

The GEO Field Guide

The GEO Field Guide is the complete method for getting your brand named by AI assistants — six layers, in the order they actually matter, with the verification step for each one. It is published here in full, ungated, because a document behind a form is invisible to the engines this guide is about.

About 45 minutes to read end to end, as a page, a PDF or Markdown — the same document either way. Also known as AEO or answer-engine optimisation.

What this covers. How assistants pick who to name, the six layers of visibility in dependency order, how to decide what to write, how to build a page that gets quoted, how to measure without fooling yourself, and how to audit a site in an afternoon.

§1 How an assistant decides who to name

§2 The visibility pyramid

§3.1 Gate: can a crawler read it?

§3.2 The entity foundation

§3.3 Structure and extractability

§3.4 Off-site corroboration

§3.5 Freshness and decay

§4 Deciding what to write

§5 Building the page

§5.6 Gating without going invisible

§6 Measuring without fooling yourself

§7 The audit, bottom-up

§10 Myths you will hear

§1 • THE CORE MODEL

How an assistant decides who to name

An assistant can only mention you through one of three routes, and every tactic in this guide maps to one of them. Diagnose the route before prescribing a fix — most wasted GEO effort is the right work aimed at the wrong channel.

CHANNEL	WHAT IT IS	HOW FAST YOU CAN MOVE IT
A • Parametric memory	What the model absorbed in training, because it appeared repeatedly and consistently across the public web. You cannot edit it directly.	Months to years, behind a training cutoff
B • Retrieval (grounding)	At query time the engine fetches a handful of passages and answers from them, usually with citations. Most "AI search" gets its facts here.	Hours to days

CHANNEL	WHAT IT IS	HOW FAST YOU CAN MOVE IT
C • Live web search	A freshness-weighted subset of B, running against an underlying search index. If you are not indexed, you cannot be retrieved.	Hours to days

The two clocks — you have to win both

A brand that wins only retrieval disappears the moment a fetch fails or the model answers from memory. A brand that wins only parametric memory is frozen at the last training cutoff and wrong the day its specs change. The slow clock rewards ubiquity, consistency and corroboration; the fast clock rewards crawlability, structure and freshness.

Indexation makes you *eligible*. It does not make you *chosen*. Large-sample citation studies keep finding that only a minority of sources cited in AI answers rank in the organic top ten — what wins the citation is being the clearest self-contained passage on the exact question.

Entities, not keywords

Old search matched strings. Models reason over entities and the relationships between them: a product connected to its category, its maker, its alternatives, the standards it meets. This is why entity work is decisive in niche and B2B categories — where there is little competing signal, whoever establishes the canonical facts first effectively writes what every model believes.

Diagnostic: the assistant says something wrong about us

SYMPTOM	CHANNEL	WHERE THE FIX LIVES
States an outdated spec or owner, confidently, with no citations	A	Publish the correction prominently, align third-party sources, update the knowledge graph. Expect lag.
Hedges — "reportedly", "around", "sources vary"	A	Contradiction in your own ecosystem. Find and remove it before anything else.
Names a competitor for your category question	A + B	Entity foundation, then the page that answers that actual prompt.

SYMPTOM	CHANNEL	WHERE THE FIX LIVES
Describes you correctly but cites someone else's page	B	Your page is less retrievable or less clear than theirs. Structure work.
Does not mention you at all when you clearly fit	B, or the gate is broken	Check crawlability first, then look for the missing page.
Right in ChatGPT, absent in Perplexity	B / C	Different indexes and crawlers. Check access per engine; never generalise from one.
Merges you with an unrelated same-named thing	A + entity	Disambiguating context in the canonical text, plus aliases and <code>sameAs</code> .

§2 · THE ORDER OF WORK

The visibility pyramid

Work bottom-up. Lower layers are prerequisites, so effort spent high is wasted while something low is broken. This ordering is the most useful thing in the guide, because it prevents the most common failure in the field: polishing structured data on a page no crawler can read.

0

Factual accuracy — one set of approved facts

A gate. If your own surfaces contradict each other, everything downstream amplifies the noise, and contradiction is the primary cause of a model hedging about you.

1

Crawlability and machine access

A gate. If an AI crawler cannot fetch and parse the page, nothing above this line matters.

2

Entity foundation

Durable, high leverage. The model must know, without contradiction, what this is, what category it belongs to, who makes it, and what it is not.

3

On-site structure and extractability

Durable principle, churning formats. Decides whether your true facts get retrieved cleanly as passages. This is the layer that wins live retrieval today.

4

Off-site authority and corroboration

Durable, high leverage. A fact stated only by you is a claim; the same fact stated by independent sources is knowledge.

5

Freshness and maintenance

Retrieval down-weights stale content and entity facts drift. Decay is real: unmaintained pages lose visibility without anything visibly breaking.

6

Measurement and iteration

Durable principle, churning tools. You cannot manage this as a one-off project, because the surface changes continuously.

Definition of done

A subject is not finished until every one of these is true. This list works well pasted straight into a ticket.

- ✓ **One canonical URL** — the undisputed home, not three competing pages.
- ✓ **A one-paragraph canonical definition** in plain language, reused verbatim across the site, the datasheet and the knowledge graph.
- ✓ **A machine-readable spec table** — a real HTML table, with units, ranges and standards.
- ✓ **Structured data on the canonical page**, with `sameAs` to authoritative profiles.
- ✓ **Crawlable to the AI agents you care about** — or a deliberate, documented decision otherwise.
- ✓ **An HTML version of every datasheet's key facts.**
- ✓ **A visible question-and-answer block** in question → direct-answer form.
- ✓ **At least three independent sources** stating the same category and the same numbers.

✓ A named owner and a review date.

\$3.1 • GATE 1

Can a crawler actually read it?

Most AI crawlers do not execute JavaScript. They read the raw HTML your server sends. Every fact you want cited has to exist as text in view-source — not only in an image, a chart, a PDF, or content that loads on click.

Googlebot is the notable exception as a rendering crawler; design for the strictest reader, not for the most capable one. A client-rendered single-page app is close to invisible to assistants however good its content is.

FORBIDDEN	FINE
Specs only in a PDF	Native <code><details></code> — the text stays in the source
A comparison shown only as an image	Server-rendered tabs with every panel in the HTML
Answers that expand by fetching over AJAX	Real <code><table></code> , <code></code> and <code><dl></code>
Any "read more" that loads text after a click	Inline SVG containing real text

Verify it — do not assume it

- ✓ **View source with JavaScript off** and search the raw HTML for the definition sentence, each headline number, each answer, each comparison value.
- ✓ **Check the status code** and that nothing sits behind a consent wall, geo-gate or login for the crawler.
- ✓ **Fetch robots.txt** and enumerate which AI agents are allowed, and whether a `Sitemap:` directive exists.
- ✓ **Confirm the sitemap** is current and lists the canonical URLs.
- ✓ **Confirm the canonical link**, and that variants point at the hub.
- ✓ **Confirm the page is indexed** in the underlying search engines — live retrieval depends on it.

Our free AI crawler access checker (mentionbeat.com/ai-access) does the robots.txt parse and the live per-agent probes for you, including the CDN and firewall blocks that robots.txt alone will not reveal.

The robots.txt decision is two decisions

Conflating them is the most expensive mistake in this section.

	SEARCH AND ANSWER AGENTS	BULK TRAINING CRAWLERS
What they do	Fetch pages to build the answer a user sees now, and to power the engine's index. Usually produce a citation.	Collect corpora for training the next model.
Blocking means	You disappear from AI answers.	Your content is less likely to enter parametric memory.
Decision type	A visibility decision — for most vendors, allow.	A licensing decision — declining is defensible.

Two further rules worth holding: robots.txt is a *request*, honoured by reputable crawlers, not an enforcement mechanism — if you need actual prevention, that is authentication or edge rules. And blocking is hard to reverse in effect, because lost presence takes time to rebuild. Ready-made rules are in the [templates pack](https://mentionbeat.com/templates) (mentionbeat.com/templates).

§3.2 · LAYER 2

The entity foundation

The model has to be able to answer, with no contradiction: what is it, what category, who makes it, what is it not, and what is it the same as. This is the highest-leverage durable work available, and it is mostly unglamorous.

- ✓ **State the full name, category and maker in the first sentence, everywhere.** The pattern is "*Name is a category made by Maker that does Y for Z*". This single repeated sentence is the most valuable string you control.
- ✓ **Resolve legacy and alias names explicitly.** After a rebrand, models trained across the transition hold conflicting facts. Record the legacy name as an alias and an `alternateName`, and let the machine-readable layer carry the bridge so the running copy stays clean.
- ✓ **Handle name collisions.** Search your name in isolation. If unrelated things share it, your canonical text needs disambiguating context so entity linking does not merge you with them.

- ✓ **One spelling and casing**, consistently, with the variants recorded as aliases.

The open knowledge graph, in order of effort

Wikidata first — structured, editable, and it feeds both training corpora and grounding systems. Give every statement a reference; unreferenced statements are low-trust and prunable. Add multilingual labels, because one item serves every language. **Wikipedia second, and slowly**: a standalone article needs notability established by independent sources you do not control, and a promotional page will be removed and can backfire. Build the independent-source base first.

Do not make covert promotional edits, cite your own marketing as a notability source, or create an article before independent coverage exists. Disclose conflicts of interest and propose sourced edits transparently. This is one of the few places in GEO where the shortcut has a real downside.

sameAs — wiring the entity together

In the canonical page's structured data, `sameAs` should point at every authoritative profile of the same entity: Wikipedia, Wikidata, official social profiles, industry registries, standards listings. It is the literal machine instruction "these all refer to the same thing". Cheap, durable, and left out more often than anything else on this page. Our schema generator (mentionbeat.com/tools/schema) emits it correctly.

S3.3 · LAYER 3

Structure and extractability

Engines retrieve passages, not pages. A page is a container; the unit that earns a citation is a paragraph that makes sense on its own.

- ✓ **Self-contained sections of roughly 50–150 words**, each answering exactly one question under a descriptive heading.
- ✓ **No orphan pronouns**. A retrieved paragraph must make sense with zero surrounding context — repeat the subject's name instead of "it" or "the platform".
- ✓ **One idea per paragraph, one claim per sentence**, conclusion first, support after.
- ✓ **Put the direct answer immediately under the question heading**, then elaborate.
- ✓ **Front-load the page**. A disproportionate share of citations come from the top, and the top is what survives truncation. The definition goes above the narrative.

Self-contained, not fragmented. There is no requirement to chop content into tiny pieces, and doing so causes cannibalisation — a supporting page out-ranking and out-citing the page it supports. The target is complete and self-contained.

Semantic HTML

One clear `<h1>` carrying the canonical name and category, with a logical heading hierarchy phrased as real user questions. Real tables for specs and comparisons, lists for steps, definition lists for term-and-definition pairs. Valid HTML, a set `lang`, a descriptive title and an answer-first meta description of about 155 characters. Native `<details>` for progressive disclosure — never a JavaScript accordion that fetches its own text.

Structured data

Structured data is the most reliable way to hand a machine unambiguous facts. It is not a ranking lever — Google states it is not required for AI features — but it is cheap, durable and it disambiguates. Mark up the primary entity, the organisation behind it, the breadcrumb trail, and any visible question-and-answer block, and make sure every marked-up fact appears in the visible text. Structured data that contradicts the page is worse than none.

Media, PDFs and datasheets

PDFs are weakly retrievable: often poorly chunked, sometimes not crawled at all, with the real numbers trapped in tables and figures. This matters enormously in B2B, where the facts live in datasheets.

- ✓ **Mirror every datasheet's key facts in HTML.** The PDF can stay the formal artifact; the HTML is what gets retrieved.
- ✓ **Never put a number only in an image.** Restate it in adjacent text, or use inline SVG with real text plus a caption carrying the values.
- ✓ **If a PDF must stand alone,** give it selectable text, a descriptive filename, a title, headings, a stable URL and an HTML landing page that summarises and links it.
- ✓ **Express every spec with its units and qualifiers inline.** "Up to 200 m, IEC 61400-12-1 classified" beats "long range", and models reproduce the qualifier when it sits beside the number.

llms.txt — the honest verdict

A proposed convention: a Markdown file at the site root pointing at your most citable pages. Google says it is not needed for Search; some assistants read it; adoption is not guaranteed and it may standardise or fade. **Cheap insurance, not a lever.** Ship it if it costs an hour, and never let it substitute for robots.txt, a sitemap, structured data or HTML hygiene. Our **llms.txt generator** (mentionbeat.com/tools/llms-txt) builds one from your sitemap.

Hub and spoke

The canonical page is the hub; focused spokes each own one question and link both ways. Engines read the internal-link graph to decide which page is the citation-worthy source, so route links and authority toward the hub with descriptive, varied anchor text — never "read more" alone. Contextual in-body links carry more signal than the same link in a related block; do both, but do the in-body one. Roughly three to five internal links per thousand words, key pages within about three clicks, no orphans.

\$3.4 · LAYER 4

Off-site corroboration

A fact stated only by you is a claim. The same fact stated by independent sources is knowledge.

Corroboration is how a claim enters parametric memory and earns retrieval trust. In specialised categories it is decisive, because whoever seeds consistent independent facts first owns the model's understanding. Ordered by trust:

1. **Standards bodies and official registers** — classifications, metrology institutes, regulators. Highest authority, durable, hard to fake.
2. **Peer-reviewed and preprint literature** that names the subject. For technical products this is the dominant channel and it ages extremely well.
3. **Independent technical press** and vertical trade publications.
4. **Review and directory platforms** — accurate third-party profiles raise citation probability.
5. **Communities and forums**, which some engines cite disproportionately. You cannot astroturf this; you can be genuinely present.

The aggregate charts are easy to misread. Published analyses find that user-generated and encyclopedic platforms dominate citations overall — but that is a statement about consumer query volume, not about your category. For a specialised B2B subject the cited sources skew heavily to standards, literature and vendor documentation. The actionable reading is about *format*: what those platforms have in common is that they are independent, question-shaped, plainly written and densely factual. That is a description of what to make your own content look like.

The durable principle is consistency. Every corroborating source should repeat the same definition and the same numbers. Divergent numbers across sources are the leading cause of model hedging on technical subjects.

§3.5 · LAYER 5

Freshness and decay

- ✓ **Keep** `dateModified` **accurate** and show a visible "last updated" date.
- ✓ **Date your facts in the prose** — "as of 2026, the current model is..." — so a stale claim is self-evident.
- ✓ **Do not fake freshness.** Bumping a date with no substantive change is a low-trust signal, and it destroys your own ability to read measurement deltas afterwards.
- ✓ **Run a propagation pass on the triggers:** a new model or spec, a rebrand or acquisition, a new engine launch, or a measured accuracy drop. Update the facts, then push them to the site, the structured data, the knowledge graph and your partners within one sprint.

§4 · PLAYBOOK B

Deciding what to write

A beautifully structured page that answers the wrong question will not get cited. Models cite the page that best matches the intent behind a prompt, so if nobody mapped the prompts, the team is guessing at intent — and whoever did map it gets cited instead.

Two governing principles. Content is evidence, not persuasion: buyers in considered purchases verify rather than get sold, and a page of adjectives loses to a page of numbers and standards. And route by stage — *learn* questions to explainer pages, *decide* questions to the product page, *act* questions to the contact route. A page serving all three usually serves none.

STEP	WHAT IT PRODUCES
1· Profiles	A specific decision-maker, not a market segment. "Offshore wind" is a market; "an offshore programme manager who buys through an integrator" is a profile — and they ask different questions.
2· Jobs	The decisions each profile is actually accountable for.
3· Questions	The real questions, harvested from sales calls, support tickets and the search box — in the buyer's words, not yours.

STEP	WHAT IT PRODUCES
4 • Prompts	Those questions rewritten the way someone would type them into an assistant. This becomes your measurement suite.
5 • Baseline test	Run the suite before you change anything. This is the step almost every team skips, and it is why most GEO efforts cannot prove anything later.
6 • Content map	Each prompt mapped to the page element that answers it — a section, a table, an FAQ entry, a new spoke.
7 • Build	Hand off to the page work below.
8 • Measure	Re-run the same versioned suite and compare.

§5 • PLAYBOOK C

Building the page

Two principles drive every choice below. **If it is not in the HTML text, it does not exist** to an assistant. And **complete but layered, not thin** — "concise" has to mean progressive disclosure, not a short page. A thin page loses three ways: it gets cannibalised by its own supporting pages, it loses to a rival's fuller page, and humans leave.

Only two things are position-sensitive: the answer-first opening and the key facts. Front-load those. Everything else you can reorder, because each section stands alone. The two real ordering mistakes are burying the definition under narrative, and writing sections that only make sense in sequence.

The twelve-point ship-it gate

Paste this into the ticket. The page is not done until every line is true.

- ✓ **Answer-first opening** — the first sentence defines the subject in plain language.
- ✓ **A key-facts block** near the top, written so a model can lift it whole.
- ✓ **Structured data present and valid**, mirroring the visible text, with zero validation errors.
- ✓ **Specs in a real HTML table** — not an image, not PDF-only.
- ✓ **One comparison table**, category-level or head-to-head if your policy allows it.
- ✓ **An honest "what this is not for" section.** It builds the trust models reward and saves bad-fit sales cycles.

- ✓ **Six to ten visible questions and answers** — the visible block is what gets cited.
- ✓ **Every in-image fact repeated in text**, with descriptive alt text.
- ✓ **No JavaScript-only or PDF-only critical content.**
- ✓ **A canonical URL** plus a fresh, visible last-updated date backed by `dateModified`.
- ✓ **Links to at least two supporting pages** with descriptive anchors, in both directions.
- ✓ **Every fact traceable** to your approved source of truth.

The invisible layer

A human never sees this; assistants depend on it. Structured data for the primary entity, the organisation and the breadcrumb, with specs as properties, `sameAs` wired to authoritative profiles, and `alternateName` carrying legacy names that should not appear in running copy. One canonical URL, an answer-first title and meta description, a set language, an accurate `dateModified`, a `robots.txt` that deliberately allows the agents you want, and a referenced sitemap.

Cross-cutting rules

- ✓ **Diagrams are text-backed, never text-replacing.**
 - ✓ **High information density everywhere** — named entities and concrete numbers in every section.
 - ✓ **Every section answers its own question** without requiring another.
 - ✓ **Plain, precise language.** Write for intent, not keywords, and name the brand and product together consistently.
 - ✓ **Avoid superlatives.** Unverifiable claims lower credibility with readers and models alike, and specific number-backed claims outperform them.
-

\$5.6 • THE CONVERSION QUESTION

Gating without going invisible

This is where GEO and demand generation collide, and the default corporate instinct is wrong. A gated PDF is invisible to crawlers: gating baseline collateral defeats the entire visibility effort *and* produces weak, noisy leads. The principle is narrower than "never gate" — **gate only when the exchange is fair and the signal is strong**. The reader must get unique value, and the act of asking must mean they are genuinely evaluating a purchase.

TIER	WHAT BELONGS HERE	TREATMENT
Open	Datasheets, brochures, general product information, spec tables, application notes, how-it-works, FAQs — anything needed to consider you at all	Ungated and crawlable. An optional "email me a copy" is fine.
Named	Recorded technical deep-dives, validation-methodology papers, raw datasets	A mild gate is fair — real extra effort, low friction.
Qualified	A quote, a demo, an evaluation unit, a site-specific proposal, a finance-grade dossier	A full form is expected and welcome. This <i>is</i> the buying signal.

Gate intent, not documents. The strongest modern move is to capture intent through tools that require someone to describe their situation in order to get value — a cost calculator, a readiness checker, a scoping builder, an evaluation request. The input *is* the lead, because you learn the site, the scale, the timeline and the standards, and the tool's own content can stay server-rendered and visible. The one document genuinely worth gating is a customer-specific dossier: asking for it is a near-certain buying signal, which is the opposite of a datasheet.

This page is the rule applied to ourselves. The whole guide is here, ungated; the file downloads without an address; the form below is a convenience, and our **scoping builder** (mentionbeat.com/plan) is where we ask for more, because that is where the exchange is fair.

\$6 • PLAYBOOK D

Measuring without fooling yourself

Model output is non-deterministic, and temperature zero does not make it deterministic. Asking ChatGPT once whether it mentions your product tells you almost nothing.

Everything below follows from that one fact.

- ✓ **Never measure once.** Sample each prompt many times across several engines and report a confidence interval, not a falsely precise single number.
- ✓ **Keep the raw answers, recompute the analysis.** Store every response so you can re-score or add a metric later without re-spending the budget. The raw log is the durable asset; the dashboard is not.

- ✓ **No brand-name leakage.** Headline metrics come from category prompts that never name the brand. If you ask "tell me about Brand X" and it answers, you have measured nothing.
- ✓ **Report grounded and remembered answers separately.** An answer carrying citations is grounded even if you did not enable search, and averaging the two measures nothing.
- ✓ **Fix the suite and version it.** Thirty to a hundred prompts spanning informational, comparative, buying-intent and troubleshooting questions. Changing the suite invalidates every historical comparison.
- ✓ **Alert only on non-overlapping intervals,** or you will chase noise every week.

The metrics that matter

METRIC	DEFINITION
Presence / mention rate	The share of suite prompts where the subject is named.
Share of voice	Your mentions divided by all options named on category prompts.
Accuracy / hallucination rate	The share of stated facts that are correct. Investigate every error individually.
Owned-citation rate	The share of answers citing one of your own URLs, as opposed to describing you correctly while crediting someone else.
Recommendation rate	The share of buying-intent prompts that shortlist you.
Framing	How you are characterised — reference standard, expensive, legacy — with limitations stated correctly.

The statistics, for whoever owns the number

Use **Wilson intervals** for proportions rather than the normal approximation, which behaves badly exactly where visibility rates live — near zero, near one, and at small samples. Use a **cluster bootstrap over prompts** for headline intervals, because repetitions of the same prompt are correlated and treating them as independent understates uncertainty. If a model scores the answers, **calibrate the judge** against human labels and correct for its error rate — an uncalibrated judge can shift a headline number by more than the change you are trying to detect. And plan the sample size before spending the budget.

Traffic, attribution, and "is this worth it?"

Expect the question, and answer it honestly. **The value is being in the answer, not the click** — when an assistant returns a shortlist, inclusion shapes the decision whether or not anyone clicks, and published analyses of assistant citations find only a small minority produce a click-through. Attribution is genuinely hard: referrers get stripped, in-app browsers hide the origin, and clickless answers leave no trace, so a large share of assistant-driven visits land as "direct". Anyone reporting precise AI-attributed revenue is over-claiming.

Track it as well as it can be tracked — assistant hostnames as a custom channel group, server-log analysis of crawler hits, and above all the prompt suite, which is the only instrument that measures presence in answers where there was never a click. On conversion quality: report the direction, refuse the magnitude. Published multipliers range from about four to twenty-three times, which is itself evidence that nobody measures it consistently.

§7 · PLAYBOOK A

The audit, bottom-up

Audit in layer order and stop to report as soon as you find a broken gate. Do not hand someone forty findings when the first two make the other thirty-eight irrelevant.

1. **Establish the subject and its ecosystem** — what exactly is being audited, what the canonical URL is meant to be, and what other pages, PDFs, partner listings and localised versions exist.
2. **Gate 0: contradictions.** Diff the facts across every surface you can reach and list every divergence with both sources. If contradictions exist, that is the finding; everything else waits.
3. **Gate 1: machine access.** Run the verification above. If a crawler cannot read the facts, that is the finding.
4. **Entity.** Search the name in isolation for collisions. Check the knowledge-graph item, the canonical definition sentence, the `sameAs` wiring, and the spelling across surfaces.
5. **Structure.** Definition first? Key facts near the top? One heading hierarchy phrased as questions? Specs in a real table? Visible FAQ? Structured data mirroring visible text? Internal links pointing at the hub?
6. **Corroboration.** Count independent sources stating the category and the key facts, and check their numbers are identical to yours.
7. **Freshness.** Visible last-updated date, accurate `dateModified`, and facts that are actually current.
8. **Measurement.** Is there a versioned suite and a baseline? If not, run one as part of the audit — it converts an opinion into evidence and it is usually the most persuasive artifact you produce.
9. **Score and prioritise** by layer depth, then effort, then impact. Gates first, then quick wins, then projects.

Report in this order: the headline — one sentence naming the single thing blocking visibility; the baseline evidence — what the assistants actually say today, verbatim, per engine, which persuades stakeholders better than anything else you can produce; then the scorecard with the gates flagged; then the prioritised findings.

Our [free page checker](https://mentionbeat.com/checker) (mentionbeat.com/checker) runs the structure and crawlability half of this automatically, and the [page checklist](https://mentionbeat.com/checklist) (mentionbeat.com/checklist) lists every check it applies.

§10 • CORRECTIONS

Myths you will hear

Correct these with the mechanism, not with authority.

“GEO replaces SEO.”

“Just add llms.txt and you’re optimised.”

“FAQ schema gives a 40% citation lift.”

“AI can read our PDFs and our JavaScript site fine.”

“Just block all the AI bots — they steal our content.”

“Longer content ranks better in AI.”

“We rank number one, so we’ll be cited.”

“We checked ChatGPT and we’re visible.”

“Our AI referral traffic is tiny, so this doesn’t matter.”

“Let’s mass-produce AI-written pages to cover every query.”

“This is a one-time project.”
