

Citation Gap Synthesizer

Closing the gap between AI Visibility and action



Visibility

tracked



Position

tracked



“Why?”

not yet answered

Mereke Dadabayeva — Product Manager, bridging strategy and engineering

Prepared for Peec AI

AI search is the new front door. Most teams can only watch it open.

Peec AI already tells growth teams how their brand performs in AI search — Visibility, Position, and Sentiment, tracked per Prompt, across every Model. That's the essential first layer. What's still missing is the next question every team asks the moment they see a loss.

● What Peec AI already shows

- Visibility, Position, and Sentiment for every tracked Prompt
- Which competitor Brand is outperforming yours, and by how much
- Whether the trend is improving or falling behind, over time

● What's left to the team, by hand

- Open the winning competitor's page and the brand's own page, side by side
- Guess which Source detail made the AI model prefer that page
- Repeat this, one Prompt at a time, across every lost query

The B2B Growth Lead

The person who already lives inside Peec AI's dashboard every week.



Growth Lead

B2B SaaS · Marketing-led growth team

GOAL

Grow AI Visibility on the Prompts that drive pipeline — and prove it to leadership.

DAILY FRICTION

Manually cross-examining competitor Sources to guess why an AI model picked them.

HOURS LOST

Once per lost Prompt, across dozens tracked weekly.



*I can see we're losing on 'best CRM for startups.'
Position 5, Visibility down. What I can't see, without
an hour of manual digging, is what to actually
change on our page.*

— a paraphrase of what growth teams tell Peec AI's own customers

A different kind of ranking signal

Legacy SEO tools optimize for a keyword index. AI models don't read one.

SIGNAL	LEGACY SEO TRACKERS	GEO / AI SEARCH
What's measured	Keyword density, backlink volume, SERP rank	Brand mention + Source citation inside an AI-generated answer
Evidence used	On-page keyword frequency	Schema markup, entity coverage, and Source-level text an AI model actually cited
How a gap is found	Manual audit against a checklist	Deterministic diffing against the competitor's live Source

One diagnostic engine, two ways to use it

Both tracks run the same deterministic scraping pipeline and zero-extrapolation diffing logic — every gap traced to text a page actually contains, nothing inferred.

● TRACK A

Single-Prompt Diagnostic

Point it at one lost Prompt. It scrapes the brand's page and the winning competitor's page, then diffs both against Schema markup, stated proof points, and topic entities.

Output: a plain-English fix, and the exact code to ship it.

● TRACK B

Portfolio Aggregator

Runs the same engine across every Prompt a Brand is tracking, then groups the gaps that repeat — so the team sees the pattern, not just one query at a time.

Output: one ranked list — fix the top item, win back the most Prompts.

Simple by default. Deep when you need it.

Every team member gets the view that matches how much detail they need.

Comparative Inspection

Your Brand

Missing FAQPage Schema

Competitor

✓ FAQPage Schema present

► Raw scraped payload

Collapsed by default — full Source text available for anyone who wants to verify a claim.

Comparative inspection

Your Brand's page and the competitor's page, aligned row by row — so the difference is visible at a glance, not buried in two open tabs.

Raw payload drawer

The scraped Source text stays one click away — technical enough for engineering, invisible for everyone else.

The hard part wasn't scraping. It was knowing what not to claim.

Live Source data fails sometimes — a page times out, a Model-facing detail changes. The dangerous failure is silent: a plausible guess quietly standing in for real evidence. This is how the engine stays trustworthy when that happens.

● FETCH RELIABILITY

Dual-reader race, not a single point of failure

Two independent readers fetch the Source concurrently on every request. Whichever returns real content first is used — raising the success rate without a fragile single dependency.

● CLAIM INTEGRITY

Every claim is tagged, never blended

✓ VERIFIED

confirmed against a live-fetched Source — never a placeholder

~ METHODOLOGY

a general best practice, explicitly not asserted as competitor evidence

“This is the same discipline Peec AI's own Visibility data has to hold: a number only matters if the team reading it can trust where it came from.”

What v1 intentionally does not do

A short, deliberate no-list — protecting engineering bandwidth so the diagnostic core ships well.



No direct CMS auto-publishing

The tool recommends the fix; a person still decides to ship it.



No automated outreach

No emailing journalists or site owners on the team's behalf.



No bypassing paywalled Sources

If a competitor's page sits behind a login, it's left out of the diff.



No multi-tenant permissions

v1 is single-workspace; role-based access is a later phase.



No fine-tuning an LLM

Diffing logic is deterministic — no model training in the loop.

One diagnosis. Two ready handoffs.

The same finding exports in the format each team already works in.

Marketing Brief · Markdown

- Why the competitor won this Prompt
- What to change, in plain English
- One click to copy into the content brief

For: Marketing & Content

PEEC-408 · Jira Ticket

Given / When / Then

Given the brand's page is live
When the fix is deployed
Then the missing Schema is present
And the stated proof point matches the competitor's

For: Engineering

Proposed North Star metrics

Layered on top of Peec AI's existing Visibility, Position, and Sentiment — as targets to work toward, not results already claimed.

TARGET: DECREASE

Time-to-Remediation

The time between a lost Prompt being flagged and the fix going live on the brand's page.

TARGET: INCREASE

Prompt-to-Action Conversion

The share of flagged gaps that turn into a shipped fix, not just a diagnosis that sits unread.

Both metrics require a live pilot to baseline — no numbers are claimed here that haven't been measured yet.

See it live, or read the source.

Both tracks run end to end in the browser — no signup, and nothing hidden from view. The source includes a parallel prototype of the same diffing logic, built independently to validate the approach a second way.

LIVE BUILD

merekedadabayeva.github.io/AI-Search-Citation-Gap

SOURCE CODE

github.com/MerekeDadabayeva/AI-Search-Citation-Gap

PORTFOLIO

mereke.me

A note on this deck

Every claim describing the shipped tool's behavior describes code that runs today. Figures in the North Star Metrics section are stated as targets, not results — in keeping with the Verified / Methodology discipline the tool itself enforces.