

The Evidence Led GEO Growth Playbook

A practical operating manual for websites and B2B teams that need AI search visibility to produce qualified demand

Version	Format	License
1.1 September 2026	Field manual and workbook	Single company internal use

Aivius

Evidence, experiments, and revenue measurement for generative engine optimization.

License and use

This purchase grants one organization a perpetual, non-transferable license to use this edition internally. The purchaser may copy the worksheets for employees and contractors working for that organization.

The license does not transfer copyright or permit resale, public redistribution, white-label publication, dataset extraction, or use as training material for a competing commercial product. Client-service firms may use the methods in their own work, but may not give clients the full manual or editable source files unless a separate license is agreed.

The manual provides an operating method, not a ranking or revenue guarantee. Search and AI systems change, and observed results depend on product fit, evidence quality, market demand, technical access, and execution. Record provider failures and unknown values honestly.

What lifetime access means

- The purchaser may keep and use version 1.x without a recurring fee.
- Minor corrections to version 1.x are included when Aivius distributes them.
- Future major editions, consulting, implementation, scans, and software access are separate products unless the order says otherwise.

Copyright 2026 Aivius. All rights reserved.

What this manual helps you build

By the end of the program, your team should have a versioned question set, a crawlable evidence library, consistent entity definitions, fair comparison content, an independent-source pipeline, a reproducible prompt test, and a measurement chain from discovery to qualified revenue.

The method treats GEO as a business system. Publishing is one activity inside that system. The team first identifies a decision worth influencing, then makes the relevant facts discoverable, earns corroboration, observes answers, and measures downstream behavior.

Use the playbook in four passes

1. Diagnose. Complete the baseline and score each of the twelve systems from zero to four.
2. Prioritize. Choose the earliest broken stage in the discovery-to-revenue chain.
3. Execute. Follow the 90-day program, assigning one accountable owner to each deliverable.
4. Learn. Rerun matched tests, compare business outcomes, and update the next experiment backlog.

Recommended team

Role	Accountability	Weekly input
Growth owner	Commercial scope and tradeoffs	90 minutes
Product or subject expert	Claims and evidence	60 minutes
Web owner	Technical publishing and instrumentation	60 minutes
Content or research lead	Question research and decision pages	4 to 8 hours
Sales or customer success	Buyer language and lead quality	45 minutes

The discovery to revenue chain

A brand cannot earn measurable value from an AI answer unless several distinct events occur. Public evidence must be accessible. The system must identify the entity and find relevant material. The answer may mention, cite, compare, or recommend the brand. The buyer may then visit, act, qualify, and eventually produce revenue.

Stage	Observable evidence	Typical intervention
Access	200 response, rendered text, allowed crawler	Technical repair
Understanding	Correct entity, category, and product facts	Entity alignment
Retrieval	Relevant owned or independent source appears	Evidence and source work
Answer	Mention, citation, comparison, or recommendation	Decision content
Action	Referral, CTA, trial, form, or demo	Conversion repair
Commercial	Qualified lead, pipeline, payment, gross profit	Sales and attribution

Diagnose the first weak stage

If the page is not accessible, more outreach cannot fix retrieval. If the brand is recommended but no visitor acts, more articles may not fix conversion. If referrals produce qualified leads but CRM IDs are lost, the immediate problem is attribution.

Six rules that prevent false precision

1. Keep an observed zero separate from a provider failure and an unmeasured value.
2. Do not call a mention a recommendation or a citation a cause.
3. Publish the sample size, date, providers, routes, locale, and retry policy beside every result.
4. Preserve raw answers and visible citations so an analyst can audit extraction decisions.
5. Use matched reruns after changes. Different prompts or models do not form a clean before-and-after comparison.
6. Report direct, assisted, modeled, and unknown revenue separately.

A score is a diagnostic summary

A weighted score can help a team triage a large portfolio, but the interface must retain the components. A 70 can mean broad mentions with no primary picks, a narrow set of strong recommendations, or a measurement gap. Those conditions require different action.

Value	Meaning	Allowed calculation
0	Observed absence in a successful answer	Include in valid denominator
Not measured	No approved test or data source	Exclude and display
Provider failure	The requested route did not return usable evidence	Exclude and report

Score the twelve systems

Score each system from zero to four using evidence, not confidence. Zero means absent. One means ad hoc activity. Two means a documented process exists. Three means the process runs on schedule with auditable records. Four means the team uses observed commercial outcomes to improve it.

System	0 to 4	Evidence	Next action
Commercial question map			
Category and entity definition			
Crawl and index foundation			
Evidence inventory			
Answer-ready content architecture			
Comparison and recommendation evidence			
Independent source strategy			
Prompt testing protocol			
Measurement and diagnosis			
Conversion and attribution			
Experiment and refresh loop			
Operating model and governance			

Commercial question map

Choose the questions where an AI answer can change a buying decision.

Operating principle

Demand research begins with the buyer's decision, not with a list of high-volume phrases. A useful GEO program separates learning, comparison, risk, implementation, and purchase questions because each requires different evidence.

Why this system comes before tactics

A broad prompt list creates activity without a decision model. High mention volume can hide complete absence from the questions that lead to demos, trials, or purchases.

Required output

A versioned question map with 30 to 80 prompts, owners, buyer stages, target pages, and test status.

Decision check

- What buyer decision will change if this system improves?
- What evidence would show that the current diagnosis is wrong?
- Which dependency must be resolved before production expands?
- Who can approve the facts, implementation, and measurement rule?

How to implement commercial question map

1. Interview sales, support, customer success, and product teams. Collect exact questions, objections, alternatives, constraints, and proof requests.
2. Group questions by decision stage and buyer role. Keep a distinct set for evaluators, budget owners, technical reviewers, and end users.
3. Select a small monitored set using commercial relevance, product fit, evidence availability, and the cost of being absent.
4. Record the wording, locale, date, intended answer type, and target action. Do not silently change the set between tests.

Measures

Measure	Definition	Data source	Review cadence
qualified-question coverage			
recommendation rate			
primary-pick rate			
assisted product action			

Release gate

Do not mark this system complete because a document exists. Mark it ready only when the responsible owner, source records, review date, and downstream decision are all visible.

Worksheet for commercial question map

Question	Working answer	Evidence or owner
What is the current observed gap?		
Which buyer or business decision does it affect?		
What is the smallest testable intervention?		
What is the primary measure?		
What could create a false positive?		
When will the matched result be reviewed?		
What decision follows each possible result?		

Owner signoff

Accountable owner	Contributors	Approval date	Review date

Category and entity definition

Make the company, product, category, audience, and use cases unambiguous across the web.

Operating principle

AI systems must resolve what an entity is before they can compare or recommend it. Repeated naming alone is weak evidence when the product is described differently across the home page, documentation, directories, reviews, and partner pages.

Why this system comes before tactics

A company can publish a large content library and remain difficult to recommend because the system cannot confidently determine what it sells or which buyer it serves.

Required output

An entity dictionary and a page-by-page consistency audit.

Decision check

- What buyer decision will change if this system improves?
- What evidence would show that the current diagnosis is wrong?
- Which dependency must be resolved before production expands?
- Who can approve the facts, implementation, and measurement rule?