

THE SCREENSHOT

*Why AI Recommends Your Competitors, and How to Fix
It*

Jason Colapietro

Johnny Suede Press • 2026

THE SCREENSHOT

Why AI Recommends Your Competitors, and How to Fix It

Jason Colapietro

ChatGPT is recommending your competitors. Here's the screenshot.

Epigraph

*“Traditional SEO gets you ranked. AI SEO gets you cited.” —
from the Suede AI SEO methodology*

About the Author

Jason Colapietro is the founder of Suede Labs, a music-technology company built around a simple conviction: creators should own their work, their rights, and their visibility. Along the way he built something adjacent and urgent. He started running founders' categories through ChatGPT, Perplexity, and Gemini, capturing timestamped screenshots of who the machines recommend, and shipping the repairs the same day. That service is Suede Scan. This book is the thinking behind it, written down so you can run the diagnosis yourself.

He is a four-time published author. This is the fifth, and the only one written because his DMs kept proving the same uncomfortable fact: most founders have never once looked at what AI says when a buyer asks about their category.

How to Read This Book

Part I makes the case that AI visibility is now a survival question, not a marketing tactic. Part II shows you how to diagnose and repair your own visibility, using the same checks I run for paying clients. Every chapter in Part I ends with a prompt you can paste into ChatGPT right now. Do not skip those. The book works because you see your own gap with your own eyes, not because I describe someone else's.

What This Book Does Not Claim

Read this before anything else, because the AI visibility space is full of people who will promise you the opposite.

Every AI answer in this book, and every answer you generate while reading it, is a point-in-time observation. AI engines change their answers between sessions, between accounts, and between days. Nobody can guarantee you a citation, a ranking, or a recommendation in any AI system, and anyone who guarantees one is selling you something they cannot control. What you can control are the inputs: whether the machines can read your site, whether your pages are structured so an engine can lift an answer from them, and whether the evidence on your pages is worth citing. This book shows you the exact answers, shows you how to ship the fixes, and shows you how to re-check and measure the delta. Evidence, shipped inputs, and deltas. Nothing else is honest.

PART I: THE STAKES

Chapter 1: The Screenshot

“The paid unit of value is the screenshot of ChatGPT, Perplexity, or Gemini recommending your competitor instead of you.” — from the Suede Scan product spec

There is a message I have sent to founders more times than I can count. It goes like this:

“Hey, I ran ‘best tool in your category for your exact customer’ through ChatGPT and Perplexity this morning. Two of your competitors are cited in both answers. Your product doesn’t appear in either. Screenshot attached.”

Nobody argues with the screenshot.

Founders argue with analytics dashboards. They argue with SEO reports full of scores and acronyms. They argue with consultants who talk about domain authority. Nobody argues with a picture of ChatGPT, asked a question their actual buyer actually asks, recommending someone else by name.

The screenshot works because it is not an opinion. It is not a projection, a trend line, or a maturity model. It is the exact answer a real buyer received when they asked a machine for a recommendation in your category. If your name is not in that answer, then for that buyer, in that moment, you did not exist. There was no second page to scroll to. There was no listing to skim past. There was an answer, the answer had names in it, and yours was not one of them.

The morning this became a business

I did not set out to write a book about AI visibility. I set out to figure out why smart founders with good products were quietly losing deals they never knew existed.

The pattern kept repeating. A founder builds something genuinely good. They do the responsible things: a clean site, some content, maybe an SEO retainer. Their traffic is fine. Their rankings are fine. And meanwhile, a growing share of their buyers have stopped searching in the way any of those numbers measure. Those buyers open ChatGPT, or Perplexity, or Google with AI Overviews turned on, and they ask a question. The machine answers with three to five names. The buyer evaluates those names. The deal happens inside that shortlist.

If you are not in the answer, you are not losing the deal. Losing implies you competed. You were never in the room.

So I started running the check for people. Ask the engines the questions their buyers ask, several different ways. Capture every prompt and every answer with timestamps. Count who gets named. Put the screenshots in front of the founder. Then, for the ones who wanted it, ship the actual repairs the same day: the structured data, the machine-readable files, the rewritten page sections that give an engine something worth citing.

That service became Suede Scan. This book exists because the diagnosis half of that work should not be a secret. You can run it yourself, today, for free, and Part II shows you exactly how. What I sell is speed and execution. What I am giving you here is the sight.

Why a screenshot and not a report

I want to be precise about why the screenshot is the unit of value, because the reasoning is the spine of this whole book.

An SEO report tells you about a system you already understand: search results, rankings, clicks. You have twenty years of intuition for what page one means. You have no intuition yet for what it means that an answer engine synthesizes one response and your company is structurally absent from it. The screenshot builds that intuition in five seconds. It converts an abstract risk into a specific, dated, undeniable event: this machine, asked this question, on this day, named these companies and not yours.

The second reason is harder to hear. Reports let you postpone. A score of 61 out of 100 feels like a project for next quarter. A screenshot of your top competitor being recommended to your exact customer, this morning, feels like a fire. It should. Chapter 5 is about what the delay actually costs, and the short version is that these answers harden over time. The competitor who owns the answer today is training thousands of buyers, and arguably the engines themselves, to treat them as the default.

One caveat, and I will repeat it throughout because the honest version of this field requires it: every screenshot is a point-in-time observation. Ask the same engine tomorrow and the answer may differ. That cuts both ways. It means a good answer today is not a permanent asset, and a bad answer today is not a permanent verdict. What persists are the inputs, and the inputs are what you control.

What this book will and will not do

By the end of Part I you will understand why this happened: where the buyers went, how answer engines choose who to name, and why absence compounds. By the end of Part II you will have run a real audit of your own AI visibility: whether the crawlers can read your site, how each major engine selects sources, whether your pages are extractable, whether your evidence is citable, and a checklist you can rerun monthly.

What this book will not do is promise you citations. Nobody can. The engines do not publish their selection logic, they change constantly, and anyone who guarantees you a spot in an AI answer is charging you for weather. What I can promise is narrower and more useful: you will know exactly what the machines say about your category right now, you will know which inputs are broken, and you will know how to fix every one of them.

That is the whole trade. Evidence, shipped inputs, and measured deltas. It is enough.

Check this yourself right now

Open ChatGPT and paste this, filled in for your business:

“What is the best [your category] for [your ideal customer]? Recommend 3 to 5 options and briefly explain each.”

Read the answer once as a founder. Then read it again as a buyer who has never heard of you. Screenshot it, note the date, and keep it. That screenshot is your baseline, and by Chapter 11 you will be measuring against it.

If your name is in the answer, good. You have something to protect, and most of Part II applies to you with equal force. If it is not, you now know precisely what this book is for.

Chapter 2: The Quiet Migration

“Free 10-second gut check first, no email required. It tells you whether AI crawlers can even read your site.” — from the Suede Scan launch campaign

Buying research used to have a shape everyone understood. A person with a problem typed words into a search box, received ten blue links, clicked a few, formed a shortlist, and bought. Every marketing discipline of the last two decades is a strategy for winning some stage of that shape. Rank higher. Write the comparison page. Earn the click. Convert the visit.

That shape is dissolving, and it is dissolving quietly, which is what makes it dangerous.

Where the buyers went

The migration has three parts, and each one removes a place where you used to be able to compete.

First, the question moved. A meaningful and growing share of buying research now starts as a conversation with an AI assistant rather than a keyword query. The buyer does not type “crm small business.” They ask, “I run a 12-person agency and our client tracking is a mess. What should we use?” That is a better question, and it gets a better answer: a synthesized recommendation with names, trade-offs, and reasoning. Notice what it does not produce. It does not produce a results page you can rank on.

Second, the answer moved. Even inside traditional search, the answer increasingly arrives before the links do. Google’s AI Overviews synthesize a response at the top of the page. The buyer

reads the synthesis. Some scroll past it. Many do not. The industry calls the broader pattern zero-click search: the query is asked, the answer is consumed, and no website receives a visit. The click you spent twenty years learning to win no longer exists for that query. I am deliberately not quoting adoption percentages at you. The widely circulated statistics in this space are mostly undated and unsourced, and this book does not traffic in unverifiable numbers. You do not need a statistic. You need Chapter 1's exercise: you watched a machine answer a buying question in your own category. That behavior is the migration, observed firsthand.

Third, the shortlist moved. This is the part founders underestimate. In the old shape, the buyer assembled their own shortlist from many sources, and a determined vendor had a dozen chances to get on it: rankings, ads, review sites, a colleague's mention. In the new shape, the engine assembles the shortlist and presents it as finished work. Three to five names, with reasons. Buyers still verify from there. But the frame is set. Everything the buyer does next happens inside a list a machine wrote in four seconds.

Why you did not notice

If this migration is real, why doesn't it show up as a crisis in your dashboards? Four reasons, and they compound.

Your analytics cannot see it. When a buyer asks ChatGPT about your category and you are not mentioned, nothing happens in any system you monitor. No impression, no lost click, no bounce. The event that mattered most, a qualified buyer receiving a shortlist without you on it, is invisible by definition. Your dashboards are not lying. They are measuring a road while traffic reroutes to one they cannot see.

Your SEO metrics still look healthy. Rankings degrade slowly, and ranking well and being cited by answer engines are related but different achievements, which is the entire subject of Chapter 4. A site can hold page one for years while being systematically absent from the AI answers layered on top of and beside those results.

The losses arrive as silence. A competitor winning a bidding war is loud. A buyer who never contacted you because a machine handed them three other names is perfectly silent. There is no lost-deal record, because there was no deal. Silence reads as normal. It is not normal. It is unmeasured.

Nobody in your company owns this. Ask your team who owns AI visibility and watch the pause. The SEO person owns rankings. Marketing owns campaigns. The question of what six different answer engines say when a buyer asks about your category sits in the gap between every job description you have written.

The gut check

This is why the first tool I ever shipped for this problem takes ten seconds and asks for nothing. At optimize.suedeai.ai you can check whether AI crawlers can even read your site. Not whether your content is brilliant. Whether the machines are able to fetch it at all.

I built the free check first because of what I kept finding: companies investing real money in content while their infrastructure silently turned AI crawlers away. A robots.txt rule written years ago for a different problem. A firewall or bot-protection layer that treats every non-Google crawler as an attacker. A site rendered so heavily in JavaScript that a text-first crawler fetches an empty shell. Each of these is invisible from inside the company. The site looks perfect in

a browser. The dashboards are green. And GPTBot, PerplexityBot, and ClaudeBot are bouncing off the front door. Chapter 6 covers how to run this diagnosis properly, bot by bot.

The migration is the reason that check matters. When buyers asked humans and search pages, being readable by machines was a technicality. Now that a layer of machines sits between you and a growing share of your buyers, being readable by machines is the new being open for business.

What this means before we go further

I want to close this chapter by defusing the two standard reactions, because both are wrong in instructive ways.

The first reaction is panic: rip up the marketing plan, chase every AI optimization trick on the internet. Wrong, because as Part II will show, much of what wins AI citations is disciplined fundamentals, done for machines as well as people, and most of the tricks are noise.

The second reaction is dismissal: this is a fad, buyers will always verify, search is not dead. Also wrong, and note that the dismissal attacks a claim I am not making. Search is not dead. Your site still matters, arguably more than ever, since it is what the engines read and cite. The claim is narrower and harder to dismiss: a new layer now sits between buyers and vendors, that layer writes shortlists, and you have almost certainly never audited what it writes about you.

The buyers moved. The dashboards stayed. This book is about closing that gap.

Check this yourself right now

Open Perplexity and ask the question a real buyer would ask at the start of research:

“I need [the problem your product solves] for [your customer type]. What are my options and how do I choose?”

Perplexity cites its sources with links. Look at the citations, not just the names. Which pages taught the machine its answer? Is any of them yours? Screenshot it, date it, save it next to the one from Chapter 1.

Chapter 3: Invisible Is the New Page Two

“Blocked silently, which is worse than an error.” — from the Suede Scan operations runbook

Every founder of the last twenty years learned one piece of search folklore: nobody looks at page two. It was the industry’s favorite dark joke, the punchline about where to hide a body. But page two, for all its deadness, had a redeeming feature that we only appreciate now that it is gone.

You could see it.

You could type your keyword, count the results above you, and know exactly where you stood. Position fourteen was bad news, but it was legible bad news. It came with a number, a trend line, and an entire industry of tools and practitioners for improving it. The system that judged you also showed you your score.

AI answers do not have a page two. They have named and not named. And the system that judges you shows you nothing unless you go ask it yourself.

The binary

When an answer engine responds to a buying question, it names a handful of companies. Perhaps three, perhaps six. Everyone else in the category is not ranked lower. They are absent from the document the buyer is reading. There is no scroll, no “see more results,” no position fourteen. The answer is the entire surface, and you are either on it or you are not.

This binary changes the economics of visibility in a way that page-based thinking cannot process. In ranked search, visibility degrades gradually: position three gets less than position one, page two gets scraps, but the curve is continuous and every improvement buys something. In answer engines the curve collapses into a step function. Inside the answer, you get evaluated. Outside it, you get nothing, and no amount of being almost included pays partially.

The nearest analogue is not SEO at all. It is retail distribution. Either your product is on the shelf when the customer walks the aisle, or it is not, and a product that is almost stocked sells exactly as well as one that does not exist. Answer engines are shelf space for recommendations, and most founders have never once walked the aisle.

Silent failure is the default

The line at the top of this chapter comes from my own operations notes, and the story behind it is worth telling because it is this whole problem in miniature.

While running client scans, I documented a failure mode with one of the major engines: logged out, its prompt box rendered normally, accepted typed text, and then, on submit, produced nothing. No answer, no error, no message. The box just went empty. My note from that day reads: blocked silently, which is worse than an error. Because an error tells you something failed. Silence lets you believe it worked.

Hold onto that sentence, because it describes almost every failure in AI visibility.

When your robots.txt blocks GPTBot, nothing warns you. Your site works in every browser. When a bot-protection layer serves AI crawlers a challenge page instead of your content, your uptime monitor stays green, because to a human visitor the site is up. When your pricing page renders beautifully for eyes but is a JavaScript shell to a text-first crawler, no tool you currently run will mention it. When an engine simply does not know your company exists, there is no notification, because notifications are sent by systems that know about you, and this one does not.

Compare this with how ranked search fails. Rankings fail loudly: traffic drops, a chart dips, Search Console emails you about coverage problems. An entire nervous system evolved to make search failures visible. AI visibility has no nervous system yet. It fails silently, by default, everywhere, and the silence is indistinguishable from health.

Nobody checks

Here is the flat, uncomfortable observation this chapter exists to deliver: most founders have never once checked what AI engines say about their category. Not monthly. Not ever. Once.

I know this because opening people's eyes to it briefly became my job. The reaction to a first screenshot is remarkably consistent: not disagreement but surprise that the question was askable. It had not occurred to them that "what does ChatGPT tell my buyers" was a thing one could go look at, this afternoon, for free.

The reason is not laziness. It is that no habit exists. Checking your Google ranking became a reflex because two decades of tools, reports, and agencies built the reflex. The equivalent reflex for

answer engines does not exist yet, and the engines themselves do not send report cards. The information is one prompt away and almost nobody asks for it.

Which means, for the moment, that checking at all is an edge. Your competitors, statistically, are as blind as you were before Chapter 1. The difference is that you are two screenshots into fixing it.

Legibility is the first repair

If the disease is silence, the first treatment is not optimization. It is instrumentation. Before you change a single page, you need the answers in front of you, dated, so that invisible becomes a measured state instead of an unexamined one.

That is what the exercises at the end of these chapters are quietly building: your first instrument panel. One question per engine, screenshots, dates. Part II will formalize this into a proper audit with a prompt set that covers your category from multiple angles, because a single question is a spot check, not a diagnosis. Chapter 11 will turn it into a monthly operating rhythm, because a point-in-time reading, as I will keep repeating, is only a point in time.

But the principle lands here: you cannot manage what fails silently. Page two was cruel and legible. This new layer is crueler precisely because it never sends the bad news. You have to go get it.

Check this yourself right now

This one takes two minutes and is the fastest instrumentation you will ever install.

First, ask ChatGPT directly: **“What do you know about [your company name]?”** The answer tells you whether you exist to the machine at all, and what it believes about you if you do. Founders

are routinely startled in both directions: total blanks for real companies, and confident descriptions that are years out of date.

Second, run the ten-second infrastructure check at

optimize.suedeai.ai. No email, no signup. It tells you whether AI crawlers can even read your site, which is the subject we take up properly in Chapter 6.

Screenshot both. Date them. Your instrument panel now has three readings on it.

Chapter 4: Ranked vs. Cited

“In AI search, a well-structured page can get cited even if it ranks on page 2 or 3.” — from the Suede AI SEO methodology

Everything you know about search visibility was built for a tournament. Pages compete for a keyword, an algorithm scores the field, and the winners are displayed in order. Twenty years of SEO is the study of winning that tournament.

Answer engines are not running a tournament. They are writing a document, on deadline, and looking for sources they can quote without embarrassment. Understanding that difference, ranking versus citation, is the single most useful mental upgrade in this book, because it explains both why your SEO success has not protected you and why the repairs in Part II look the way they do.

What an answer engine is actually doing

Strip the mystique away and the mechanics are almost mundane. When a buyer asks an assistant for a recommendation, the engine typically runs searches of its own, retrieves a set of candidate pages, reads them, and synthesizes an answer, naming companies and, on several platforms, citing the pages it drew from. Some of what it says also comes from its training data: the accumulated text of the public web as of some cutoff, which is why engines sometimes describe companies in confidently outdated terms.

Now put yourself in the engine’s position at the retrieval step. You have a few seconds, a stack of candidate pages, and a paragraph to write. Which sources do you actually use?

You use the page that answers the question directly, in a self-contained passage you can lift cleanly. You use the page whose claims come with evidence: names, dates, specifics. You skip the page that takes four hundred words of throat-clearing to approach its point, the page whose key facts live inside a JavaScript widget you never rendered, and the page you could not fetch at all.

That is citation logic. Notice what it is not: it is not a popularity ranking. Which produces the strangest and most hopeful fact in this field.

The decoupling

Ranking and citation correlate, but they are not the same contest, and the gap between them is where opportunity lives.

The line at the top of this chapter is from my own methodology notes, and it is the crux: a well-structured page can get cited even if it ranks on page two or three. The engines that search the live web select sources for quotability, not just for rank position. A page that states the answer plainly, in extractable blocks, with real evidence, can be quoted over a page that outranks it but buries its substance. Meanwhile the reverse also happens constantly: a site that dominates page one while being structurally useless to quote, and therefore absent from the answers written on top of its own rankings.

The decoupling is uneven across platforms, and the differences matter operationally. Google's AI features are explicitly rooted in its core search and quality systems, so there, traditional ranking strength carries the most weight, and Google's own guidance says no special markup or AI-specific files are required. Engines like ChatGPT, Perplexity, and Copilot draw from a wider pool and lean

harder on structure, freshness, and authority signals, and they cite third-party surfaces, review sites, community threads, comparison pages, more heavily than top-ranked vendor pages. Chapter 7 walks the engines one by one. The point here is the pattern: rank is one input into citation, not a synonym for it.

For a small company this is unambiguously good news. You may never outrank an incumbent with a decade of accumulated links. Out-structuring them, and out-evidencing them, is a Tuesday.

One more decoupling: your page is not the only door

The tournament model carried a hidden assumption: the way to be found is your own website. Citation logic breaks that assumption too.

Watch which sources answer engines actually cite for buying questions and a pattern appears: comparison articles, review platforms, community discussions, reference pages. The engines are trying to recommend responsibly, and third-party accounts of you read as evidence in a way your own homepage cannot. Your homepage says what you claim about yourself. A review site or a practitioner's comparison says what the world has registered about you.

This means AI visibility is a property of your whole footprint, not just your domain: whether the places engines trust for your category have anything accurate to say about you. It also previews Chapter 9's argument about receipts. An engine deciding whether to name you is asking, in machine form, the same question a skeptical buyer asks: who besides you says you are real?

Same product, opposite fates

Hold the two logics side by side and the strategy difference becomes concrete.

The ranking playbook optimizes for the click: win position, earn the visit, convert on your page. Its content grew long and comprehensive because comprehensiveness won rankings, and it could tolerate burying the answer, since the human, once landed, would find it.

The citation playbook optimizes for the quote: answer the question in the first breath, structure every key claim so it survives being lifted out of context, attach evidence to everything, and make sure the machines can fetch it all. It cannot tolerate burying the answer, because the machine does not dig. It quotes what is quotable and moves on.

Here is the relief in all this: the two playbooks barely conflict. Google's stance, write for people, organize with normal headings, no separate content for AI, is also simply good writing, and the extractable structure the other engines reward does not hurt your rankings. You are not choosing between audiences. You are removing the barriers that kept one of them from using what you already built. Almost everything in Part II follows from that one sentence.

Check this yourself right now

Take the most important buying question in your category from your earlier exercises. Search it in ordinary Google and note the top three organic results. Then ask the same question in ChatGPT and in Perplexity, and note which companies are named and, in Perplexity, which pages are cited.

Now compare the lists. Ranked but not cited, cited but not ranked, or both? Whatever pattern you find, you have just observed the decoupling firsthand, in your own category. That is the pattern Part II teaches you to exploit.

Chapter 5: Compounding Absence

“Impressions are vanity; DMs are signal.” — from the Suede Scan launch campaign

Every founder understands compounding when it works for them. Content compounds. Reputation compounds. Distribution compounds. This chapter is about the version nobody budgets for: absence compounds too.

Being missing from AI answers is not a static condition, like a billboard you have not bought yet. It is a position that deteriorates while you wait, through three loops that feed each other. None of them requires the engines to be malicious or even particularly smart. They just require time.

Loop one: the buyer’s default hardens

Start with what an AI answer actually does to the person reading it. It does not just inform them. It frames the category. When a thousand buyers ask about your space this quarter and the same two competitors appear in most of the answers, those companies are not merely getting leads you are not getting. They are being installed as the reference points, the names against which everything else in the category gets compared.

Buyers repeat what machines tell them. The founder who asked ChatGPT mentions the shortlist in Slack. The consultant pastes it into a client deck. The junior analyst’s market summary is one-third synthesized answer. Each repetition is invisible to you and cumulative for them. By the time you finally meet this buyer,

months from now, you are not a candidate. You are an unfamiliar name being measured against their defaults, and the defaults were written by a machine you never audited.

Sales teams have a phrase for the vendor who arrives after the frame is set: column fodder. The new mechanism producing that old outcome is that the frame now gets set in four seconds, before you have ever heard of the deal.

Loop two: the evidence trail thickens on one side

Chapter 4 established that engines lean on third-party evidence: reviews, comparisons, discussions, reference pages. Now watch what happens to that evidence over time when one company is in the answers and another is not.

The cited company gets the buyers, which means it gets the users, which means it gets the reviews, the community threads, the “we switched to X” posts, the comparison articles that include it by default. Every one of those artifacts is a new page that tomorrow’s retrieval can find and cite. The absent company generates none of this, not because its product is worse but because the flywheel that produces public evidence starts with buyers, and the buyers are being routed elsewhere.

This is the loop that should genuinely worry you, because it operates on the inputs rather than the outputs. A stale answer can flip tomorrow. A category’s entire evidence trail, tilted three years toward your competitor, cannot. The longer you wait, the more the public record of your category is written by people who have never used your product, describing a market that does not include you.

Loop three: today's answers are tomorrow's training data

The third loop is slower and quieter. The text of the public web is what these systems learn from. The answers being generated today, and the articles, posts, and summaries humans write downstream of those answers, become part of the corpus future systems train on. A category narrative that hardens in this era gets inherited by the next one.

I will not overclaim here, because this loop is the least measurable of the three and this book does not deal in unverifiable numbers. Nobody outside the labs can tell you precisely how much today's answer shapes next year's model. But the direction is not in serious doubt: these systems learn from the written record, the written record is being written now, and you are either in it or you are not. Absence, left alone, archives itself.

Why waiting feels safe and is not

Put the three loops together and the cost of waiting stops being abstract. A quarter of delay is not a quarter of missed leads. It is a quarter of buyers trained on a shortlist without you, a quarter of evidence accruing to the companies on it, and a quarter of the public record hardening in a shape you will later have to argue with.

Yet waiting feels safe, and it is worth naming exactly why. Nothing hurts. The loops run in the silence Chapter 3 described: no alert fires, no chart dips, no deal is visibly lost. The line at the top of this chapter is from my own launch notes, where it meant something narrow about marketing metrics: impressions are vanity, DMs are signal. It generalizes into the operating principle for this whole problem. The numbers that feel good and arrive automatically are vanity. Signal is what a real buyer, or a real buying machine,

actually does. And the signal here, the answers themselves, never arrives on its own. You have to go collect it, which by now you have done three times.

Be suspicious, too, of the comfort on the other side. If you ran the exercises and found yourself in the answers, the loops are running for you, which is leverage, not safety. Point-in-time, as always: answers move between days and sessions. An incumbent who stops tending the inputs is exactly the kind of incumbent the decoupling in Chapter 4 exists to punish.

The turn

Here is where Part I lands. The buyers migrated to a layer you were not watching. That layer writes shortlists, binary and silent. It selects for citation logic, not ranking logic. And its verdicts, left alone, compound.

Everything in that paragraph is outside your control except one thing: the inputs the machines encounter when they come looking. Whether they can read your site. What your pages give them to quote. What evidence exists that you are real. Those are concrete, inspectable, fixable properties, and fixing them is not a dark art. It is a checklist, and it is the entire second half of this book.

The stakes were the hard part. The fix is work. Let's work.

Check this yourself right now

One last stakes exercise, and it is the one that reliably ends the debate inside a company. Ask ChatGPT:

“Compare [your company] and [the competitor you most often lose to] for [your ideal customer]. Which would you recommend and why?”

Read what the machine believes the comparison is. Founders regularly discover the engine misstates their pricing, their features, or their category, or politely declines to say much about them at all while speaking fluently about the competitor. Screenshot it, date it, add it to the panel. Then turn the page, because everything from here forward is repair.

PART II: THE FIX

Chapter 6: Can the Machines Even Read You?

“Fetch the file and read the rules. Do not assume.” — from the Suede AI SEO methodology

Every repair in this book depends on one precondition: when an AI crawler shows up at your site, it gets your content. Not a block, not a challenge page, not an empty JavaScript shell. Your actual words.

This is the least glamorous chapter in the book and the first one for a reason. In the scans I run, access problems are the most common severe finding, and they are always invisible from inside the company. The site looks perfect in every browser anyone has ever opened. Meanwhile some or all of the AI crawlers are being turned away at the door, and nothing anywhere records the refusal. Silent failure, exactly as Chapter 3 promised.

The good news: access is the most mechanically checkable thing in this entire field. No judgment calls, no engine mystique. A file either allows a bot or it does not. A URL either returns your content or it does not. You can verify every piece of this yourself in under fifteen minutes.

Know the bots by name

Each AI platform sends its own crawler, identified by name, and blocking a platform’s bot generally means that platform cannot fetch and cite your pages. The names to know:

- **GPTBot** and **ChatGPT-User**, from OpenAI. The first crawls; the second fetches pages live during ChatGPT browsing.

- **PerplexityBot**, from Perplexity.
- **ClaudeBot** and **anthropic-ai**, from Anthropic.
- **Google-Extended**, Google's control for its Gemini models.
Note that Google's AI Overviews ride on ordinary Google Search crawling, so your regular Googlebot access matters there.
- **Bingbot**, which feeds Microsoft Copilot through the Bing index.
- **CCBot**, from Common Crawl, a dataset many models train on.

These are different doors. It is entirely possible, and depressingly common, to be open to Google, closed to OpenAI, and challenged by everything else, without anyone in the company having decided any of it.

Check one: robots.txt, actually read

Your robots.txt file, at yourdomain.com/robots.txt, is the posted policy for crawlers. My methodology note for this check is four words long: fetch the file and read the rules, do not assume. Assumption is how these blocks survive for years.

Open the file in a browser. Reading it takes one rule: a Disallow line belongs to the User-agent line above it, and a `User-agent: *` block applies to every bot that does not have its own block. So `User-agent: *` followed by `Disallow: /` blocks everyone from everything, including all the bots above. A specific block like `User-agent: GPTBot / Disallow: /` shuts out exactly one platform.

What you are looking for: any of the named bots blocked, or a blanket rule doing it wholesale. If the file fails to load at all, that is a finding too. Treat access as unverified, not as open, and find out why.

One nuance before you reach for the delete key. Blocking AI bots is a legitimate business decision for some companies; publishers with licensing concerns, for instance, block training crawlers on purpose. The problem is not blocking. The problem is blocking by accident, inherited from an old contractor's template. There is also a middle position worth knowing: block the training-only crawler, CCBot, while allowing the search-and-answer bots, keeping your content out of bulk training sets while remaining citable. Whatever you choose, the fix for an unintended block is one line of text removed. It may be the highest-leverage single edit in this book.

Check two: what the bot actually receives

Robots.txt is the policy. Now verify the practice, because plenty of sites say “allowed” in the policy and serve something else in fact.

The common offenders sit in front of your site: CDN bot protection, web application firewalls, DDoS shields. These layers score visitors, and crawlers that are not Googlebot often score badly, receiving a challenge page, an error status, or a refused connection. Nobody configured this on purpose. It shipped as a default setting labeled something reassuring like “bot fight mode.”

The other offender is your own rendering. Text-first crawlers do best with content present in the initial HTML response. If your pages arrive as a nearly empty shell that assembles itself in the visitor's browser via JavaScript, a crawler that does not execute your scripts fetches the shell. Quick test: view your key page's source, the raw source, not the browser's rendered inspector, and search for a sentence of your actual copy. If your pricing, your product description, and your answers are not in that raw response, the machines may not be reading the page you think you published.

If you want the fifteen minutes done for you, this is exactly what the free check at **optimize.suedeai.ai** does: ten seconds, no email, and it tells you whether AI crawlers can even read your site. It is the front door of the same diagnosis this chapter just taught you to run by hand.

Check three: can the machines find everything

Access is not only the front door. A crawler that can read you still needs to discover the pages that matter. Confirm you have a sitemap listed in robots.txt, confirm your important pages are in it, and confirm those pages do not carry stray noindex tags or point their canonical URLs somewhere unintended. These are classic SEO hygiene items, and Chapter 4 explained why they now pay double: the same crawl that feeds your rankings feeds the answer layer built on top of them.

The finding, written down

Run all three checks and write down the result per bot: allowed, blocked, or unverified, with the reason. That per-bot line is the professional standard for this diagnosis, and “unverified” is an honest and common answer; a fetch that fails tells you less than a rule that says Disallow, and the two should never be reported as the same thing.

If everything came back open, congratulations: your problem is upstream, in structure and evidence, which is where the next chapters live. If you found blocks, fix them before touching anything else in this book. Every hour spent on content while the crawlers bounce off your firewall is an hour spent decorating a room the machines cannot enter.

Check this yourself right now

Open **yourdomain.com/robots.txt** and read it against the bot list above, block by block. Then run **optimize.suedeai.ai** and compare its result with your reading. Write the per-bot verdict into your notes: allowed, blocked, or unverified. That one line of notes is the foundation the next four chapters build on.

Chapter 7: The Six Engines and Who Each One Trusts

“I audit how ChatGPT, Perplexity, and Gemini answer your category, then ship the fix the same day.” — Jason Colapietro, @johnnysuede bio

Founders talk about “AI” as if it were one place, the way people once said “the internet.” Operationally, there is no such place. There are six engines that matter for buying questions, each with its own way of finding sources, its own trust preferences, and its own failure modes. Treating them as one thing produces the classic mistake of optimizing hard for a behavior only one engine has.

This chapter is the field guide. It is deliberately practical: for each engine, how it selects, and what that means for your pages. Two honest caveats first. These systems change fast, so treat this as the map I would draw today, not scripture; the verification habit from Part I is what keeps your map current. And none of this is inside knowledge of proprietary ranking systems. It is the published guidance plus what the engines observably do when you run real category questions through them, which you now know how to do.

Google AI Overviews and AI Mode

The synthesized answers inside Google Search are, by Google’s own explicit statement, rooted in its core search ranking and quality systems. That single sentence sets your whole strategy for this engine: whatever earns you strong ordinary rankings is what earns you presence in the answers assembled above them.

Google is unusually direct about what not to do. No special markup or files are required. Do not chunk your content into artificial fragments for AI. Do not write separate content for machines; that path risks tripping spam policies about scaled content abuse. Write helpful, people-first content with normal headings and paragraphs, demonstrate real experience and expertise, and keep your indexability clean.

One Google behavior deserves special attention: query fan-out. Google's AI features do not answer only the literal question asked. They generate related queries under the hood, retrieve for each, and synthesize across the set. A user asking about fixing lawns triggers hidden retrievals about herbicides, prevention, chemical-free methods. The implication for you is significant: covering your topic comprehensively, the parent question and its natural sub-questions, makes you retrievable across the fan-out. Ten shallow pages targeting ten keywords are worth less than one page, or one connected cluster, that genuinely covers the territory.

ChatGPT

ChatGPT answers from two layers: its training data, the accumulated public web as of a cutoff, and live web search when it browses. The two layers fail differently. Training data is why it can describe your company confidently and wrongly, using facts from two years ago. Live search, via GPTBot and ChatGPT-User, is where today's pages compete for citation.

What ChatGPT observably rewards is extractable structure: passages that answer a question in one self-contained block, FAQs, comparison tables, definitions that stand alone. It draws from a wider pool than the top of Google's rankings, which makes it one of the friendliest arenas for the decoupling in Chapter 4: a modest-

ranking site with quotable structure gets named. It also leans noticeably on third-party surfaces, review sites, comparison articles, community threads, when recommending in a category.

Perplexity

Perplexity is the transparency engine: always searching, always citing, links visible on every answer. That transparency makes it your best diagnostic instrument, the engine where you can see exactly which pages taught the machine its answer, which you exploited in Chapter 2.

It favors authoritative, recent, well-structured content, and its freshness preference is real enough to act on: stale pages lose citations here first. Keep your key pages visibly current, dates included, and Perplexity is winnable structure-first territory. PerplexityBot must be able to reach you, which Chapter 6 already had you verify.

Gemini

Google's assistant draws on the Google index and, critically, the Knowledge Graph, Google's structured understanding of entities: companies, products, people, and how they relate. Gemini rewards being a well-defined entity, not just a well-ranked page. Is your company unambiguous to Google: consistent name, consistent description, structured data connecting your organization to your site, coherent presence on the surfaces Google trusts? Chapter 8 covers the schema markup that feeds this. Access-wise, Google-Extended is the switch; you checked it in Chapter 6.

A field note from my own scan runbook, as a reminder that these engines are living systems: during one client-scan period, Gemini's logged-out interface accepted a typed prompt and returned nothing

at all. No error, just an empty box. Blocked silently, worse than an error. Verify your engines are actually responding when you run your checks, and note the conditions, logged in or out, which account, what day.

Microsoft Copilot

Copilot is Bing-powered, which makes it the engine founders most consistently forget, because founders forgot Bing. The Bing index is its retrieval pool, Bingbot is its crawler, and classic Bing hygiene, indexability and Bing Webmaster Tools, its verification tooling, is the unglamorous work. It matters more than its mindshare suggests because Copilot ships inside Windows, Edge, and Microsoft 365, surfaces where a large share of B2B buyers spend their working day. If your buyer is a corporate employee, there is a real chance their first AI answer about your category comes from Copilot.

Claude

Anthropic's Claude answers primarily from training data, plus web search where enabled, drawing on an external search index. The operational notes are simple: ClaudeBot and anthropic-ai are the crawlers to allow, and your durable public footprint, the evidence trail from Chapters 5 and 9, is what a training-data-weighted engine most reflects. Claude also matters for a reason beyond its chat interface: it is widely embedded inside other products and agent workflows, answering category questions in places you will never see.

Reading the table

Six engines, three families of trust. Google's surfaces, AI Overviews and Gemini, trust their own ranking and entity systems: win them with fundamentals and structured entity clarity. The live searchers, ChatGPT, Perplexity, Copilot, trust what they can fetch and quote today: win them with access, extractable structure, and freshness. The training-data-weighted layer, Claude, and every engine's offline knowledge, trusts the durable public record: win it with the evidence trail, which is Chapter 9's subject.

Notice what every column shares: readable access, clear structure, real evidence. That is why this book's repairs are not six playbooks. They are one playbook with engine-specific accents, and the next chapter starts executing it.

Check this yourself right now

Take your single most important buying question and run it through three engines you have not yet tested as a set: Gemini, Copilot, and Claude. Same question, word for word. Add the results to your panel with dates and login conditions noted.

You now have readings across the major engines. Look at the spread. Named in some and absent in others is the normal finding, and it is good news: it tells you which family of trust you already satisfy and which repairs, structure, entity, or evidence, the next three chapters should get first.

Chapter 8: Extractable or Invisible

“AI systems extract passages, not pages.” — from the Suede AI SEO methodology

The machines can read you. You know which engines trust what. Now comes the chapter where you change your site, and it starts with the five-word sentence at the top, which is the closest thing this field has to a law of physics.

AI systems extract passages, not pages. When an engine cites you, it does not present your page. It lifts a passage, a sentence to a short paragraph, and builds its answer from that. Which means the unit of AI visibility is not the page you have been polishing. It is the passage, and most sites, including most well-written ones, contain almost no passages that survive extraction.

The extraction test

Take any paragraph from your site and apply one test: pulled out alone, with no surrounding context, does it still say something true, complete, and attributable?

Most marketing copy fails instantly. “We take a fundamentally different approach” extracts to nothing; different from what, at what? “That’s why thousands of teams trust us” fails on both attribution and content; who is we, trusted to do what? These sentences work on the page, where headers and context carry the meaning. Extracted, they are fog.

Now the passing version: “Acme is project management software for construction subcontractors; it tracks change orders, lien deadlines, and payment schedules in one dashboard, with plans starting at \$49 per month.” Ripped from the page and dropped into an AI answer, that passage still works. Every claim is self-contained: who, for whom, what, how much. That is an extractable passage, and building them is a discipline of one rule: write every key claim so it survives alone.

The blocks that get lifted

You do not have to guess which shapes engines quote. Watch the answers, as you have been doing since Chapter 1, and the same few blocks appear over and over. Build these, in normal prose, on the pages that matter:

Definition blocks. For “what is” questions. One tight paragraph, early in the page, naming the thing and what it does. Your product page should define your product in its first breath, not after the hero video.

Step blocks. For “how to” questions. Numbered, each step a verb with a result.

Comparison tables. For “X vs Y” questions, which are pure buying intent. A real table with honest rows. Engines lift tables gratefully, and so do buyers.

FAQ blocks. Real questions phrased the way buyers ask them, each with a direct answer in the first sentence. This is also where Chapter 7’s fan-out pays off: the sub-questions Google retrieves for are exactly what a good FAQ covers.

Evidence blocks. A specific claim with its source and date attached. The next chapter is entirely about why these outperform adjectives.

Aim answers at roughly the 40-to-60-word mark: a complete thought, quotable whole. And answer first, then elaborate. The machine does not dig, and increasingly, neither does the human.

Structure is signposting, not decoration

Around the passages goes the page's skeleton, and here the rule is boring on purpose: one H1 that says what the page is about, H2s and H3s that name their sections honestly, headings that match how buyers phrase questions. A heading is a signpost telling a machine "the answer to this lives here." "How much does Acme cost" is a signpost. "Flexible plans for every journey" is interior decorating.

Recall Google's guidance from Chapter 7, because it draws the line perfectly: do not chunk content into artificial fragments for AI, do not write separate machine-facing content. Write for people, organize for clarity. Everything in this chapter lives on the right side of that line. An extractable passage is just a clear paragraph. A signpost heading is just an honest one. You are not gaming anything; you are removing the fog that kept machines, and skimming humans, from using what you wrote.

Schema: say it in the machines' grammar

Structured data, schema markup, is a block of JSON-LD in your page that states facts in a standard vocabulary: this page describes an Organization with this name and logo; this is a Product with this price; these are FAQ questions with these answers. It feeds Google's

entity understanding, the Knowledge Graph plumbing behind Gemini from Chapter 7, and removes ambiguity about who and what you are.

Three rules keep you out of trouble. Mark up only what is visibly true on the page; schema that contradicts your content is worse than none. Use the boring, supported types: Organization, Product, FAQPage, Article, HowTo, BreadcrumbList. And keep your identity consistent: your organization should be one entity across your site, not a slightly different name and description on every page. Validate what you ship with Google's Rich Results Test rather than trusting that it parses.

llms.txt: what is real and what is hype

You will hear about llms.txt: a proposed convention, a plain-text-markdown file at your domain root that gives AI systems a curated map of your site, your key pages and what they contain, in a format built for machine reading.

The honest status report: it is a proposal, not a standard. Google has said plainly that no special files are required for its AI features. Some other engines parse llms.txt and similar machine-readable files when present; support is uneven and shifting. So the claim I will make is deliberately modest: it is cheap, it cannot hurt, and for the engines that do read it, it hands them your site's map in their native grammar. I ship one in most repair engagements, in about twenty minutes, and I would never claim it guarantees anything. That modest claim generalizes: anyone selling you a secret file, tag, or trick that "gets you into AI answers" is selling weather again. Access, passages, structure, schema, evidence. There is no sixth secret.

The one-page drill

Do not renovate the whole site this week. Take one page, the one that should answer your category's biggest buying question, and run the drill: define the thing in the first paragraph. Convert the key claims into passages that pass the extraction test. Add the comparison table and the FAQ if the question warrants them. Fix the headings into signposts. Add the supported schema. Then reread it as a human and confirm it got better for them too. It will have. That is the tell that you did this right.

Check this yourself right now

Open your most important page and view the raw source. Find the paragraph that is supposed to answer your buyer's main question. Copy it out, paste it somewhere blank, alone, and apply the test: true, complete, attributable, with no help from the rest of the page?

Then rewrite it until it passes, and ship that one paragraph. One passing passage on your most-asked question is the smallest real repair in this book, and you can have it live within the hour.

Chapter 9: Receipts Beat Claims

“A founder buys from a named human with receipts, not a brand account with no face.” — from the Suede Scan launch playbook

The line above is from my own launch playbook, where it decided something small: which account posts. Everything for Suede Scan ships from my personal account, under my name, with evidence attached. The brand account amplifies; it never originates. The reasoning was simple: a founder buys from a named human with receipts, not a brand account with no face.

I did not expect that sentence to double as technical SEO advice. Then I watched what the engines cite, and it turns out machines evaluate trust the way skeptical buyers do. This chapter is about making your pages carry receipts, because in a world of infinite generated text, receipts are what remains scarce.

Why machines rediscovered trust

Put yourself back in the engine’s position from Chapter 4, with one new pressure added: the web is now flooded with plausible, machine-written content, and citing something wrong or fake embarrasses the engine in front of its user. Every major engine is therefore tuned, and continuously re-tuned, to prefer sources that look like they know what they are talking about.

What does that look like, mechanically? It looks like E-E-A-T, the framework from Google’s own quality guidelines: Experience, Expertise, Authoritativeness, Trustworthiness. Founders hear that acronym as branding mush. Read it instead as a checklist of

forgeable versus unforgeable signals. Adjectives are free and infinitely forgeable; any text generator can produce “industry-leading.” Receipts are expensive to fake: a named author who verifiably exists, a specific number with a date and a methodology, a screenshot of a real result, a client who says so in public, a document trail that third parties corroborate. Engines lean toward the expensive signals for the same reason buyers do. That is the entire theory of this chapter. Now the practice.

Receipt one: a named human

Anonymous content is the weakest content on the modern web, because anonymity is what generated filler looks like. Put names on your work. Your key pages and articles should carry a real author with a real bio: who they are, why they are qualified on this subject, where else they exist, their site, their profiles, their history. Mark it up with Person schema from Chapter 8 so the entity plumbing connects the name to the work.

This is Experience and Expertise made machine-legible, and it is also why the faceless-brand pattern fails twice: buyers do not trust it, and now neither do the machines. If you are a founder, you are sitting on an unfair advantage here. Your story, your reasons, your firsthand account of the problem you solve, none of it can be generated by a competitor’s content vendor.

Receipt two: specifics with dates and sources

Go through your key pages and find every claim shaped like “faster,” “trusted by thousands,” “best in class.” Each is a receipt-shaped hole. The repair is the evidence block from Chapter 8: the specific number, with its date, source, or methodology attached, written as an extractable passage.

“Significantly faster imports” becomes “In our March benchmark, importing a 10,000-row catalog took 41 seconds, down from 9 minutes in the previous version; methodology at the link.” One of these is quotable by a machine trying not to embarrass itself. The other is fog. And when you cite numbers that are not yours, name where they came from and when. An engine, like a careful reader, treats a sourced claim and a floating claim as different species.

The discipline has a hard edge that most marketing resists: if you do not have the receipt, do not make the claim. An unsupported superlative is not neutral filler; it is a signal about everything else on the page. Write around the gap or go earn the number.

Receipt three: other people’s words

Chapters 4 and 5 established that engines lean heavily on third-party surfaces for buying questions: review platforms, comparison articles, community threads. That is trust logic too. Your page says what you claim; the outside record says what the world corroborates. Which makes part of AI visibility work happen off your site entirely.

The playbook is patient and honest. Be present and accurate on the review platforms that matter in your category, and actually ask your happy customers to write there; most silence is unasked, not unhappy. Be genuinely useful in the communities where your buyers ask questions, which produces threads that mention you for real reasons. Publish reference material accurate enough that comparison writers use it, since they are drafting tomorrow’s citations. None of this is manufacturable in a week, and that is precisely why it works. It is the evidence flywheel from Chapter 5, deliberately spun in your direction.

The disclosure edge

Here is the counterintuitive move that ties this chapter together: disclose what promoters hide.

My own testimonial protocol says every quote from a founding-rate client ships with a disclosure line attached, verbatim, stating that the client got a reduced founding rate and that the rate was not conditioned on what they said. My delivery emails carry a point-in-time disclaimer and a refund policy on every offer. Say plainly what your product does not do; state the conditions under which your numbers were measured; put the caveat in the copy.

The reflex says this weakens the pitch. Watch what it actually does. A page that discloses reads as a page with nothing to hide, to a buyer and, increasingly, to systems tuned to detect the texture of trustworthiness. One honest limitation makes every adjacent claim more credible. In a sea of breathless generated copy, the disclosed caveat is the strongest scarcity signal you can print, and unlike every other receipt in this chapter, it costs nothing but nerve.

Check this yourself right now

Two audits, ten minutes each.

Onsite: open your most important page and count claims versus receipts. Every “leading,” “trusted,” and “fastest” goes in one column; every named author, dated number, sourced stat, and disclosed caveat in the other. Most pages run heavily claims-first. Convert two: one adjective into an evidence block, one anonymous section into a named one.

Offsite: ask Perplexity, **“What do people say about [your company]? Include sources.”** Read the citations. That is your third-party evidence trail as a machine sees it, and if it is thin, you

have just met the patient work this chapter assigned you.

Chapter 10: The Founder's Visibility Audit

“Run every check, and state explicitly when a check was skipped and why.” — from the Suede SEO audit discipline

Everything so far has been one repair at a time. This chapter assembles the whole diagnosis into a single sitting: the founder's edition of the audit I run professionally, reorganized so one non-technical person can complete it in an afternoon with a browser, a notes file, and the screenshots you have been collecting since Chapter 1.

Two rules from my professional discipline come with it. First, run every check; the audit's value is completeness, because the checks you skip are statistically where the problem lives. Second, when you do skip one, write down that you skipped it and why. An audit with honest gaps is a document you can trust and rerun. An audit with silent gaps is a false clean bill of health, and false health is exactly the disease this book keeps diagnosing.

Score each lane the obvious way: pass, partial, or fail, with one line of evidence for the score. No twenty-point scales. You are building a punch list, not a dashboard.

Lane 1: The answers themselves

You have most of this from the chapter exercises. Now do it systematically. Build a prompt set of 10 to 20 questions covering your category from the angles buyers actually use: what is this category, best option for your specific customer, your company versus your main competitor, how to solve the underlying problem,

and pricing. Run the set across the engines from Chapter 7. For every prompt, record: engine, date, whether you were named, who was named, and, where the engine shows sources, which pages were cited.

That grid is your visibility baseline. The score writes itself: how often are you present, and who owns the answers you are absent from? Pass is present in a majority of your money questions. Most first audits are not a pass. That is why you are here.

Lane 2: Access

Chapter 6, executed and written down per bot: GPTBot, ChatGPT-User, PerplexityBot, ClaudeBot, anthropic-ai, Google-Extended, Bingbot, and your verdict for each, allowed, blocked, or unverified with the reason. Robots.txt read block by block; the ten-second check at optimize.suedesai.ai; the raw-source test on your key pages, is your actual copy in the initial HTML; sitemap present and listing the pages that matter; no stray noindex or misaimed canonicals on money pages. Any fail in this lane outranks every finding in every other lane. Fix access first, always.

Lane 3: The money pages

Pick the three to five pages that should be earning citations: your product page, your category explainer, your comparison page, your pricing page. For each, run the Chapter 8 drill as a checklist. Definition in the first paragraph. Key claims pass the extraction test. Comparison table where the buying question calls for one. FAQ with answers-first phrasing. Headings that signpost real questions. Supported schema present, validated, and matching the visible content. Freshness visible, a real updated date, and true.

Partial credit is normal here. Score each page, then rank the fixes by one question: which page is closest to the money question you lose most often?

Lane 4: Evidence

Chapter 9, as a count. Onsite: claims versus receipts on the money pages, named authors versus anonymous sections, dated and sourced numbers versus floating adjectives, disclosures present where a skeptic would want them. Offsite: the Perplexity sources exercise, what third-party record exists, review platforms, comparison articles, community threads, and whether what it says is accurate. Score onsite and offsite separately. They fail independently and get fixed on different timescales: onsite in a week, offsite over quarters, which is exactly why offsite starts now.

Lane 5: Entity clarity

The Gemini lane. Is your company unambiguous to machines? One consistent organization name and description across your site and profiles. Organization and Person schema tying your people to your pages. Consistent presence on the surfaces that define entities: your domain, your professional profiles, the registries and directories of your industry. The quick test is Chapter 3's: ask the engines what they know about your company, and grade the coherence of what comes back.

Reading the results: the severity order

You now have five scored lanes and a pile of findings. The order of operations is fixed, and it is the same order I use on paid engagements:

1. **Access blocks.** Hours to fix, gates everything else.

2. **The single most-lost money question.** Take the one prompt from Lane 1 whose answer most often names a competitor and not you, and aim Lanes 3 and 4 repairs at that page specifically: extraction rewrite, table, FAQ, schema, receipts. One question, fixed end to end, beats ten questions half-fixed.
3. **Entity basics.** Schema and consistency, a day of unglamorous work.
4. **The evidence flywheel.** Reviews, community, reference material. Start the asks this week precisely because they pay next quarter.

Then rerun the prompt set and compare against your baseline, which is Chapter 11's whole subject: deltas, not vibes.

Do it yourself, or have it done

Everything above is genuinely doable by one founder in an afternoon, and the point of this book is that you can. It is also, fairly, several focused hours plus repairs, and some readers checked the price of having it done before finishing Part I.

So here is the plain version of what I sell, disclosed the way Chapter 9 says everything should be. At **scan.suedeai.ai**: a 48-hour Visibility Snapshot, the Lane 1 grid run professionally with timestamped screenshots. A Full Audit, all five lanes with a prioritized punch list. A Fix Sprint, where the repairs ship as pull requests and CMS edits with a before-and-after re-scan. A monthly re-scan that keeps Chapter 11 running without you. Scope and price are quoted by reply, because both depend on the site. At **seo.suedeai.ai**, the same five lanes run as an ongoing retainer with PR, entity and reputation work beside them, for a small number of companies. Every engagement carries the point-in-time

disclaimer and a refund policy, and none of them promises you a citation, because Chapter 1 told you what to think of anyone who does. The free ten-second access check at **optimize.suedeai.ai** stays free either way.

That is the whole pitch, and it is the only one in this book. The audit is yours now. Run it or hand it off, but get it run.

Check this yourself right now

Open your calendar and block the afternoon, this week, named “AI visibility audit.” Then create the notes file and paste in five headers: Answers, Access, Money Pages, Evidence, Entity. The screenshots from Chapters 1 through 9 already seed the first lane.

An audit that exists as a calendar block and a file with headers gets finished. One that exists as a good intention joins the silence this book has been teaching you to distrust.

THE CLOSE

Chapter 11: Measure Like an Operator

“Never claim an outcome. Evidence, shipped inputs, and deltas only.” — from the Suede Scan operations runbook

The sentence at the top of this chapter is the rule I wrote for myself before I let myself sell any of this. It was a promise-discipline for client work: never tell a founder “you’ll get cited.” Show the evidence, ship the inputs, measure the delta. But read it again as a founder running your own program, because it is also the complete theory of measurement for AI visibility, in nine words.

You cannot control outcomes here. The engines are moving targets: answers shift between days, sessions, and accounts, and no one outside those companies decides who gets named. What you control, entirely, are the inputs, and what you can observe, honestly, are the deltas. An operator builds their measurement on exactly those two things and refuses to be graded, or grade anyone else, on the weather in between.

The monthly re-scan

The instrument is the one you built in Chapter 10: the prompt set, 10 to 20 buying questions across the engines, screenshots, dates, a grid of who got named and who got cited. The operating rhythm is monthly, same prompts, same conditions, logged in or out noted, results into the same file next to last month’s.

Monthly is deliberate. Weekly makes you a day trader of answer noise; engines wobble between sessions, and you will chase ghosts. Quarterly lets a silent regression, a replatform that broke crawler

access, a competitor's content push, run for ninety days unobserved. Monthly matches the speed at which inputs actually take effect and keeps the habit cheap enough to survive. Thirty minutes, once a month. Put it on the calendar next to the audit block, or let a re-scan subscription do it for you; either way, the reading happens on schedule, not on curiosity.

Read deltas, not weather

One reading is weather. The comparison between readings is climate, and climate is the only thing worth reacting to.

Three shapes matter. **Trend:** your presence rate across the prompt set, this month versus the baseline. Moving up over two or three readings after you shipped repairs is signal; a single month's wobble in either direction is not. **Displacement:** who newly appears and who drops in your money answers. A competitor entering three answers they were absent from last quarter tells you where their inputs are improving, and Perplexity's visible citations will usually show you exactly which page did the work; go read it. **Correlation with shipped inputs:** the operator's discipline of writing down what you changed and when, the access fix on the 3rd, the rewrite on the 14th, the schema on the 20th, so that when the delta arrives you can connect it to something. Not proof of causation, and never claim it as proof. But shipped inputs plus a dated delta is an honest evidence trail, and an honest evidence trail is what lets you decide what to do more of.

And when a delta shows up before an input did, ask why. Sometimes the answer is that the engines moved. Sometimes it is Chapter 5's flywheel: a review, a thread, a comparison article you never commissioned, doing its patient offsite work.

The day-6 lesson

One habit from my runbook transfers whole, and I want to hand it over explicitly because it is the most operator-grade sentence in this book: verify the inputs are live before you measure the delta.

In client Fix Sprints, the day before any re-scan, I run a five-minute check: fetch the llms.txt, view source on the repaired page, confirm the schema block is actually in production. Because deploys get reverted, CMS edits sit unpublished, and a re-scan of an unpatched site measures nothing but noise, while looking exactly like a real result. If the inputs are not live, the re-scan waits.

Run your own version. Before each monthly reading, spend five minutes confirming that what you think you shipped is what the machines can currently fetch. The audit taught you every check involved. It is the difference between measuring your work and measuring your assumption that the work happened.

The operator's stance

Step back far enough and the whole book compresses into a stance.

The buyers migrated to a layer that writes shortlists. That layer is binary, silent, and compounding, which is why waiting is the one indefensible strategy. It selects sources by citation logic: access, structure, evidence, entity clarity, every one of them an input you control and can inspect. So the operator's loop is: baseline the answers, fix the inputs in severity order, verify the inputs are live, re-scan on a rhythm, read the deltas, repeat. No step requires permission, a budget committee, or a guru. The first pass costs an afternoon, and every pass after costs half an hour a month.

Most of your competitors will not do this. Not because it is hard, but because it is silent, and silence never makes it onto a roadmap. You have watched the machines answer your category's questions with your own eyes, which puts you, permanently, in the minority that has looked.

Stay in that minority. Check monthly. Ship inputs. Read deltas. The machines are already talking about your category, every day, to your buyers.

Make sure they have something true, fetchable, and worth quoting when they talk about you.

Check this yourself right now

The last exercise is the smallest. Open your calendar and create the recurring event: monthly, thirty minutes, titled "AI visibility re-scan," with a link to your prompt-set file in the description. First occurrence within thirty days of your Chapter 10 audit.

Then go run your baseline, if the afternoon block from Chapter 10 has not happened yet. If you want the baseline done for you instead, the Snapshot at **scan.suedeai.ai** is the same grid with professional screenshots, and the free check at **optimize.suedeai.ai** takes ten seconds today.

Evidence, shipped inputs, and deltas. That is the whole discipline. Go look at what the machines are saying.

Appendix A: The Founder's AI Visibility Checklist

Print this. It is the whole book in one page, in severity order.

Baseline (Chapter 10, Lane 1) - ☐ Prompt set written: 10 to 20 buying questions (category, best-for, versus, how-to, pricing) - ☐ Run across ChatGPT, Perplexity, Gemini, AI Overviews, Copilot, Claude - ☐ Screenshots dated; named/not-named and citations recorded per engine - ☐ Login conditions noted

Access (Chapter 6): fix before everything else - ☐ robots.txt fetched and read block by block - ☐ Per-bot verdict written: GPTBot, ChatGPT-User, PerplexityBot, ClaudeBot, anthropic-ai, Google-Extended, Bingbot: allowed / blocked / unverified with reason - ☐ Bot-protection and firewall layers checked for AI-crawler challenges - ☐ Raw-source test: real copy present in initial HTML of money pages - ☐ Sitemap present, listed in robots.txt, contains money pages - ☐ No stray noindex or misaimed canonicals on money pages - ☐ Ten-second check run at optimize.suedeai.ai

Money pages (Chapter 8) - ☐ Definition in the first paragraph of each money page - ☐ Key claims pass the extraction test (true, complete, attributable, alone) - ☐ Answers-first FAQ, phrased as buyers ask - ☐ Comparison table where the buying question calls for one - ☐ Headings are signposts, not slogans - ☐ Supported schema only (Organization, Product, FAQPage, Article, HowTo), validated, matching visible content - ☐ Visible, truthful updated dates - ☐ llms.txt shipped, claims kept modest

Evidence (Chapter 9) - [] Claims-versus-receipts count run on money pages - [] Named authors with real bios and Person schema on key content - [] Numbers carry dates, sources, or methodology - [] Disclosures present where a skeptic would want them - [] Review-platform presence accurate; happy customers actually asked - [] Offsite trail checked: “What do people say about [company]?” Include sources.”

Entity (Chapters 7, 10) - [] One consistent organization name and description everywhere - [] Organization and Person schema connect people, product, and site - [] “What do you know about [company]?” asked across engines; coherence graded

Rhythm (Chapter 11) - [] Repairs logged with ship dates - [] Inputs verified live before each re-scan (the day-6 lesson) - [] Monthly 30-minute re-scan on the calendar, same prompts, same file - [] Deltas read as trend, displacement, and input correlation; single-month wobble ignored

Appendix B: Glossary

AI Overviews. Google's synthesized answers at the top of search results, built on its core ranking and quality systems.

Answer engine. Any system that responds to a question with a synthesized answer naming sources or companies, rather than a list of links.

Citation. An engine using, and often linking, a specific page as a source for its answer. The unit of victory in this book.

Crawler / bot. Software an engine sends to fetch web pages. Each platform's crawler has a name your robots.txt can allow or block (GPTBot, PerplexityBot, ClaudeBot, Google-Extended, Bingbot, CCBot).

Delta. The change between two dated readings of the same prompt set. The only measurement this book trusts.

E-E-A-T. Experience, Expertise, Authoritativeness, Trustworthiness: the quality framework behind Google's guidelines, read in this book as a checklist of unforgeable signals.

Entity. A thing machines can identify unambiguously: a company, person, or product. Entity clarity is being one consistent thing everywhere.

Extraction test. Pulled out alone, does the passage still say something true, complete, and attributable?

GEO. Generative engine optimization: the practice of earning presence in AI-generated answers. Also called AI SEO, AEO (answer engine optimization), and LLMO. All one discipline: access, structure, evidence, entity, rhythm.

Knowledge Graph. Google's structured map of entities and relationships, feeding Gemini and the entity lane of the audit.

llms.txt. A proposed convention: a machine-readable file at your domain root mapping your key pages for AI systems. Cheap, modest, unevenly supported.

Point-in-time. The status of every AI answer ever captured. Answers move between days, sessions, and accounts; screenshots are dated for a reason.

Prompt set. Your standing list of 10 to 20 buying questions, run identically each month.

Query fan-out. Google's AI generating related queries under the hood and synthesizing across them; the reason topical coverage beats keyword sniping.

Schema / structured data. JSON-LD in your pages stating facts in a standard vocabulary machines parse directly.

Zero-click search. A query answered on the results surface itself, with no visit to any website.

Appendix C: The Tools

Everything in this book can be done by hand. These are the shortcuts, disclosed plainly per Chapter 9.

optimize.suedeai.ai: free, ten seconds, no email. Checks whether AI crawlers can read your site: the access lane's front door.

scan.suedeai.ai: the done-for-you ladder. The 48-Hour Visibility Snapshot is your prompt-set grid, run professionally, with timestamped screenshots. The Full Audit is all five lanes with a prioritized punch list. The Fix Sprint ships the repairs as pull requests and CMS edits, with a before-and-after re-scan. The Monthly Re-Scan is Chapter 11 on autopilot. Scope and price are quoted by reply.

seo.suedeai.ai: the retainer practice. The same five lanes run continuously, with PR, entity and reputation work beside them, for a small number of companies. Quoted by reply. This book, in both editions, lives at seo.suedeai.ai/book.

Every engagement carries a point-in-time disclaimer and a refund policy. None of them guarantees a citation, a ranking, or a recommendation, because nobody honest can.

The Screenshot is a Johnny Suede Press book by Jason Colapietro, founder of Suede Labs. All engine behaviors described were observed as of this edition's writing and will drift; the discipline is built to outlast the details. Trademarks belong to their owners; the engines named here are products of their respective companies, and nothing in this book implies their endorsement.