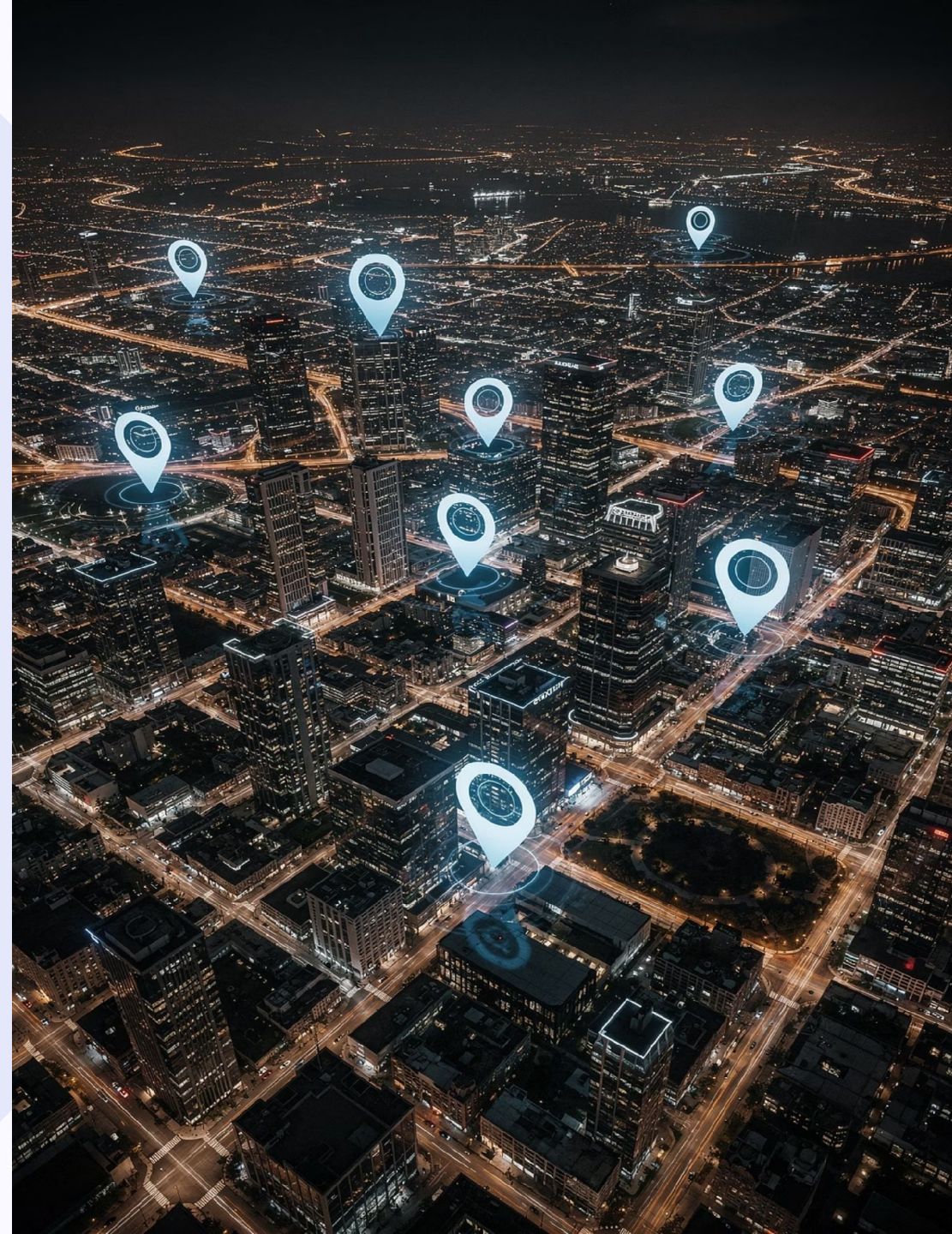


AI Workflows for Local SEO

What to Delegate, What to Own

Amanda Jordan · Senior SEO Specialist at Owner.com



Who Is Amanda Jordan

Current Role

Senior SEO Specialist at Owner.com — local SEO for restaurant brands at scale

AI Experience

Built AI-assisted workflows deployed across thousands of locations

Background

Former agency: multi-location franchise SEO, currently supporting 20,000 accounts

AI is everywhere in local SEO right now.

But most teams are delegating the wrong tasks — and getting burned by it.

What We're Covering Today

01

The Filter

Automation vs. AI vs. strategist
— getting this wrong wastes
money

02

Guardrails

Setting AI guardrails for
consistent, high-quality output

03

Competitive Analysis

Advanced competitive analysis
using AI

04

Action Plans

Creating GBP and location page action plans from
competitor gaps

05

Test Backlogs

Building performance-based test backlogs with
clear success criteria

01 · THE FILTER

Automation, AI, or Strategist?

Getting this wrong wastes money and produces bad SEO.

The Question Isn't "Human or AI?"

It's: automation, AI, or strategist?

Each has a different cost, capability, and appropriate use case. Mixing them up is expensive — in tokens, in time, and in SEO quality.

OPTION 1

Pure Automation

Logic-based rules that execute without thinking — **no AI tokens required.**

Use When

- Task follows clear if/then logic
- No interpretation needed
- Same rule applies every time

Examples

- Flag duplicate listings by name/address/phone
- Auto-route tickets by keyword category
- Alert when GBP field is blank

OPTION 1

Why Pure Automation Matters



Costs \$0 in Tokens

Logic tasks never need an LLM



Faster & More Reliable

No hallucination risk on rule-based tasks



Built-in Guardrails

Logic itself is the quality control



Scales Without Drift

Quality doesn't degrade at volume

OPTION 2

AI/Agent

Pattern recognition and generation from structured, specific inputs — **with guardrails.**

Use When

- Task requires pattern recognition across data
- Output can be generated from structured input
- Errors are detectable and non-catastrophic

Examples

- GBP post content from menu/category data
- Review response drafts from review text
- Competitor category pattern analysis

AI/Agent Requirements

Single-Focus Prompt or Chaining

Clear constraints — one job per prompt, not everything at once

Structured Data Provided

Not generic context — the specific data for this specific task

Validation Rules

Output guardrails built into the prompt design

Spot-Checking or AI Feedback Loop

Regular review to catch quality drift before it compounds

OPTION 3

Strategist

Anything requiring your SEO philosophy, cross-dataset analysis, and contextual judgment.

Use When

- Requires holding context across many factors
- Outcome depends on business-specific knowledge
- A wrong call has significant SEO consequences

Examples

- 6-month strategy from performance + competitors
- Primary GBP category for a specific business
- Deciding when to build vs. consolidate location pages

Why AI Can't Replace the Strategist

→ Context Loss

Can't hold all context without losing guardrails

→ Assumption Filling

Makes assumptions when context is incomplete and is very confident while being wrong

→ No SEO Philosophy

Doesn't understand your approach or priorities

→ Silent Errors

Judgment errors compound quietly — no error thrown

The Three Options at a Glance

Factor	Pure Automation	AI/Agent	Strategist
Thinking required?	None	Some	Yes — essential
Token cost	\$0	Per task	Your time
Error risk	Low — logic-based	Medium — hallucination	Low — you're the expert
Scales automatically?	Yes	Yes, with guardrails	No

Token Costs Are Real

AI Tokens Are Not Free

Companies are already seeing real cumulative costs — many from tasks that didn't need AI at all.

Logic Tasks Cost \$0

Duplicate detection, field validation, ticket routing — if it's IF/THEN logic, skip the LLM entirely.

Costs Multiply at Scale

One prompt at 4,000 tokens feels cheap. That same workflow across thousands of locations adds up fast.

- ❏ Automation first, AI second. Before sending any task to AI, ask: could I resolve this with a case statement?

The Teachability Test

"Could you explain this task to someone who has never done SEO, and have them do it correctly?"

YES → Automation or AI

The task has clear rules or patterns that don't require your accumulated SEO judgment to get right.

NO → Strategist Territory

It requires your SEO philosophy. The context that makes it work lives in your head — not in a prompt.

The benchmark isn't "good enough."

It's your quality.

If you cannot replicate the quality you would produce at least **80% of the time**, the tool is not ready to use.

What Happens When AI Gets Judgment Calls

Judgment errors are invisible. They don't throw an error — they quietly underperform for months.

Wrong GBP Category

AI picks the most common option, not the most accurate one. Category accuracy directly affects discoverability.

Wrong Geo Targeting

Target market selection requires knowing the actual service radius — no prompt can capture that from outside.

Poor Content

AI defaults to what it sees most often. Content built with no real guardrails or guidance is generic slop.

02 · GUARDRAILS

Guardrails

AI executes what you tell it to. The quality of your outputs is entirely a function of your constraints.



Guardrails Are Quality Standards

A guardrail tells AI: here is what I expect, here is what I will reject, here is when to stop and ask.

1

Input Guardrails

What context must exist
before AI touches this task

2

Logic Guardrails

IF/THEN rules that catch
predictable failure modes —
no tokens needed

3

Validation Thresholds

When does AI proceed vs.
flag for human review

Data Is Your Most Powerful Guardrail

"Use only the data I provide. Do not draw on outside knowledge."

Helps Prevent Hallucination

Reducing the amount of outside knowledge the model needs makes hallucinations easier to prevent.

Makes Errors Visible

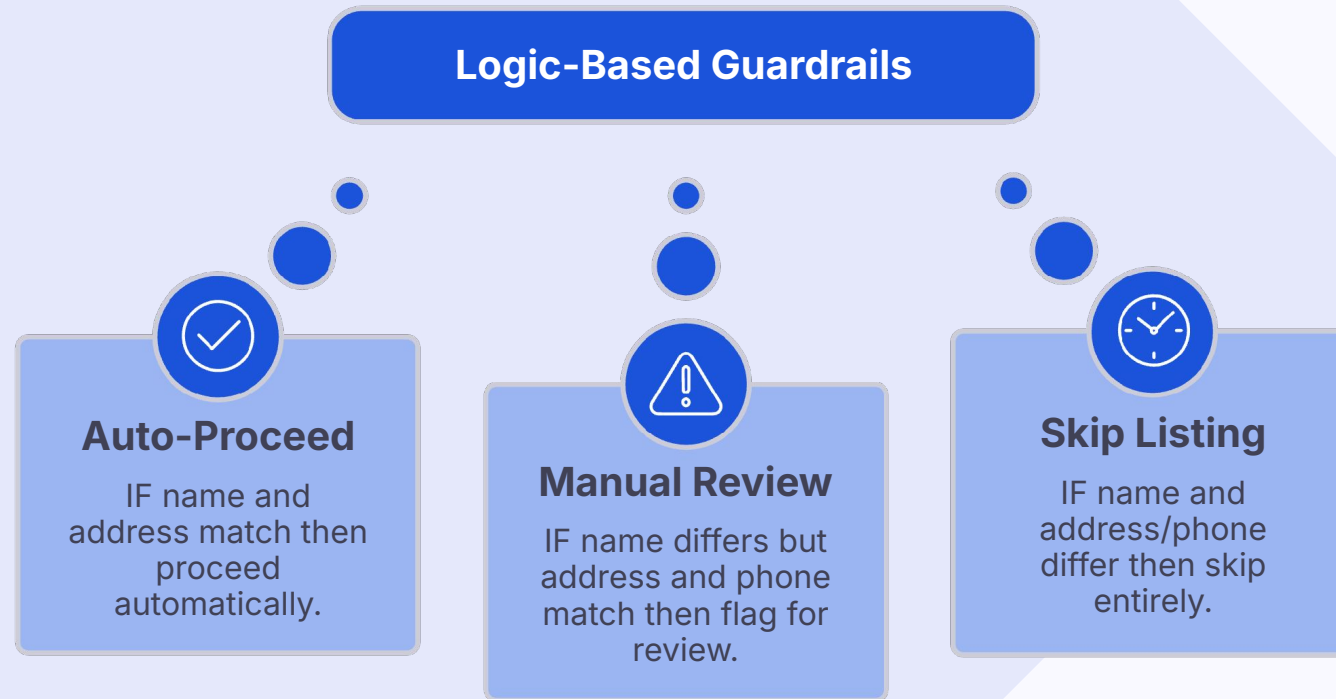
If the data doesn't contain an answer, the model must flag it — not invent one.

Applies to Every Task

Categories: give it the menu. Geo: give it the address. Content: give it the competitor page text.

Logic Guardrails: No AI Required

Some of your best guardrails are just IF/THEN logic.



Zero AI tokens. Zero hallucination risk. Built entirely with logic.

Validation Rules: Teaching AI When to Ask

Content Generation

Flag if: Output repeats the same word or phrase several times — a hallucination and low context signal

Category Recommendation

Flag if: Recommended category doesn't appear in the menu or service list provided

Service Area Targeting

Flag if: Recommended city is more than 30 miles from business address

Review Response

Flag if: Response mentions a product or detail not found in the review or GBP data

Know Your Edge Cases Before You Build

Every task has predictable failure modes. List them before you deploy — then build the catch for each one.

- 1 What are the 3 most likely ways this output could be wrong?**
- 2 What would a wrong output look like — and would I catch it?**
- 3 What data could be missing that the model might fill with a guess?**
- 4 What's the worst-case outcome if this runs incorrectly for 30 days?**
- 5 Have I built a catch for each of those failure modes?**

Build a Review Gate

Define which outputs go through a human — and which ones don't.

Always Review

- Hard-to-reverse changes (category, geo, homepage)
- Anything touching core local signals
- Workflows not yet run at scale

Spot-Check Only

- Content following a tested template
- Tasks with outputs that may require a validation check.
- Workflows passing QA for 30+ days

Auto-Proceed

- Pure logic tasks (no AI involved)
- High-confidence outputs from tested workflows
- Tasks where errors self-correct on next run

The Danger of Not Checking: False Confidence

LLMs Are Not Consistent

They're generative — they can give a different response for the same prompt every time. What passed QA in week 1 may not hold in week 8.

Quality Drift Is Gradual

It doesn't break all at once. It degrades output by output — and if no one is checking, no one sees it until customers do.



Set a review cadence. Weekly for new workflows. Monthly minimum for established ones.

You're Responsible for the Tools You Build

If the tool is inaccurate and you know it, you shouldn't be using it. That standard applies to the tools you create as much as the ones you buy.

03 · COMPETITIVE ANALYSIS

Competitive Analysis

AI can find patterns across hundreds of competitors in minutes. But only if you give it one job at a time.



One Agent. One Job.

The more context you load into a single conversation, the more the important instructions get buried. Specialists outperform generalists — in humans and in AI.

What Fails

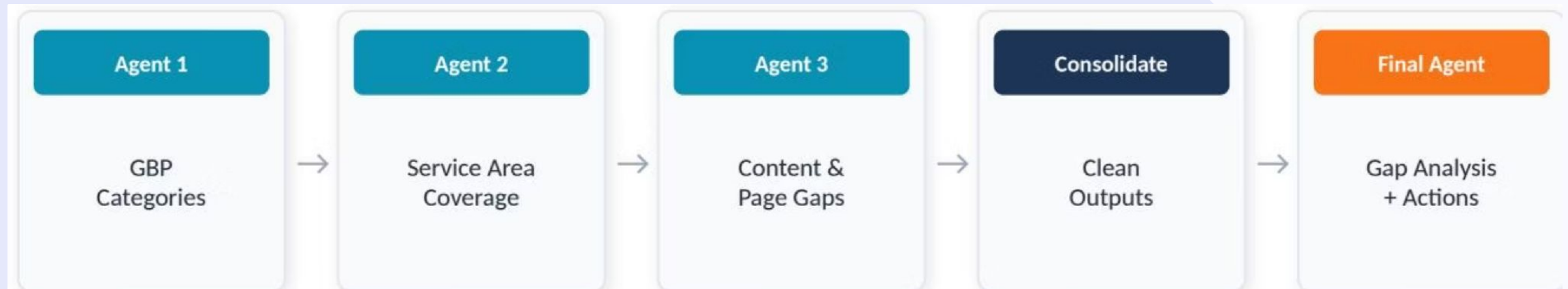
Give one agent all data — GBP, website, rankings, reviews — and ask it to find everything wrong with your SEO. Result: vague advice, missed patterns, hallucinated insights.

What Works

Agent 1: GBP categories only. Agent 2: service area gaps only. Agent 3: content gaps only. Final agent: consolidate clean outputs.

The Chain Approach

Chunk → Consolidate → Final Pass



The key: only pass **clean, consolidated data** to the final agent — not the raw inputs. The final agent performs better because it has less to process and more to work with.

GBP Competitive Data: What to Collect

Pull this data first, then give it to AI. Garbage in, garbage out.

Location & Proximity

- Business name, address
- Distance from your location & city center
- City population

GBP Signals

- Primary + all additional categories
- Features enabled (menu, booking)
- Mobile local panel visibility

Review Signals

- Rating, total count, recency
- Keywords in reviews

Content & Geo

- Menu or service highlights
- Nearby area rankings
- Page linked from GBP profile

Location Page Competitive Data

For every competitor page, pull the same data points. Consistent input → comparable output.

Page Type

Homepage, dedicated location page, or service page? Signals how they're thinking about local intent.

Internal Link Structure

How many pages link TO this page? Internal equity signals how much Google should trust it.

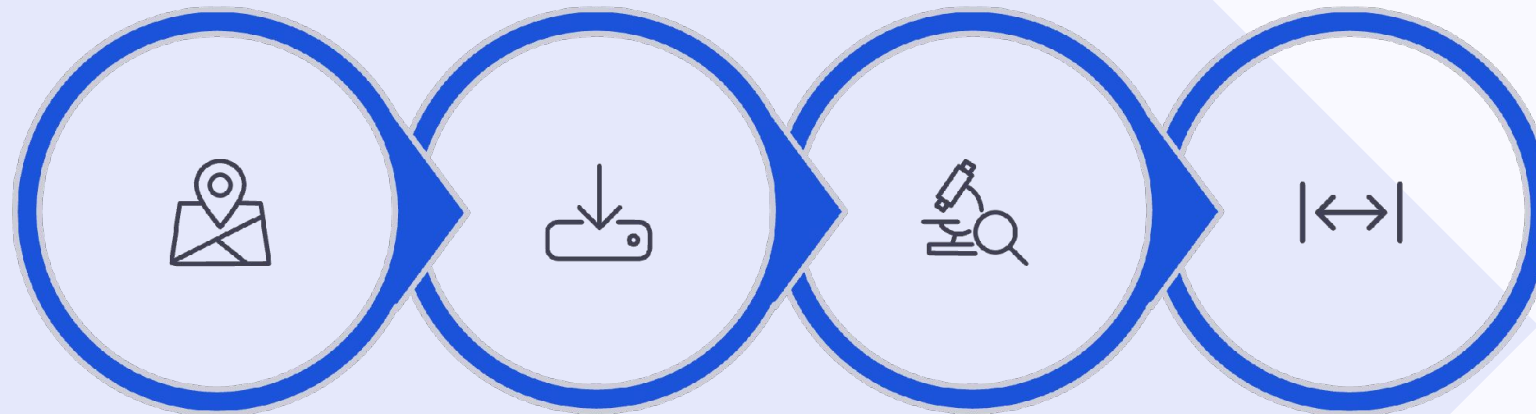
Content Structure & Quality

What topics appear consistently across top performers that yours doesn't have? Is content location-specific or generic?

Historical Organic Performance

High-performing organic pages reinforce local signals — organic and GBP visibility are connected.

The Competitive Analysis Process



Define
Dataset

Pull Data

Run Patterns

Identify Gaps

Who actually wins the local pack for your target queries in your target geo? Not your brand competitors — your **map pack rivals**.

Prompting for Competitive Analysis

"You are analyzing local SEO signals for [business type] in [city]. Use ONLY the data provided below. Do not draw on outside knowledge. Identify: (1) the most common primary category, (2) secondary categories in 3+ competitors, (3) any category patterns in the top 3 that differ from others. Return a structured table. Flag any business with a significantly different pattern."

- Single job only — categories, not everything at once
- Data-as-guardrail instruction built in
- Structured output requested with outlier flagging

What Patterns Reveal



Category Consensus

Reveals how Google and top competitors classify the market. Consensus can surface patterns worth investigating, but every category recommendation must still accurately describe the business.



Content Presence

Topics appearing consistently across top-ranking pages that yours are missing. These are the content gaps that matter most.



Geo Coverage Gaps

Cities top performers target that you're invisible in. Addressable with dedicated pages.



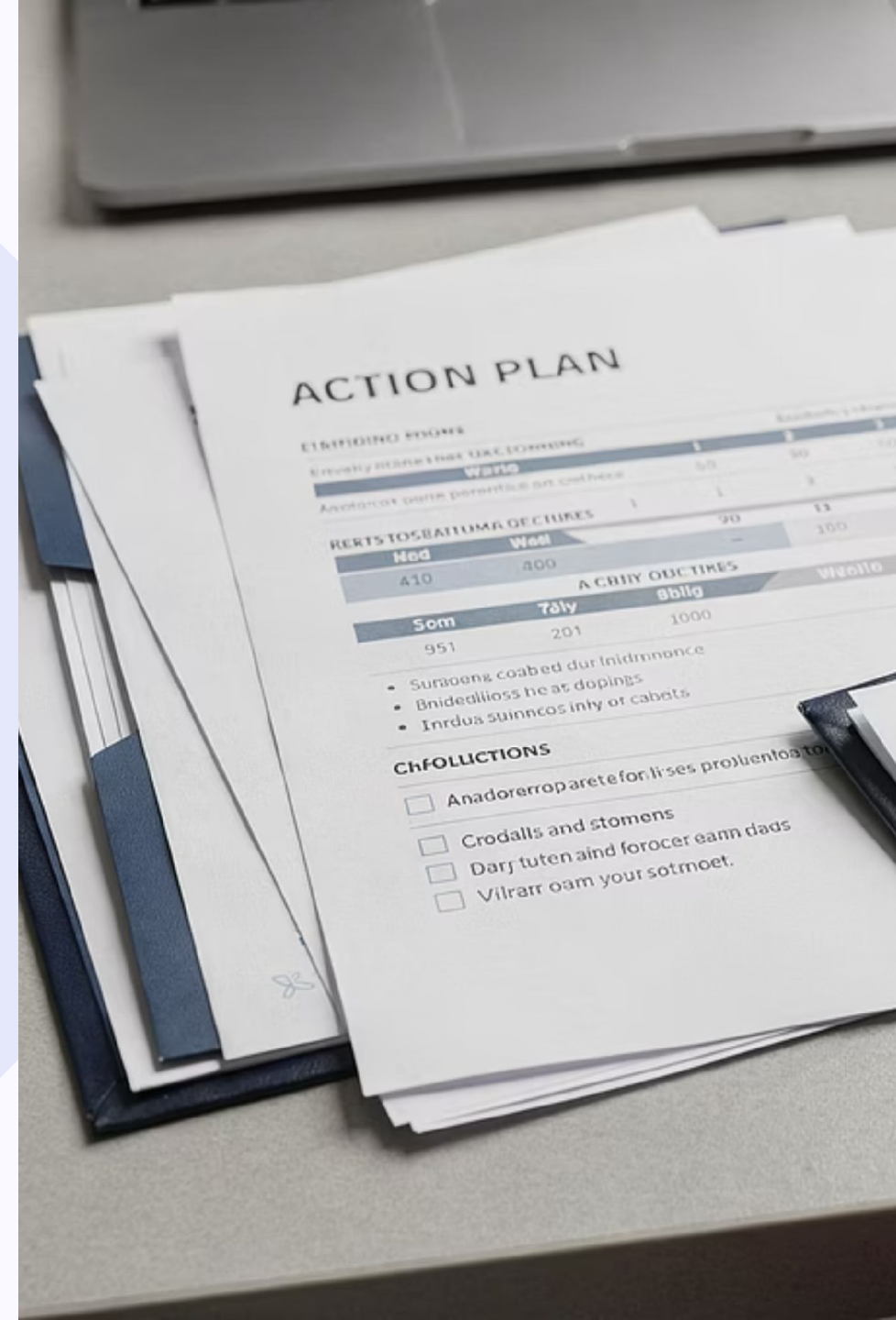
Review Velocity

Not just volume — recency. Recent review activity signals active customers to Google and to searchers.

04 · ACTION PLANS

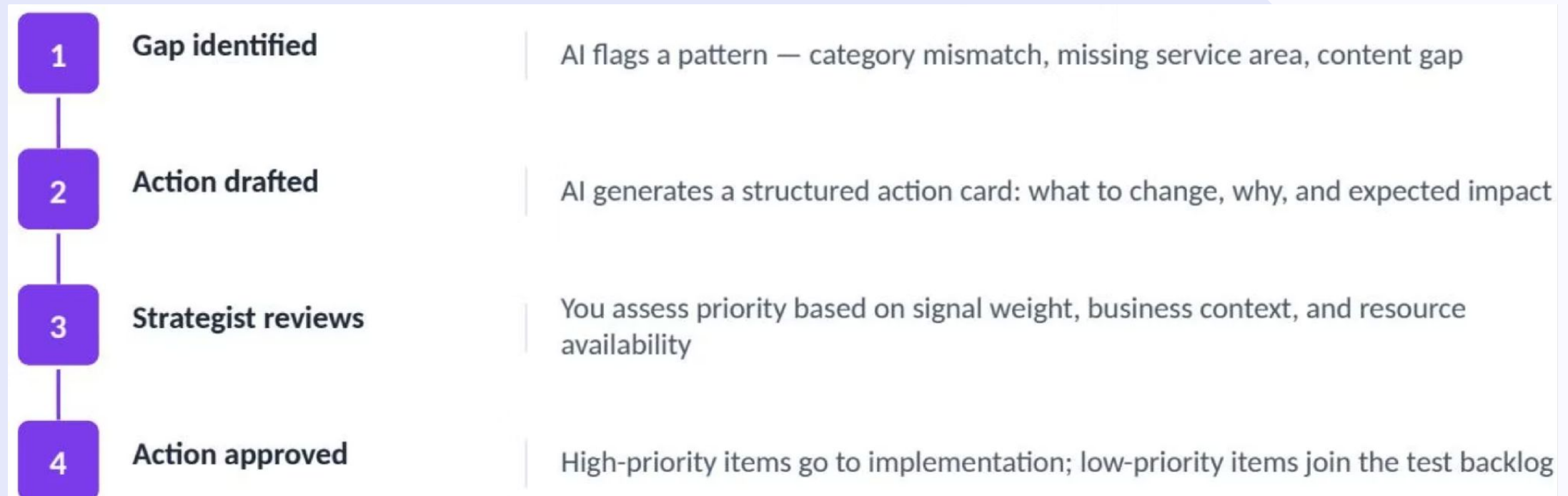
Action Plans

A gap with no action plan is just a list of problems.



From Gap to Action

AI is useful for turning structured gap data into draft action cards. The strategist reviews priorities and approves changes.



GBP Action Plan Template

Signal	Our Current Status	Competitor Benchmark	Recommended Action	Priority
Business Hours	Inaccurate listings found	Consistent hours across major listings	Update hours on inaccurate listings	HIGH
Primary Category	Sandwich Shop	Restaurant	Validate the category	HIGH
Service Area	Home city only	3 nearby cities	Add target cities	MED
Photos	5 total	Avg. 30+	Add 25+ quality photos	MED

Location Page Action Plan

Page Intent

- What topic is the page currently built to win?
- Is that the highest-value local query for this business?
- Does the page answer it better than competitors?

Secondary Market Pages

- Does this location serve cities beyond its home city?
- Does each secondary market have dedicated content?
- Are those pages indexed and internally linked?

Content Gaps vs. Competitors

- What topics appear on top competitors' pages but not yours?
- Do your pages use the same language customers use?
- Are CTAs aligned to discovery (directions, call)?
- Do competitors have features your site doesn't?

Not All Gaps Are Equal

Prioritize by the intersection of three things:



Signal Weight

How much does this factor influence local ranking?
Category and geo targeting are not equal to photo count.



Competitive Gap Size

A small gap in a high-signal area matters more than a large gap in a low-signal one.



Implementation Ease

A high-impact change you can make today beats a theoretical improvement requiring a full site rebuild.

Hold AI Output to Your Own Standard

If you would edit it significantly before sending it to a client, it's not ready. Full stop.

GBP Category Selection

Does this match what's actually on their menu and data — not just what sounds right?

Business Description

Is it accurate, local, and specific — or generic language that could apply to any competitor?

Location Page Content

Would I publish this without significant editing?
Does it reflect the business's actual voice?

Review Response

Does this sound like it came from this specific business — or is it clearly a template?

05 · TEST BACKLOGS

Test Backlogs

Use what you've learned to make future iterations better.

Why Most Test Backlogs Fail Before They Start

No Success Criteria

You can't tell if the test worked — so everything looks like it worked

Wrong Metric

Watching traffic when you should be watching GBP actions — or vice versa

No Baseline

You don't know where you started, so you can't measure how far you've moved

Window Too Short

Local SEO changes take time to propagate. 2 weeks is not enough.

No Rollback Plan

If the test fails, you don't know what to revert — or when to admit it failed

If your AI tool cannot replicate your quality at least 80% of the time, it is not ready to use.

The remaining 20% should be either non-detrimental to SEO — or flagged for your manual review. That's not a workaround. That's the design.

If You're Not Hitting 80%: Diagnose It

Asking Too Much in One Prompt

Fix: Split the task. One agent, one job. Chain the outputs.

Not Enough Context or Data

Fix: Provide specific data for this task. Restrict the model to only what you give it.

Prompt Is Unclear or Open-Ended

Fix: Give your current prompt + a bad output + "here is what's wrong" to another LLM. Ask it to suggest improvements.

No Scoring Rubric

Fix: Define what good looks like before you ask for it. Have the model evaluate its own output against it.

The 20% that AI can't handle is not failure.

It's the feedback loop.

The cases AI flags — the ones you resolve yourself — are your R&D. What you used to fix them informs the next version of the workflow.

Anatomy of a Good Local SEO Test

What

The specific change — not "improve SEO" but "add a Fort Lauderdale page"

Why

The hypothesis: if we do X, we expect Y because Z

Success Metric

The exact metric that confirms success — defined BEFORE the change is made

Measurement Window

30 days? 60? 90? Define it upfront and don't move the goalpost

Baseline + Rollback

Document current state before any change. Define the reversal path if the test underperforms.

Success Criteria: What to Actually Measure

DO Measure

- GBP direction requests (30/60/90 day windows)
- GBP actions by location
- GBP Impressions for target city queries
- Location page traffic from organic
- Conversions originating from the location page

DON'T Rely on Alone

- Overall organic traffic (too many variables)
- Keyword ranking position
- Aggregate metrics that blend all locations
- "Visibility" metrics without GBP action correlation

Building Your Test Backlog

Every item must answer all of these before it moves to "in progress":

Change	Hypothesis	Success Metric	Window	Priority
Update primary category	Better query match	GBP direction requests +15%	60 days	P1
Add 3 location pages	More discovery impressions	Impressions in target cities +20%	90 days	P1
Add 25 location photos	Below category average	GBP photo views +30%	30 days	P2
Add FAQ to location page	Content gap vs. competitors	Location page sessions +10%	60 days	P2

Your 20% Is Your R&D Budget

The pattern of what keeps getting flagged tells you exactly what to improve next.

01

Track What Gets Flagged

Log every case your workflow flags for manual review. Look for patterns — not one-offs.

03

Identify the Edge Case

What specifically made this a case the AI couldn't handle? Missing data? Logic gap? Prompt gap?

02

Resolve Them Yourself

Make the call. Document what you did and why. That resolution is your training data.

04

Build It Into the Next Version

Add a guardrail, expand the data input, update the prompt. Test again. The 20% floor should drop.

When Quality Is Bad: How to Improve Your Prompt

1

Be Specific About What's Wrong

Don't say "the output is bad." Say: it repeated "cozy atmosphere" 8 times, it recommended a category not on their menu.

2

Give It to Another LLM to Fix

Paste your prompt + bad output + explanation. Ask: "How should I change this prompt to prevent this?"

3

Build a Scoring Rubric

Define what a good output looks like as a numbered checklist. Have the model score its own output against it.

4

Test at Scale Before Deploying

Run the improved prompt on 20+ test cases before putting it back into production. Look for edge cases.

The Full AI-Assisted Workflow



The Rule

Automate execution.

Own the judgment.

AI handles the scale. You handle the decisions that determine whether that scale helps or hurts.

Five Things to Take With You

TAKEAWAY 1

Use the Right Level of Expertise for the Job

Automation, AI, or strategist. Use the right one for the task.

Each has a different cost, capability, and appropriate use case. Mixing them up is expensive — in tokens, in time, and in SEO quality.

TAKEAWAY 2

Data Is Your Most Powerful Guardrail

Restrict the model to only what you provide.

"Use only the data I provide. Do not draw on outside knowledge." That one instruction eliminates an entire category of errors.

TAKEAWAY 3

One Agent, One Job

Chain the outputs. Never load everything into one conversation.

Specialists outperform generalists — in humans and in AI. Break analysis into focused pieces. Each agent handles one dimension, then consolidate.

TAKEAWAY 4

80% Is the Floor, Not the Goal

If you're not there yet, diagnose - don't deploy.

If your AI tool cannot replicate your quality at least 80% of the time, it is not ready to use. The remaining 20% should be flagged for your manual review.

TAKEAWAY 5

You're the One Being Trained

Get better at prompting and the tools get better too.

The cases AI flags — the ones you resolve yourself — are your R&D. What you used to fix them informs the next version of the workflow. The 20% floor should drop over time.

Scalable Workflow In Action

How strategy, AI, and automation work as one system

One System. Three Stages.

1

01 · Strategist Creates

Sets direction. Decides what requires judgment and what gets delegated.

2

02 · AI + Automation Execute

Handles repeatable, pattern-based, logic-driven tasks at scale.

3

03 · Sharpen Your Tools

Edge cases return for manual resolution. Learnings sharpen the tools.

Use learnings from the 20% to sharpen prompts, add logic rules, and raise the bar before the next strategy cycle.

It Starts with the Strategist

Before you can delegate anything to AI or automation, someone has to decide **what** to delegate. That requires thinking. That's the strategist.

In Practice: 6-Month SEO Strategy

What Goes In

- Current performance data (rankings, GBP actions, traffic)
- Competitor GBP + location page analysis
- Recent Google algorithm behavior in the vertical
- Business-specific context (priorities, constraints, goals)

What Comes Out

- Prioritized 6-month roadmap by SEO impact
- GBP + location page action priorities per business
- Decisions on what to build, consolidate, or leave alone
- Clear rationale tied to performance and competition

❏ Holding performance + competitor + algorithm context simultaneously — without losing your philosophy — is the strategist's job. AI can't own this.

The Strategy Decides What Gets Delegated

Stays With You

- Primary GBP category selection
- Homepage intent and page structure
- Which nearby areas to target
- Reading cross-market patterns
- Setting the 6-month direction

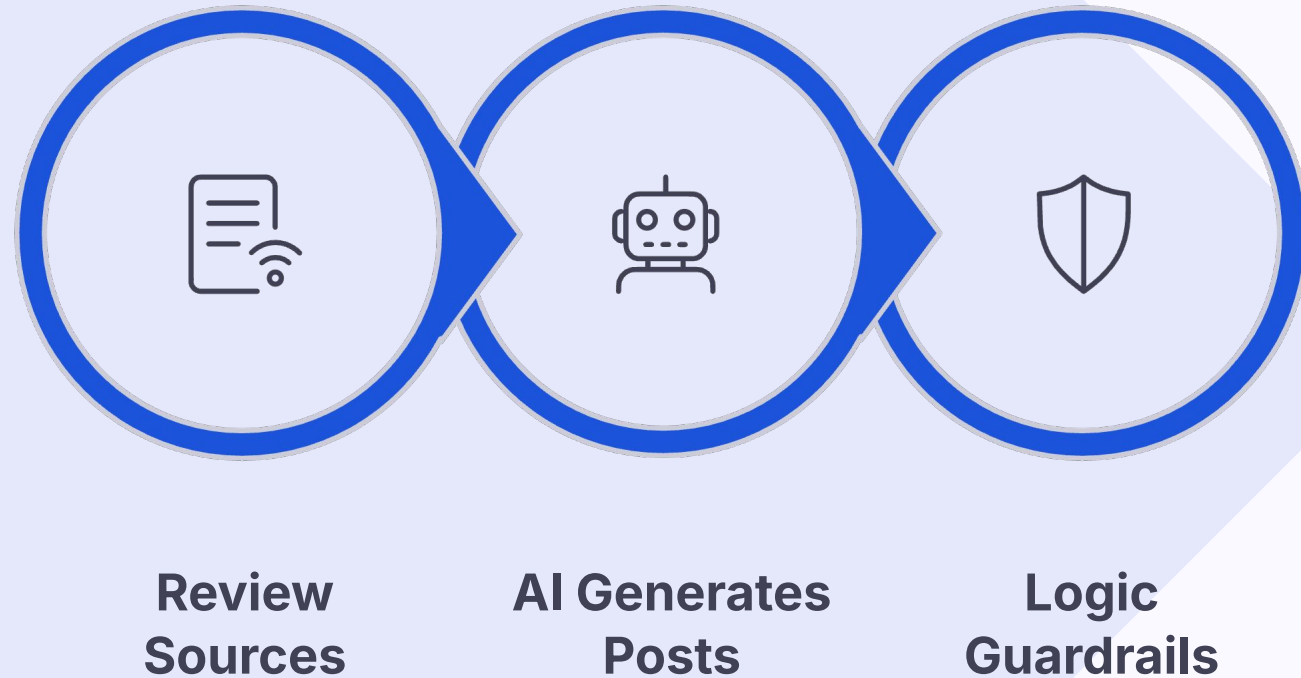
Gets Delegated

- GBP post creation from social + specials data
- Review response drafts
- Listing accuracy checking across directories
- Duplicate detection by name + address + phone
- Support ticket routing by keyword

AI and Automation: Execute at Scale

Not because you're offloading responsibility — **because the tools run at your quality standard.**

In Practice: GBP Post Creation Workflow



Sources feed AI generation, which is filtered by guardrails before any post goes live.

GBP Post Workflow: System Prompt + Guardrails

"You are a local SEO specialist responsible for finding source content for GBP posts and publishing them to an account. Please review sources to create GBP posts based on brand guidelines and previous posts. Prioritize deals/specials and include end dates."

Guardrail: Stale Promos

A separate call checks each source item. If it contains a promotion older than X days, it is excluded before generation.

Guardrail: Photo Match

The photo used must be the same image that accompanied the content on social media or the restaurant website.

In Practice: Listing Accuracy Checker

No AI needed — pure logic. Faster, cheaper, and more reliable for this task.

Check All Fields

Name, Address, Phone

All Fields Match Exactly?

YES → ACCURATE ✓

NO → Continue

Name Differs Only by Stop Words?

e.g., "the", "&", "LLC"

YES → ACCURATE ✓

NO → Continue

Address Differs Only by Abbreviation?

e.g., Rd/Road, Ste/Suite, #

YES → ACCURATE ✓

NO → INACCURATE × — Label the wrong field

Sample Output

Name	Status	Issue
Joe's Pizza	Accurate	—
Joe Pizza Co.	Accurate	—
Joe's Bistro	Inaccurate	Name
Joe's Pizza	Inaccurate	Phone

The 20% is not the failure rate.

It's the signal.

The cases AI and automation can't confidently handle tell you exactly where to improve.

What You Actually Do with the 20%

01

Resolve It Manually

You're the specialist. The 20% is exactly the kind of case that requires your judgment.

03

Fix the Tool

Improve the prompt, add a logic rule, or mark this task type as strategist territory. Document the fix.

02

Identify Why It Failed

Was the prompt missing context? Did the logic rule miss this edge case? Was it a strategist task all along?

04

Update Before the Next Cycle

Whatever fixed the 20% gets baked in before the next strategy cycle runs. Tools compound in quality.

Build Systems That Learn from You

Not just systems that work for you.

That's what separates AI tools from AI workflows. The goal: spend more time on SEO that requires deeper thinking and experimentation.

AI Workflows for Local SEO

Automate execution

Rules-based ops at zero
token cost

Own the judgment

Strategy, context, and key
decisions stay with you

Build the muscle

Better prompting = better
tools at your quality level

Thanks!