

Multi-Model + Agentic AI for Dental Practices

Blake McClellan + Dr. Peter Boulden

August 9, 2026 — Session B Workshop Kit

CONTENTS

Multi-Model + Agentic AI for Dental Practices

Table of Contents

How to Use This Guide

What It Is

Worked Example — Associate Hiring Decision

The Total Loop

The Universal Rules

The 10 Dental Scenarios

When to Deviate from the Table

Model Subscription Recommendations

Card Layout

How to Fill It Out — Live Example

Sample Filled-Out Card

What to Do After You Fill the Card

After a Month

What “Agentic” Means

The 6 Patterns — At a Glance

Scheduled Jobs

Slack Access (Team-based)

Trigger-Based
Monitor + Alert
Human Review Gate — The Safety Layer
Multi-Agent (Advanced)
The Honest Prioritization Order
When This Pattern Makes Sense
Setup Sequence (High Level)
Realistic Timeline for You
What Blake’s Bot Actually Costs
Why This Pattern Works for Solo Owners
What Blake Runs (For Transparency)
The Adaptation for a Dental Owner
How to Build a Solo-Owner Version
Realistic Costs for a Solo Owner
When NOT to Build This Yet

Multi-Model + Agentic AI for Dental Practices

Session B Workshop Kit Blake McClellan + Dr. Peter Boulden Bulletproof Dental CE Summit · August 9, 2026

Your take-home cookbook. Use it Monday. Refer back all year.

Table of Contents

1. **The 4-Step Multi-Model Decision Workflow** — the daily-use workflow. 100% self-serve. Do Monday.
 2. **Which Model to Use When** — decision tree + 10 dental scenarios mapped model-by-model.
 3. **The Multi-Model Decision Card** — how to fill out your physical card on Monday.
 4. **Agentic Patterns for a Dental Practice** — 6 patterns, honest cost, honest timeline.
 5. **Slack Access Setup for Team-Based Practices** — the @claude-aipg pattern, adapted.
 6. **The Overnight Cron Pattern for Solo Owners** — Blake's Learning Queue pattern, adapted.
-

How to Use This Guide

Six chapters. Read them in order the first time. After that, jump to whatever chapter matches what you're building.

HONEST FRAMING THROUGHOUT

Where a pattern is achievable Monday, we say so. Where it needs \$500–2,000 of freelance help, we say so. Where it's not worth building this month, we say so. This guide is **not** a sales pitch for AIPG. It's a shopping list you can hand to your own tech person.

CHAPTER 1

The 4-Step Multi-Model Decision Workflow

The core workflow. Every dental owner in the room can start running this Monday. No developer required. No BAA required. No PHI touched.

WHAT IT IS

The 4-step multi-model workflow is a discipline for making better business decisions using three AI tools that already exist as consumer products: Claude, ChatGPT, and Grok. Each of the four steps has a specific job. Each step takes about 60–90 seconds. The whole loop takes about 5 minutes. You end with a decision that’s stress-tested by three different AI perspectives — meaningfully better than a decision made from a single AI’s output.

The four steps:

STEP	MODEL	TIME	OUTPUT
1. Ideate	● Claude (default)	60–90 sec	First draft framework
2. Critique	● ChatGPT (switch models)	60–90 sec	Blind spots surfaced
3. Refine	● Claude (back to primary)	60–90 sec	Version 2 draft
4. Decide	You (the operator)	30–60 sec	The call
4.5. Adversarial	● Grok — for \$5K+ decisions only	60–90 sec	Worst-case take

RULE OF THUMB

If the decision costs more than
\$5,000

to undo, run all three models. Under \$5,000, Steps 1–4 are enough.

WORKED EXAMPLE — ASSOCIATE HIRING DECISION

The exact decision Blake and Boulden ran live on stage. Full transparency on every prompt, what to expect at each step, and what the output looks like.

Setup. You are considering hiring an associate dentist. Two paths:

- **Path A:** Full partnership path — slower equity build, retain talent long-term
- **Path B:** Employed associate — faster cash flow, higher turnover risk

Cost of choosing wrong: \$50,000–\$200,000 in lost production over 3 years. High stakes. Run all four steps + Grok adversarial pass.

STEP 1 — IDEATE (CLAUDE)

Open: Claude Desktop or claude.ai

Type this prompt:

```
I'm a dental practice owner considering hiring an associate dentist.
I need to decide between two paths:

PATH A: Full partnership path (slower equity build, retain talent long-term)
PATH B: Employed associate (faster cash flow, higher turnover risk)

Build me a comprehensive decision framework. Include:
- Financial comparison (5-year projection for each path)
- Talent retention metrics for each path
- Pros, cons, and red flags for each
- Non-financial factors (culture, patient relationships, leadership burden)
- Decision triggers: when should I choose each path?

Do not include patient names or clinical protocols. Business decision only.
```

Time to output: 30–60 seconds.

What you'll see: Structured framework. Financial projections (often with placeholder numbers you'll need to adjust — Claude doesn't know your production). Pros/cons balanced across both paths. Decision triggers listed.

What you should notice: Claude will give equal weight to both paths. That's the Claude "helpful" tendency. It will not tell you one path is worse than the other unless you push it. That's fine for Step 1. That's exactly what the critique step is for.

Don't do this at Step 1: Don't argue with Claude. Don't add more constraints. Don't try to make the framework "perfect." Just accept the first draft and move to Step 2.

STEP 2 — CRITIQUE (CHATGPT)

Open: ChatGPT — chat.openai.com or the desktop/mobile app

Copy Claude's full output to your clipboard.

Type this prompt into ChatGPT (do not paste Claude's output yet):

```
I'm going to share a decision framework an AI built for me about hiring an associate dentist (partnership path vs. employed). Critique this framework. Push back on the assumptions. Tell me what I'm NOT seeing. What are the hidden risks? What are the emotional factors a financial model misses? What would a dental practice owner who made the WRONG choice look back and say they should have known?
```

```
[Then paste Claude's full output below this prompt in the same message.]
```

Time to output: 60–90 seconds.

What you'll see. A different tone entirely. ChatGPT will find gaps. Common critiques for this scenario:

- "The financial model doesn't account for the associate-departure cascade — top hygiene patients following the associate out"

- "You haven't factored in the golden-handcuffs problem — where a partnership commitment locks both parties in and one wants out"
- "The compliance risk of an employed associate practicing without full ownership investment"
- "The team-morale impact when senior hygienists learn a new dentist is being fast-tracked to ownership"
- "The referral network impact if the associate leaves in year 3"

What you should notice: At least 3 blind spots should be surfaced. If ChatGPT only finds 1–2, sharpen with a follow-up: *"That's a start. Push harder. Pretend you've hired ten associates and every one had a specific problem. What are those problems?"*

STEP 3 — REFINE (BACK TO CLAUDE)

Go back to Claude — same conversation you started in Step 1.

Copy ChatGPT's full critique to your clipboard.

Type this prompt in Claude:

Here's a critique of the framework you built. Address every gap the critique identified. Keep the structure that worked. Produce version 2 that a real practice owner would sign off on. Make decision triggers more specific — I need actual "if X then Path A, if Y then Path B" guidance, not general considerations.

[Then paste ChatGPT's full critique below this prompt.]

Time to output: 60–90 seconds.

What you'll see: Version 2 will be visibly tighter. The golden-handcuffs risk will be explicit and mitigated. The associate-departure cascade will have a contingency. Decision triggers will be sharper: *"Choose partnership if [three specific conditions]. Choose employed if [three specific conditions]."*

What you should notice: Claude has been forced to hold itself to a higher standard. The "eager to please" tendency is gone. You now have a framework a real operator could defend.

STEP 4 — DECIDE

Read version 2. Adapt to your practice reality. Ship.

Before you ship, ask yourself: *Would I defend this decision to my team, to my patients, to my accountant?*

- **If yes:** ship. Make the decision. Move.
- **If no:** what's still bothering you? Back to Step 3 with the specific concern, or add Step 4.5.

NON-NEGOTIABLE

Do NOT let AI make the decision. The AI gave you the analysis. You give the decision. Every time.

STEP 4.5 — ADVERSARIAL (GROK, FOR \$5K+ DECISIONS)

Open: grok.com or Grok in the X app

Type this prompt:

I'm a dental practice owner deciding between partnership path and employed associate for an associate dentist. I need the brutal truth.

For each path, tell me:

- (a) What's the single worst-case scenario?
- (b) What did owners who chose this path later regret?
- (c) What are they proud of?
- (d) If you could give one piece of advice to someone making this decision RIGHT NOW, what would it be?

Do not soften. Do not hedge. Direct language only.

Time to output: 30–60 seconds.

What you'll see: Direct, opinionated, sometimes uncomfortable. This is what you're paying for.

How to use it: Read Grok’s take. Ask yourself: does my decision from Step 4 hold up against these worst-case scenarios? If yes, ship. If no, either mitigate the worst-case explicitly in your implementation plan, or reconsider the decision.

THE TOTAL LOOP

WHAT IT COSTS

Time: about

6 minutes

. Cost:

\$0

on free tiers.

\$20–40/mo

if you subscribe to Claude Pro + ChatGPT Plus (recommended for daily use). Output: a decision stress-tested by three AI perspectives, refined against critique, and stress-tested against worst-case.

MONDAY This is Monday-morning implementable. 100% self-serve. Do this on your first big decision this week.

CHAPTER 2

Which Model to Use When

The decision tree for picking your models. Use this reference when you're unsure which model to reach for.

THE UNIVERSAL RULES

- **Ideate with the model best at long-form structure.** For dental owners, that's Claude.
- **Critique with a different model than you ideated with.** Different training data, different blind spots. Switching is the whole point.
- **Adversarial only when the decision costs \$5K+ to undo.** Otherwise skip. Save the time.
- **Never let AI make the decision.** All four steps end with you.

THE 10 DENTAL SCENARIOS

Copy this table to your desk. Reference it every time you face a decision.

#	SCENARIO	IDEATE	CRITIQUE	ADVERSARIAL
1	Hiring an associate	● Claude	● ChatGPT	● Grok
2	Hiring a hygienist / front-desk / TC	● Claude	● ChatGPT	— (skip)
3	Marketing copy — website, ads, social	● Claude	● ChatGPT	— (skip for low stakes)
4	Vendor evaluation (imaging, PM software, etc)	● ChatGPT (fast for comparison)	● Claude	● Grok (worst-case)
5	Ad creative review (existing ad set)	● ChatGPT (image analysis)	● Claude (tone rewrite)	● Grok (current market)
6	SOP drafting (new office procedure)	● Claude	● ChatGPT	— (skip)
7	Review response drafts (Google reviews)	● Claude	● ChatGPT	— (skip)
8	Strategic pivot (add specialty, second location)	● Claude	● ChatGPT	● Grok
9	Financial modeling (real estate, big equipment)	● Claude	● ChatGPT	● Grok
10	Leadership scenario (team conflict, associate mgmt)	● Claude	● ChatGPT	● Grok (for high-stakes)

WHEN TO DEVIATE FROM THE TABLE

Sometimes you’ll want to swap the pattern. Here’s when:

- **Ideate with ChatGPT instead of Claude when:** You need image analysis in the first pass. Example: “Compare these three vendor demo pages side by side.” ChatGPT’s vision is stronger; start there.

- **Skip the critique step when:** You already know the domain deeply and the AI's job is just to save you drafting time. Example: you've written 20 review responses this year — Claude's draft plus your edit is enough. You don't need a second model on it.
- **Add a second critique when:** The decision is enormous and you want two independent second opinions. Example: buying a \$2M practice. Run Claude → ChatGPT critique → Grok critique → Claude refine → Grok adversarial. The extra 90 seconds is worth it.
- **Use Grok as the primary when:** You need current information the other models don't have. Example: "What's happening in dental M&A in Q2 this year?" Grok has real-time search where Claude and ChatGPT have knowledge cutoffs.

MODEL SUBSCRIPTION RECOMMENDATIONS

For a practice owner using this workflow daily:

CLAUDE PRO — \$20/mo

Worth it if you use Claude more than 3×/week. Higher usage limits, GPT-5 class quality, project workspaces.

CHATGPT PLUS — \$20/mo

Worth it if you need image analysis, GPT-5, or higher usage. Skip if you rarely need those features.

GROK — Free at grok.com

Free tier works for the adversarial passes. X Premium+ if you want expanded features. Optional.

Total monthly cost for full workflow: \$40 (Claude + ChatGPT). Compare to hiring a \$75/hr consultant to help you think through decisions — the workflow pays for itself in one decision per month.

CHAPTER 3

The Multi-Model Decision Card

You have a physical Multi-Model Decision Card in your Bulletproof program. This chapter walks you through filling it out Monday morning on your first real decision.

CARD LAYOUT

The card has six fields:

1. **Decision** — the actual choice you’re facing (in one sentence)
2. **Stakes** — dollar cost if you get it wrong
3. **Ideate Model** — usually Claude
4. **Critique Model** — usually ChatGPT
5. **Adversarial Model** — Grok, only if stakes are \$5K+
6. **Deadline** — when you’ll ship the decision

HOW TO FILL IT OUT — LIVE EXAMPLE

Take your card. Follow along.

FIELD 1 — DECISION

Write your decision as a decision, not a problem. This distinction matters.

- **Problem:** “Hygiene chair is over-booked.”
- **Decision:** “Should we hire a second full-time hygienist by Sept 1?”

A decision has paths. A problem has complaints. If you can’t state your decision as a choice between two or more paths, you have a problem, not a decision. Rewrite it.

Good decision examples for the card:

- “Should we hire a treatment coordinator or expand front-desk hours?”

- "Should we invest in a second cone-beam or upgrade our current one?"
- "Should we run Facebook ads or double down on Google Business Profile optimization?"
- "Should we lease the space next door for a specialty suite or wait 12 months?"

FIELD 2 — STAKES

Write the dollar cost of getting it wrong. Rough estimate is fine — this determines whether you run Grok's adversarial pass.

- **Under \$5,000:** skip Grok
- **\$5,000–\$50,000:** run all four steps + Grok
- **Over \$50,000:** run all four steps + Grok + sleep on it + repeat the loop with a different framing

FIELD 3 — IDEATE MODEL

Almost always Claude. Only swap if you need image analysis in the first pass (rare for decisions, common for ad creative reviews).

FIELD 4 — CRITIQUE MODEL

Almost always ChatGPT. The critique model must be different from the ideate model. That's the whole point.

FIELD 5 — ADVERSARIAL MODEL

Grok if stakes are \$5K+. Empty if not.

FIELD 6 — DEADLINE

When you'll actually ship the decision.

Deadlines make decisions real. Without a deadline, the loop turns into research paralysis. Common deadlines that work:

- **Simple decision, under \$5K:** 24–48 hours
- **Medium decision, \$5–50K:** 3–7 days
- **Big decision, over \$50K:** 1–2 weeks max — after that, more analysis doesn't help you

SAMPLE FILLED-OUT CARD

DECISION: Should we hire a second full-time hygienist by Sept 1?
STAKES: \$180K over 24 months (salary + benefits + lost prod if wrong hire)
IDEATE: Claude
CRITIQUE: ChatGPT
ADVERSARIAL: Grok
DEADLINE: Friday Aug 15 (this Friday)

Notes:

- Run the loop tonight after 8 PM
- Sleep on it, revisit Wednesday
- If clean, ship by Friday

WHAT TO DO AFTER YOU FILL THE CARD

1. **Tonight (or Monday, latest):** Run the 4-step loop on your card's decision.
2. **Track:** After Step 4, write down what changed between version 1 and version 2. The delta is the value of the critique step.
3. **Ship:** By your deadline, make the call. Not "I need to think more." Make the call.
4. **Log:** In a week, write one sentence about how the decision is going. This is your feedback loop that makes the workflow better over time.

AFTER A MONTH

You should have 8–12 filled cards. Look at them. Ask:

- Which decisions changed because of the critique step?
- Which decisions did NOT change even after Grok's adversarial pass?
- Where is the workflow paying you the most?

The answer to that last question is your leverage point. Double down there.

CHAPTER 4

Agentic Patterns for a Dental Practice

The 4-step workflow you just learned is what you do at the keyboard. The agentic layer is what runs while you sleep. Six patterns. Full transparency on each.

WHAT “AGENTIC” MEANS

Agentic = AI that executes work on a schedule or a trigger, then reports back to a human for review.

CRITICAL PRINCIPLE

Every agentic pattern has a **human review gate** before anything external ships. AI drafts. Human approves. Then the message, the post, the response goes out.

Never skip the human gate for patient-facing work. Non-negotiable.

THE 6 PATTERNS — AT A GLANCE

#	PATTERN	WHAT IT IS	BLAKE RUNS?	YOU SHOULD BUILD?
1	Scheduled Jobs	AI runs on cron	Yes	Q3
2	Slack Access	Team @-mentions AI in Slack	Yes	Q3
3	Trigger-Based	Form/event fires AI	Yes	Q4
4	Monitor + Alert	AI watches for anomalies	Yes	Q4
5	Human Review Gate	Approvals dashboard	Yes	Q3-Q4 (critical)
6	Multi-Agent	Multiple AIs coordinate	Yes	2027

1 SCHEDULED JOBS

Q3 BUILD

What it is: A script runs on a schedule (usually cron on Linux/Mac, Task Scheduler on Windows, or a cloud service). At the scheduled time, it invokes an AI with a preset prompt, saves the output, and posts a summary somewhere you'll see it.

What Blake has running:

- **File path:** `~/Library/Logs/hulk/overnight/scripts/job10-learning-queue.sh`
- **Schedule:** Every Sunday 8 AM
- **What it does:** Pops the top research topic from a queue file, sends it to Claude via subprocess, produces a knowledge card (600–1,200 words), writes to a Markdown file, appends to an index, posts a 5-line digest to Slack + Telegram.
- **How Blake reviews:** Slack digest with the topic + hook. Blake opens the full card on his phone (Obsidian mobile) or laptop when he has time.

What you'd build in your practice:

Example — Weekly Google Business Profile review scan.

- **Schedule:** Every Sunday 8 PM.
- **What it does:** Fetches your recent Google reviews (last 7 days) via your review-monitoring tool or Google Business API. Sends to Claude with the prompt: *"Summarize the sentiment. Flag any reviews that need a personal response. Draft response templates for the flagged ones."*
- **Output:** Google Doc summary + Slack notification.
- **Human review:** Monday morning, you (or your office manager) reviews the drafted responses and posts them.

COST TO BUILD

\$500–\$800 freelancer (5–10 hours). Fiverr or Upwork. 1–2 days elapsed. You spend 1–2 hours on spec + review.

COST TO RUN

Effectively \$0 beyond your Claude subscription. Cron runs on your Mac or a \$5/mo VPS.

REALISTIC TIMELINE

Q3 2026 build. Send the spec to a freelancer by mid-Aug.

SPEC YOU'D SEND TO A FREELANCER

I want a shell script that:

1. Runs every Sunday at 8 PM local time
2. Fetches Google Business Profile reviews from the last 7 days
3. Sends them to Claude (via anthropic API or Claude CLI subprocess) with a
fixed prompt (I will provide) for sentiment analysis + response drafts
4. Writes the output to a Google Doc (via Docs API) in a shared folder
5. Posts a 5-line summary to a specific Slack channel
6. Logs all runs to a local file
7. Handles errors gracefully (retry once, then alert me)

Please use bash + cron, or Python + APScheduler, whichever is simpler.

Budget: \$500-800. Timeline: 2 weeks.

HONEST FRAMING

This is a real project, not a weekend hobby. Budget for it. But once built, it runs forever with essentially zero maintenance.

2 SLACK ACCESS (TEAM-BASED)

Q3–Q4 BUILD

What it is: Your team @-mentions an AI bot in a Slack channel. The bot responds in-thread. Same brain as a Claude window — just accessible from within your team’s chat, without anyone having to switch apps.

What Blake has running:

- **File path:** `~/hermes/bin/claude-aipg-responder/responder.py`
- **What it is:** 200-line Python bot using `slack_bolt` (Slack’s official SDK) in Socket Mode. When @-mentioned, strips the mention, sends the message text to `claude -p` (the Claude CLI subprocess), posts the response back to the thread.
- **Auth model:** Two Slack tokens (app-level Socket Mode token + bot user OAuth token). Claude auth inherits from Blake’s Claude Pro plan via the CLI — no separate API key needed.
- **Logs:** `~/Library/Logs/hulk/claude-aipg/responder-YYYY-MM-DD.log`
- **Who uses it:** Blake’s team at AIPG. Blake, Shane, Travis. They @-mention the bot in project channels; it responds in-thread using shared context.

What you’d build in your practice:

Example — Front-desk AI assistant in Slack.

- **Setup:** Your practice has a Slack workspace with channels like `#front-desk`, `#treatment-coordination`, `#ops`.
- **Bot:** `@ai` bot lives in the channels.
- **Use case 1:** “@ai draft a follow-up email for a patient who was here today for the crown consultation” — bot drafts the email in the thread. Front desk edits, sends via their normal patient email system.
- **Use case 2:** “@ai what’s the standard verbiage for our financing policy?” — bot pulls from a knowledge base and answers.
- **Use case 3:** “@ai draft a Google review response for the review in the screenshot” — bot drafts, front desk reviews, front desk posts.

COST TO BUILD

\$800–\$1,500 freelancer. Slightly harder than Pattern 1 because of Slack app configuration and OAuth setup. 2–3 days elapsed. You spend 3–4 hours on spec + review + Slack workspace setup.

COST TO RUN

\$0–5/mo. Runs on your Mac or a \$5/mo VPS. Slack is free at the workspace level.

REALISTIC TIMELINE

Q3–Q4 2026. Do NOT try to build this Monday.

SPEC YOU'D SEND TO A FREELANCER

I want a Slack bot that:

1. Lives in specified channels (#front-desk, #treatment-coordination, #ops)
2. Responds when @-mentioned by any team member
3. Uses Claude (via Anthropic API or Claude CLI) as the underlying brain
4. Has a system prompt I'll provide that establishes the bot as our practice assistant, with our voice, our policies, our compliance boundaries (PHI-free – never asks for or accepts patient names)
5. Threads responses so channel history stays clean
6. Logs all interactions to a local file for review

Setup includes:

- Slack app configuration (Socket Mode, bot user, OAuth scopes, event subs)
- Server that runs the bot (I can provide a Mac or VPS)
- Documentation on how to update the system prompt

Budget: \$800–1,500. Timeline: 3 weeks.

HONEST FRAMING

This is the pattern that makes AI feel “always there” for a team. But it’s a real project. And you have to train your team on the “PHI-free zone” rule — no patient names in the bot chat, ever. Set that boundary before you launch.

3 TRIGGER-BASED

Q4 BUILD

What it is: An external event fires (a form submits, a new patient books, a lead comes in) → your system captures the event → an AI drafts a personalized response → a human reviews and sends.

What Blake has running: GHL webhooks + Vercel serverless functions + Supabase. When a new lead submits a form on a client's site, the webhook fires, the function drafts a personalized follow-up, and the draft lands in an approval queue (Pattern 5).

What you'd build in your practice:

Example 1 — New patient welcome email.

- **Trigger:** New patient books their first appointment through your website or calls in.
- **What happens:** Your system (booking software webhook, or a manually-fired script) sends the new patient's public info (name, appointment type, date — but NOT clinical data) to an AI-drafting endpoint.
- **Draft:** AI drafts a personalized welcome email based on the appointment type.
- **Human review:** Front desk reviews the draft in an approvals dashboard (Pattern 5). Clicks approve. Email sends via your normal email system.

Example 2 — Missed-call callback draft.

- **Trigger:** Someone calls the practice, hangs up before speaking to reception.
- **What happens:** Call log fires a webhook. Script drafts a callback text.
- **Human review:** Front desk reviews and sends.

COST TO BUILD

\$1,000–\$2,500 freelancer. More complex than Pattern 1 because you need the trigger integration. 3–5 days elapsed.

COST TO RUN

\$10–30/mo (hosting for the webhook receiver + AI API calls).

REALISTIC TIMELINE

Q4 2026 or Q1 2027.

HONEST FRAMING

This is where you start to feel real leverage — AI drafts every routine communication, front desk becomes an editor rather than a writer. But it's also where PHI risk starts to creep in. Make sure your PHI-free doctrine is explicit: patient names are OK in the draft prompt, clinical data is NOT.

4 MONITOR + ALERT

Q4 BUILD

What it is: A script watches something at regular intervals (an inbox, a metric, a signal). When something crosses a threshold or looks anomalous, it alerts a human.

What Blake has running:

- **File path:** `~/aipg-os/scripts/fleet-health-check.sh`
- **Schedule:** Weekly (via launchd)
- **What it does:** Reads status signals across Blake's AI infrastructure (Hermes gateway, Hulk local model, Tailscale network, Git repos). Writes a dated health report to a vault file. Sends a Telegram alert to Blake with the verdict + any red or yellow flags.
- **Cost:** \$0 (no LLM in the check — deterministic bash script).

What you'd build in your practice:

Example 1 — Google review sentiment monitor.

- **Watches:** Your Google Business Profile reviews.
- **Schedule:** Every 6 hours.
- **What it does:** Fetches recent reviews. Sends any new ones to Claude with the prompt: *"Rate this review's sentiment 1–10. Flag any that mention: staff, wait time, billing, or clinical concerns. Rate urgency for response."*
- **Alert:** If a review is rated below 6 (negative) or flagged urgent, sends a Slack + text alert to the owner + office manager.

Example 2 — Ad spend anomaly monitor.

- **Watches:** Your Google Ads or Meta Ads spend + cost per lead.
- **Schedule:** Daily at 9 AM.
- **What it does:** Compares yesterday's cost per lead to the 7-day average. Flags if variance is >30%.
- **Alert:** Slack + text to owner / marketing lead if flag triggered.

COST TO BUILD

\$500–\$1,200 freelancer. 2–4 days elapsed.

COST TO RUN

\$0–10/mo depending on API costs.

REALISTIC TIMELINE

Q4 2026.

HONEST FRAMING

This pattern's value is proportional to how much you already ignore the signals. If you check your Google reviews every day, you don't need this. If reviews go unread for a week, this pattern is high-value.

5 HUMAN REVIEW GATE — THE SAFETY LAYER

Q3–Q4 · CRITICAL

What it is: An approvals dashboard where every AI-drafted patient-facing communication queues up, waiting for a human to click approve before it goes out.

What Blake has running:

- **File path:** `~/aipg-os/apps/aipg-console/`
- **Live URL:** `aipg-console.vercel.app`
- **What it is:** A Next.js + Supabase dashboard, mobile-first, that Blake and his team use to approve AI-drafted work orders and client deliverables. Every AI-generated draft that would go to a client (email, ad copy, report, review response) queues here. Blake or a team member reviews on desktop or phone, clicks approve, and only then does the draft ship.

WHY THIS PATTERN IS NON-NEGOTIABLE

Every other agentic pattern in this list is safe to run **only if Pattern 5 is running behind it**. If you're going to have AI drafting review responses, patient emails, or ad copy — a human must approve before publish. Not “usually approve.” Not “spot-check.” **Every single one.**

This is the pattern that separates AI-augmented practices from AI-embarrassed practices.

What you'd build in your practice:

Example — Communications approval queue.

Every AI-drafted piece of patient-facing communication queues in the dashboard:

- Review response drafts
- Follow-up email drafts
- Text callback drafts
- New patient welcome email drafts
- Missed-call outreach drafts

Dashboard shows: draft text, context (which patient interaction generated it), source model, timestamp queued. One click to approve (sends via your normal comm system). One click to reject or edit. Log of approve/reject decisions kept for training/audit.

COST TO BUILD

\$3,000–\$8,000 for a solid MVP (full-stack developer). More if you want mobile-native experience. 3–6 weeks elapsed.

COST TO RUN

\$20–100/mo (hosting + database).

REALISTIC TIMELINE

Q4 2026 or Q1 2027 — build this BEFORE you build Patterns 3 and 4 in production. Approvals dashboard first, agentic senders second.

HONEST FRAMING

This is the biggest single investment in the six patterns. It's also the one that makes everything else possible. If you can't afford this yet, don't build Patterns 3 and 4 for patient-facing work. Use them only for internal drafts that a human copy-pastes into the sending system manually.

6 MULTI-AGENT (ADVANCED)

2027 ASPIRATION

What it is: Multiple AIs coordinating. One agent drafts, another critiques, a third publishes to different surfaces.

What Blake has running:

- Blake's fleet: **Claude Code** (orchestrator, planning), **Jarvis** (GPT-5.5-based worker, execution), **Hulk** (local Qwen 3.6 model, overnight knowledge work), **@claude-aipg** (Slack-side responder).
- These agents coordinate via a shared vault (`~/HulkBrain/`), a handoffs directory (`08 - Operations/_handoffs/`), and cross-machine sync (Obsidian Sync).
- Blake reviews their coordinated output as an orchestrator.

What you'd build in your practice:

Honestly? Nothing in 2026. This is a 2027-and-beyond aspiration.

The pattern requires: multiple LLM subscriptions, orchestration infrastructure, review dashboards for each agent's output, and a lot of ongoing tuning. For a solo practice or a group of 2–5 dentists, the ROI isn't there yet.

Where you might get to it: if you scale into a DSO or a multi-location group where the marketing + ops workload is too big for single-agent coordination, then Pattern 6 becomes interesting. Until then — Patterns 1–5 give you 90% of the value.

HONEST FRAMING

Do NOT let a vendor sell you a multi-agent solution in 2026. The technology is real but the packaging isn't there yet. Wait 12–18 months. If someone pitches you multi-agent for your practice today, ask them what their approvals gate is. If they don't have Pattern 5 baked in, walk away.

THE HONEST PRIORITIZATION ORDER

If you have a technical friend or a budget for a freelancer, build in this order:

1. **First** — **Q3 2026** — Pattern 5 (Human Review Gate). Basic version. Even a Google Sheet with columns “Draft Text | Context | Approve/Reject” works as a v1.
2. **Second** — **Q3 2026** — Pattern 1 (Scheduled Jobs). Start with one non-critical use case (weekly review summary).
3. **Third** — **Q3–Q4 2026** — Pattern 2 (Slack Access), if your team already uses Slack. Big adoption boost.
4. **Fourth** — **Q4 2026** — Pattern 4 (Monitor + Alert). Add one signal you’re already ignoring.
5. **Fifth** — **Q4 2026 or Q1 2027** — Pattern 3 (Trigger-Based), once Pattern 5 is solid.
6. **Never in 2026** — Pattern 6 — wait.

CHAPTER 5

Slack Access Setup for Team-Based Practices

Detailed walkthrough for Pattern 2. If you have a team on Slack, this pattern gives you the biggest single feel-good adoption moment: AI feels “always there” without anyone switching apps.

WHEN THIS PATTERN MAKES SENSE

- Your team already uses Slack daily.
- You have at least 3–5 team members who would use an AI assistant in Slack.
- You have someone (owner, office manager, or hired freelancer) willing to invest 3–4 hours in the initial setup.

If Slack isn’t already a habit for your team — don’t force it just to have an AI bot there. Adopt Slack first. Then add the bot 60 days later.

SETUP SEQUENCE (HIGH LEVEL)

You are not building this yourself. This is a hire-a-freelancer walkthrough so you know what they’re doing on your behalf.

Prerequisites:

- A Slack workspace (free or paid)
- A Claude Pro subscription
- A computer that’s on 24/7 (a Mac in your office, or a \$5/mo Digital Ocean droplet)
- 3–4 hours of your time to review, test, and roll out

Step 1 — Create the Slack App. Freelancer creates a Slack app in your workspace, configures Socket Mode, sets OAuth scopes.

Step 2 — Install the Bot. Freelancer installs the bot to your workspace. You add it to channels where you want it available.

Step 3 — Configure the Brain. Freelancer sets up a system prompt that establishes the bot's role. Yours might look like:

You are the AI assistant for Practice Name, a dental practice with 12 team members across Front Desk, Treatment Coordination, Clinical, and Ops.

Your job:

- Draft communication (patient emails, review responses, callback texts)
- Answer questions about our SOPs (which are in a shared knowledge doc)
- Help team members think through decisions using the multi-model workflow

Rules:

- NEVER accept or ask for patient names, dates of birth, clinical details, or any PHI. Segment-level language only ("the crown patient we discussed" is fine, "John Smith DOB 3/15/1974" is not).
- Match our practice voice: warm, competent, direct. Avoid corporate marketing language.
- If someone asks something outside your scope (clinical protocol, HR issue, financial decision), tell them to loop in Blake or Kim directly.
- Always sign off drafts with "- draft, please review before sending."

Step 4 — Test in a Sandbox Channel. Freelancer creates a `#ai-testing` channel. You and one other team member test the bot for 3–5 days: prompts, edge cases, PHI-doctrine enforcement. Iterate on the system prompt.

Step 5 — Roll Out to Real Channels. Once you're comfortable, add the bot to real work channels. Announce to the team: *"@ai is now available in these channels. Here's what you can use it for. Here's the PHI rule."*

Step 6 — Monitor + Iterate. For the first month, review the bot’s interactions weekly. Look for: PHI slips (correct with a system-prompt tweak), voice drift (retune the prompt), missed opportunities (adjust the “job” section).

REALISTIC TIMELINE FOR YOU

- **Now → Week 2:** Get quotes from freelancers. Post on Fiverr or Upwork: *“Need a Slack bot that responds using Claude, ~200 lines of Python, based on public docs from Anthropic + Slack Bolt SDK.”* Budget \$800–1,500.
- **Weeks 3–5:** Build phase. Freelancer builds. You review, iterate on system prompt.
- **Weeks 6–7:** Sandbox testing.
- **Week 8:** Real rollout, one channel at a time.

Realistic elapsed time: 2 months. **Realistic setup effort from you:** 8–10 hours across those 2 months. **Post-rollout:** 1 hour/week for the first month, then minimal.

WHAT BLAKE’S BOT ACTUALLY COSTS

DEVELOPMENT

Blake wrote his own bot; if he’d hired it out, ~\$1,000.

HOSTING

Runs on Blake’s Mac Mini. Moved to a VPS: \$5/mo.

CLAUDE API COST

Zero, because the bot uses Blake’s Claude Pro subscription via the CLI subprocess pattern. If you use the Anthropic API instead, budget \$30–100/mo depending on usage.

HONEST REALITY CHECK

This is a real project. Not a Monday-morning install. Set expectations with your team accordingly. The reward is that once it’s running, it feels like magic. But you’re 6–8 weeks from magic, not 6–8 hours.

CHAPTER 6

The Overnight Cron Pattern for Solo Owners

The exact pattern Blake uses for his Learning Queue — adapted for a solo dental owner.

WHY THIS PATTERN WORKS FOR SOLO OWNERS

You don't have a team on Slack. You don't need a bot in a chat channel. What you DO need: **AI doing homework overnight so you have better information in the morning.**

The cron pattern (short for “scheduled job”) is the solo-owner's leverage move.

WHAT BLAKE RUNS (FOR TRANSPARENCY)

- **File path:** `~/Library/Logs/hulk/overnight/scripts/job10-learning-queue.sh`

How it works:

1. Blake maintains a text file (`~/HulkBrain/04 - Domain Knowledge/hulk-operations/learning-queue.md`) with a list of research topics he wants deep-dive briefs on. Format: `- [ ] Topic name → category/`
2. Every Sunday 8 AM, the shell script fires (via launchd, macOS's cron equivalent).
3. Script picks the top unchecked item, sends the topic to Claude via a CLI subprocess with a specific research prompt.
4. Claude produces a 600–1,200-word knowledge card and writes it to a Markdown file.
5. Script appends the card to an index file and marks the queue item as done.
6. Script posts a 5-line digest to Slack + Telegram.

7. Blake reads the digest on his phone Sunday morning. If interesting, he opens the full card in Obsidian mobile.

What Blake learns from it: Every week, one deep-research knowledge card lands in his vault without him spending 5–15 minutes writing it. Over a year, that's 52 deep briefs.

THE ADAPTATION FOR A DENTAL OWNER

Same pattern, different content queue. Some ideas:

IDEA 1 — NIGHTLY REVIEW SENTIMENT SCAN

- **What it does:** Every night at 11 PM, fetches your Google reviews from the last 24 hours. Sends them to Claude: *"For each review, rate sentiment 1–10, extract themes, flag any that mention staff, wait time, billing, or clinical concerns."* Writes output to a text file in a shared Google Drive folder. Emails you a 3-line summary in the morning.
- **Value:** You start every morning knowing exactly what patients are saying about your practice — without opening Google Business Profile manually.

IDEA 2 — WEEKLY COMPETITOR CONTENT SCAN

- **What it does:** Every Sunday 8 PM, fetches the 3 nearest competing practices' Google Business Profiles + Instagram feeds. Sends to Claude: *"What did they post this week? What's their pricing showing? Any promotions? Anything worth responding to?"* Writes a 400-word summary to your Google Drive. Slack or email digest Monday morning.
- **Value:** You know what your competitors are doing before you sit down for your Monday morning planning.

IDEA 3 — MONTHLY AD PERFORMANCE DEEP DIVE

- **What it does:** On the 1st of each month, fetches last month's Google Ads + Meta Ads data. Sends to Claude: *"Analyze cost per lead by campaign. Flag anomalies. Recommend 3 optimizations."* Writes a report to Google Drive. Emails you.
- **Value:** You get a data-driven ad performance report every month without asking your marketing agency for one.

HOW TO BUILD A SOLO-OWNER VERSION

You are not writing bash scripts yourself. This is a spec-to-freelancer walkthrough.

Step 1 — Pick Your Use Case. Start with ONE. Not three. Pick the one that answers: *“What information would I love to have every morning that I don’t have today?”* Best starting picks for most dental owners: Idea 1 (review sentiment) or Idea 2 (competitor scan). Both are self-contained, low-risk, high-signal.

Step 2 — Write the Spec. Take the template below and fill in the blanks:

SOLO-OWNER CRON SPEC TEMPLATE

I want a scheduled script that:

1. Runs [WHEN – e.g., every Sunday at 8 PM]
2. Fetches [WHAT – e.g., recent Google reviews for my practice via Google Business Profile API]
3. Sends the data to Claude (via Anthropic API OR the Claude CLI subprocess) with a fixed prompt I'll provide (attached)
4. Writes the output to [WHERE – e.g., a Google Doc in Drive folder XYZ]
5. Sends me a summary via [HOW – e.g., email, Slack DM, Telegram]
6. Logs all runs to a local file
7. Handles errors gracefully (retries once, then alerts me if it still fails)

Please build with bash + cron OR Python + APScheduler, whichever is simpler for you to maintain.

Delivery: I want a README explaining how to modify the prompt later without your help.

Budget: \$500–1,000.

Timeline: 2 weeks.

Attached: the exact Claude prompt I want the script to use.

Step 3 — Post the Spec. Fiverr or Upwork. Search terms: *"Python API scheduler"* or *"shell script automation."* Look for freelancers with:

- 4.8+ rating
- 20+ jobs completed
- English fluency

Get 3 quotes. Pick the one whose portfolio shows real automation work, not just "I know Python."

Step 4 — Provide the Prompt. The prompt is where YOUR domain expertise matters. The freelancer builds the plumbing; you write the prompt.

Example prompt for a review sentiment scanner:

You are an assistant helping a dental practice owner review recent Google reviews. Analyze each review below and produce a summary in this format:

For each review:

- Sentiment score (1-10, where 10 is glowing)
- Topics mentioned (staff, wait time, billing, clinical, other)
- Urgency for response (immediate / this week / no action needed)
- Suggested response tone (grateful / empathetic-apologetic / neutral)

Then a 3-line overall summary:

- Total reviews this period
- Sentiment average
- Any flagged patterns

Do NOT draft the response itself. Just the analysis.

Reviews:

[REVIEWS INSERTED HERE]

Notice: this prompt does NOT ask AI to send anything. It only produces analysis. That's the safe pattern for solo-owner automation — AI produces info, human takes action.

Step 5 — Test + Iterate. For the first 2 weeks after the script goes live: read every output. Refine the prompt when the output isn't quite right. Retire the script if it isn't earning its keep. The freelancer's README should tell you how to modify the prompt yourself without re-hiring them.

Step 6 — Add Pattern 2 (Optional). After 30–60 days of running one script and getting value, add a second script. Never build all three at once — you'll get overwhelmed and stop reading the outputs.

REALISTIC COSTS FOR A SOLO OWNER

SETUP (FREELANCER)

\$500–\$1,000 per script.

ONGOING HOSTING

\$5–10/mo per script (or \$0 if you have a Mac that's always on).

AI API COST

\$5–30/mo depending on how much data you feed it.

YOUR TIME

4–6 hours to spec + review + iterate on the first script. 1–2 hours for each additional script.

Compare to: hiring a virtual assistant to do the same job manually. VA cost: \$500–1,000/mo minimum. And they don't scale.

WHEN NOT TO BUILD THIS YET

Skip the cron pattern if:

- You're not yet reading the outputs of the manual info you already have. If you don't read your reviews now, adding a summary doesn't help.
- You don't have a computer that's on 24/7 or a VPS budget of \$10/mo.
- You're not willing to spec + review a freelancer's work (this is a \$500 gamble if you set no expectations).

Come back to this pattern once you've done the 4-step workflow for 30 days and know you'll actually use the outputs.

FINAL NOTE

Every pattern in this guide has one gate: **a human, at a moment of clarity, decides what ships.** Never let a script send a message to a patient without you or a trusted team member seeing it first. The workflow is powerful because it saves you time on the drafting, not because it removes you from the sending.

Blake runs six of these patterns. His AI fleet coordinates work across three machines. He still reviews everything patient-facing personally. **The gate never disappears. The gate is the practice.**

Monday morning. Claude Desktop. Your first decision. Go.

Bulletproof Dental CE Summit · Session B Blake McClellan + Dr. Peter Boulden
Aug 9, 2026

Questions? blake@aipg.com · aipg.com