 plant items, 7 labour codes, wet hire, dry hire, mobilisation zones, and overtime rates.
Step 3	Tool 7 — Phase 3 (V2)	The daily rates schedule was attached to the Phase 3 prompt as an input document. Qwen now had confirmed figures and produced single-figure rates throughout. The financial QA pass rate improved from "draft grade" to "submission-ready."

CHAINING PATTERN — ATTACHMENT AS INPUT

```
# V1 prompt (no attachment)
Produce the Schedule of Rates for this tender using NQ market benchmarks.

# V2 prompt (with attachment)
Produce the Schedule of Rates using the attached RC-RATES-2026-01.
Use confirmed single figures only. Do not use market benchmark ranges.
# Result: confirmed figures replace estimated ranges throughout
```

Mid-Session Discovery: The ISO 45003 Gap

During the WHS deep research session, analysis revealed that the Prompt Assembler's WHS prompts (PT[6] and PT[10]) contained no reference to ISO 45003:2021 (psychosocial risk) or the Queensland WHS Amendment Regulation 2022. This was a legally material omission — the QLD 2023 regulation amendments specifically mandate psychosocial hazard management.

The response was a structured patch: PT[6] and PT[10] were updated with psychosocial hazard rows and ISO 45003 references, a new prompt PT[11] (Psychosocial Risk Review) was added, and the Chaining Workflows C2 WHS Policy Framework prompts were updated to match. The Business DNA WHS status was then updated from Code A to Code C. Every subsequent document produced in the session reflected the corrected status.

PATCH PATTERN — MID-SESSION TOOL UPDATE

```
# Before patch: PT[6] output
What to Build Next: WHS Policy Statement, SWMS suite

# Missing: No psychosocial hazard reference

# After patch: PT[6] output
What to Build Next: WHS Policy Statement, SWMS suite,
Psychosocial Risk Register (ISO 45003:2021 – Run PT[11])
```

Gap identified → prompt patched → DNA updated → all future docs correct

QUALITY

What Drove a 4.87/5 Average Across 52 Prompts

Driver	Mechanism	Where It Showed Up
Rich context block	Full Business DNA pasted into every prompt. AI never guessed missing details — it found them in the context.	All 52 prompts. Quality dropped when context was abbreviated.
Specific output format	Every prompt specified exact document structure, word counts, section headers, and tone requirements.	WHS Policy (8 policy areas), Tender (4 phases), Case studies (STAR format).
Three-phase workflow	Brainstorm/Outline → Draft → Polish/QA. No single-shot submissions. Quality improved each phase.	C2 WHS Policy Statement, B6 Grant Application, A4 Tender Application.
Financial QA in Qwen	Separate financial review pass caught: rate ranges (unsubmittable), co-contribution mislabelling, payment terms misalignment.	B6 Grant (\$15K internal cost), A4 Tender (rate ranges → single figures after rates schedule attachment).
Iteration on failure	Layout issues (CV header overflow, table column widths) were diagnosed and fixed systematically. Each fix was captured as a reusable pattern.	Staff CVs (header overlap), Insurance Register (column width), Checklist (nested table overflow).

Six Prompt Engineering Lessons

1	Context compounds — update it. The Business DNA block is not a one-time input. Every time a deliverable was completed (WHS documents, insurance update, LinkedIn rebuild), the DNA block was updated. This made every subsequent prompt smarter without any additional effort per prompt.
2	Route by data sensitivity, not by task complexity. The rule is simple: if the prompt contains real financial figures, actual revenue, or rate structures — it runs in Qwen locally. If it's narrative, strategy, or compliance framing — cloud AI is fine. This is not about trust; it's about data governance.
3	Attachments are inputs, not references. Attaching the daily rates schedule to the Phase 3 tender prompt transformed the output from benchmark estimates to confirmed figures. The model doesn't "reference" attachments passively — it ingests them as authoritative context that overrides its training data.
4	Name the gap, then patch the tool. When ISO 45003 was missing from the WHS prompts, the correct response was to patch the tool, not just fix the document. A gap found in one output is almost certainly present in every output that tool produces.
5	Three phases beats one. Every major document ran through Brainstorm → Draft → Polish. Single-shot prompts produced usable output ~60% of the time. Three-phase runs hit 4.87/5 average quality. The extra two prompts per document are always worth it.
6	Specify the format or inherit the default. Prompts that specified exact word counts, named sections, and tone requirements produced documents that matched the brief. Prompts without format specifications produced competent but generic output that required more revision iterations.

Explore the prompts and outputs

Full case study resources at consultancydd.com/results/redstone-civil/
 Prompt library and AI workflow guides at expertaiprompts.com/
 The 7 Tools system at consultancydd.com/

Case Study Ref	CS-EAP-RC-2026-01	Publisher	Expert AI Prompts expertaiprompts.com
Case Company	Redstone Civil Pty Ltd (fictional practice company)	CDD Tools	consultancydd.com

Redstone Civil Pty Ltd is a fictional practice company for CDD 7 Tools training. All prompt outputs, quality scores, and workflow details are illustrative. expertaiprompts.com · consultancydd.com