



GRAND HAVEN CONFERENCE EDITION

CONFERENCE GUIDEBOOK

Working With AI the Right Way

A Practical Guidebook for Business Owners and Operators

Event: Grand Haven Conference - Thursday & Friday

Presented by: Prospectr Digital

Website: prospectrdigital.com

Every Channel. One Team. Engineered for Performance.

Why This Guide Exists

AI is no longer an experiment. It is a workforce. The businesses pulling ahead right now are not the ones with the smartest AI - everyone has access to the same frontier models. The businesses pulling ahead are the ones who have learned how to **work with AI correctly**: how to ask, how to constrain, how to document, how to secure, and how to architect systems that remember.

This guidebook is the playbook we use every day. Over the next two days, we will work through it together. By the end, you will understand how to move from "chatting with a bot" to **operating an AI system** - one you own, one that compounds in value, and one that gets cheaper to run over time.

WHAT YOU WILL TAKE HOME

Nine structured parts covering how to write requests that actually work, how to secure AI in your business, how to document skills in plain text, how to give AI memory that survives session resets, how to orchestrate multi-agent systems, how to route tasks by cost tier, and how to think about AI ownership. Plus real prompt patterns from daily production use, integration examples, a skill library reference, and a live worksheet to complete during the conference.

Clear and Specific Requests

The Foundation of Everything

The single biggest reason people get bad results from AI is not the model. It is the request.

AI is a brilliant employee with zero context. It does not know your business, your customer, your voice, or your goal - unless you tell it. Vague in, vague out.

The anatomy of a good request

1

ROLE & CONTEXT

Who should the AI act as, and what does it need to know? Include industry, customer type, and key facts.

2

TASK

What exactly do you want produced? "Write a 3-email sequence" not "write emails."

3

CONSTRAINTS

What are the rules? Word limits, tone, what to exclude, what never to do.

4

FORMAT

What should the output look like? Table, list, numbered steps, HTML, plain text.

Weak vs. strong

WEAK REQUEST

"Write me a sales email."

STRONG REQUEST

"You are writing on behalf of a 20-year commercial cleaning company. Write a first-touch cold email to a facilities manager at a medical office building. Under 80 words. Problem-first opening. One specific proof point. One low-friction call to action (a 10-minute call). No pricing. Professional but human tone."

The strong version takes 60 seconds longer to write and saves you five rounds of revisions.

The revision principle

Never rewrite AI output by hand when it misses. **Tell it what was wrong and why.** "Too formal - this reads like a lawyer. Our clients are blue-collar owners. Rewrite at an 8th-grade level with shorter sentences." The AI improves in one pass, and you have just taught it your standard - which matters enormously once your system has memory (Part 4).

Responsible AI and Security Guardrails

The moment AI touches your customer data, your inbox, your CRM, or your money, it stops being a toy and becomes an operational and legal surface area. Guardrails are not optional.

The core principle: least privilege

An AI agent should have access to exactly what it needs for its job and nothing more. Your appointment-setting agent does not need access to payroll. Your content agent does not need database write access. Scope every credential, every API key, and every integration to the narrowest permission that still gets the job done.

Human-in-the-loop for anything irreversible

There is a simple rule we run in our own operation: **AI drafts, humans approve, then it ships.** No AI in our system sends a lead to a client, sends money, deletes data, or emails a customer without a named human approving it first. Autonomy is earned per-task, not granted globally. Start every new automated workflow with an approval gate, and only remove the gate after the workflow has proven itself over volume.

Compliance is your job, not the AI's

AI will happily do things that violate messaging compliance law, FTC endorsement rules, state licensing law, or a platform's terms of service - because it does not carry your liability. You do. Before any AI workflow touches consumer data, outbound messaging, reviews, or lead sourcing, run it past the same compliance lens you would apply to a human employee.

Practical guardrail checklist

- ☐ Every agent has its own credentials (never share your master login with an automation)
- ☐ Write / send / delete actions require human approval until proven
- ☐ Customer PII is minimized - the agent sees only the fields it needs
- ☐ A written Incident Response Plan exists (what happens if an agent misfires at scale?)
- ☐ Cyber insurance application reflects your actual AI usage

- ☐ An audit trail exists - every agent action is logged with a timestamp
- ☐ A kill switch exists - you can stop every agent in under 60 seconds

Skills and Markdown

Documentation Is the Product

THE MINDSET SHIFT

When you work with AI, your documentation is your software. A well-written instruction file is not a record of work done - it is the capability itself.

What is a "skill"?

A skill is a packaged, repeatable capability you have taught your AI system - "generate a client report card," "audit a website before launch," "classify an inbound lead." Each skill is defined by a plain-text instruction document, written in **Markdown**, that tells the AI exactly how to perform that job every single time.

Why Markdown?

- **Readable by humans** - anyone on your team can open and edit it
- **Readable by every AI model** - it is the native language of AI instructions
- **Version-controllable** - you can track every change, forever
- **Portable** - works whether your AI runs on any frontier provider or a local model in your own cloud

The rule: every skill gets a Markdown document

If a task is worth doing twice, it is worth documenting once. Every skill should contain:

1. **Purpose** - what this skill does and when to use it
2. **Inputs** - what information the skill needs to run
3. **Process** - step-by-step, exactly how the job is performed
4. **Quality standard** - what "done and correct" looks like, with pass/fail criteria
5. **Output format** - the exact structure of the deliverable
6. **Examples** - at least one real example of a correct output

REAL EXAMPLE

Before any client website goes live, an AI agent runs a 12-criteria readiness check - copy integrity, business identity consistency, schema markup, redirect mapping, analytics, link integrity, and more. That entire quality standard lives in one Markdown document. Any agent, any model, any day of the week produces the same audit to the same standard. That is what a documented skill buys you: **consistency at scale.**

Persistent Memory

AI Forgets Itself - Here Is the Fix

THE UNCOMFORTABLE TRUTH

Every AI conversation is amnesia by default. When a session ends - when you close the window, when the app crashes, when the context fills up - the AI forgets everything. It does not remember your last project, your standards, your client, or the correction you gave it yesterday. Tomorrow, it is a stranger again.

Worse: even within a long session, AI has a finite "context window." When it fills up, the earliest material silently falls away. The AI can forget the instructions you gave it an hour ago while you are still talking to it.

If you do not solve this, you will spend your life re-explaining your business to a very smart goldfish.

The solution: architect memory, do not assume it

Memory is not a feature you wait for. It is a system you **build** - out of plain Markdown files, checkpoints, and discipline. The AI writes its own notes; the notes survive the session; the next session begins by reading the notes.

A correctly architected memory system

```
/business-system/  
|-- OPERATIONS.md <-- The constitution. Who we are, standing rules,  
| approval requirements, tone, non-negotiables.  
| Every new session reads this FIRST.  
|-- MEMORY.md <-- Long-term memory. Key decisions, client facts,  
| lessons learned. Curated and trimmed regularly.  
/skills/  
| |-- report-card/SKILL.md  
| |-- site-audit/SKILL.md  
| +-- lead-classify/SKILL.md  
/logs/  
+-- 2026-07-14.md <-- Daily working log. What was done, what is pending.  
.checkpoint.json <-- Crash recovery. Current task state, saved continuously.
```

How it works in practice

1. **Session start:** the agent reads `OPERATIONS.md` and `MEMORY.md`. It now knows the business, the rules, and the history - before you type a single word.
2. **During work:** the agent maintains a task manifest and writes checkpoints. If the session crashes mid-job, you start a fresh session, point it at the checkpoint, and it resumes rather than restarts.
3. **Session end:** the agent writes a summary to the daily log and promotes anything important to `MEMORY.md`.
4. **Weekly:** a human prunes `MEMORY.md` so it stays sharp. Memory that grows forever becomes noise.

THE KEY INSIGHT

Persistent instructions live in files, not in conversations. If a rule matters, it goes in `OPERATIONS.md` - because anything that only exists in a chat will eventually be forgotten.

Agentic Orchestration

Agents Managing Agents

Once you have documented skills and persistent memory, you can graduate from "one AI assistant" to an **orchestrated system**: a manager agent that delegates to specialist agents.

The org-chart model

Role	What it does	Assigned by
Orchestrator (manager)	Receives the goal, breaks it into tasks, assigns each to the right specialist, reviews the work, assembles the final deliverable	The human owner
Specialist agents (staff)	Each runs one documented skill - one writes copy, one audits websites, one classifies leads, one builds reports	The orchestrator
Human (owner)	Sets the goal, approves anything irreversible, handles exceptions	N/A

Why orchestration beats one mega-agent

1. **Focus** - an agent with one job dramatically outperforms an agent juggling ten
2. **Auditability** - when something goes wrong, you know exactly which agent and which skill to fix
3. **Cost control** - specialists can run on smaller, cheaper models (Part 6)
4. **Parallelism** - five specialists work simultaneously; one assistant works serially

The delegation contract

An orchestrator hands a specialist the same thing you would hand a contractor: a written work order. Task, inputs, quality standard, output format, deadline. The specialist returns the work; the orchestrator checks it against the skill's pass/fail criteria before accepting. **Never let one agent's unchecked output become another agent's input.** That is how errors compound silently.

The Model Waterfall

Cost Efficiency by Design

Not every task deserves your most expensive model. Frontier models are brilliant and priced accordingly; using one to classify an email is like sending your senior partner to sort the mail.

Tier	Model class	Cost profile	Best for
Tier 1	Small / local model	Near-zero per task	Classification, extraction, formatting, routing, yes/no checks
Tier 2	Mid-tier frontier model	Low	Drafting, summarization, standard client work, most skills
Tier 3	Top frontier model	Premium	Strategy, complex reasoning, novel problems, high-stakes final review

Route every task to the **cheapest model that can pass the quality bar**, and escalate only on failure. Over time, your logs tell you exactly which skills need which tier - and your average cost per task falls month over month.

THE ECONOMICS

The majority of task volume (lead classification, data extraction, formatting) runs at Tier 1 for pennies, while the minority of task value (strategy, client-facing deliverables) runs at Tier 3. That inversion - high volume cheap, high value premium - is the entire economics of running AI profitably.

The Emerging Landscape

Frontier Models, Local LLMs, and Owning Your Bot

The two worlds

Frontier models (major AI providers - Claude, GPT, Gemini) are the most capable AI on earth, accessed via API, priced per token. They improve every few months without you lifting a finger.

Local / open-weight LLMs are models you download and run on hardware **you control** - your own cloud servers or even on-premise. No per-token fee to a vendor. Your data never leaves your environment.

Two years ago, local models were toys. Today, the best open-weight models rival last year's frontier models - and for Tier 1 and much of Tier 2 work in the waterfall, they are already good enough, which is the only bar that matters.

The ownership model: your bot, your cloud

The client owns the bot, deployed inside their own cloud account. Not ours. Theirs.

Why it matters	What it means for you
Data sovereignty	Your customer data, call recordings, and CRM never live on a vendor's infrastructure. Decisive for regulated or trust-sensitive industries.
No vendor lock-in	If any relationship ends, the system keeps running. You hold the keys.
Cost architecture	Local LLMs run at flat infrastructure cost alongside frontier API access. Tier 1 traffic stops hitting a per-token meter.
Compounding asset	Memory files, skill library, and logs accumulate in your environment. Every month the system gets smarter and your knowledge moat gets deeper.

DIRECTION OF TRAVEL

AI is becoming infrastructure you own, not a subscription you rent. The businesses setting this up now will spend the next decade paying a fraction of what their competitors pay per task - while owning every byte of the intelligence they generate.

Dispatch in Practice

Real Prompts From Our Operation

"Dispatch" is what we call our production AI operations system - the orchestrated, memory-equipped, skill-documented setup described in this guide, running live on our own business. Below are real prompt patterns pulled from daily use. Study the structure, then adapt the content to your business.

Prompt pattern 1: The Intelligence Brief

```
Run the Intelligence Brief skill for [Client], [website]. Industry: commercial cleaning. Pull their service area, target customer, differentiators, and existing testimonials from the site. Output the brief in our standard format and flag anything missing that I need to get from the client before campaign launch.
```

Why it works: Names the skill, provides the inputs, defines the output, and asks the AI to surface gaps instead of guessing.

Prompt pattern 2: The Report Card

```
Generate an SEO/SEO report card for [domain] using our D-through-A grading framework. Check schema markup, business identity consistency, review markup, AI-crawler accessibility, and answer-first content structure. Output as branded HTML and PDF. Do not inflate grades - a D that motivates action is worth more than a B that flatters.
```

Why it works: Explicit quality standard, explicit format, and an instruction that shapes judgment ('do not inflate grades').

Prompt pattern 3: The Pre-Launch Audit

```
Run the full 12-criteria site cutover readiness check on [staging URL] before we point DNS. Pass/fail each criterion with evidence. Anything failing blocks launch - list the exact fix required for each failure. Do not mark the site ready if any criterion fails.
```

Why it works: Binary pass/fail against a documented standard, evidence required, and a hard rule that removes wiggle room.

Prompt pattern 4: The Campaign Refresh

Client [X] has produced 0-1 leads in 60 days. Pull their current sequence, diagnose why it is underperforming against our standard brief criteria, then write a replacement 3-step sequence plus a follow-up. Keep the industry-specific angle. Stage it for review - do not load anything live until approved.

Why it works: Diagnosis before rewrite, references the documented standard, and enforces the human approval gate.

Prompt pattern 5: The Resumable Work Order

Before starting, write a task manifest for this job and checkpoint after each completed section. If the session fails at any point, output a resume prompt I can paste into a fresh session to continue exactly where we left off.

Why it works: Memory architecture applied to a single job. Crashes lose minutes instead of days.

Prompt pattern 6: Inbox Triage and Classification

Triage today's inbox. For each thread: classify it as Action-Required, FYI, Automated-Noise, or Needs-Reply. For Action-Required, draft the reply and flag the deadline. For Needs-Reply, write a one-sentence proposed response. Group by sender type: client, vendor, prospect, internal. Return a prioritized list with the highest-urgency item at the top.

Why it works: Structured classification taxonomy forces the AI into consistent buckets rather than vague summaries. The grouping step surfaces patterns a human would miss scanning one-by-one.

Prompt pattern 7: Lead Routing and Qualification

New inbound reply from [prospect name and company]. Their message: [paste message]. Using our qualification criteria - budget, authority, need, timeline - score this lead 1-5 and recommend a routing action: hot (immediate owner outreach), warm (nurture sequence), or cold (archive). Include a one-paragraph summary of why, and draft the first reply if hot or warm.

Why it works: Lead routing with explicit criteria produces defensible, consistent decisions. The draft reply removes friction from the next human action.

Prompt pattern 8: Data-Pull Plan Confirmation

We want to pull data for [Client] in [geographic territory]. Proposed method: [describe approach]. Target industries: [list]. Job titles: [list]. Record cap: [N]. Before I authorize this, flag any data-quality risks, confirm deduplication approach against their existing contacts, and estimate what percentage of the target universe we are likely to reach. Return a one-page pull plan I can approve or modify.

Why it works: Forces a pre-authorization review instead of diving straight into execution. Surfaces cost, quality, and overlap issues before the order runs.

Prompt pattern 9: Weekly Client Report

Generate the weekly client report for [Client]. Pull performance across all active channels for the period [date range]. For each channel, show: sends or impressions, response or open rate, leads generated, and a one-line interpretation. Highlight the top win and the one metric that needs attention. Format as a concise email I can send directly.

Why it works: Channel-by-channel structure ensures nothing is missed. The 'top win + one concern' format trains the AI to be editorial rather than just descriptive.

Prompt pattern 10: Testimonial Extraction

Review the testimonials and case studies in this document: [paste or attach]. Extract the three strongest proof points - real numbers, specific outcomes, named results. Then write a testimonial block (3-4 sentences) in the client's voice suitable for use in outbound copy. Keep it first-person and outcome-focused. Do not invent anything not in the source material.

Why it works: The 'do not invent' constraint is essential. Testimonial fabrication is both a compliance risk and a trust risk. Anchoring to source material keeps the output defensible.

Prompt pattern 11: Scope of Work Drafting

Draft a scope of work proposal for [prospect company] based on the intake notes: [paste notes]. Include: situation summary, proposed engagement, deliverables, timeline (phased if multi-month), and investment. Use our standard pricing tiers. Flag any scope gaps I need to resolve before sending. Stage the draft - do not send.

Why it works: The 'flag scope gaps' instruction is the key. It turns a passive drafting task into active quality control - the AI identifies what is missing, not just what was provided.

Prompt pattern 12: Memory Checkpoint

We are approaching the end of this working session. Before we close: write a checkpoint summary covering what was completed, what is still pending, any decisions made, and a ready-to-paste resume prompt for the next session. Save the summary as a dated log entry. This is mandatory before any session ends on an in-progress task.

Why it works: Memory architecture made into a standing work order. The 'mandatory' framing prevents the AI from skipping the step when a session feels rushed.

Prompt pattern 13: Model-Waterfall Routing Decision

This task requires a research-and-synthesis output. Route the initial research and extraction work to the fast and cheaper model tier, then hand the structured data to me for a quality check. If I approve, route the final synthesis and writing to the premium tier. Do not escalate to the premium tier until I have confirmed the research output is accurate.

Why it works: Model-waterfall routing applied explicitly to a multi-step task. Separating research from synthesis controls cost without sacrificing output quality on the step that matters.

THE PATTERN BEHIND THE PATTERNS

Every one of these prompts does five things: **names a documented skill, supplies the inputs, states the quality bar, defines the output format, and preserves the human approval gate.** That is the whole art. Master those five and you will outperform 95% of AI users regardless of which model you use.

How Skills Connect to Your Tools

A skill document is not just instructions for the AI - it is also a connector specification. Because skills are plain text, they can be wired to any tool your business already uses: your CRM, your sending platform, your cloud storage, your calendar. Here are examples of how this works in practice.

CRM INTEGRATION

New-Contact First-Touch

A skill monitors new contacts added to any CRM platform. When a contact is created, the skill pulls their industry and role, drafts a personalized first-touch email in the owner's voice, and stages it as a draft. The human reviews and sends. Zero templates. Real personalization at scale.

CLOUD STORAGE

Inbound Lead Filing

A skill watches a shared inbox folder. When a new lead email arrives, it extracts the key fields (name, company, request, source), creates a structured summary file in cloud storage, and sends a routing notification to the appropriate team member. Every lead is filed and traceable from day one.

EMAIL / SENDING PLATFORM

Campaign Performance Monitor

A skill pulls open rates, reply rates, and lead counts from your sending platform on a weekly cadence. It compares them to your documented benchmarks, flags any campaign below threshold, and drafts a diagnostic summary with a recommended next action. No spreadsheet-pulling required.

CALENDAR

Pre-Meeting Brief

A skill monitors your calendar for client meetings scheduled within 24 hours. It pulls the client record, recent activity, and open items, then produces a one-page meeting brief: context, talking points, and any decisions needed. You arrive prepared every time with zero manual prep.

CRM + CLOUD STORAGE

Weekly Account Health Report

A skill reads activity data from your CRM and recent files from cloud storage, then produces a graded health report for each active account. Accounts below threshold trigger a flag with a suggested intervention. The owner sees the full picture in a single digest each Monday morning.

ANY PLATFORM

The General Pattern

Any skill can be triggered by: a new record, a scheduled time, a webhook, or a manual command. It reads input from one system, applies a documented process, and writes output to another. The skill document is the contract between your AI and your tools - change the skill, change the behavior everywhere instantly.

Skill Library Examples

Here is what a mature skill library looks like. These are representative examples from our production operation, described generically so you can map them to your own business needs.

Skill name	What it does	Trigger
site-audit	12-criteria pre-launch website readiness check. Pass/fail per criterion with required fix for each failure. Blocks launch if any criterion fails.	Manual (before DNS cutover)
report-card	Multi-channel graded client scorecard. D-through-A grades across key performance areas. Branded HTML and PDF output. No grade inflation allowed.	Weekly or on-demand
lead-classify	Inbound reply classification and routing. Scores each reply on qualification criteria, assigns a routing action (hot/warm/cold), drafts the first reply for hot and warm leads.	On new inbound email
intelligence-brief	Structured client research brief. Pulls service area, differentiators, target customer, and testimonials. Surfaces gaps before campaign launch.	Manual (new client onboard)
brand-kit	Single source of truth for visual branding. Colors, fonts, logo rules, gradient specs, document templates, email signatures. Loaded before any client-facing deliverable.	Automatic (every deliverable)
email-triage	Inbox zero plus action-item extraction. Classifies every thread, surfaces Action-Required items, drafts replies for approval, archives automated noise. Runs on a schedule.	Scheduled (daily)

Your Next Steps

Work through these in order after the conference. Each step builds on the last.

WEEK 1 - FOUNDATIONS

1. Pick the three tasks in your business you repeat most often
2. Write your first request for each using the four-part anatomy (role, task, constraints, format)
3. Iterate until the output meets your standard - then save the winning prompt

WEEK 2 - DOCUMENTATION

1. Convert each winning prompt into a Markdown skill document (purpose, inputs, process, quality standard, output format, example)
2. Draft your OPERATIONS.md - one page: who you are, your voice, your non-negotiables, what always requires human approval

WEEK 3 - GUARDRAILS

1. Run the security guardrail checklist from Part 2 against everything you have built
2. Establish your approval gates: list every action AI may draft but never send without a human

WEEK 4 - ARCHITECTURE

1. Set up your memory structure (the file tree from Part 4) - even a simple folder is a massive upgrade over amnesia
2. Identify which of your skills could run on a cheaper model tier - begin your waterfall

READY TO GO FURTHER?

If you want the full ownership model - your bot, your cloud, your data, local models cutting your costs - that is exactly the system we build and deploy for clients. Find us during the conference or reach out afterward at **prospectrdigital.com**. Bring this guidebook; we will mark it up together.

Your Next Step Worksheet

Work through this live during the conference. Take it home and do not lose it.

Complete sections A, B, and C during the conference sessions. Section C is your personal approval-gate list - the foundation of your human-in-the-loop policy.

SECTION A

My 3 Most-Repeated Tasks

Write down the three tasks you do most often that feel repetitive. These are your first candidates for AI skills.

Task 1: _____

Task 2: _____

Task 3: _____

SECTION B

The First Skill I Will Document

Pick one task from Section A. Document it as a skill.

Skill title: _____

One-line purpose: _____

Process (write out the steps as if explaining to a new hire):

Quality standard - what does "done right" look like? _____

SECTION C

My Approval-Gate List

List every action your AI system may **draft** but must **never send or execute** without a human reviewing and approving first. Check the box when you have confirmed the gate is in place.

☐	
☐	
☐	
☐	
☐	

Common examples: all outbound emails to clients or prospects - any financial transaction or invoice - data deletes or record edits - public social media posts - any message sent on behalf of the owner.

Prospectr Digital - prospectrdigital.com - Every Channel. One Team. Engineered for Performance.