

EXECUTIVE REPORT

Does your clinic make the AI shortlist?

An illustrative benchmark for a fictional Istanbul hair transplant clinic targeting UK patients.

SAMPLE SCOPE			
MARKET	CATEGORY	PLATFORM	DESIGN
United Kingdom	Hair transplant	ChatGPT (illustrative)	30 queries x 3 runs

ILLUSTRATIVE SAMPLE - NOT A CLIENT RESULT

No real clinic, patient, provider output, or commercial result is represented in this document.

BEFORE THE FINDINGS

A demonstration of the deliverable - not evidence of performance.

This report shows the structure, analytical depth, and decision logic a clinic would receive after a completed benchmark.

WHAT IS FICTIONAL

- Every clinic name, percentage, query result, source count, and diagnosis
- The observation window and the example competitive landscape
- The priority roadmap and effort estimates

What a real engagement would replace

01

Agreed scope

Named clinic, market, language, treatment lines, platforms, and competitors.

02

Preserved evidence

Dated provider outputs, run identifiers, classification decisions, and review notes.

03

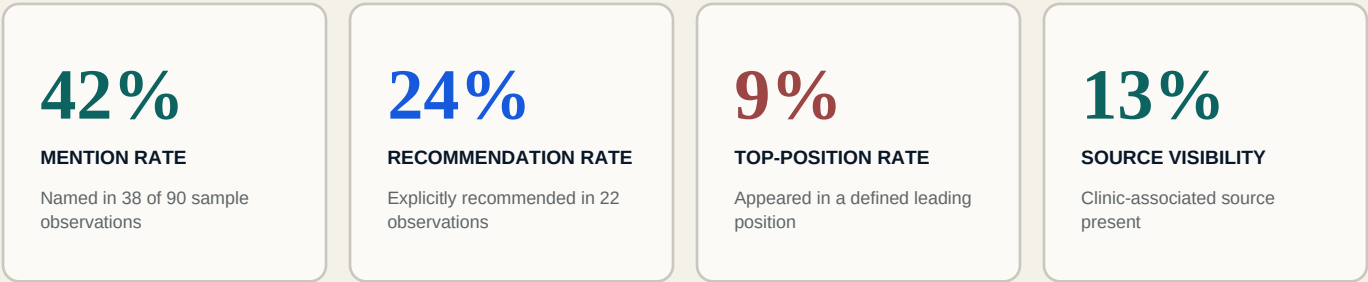
Clinic-specific diagnosis

Query losses connected to content, entities, authority, and source visibility.

SAMPLE EXECUTIVE VIEW

The clinic is known, but rarely earns a leading recommendation.

Illustrative data suggests the clinic is occasionally mentioned, yet competitors dominate explicit recommendations and supporting sources.



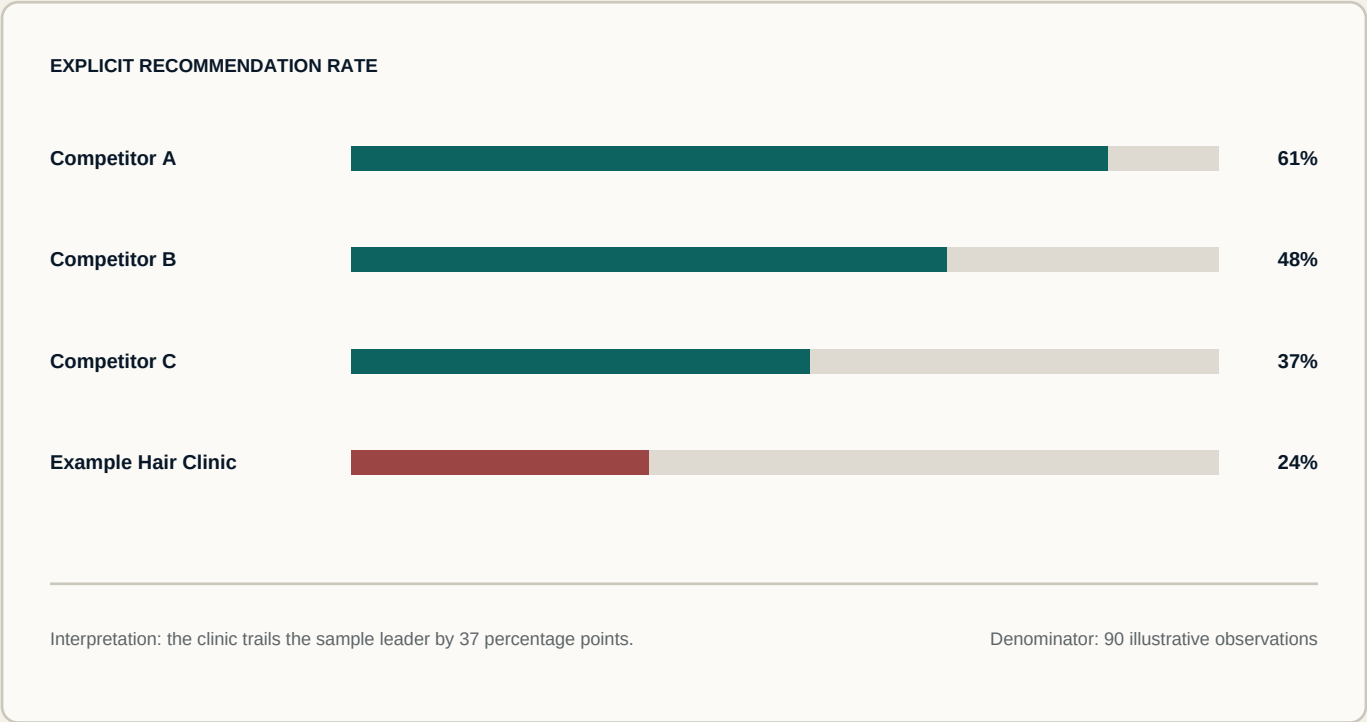
What leadership should take from the baseline

- 01 Awareness is not the main constraint**
The clinic is named often enough to confirm entity recognition. The larger gap appears after recognition: explicit recommendation and position.
- 02 Commercial queries expose the weakness**
Losses cluster around trust, physician involvement, procedure choice, aftercare, and UK-specific decision questions.
- 03 The evidence environment favors competitors**
Competitors appear beside clearer clinician evidence and more consistent third-party sources in this fictional sample.

ILLUSTRATIVE COMPARISON

Competitors win the recommendation layer, not just the mention layer.

The gap becomes clearer when the same entities are compared across separate outcomes.



Why one percentage is not enough

42%

MENTION

The clinic is recognized.

24%

RECOMMENDATION

Recognition converts inconsistently.

9%

TOP POSITION

The clinic rarely leads the shortlist.

13%

SOURCE

Supporting evidence is limited.

WHERE THE GAP CONCENTRATES

Discovery holds up. Decision and trust visibility do not.

A category view turns a single headline rate into an implementation backlog.

QUERY CATEGORY	YOUR CLINIC	BEST COMPETITOR	PRIORITY
General discovery	46%	62%	MONITOR
Procedure choice	31%	58%	PRIORITY
Physician involvement	18%	54%	CRITICAL
Trust and safety	21%	49%	CRITICAL
Aftercare	16%	45%	CRITICAL
UK patient decisions	12%	43%	CRITICAL

DIAGNOSIS

Visibility breaks down as the question becomes more specific.

The fictional pattern suggests a need for clearer clinician roles, procedure evidence, aftercare detail, and market-specific proof - not simply more generic pages.

COMMERCIALLY RELEVANT ABSENCES

A lost query is useful only when it points to a fixable evidence gap.

The sample below connects each fictional query loss to an observed competitor pattern and an implementation hypothesis.

EXAMPLE LOST QUERY	WINNERS	IMPLEMENTATION HYPOTHESIS
Best hair transplant clinic in Istanbul for UK patients	A, B	Market-specific trust evidence
Which clinics have a doctor involved in every procedure?	A, C	Clinician role and responsibility
FUE or DHI for a receding hairline?	A, B	Procedure-choice evidence
Hair transplant clinics with strong aftercare	B, C	Aftercare pathway and continuity
Safe Istanbul hair transplant clinic for first-time patients	A, C	Governed trust and safety content
What should a UK patient verify before booking?	A, B, C	Decision support and corroboration

These are fictional examples. A real report would link every row to preserved provider evidence and review notes.

FROM ANSWER TO FINDING

Every recommendation metric should be traceable to a dated response.

This fictional evidence card shows how a single observation can be reviewed without treating one answer as definitive.

OBSERVATION Q17 / RUN 02

ILLUSTRATIVE

QUERY

What are the best hair transplant clinics in Istanbul for a UK patient?

OBSERVED SHORTLIST

01	Competitor Clinic A	RECOMMENDED
02	Competitor Clinic B	RECOMMENDED
03	Competitor Clinic C	MENTIONED

EXAMPLE HAIR CLINIC WAS NOT EXPLICITLY RECOMMENDED

Classification: lost query

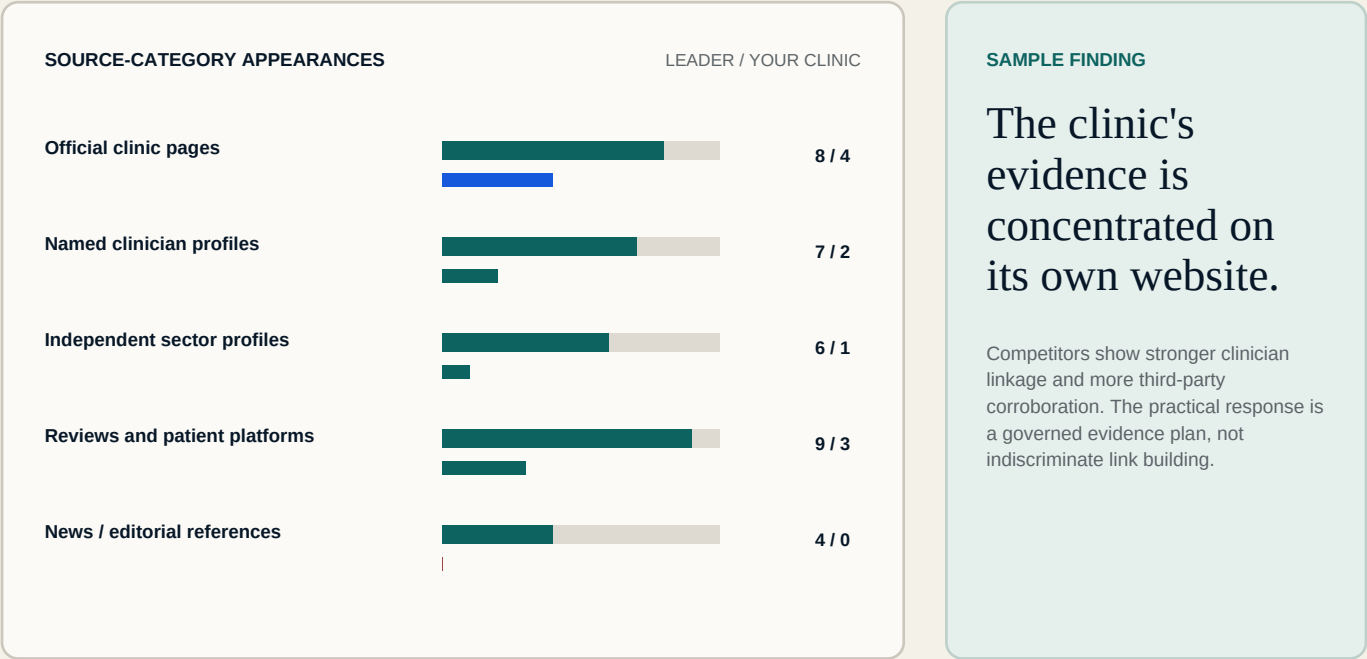
Review rule

A single observation contributes to the denominator; it never becomes the conclusion on its own. The report aggregates comparable runs and preserves exceptions.

WHAT APPEARS AROUND THE RECOMMENDATION

Competitors are supported by a broader and more consistent source environment.

This fictional source view separates clinic-owned evidence from independent corroboration.



Evidence quality controls

- Relevance to the exact clinic entity
- Named clinician and procedure consistency
- Independence and editorial context
- Date, market, and language alignment

WHY THE CLINIC LOSES

Three connected gaps explain the fictional recommendation pattern.

The benchmark ends with a diagnosis that can be converted into owned workstreams and review gates.

01

Expertise evidence

Procedure pages do not clearly connect named clinicians, roles, suitability, process, and aftercare.

PRIMARY OWNER

Content + medical review

02

Entity consistency

Clinic, doctor, procedure, and location information varies across pages, profiles, and structured data.

PRIMARY OWNER

Web + technical

03

Source authority

External profiles and independent references do not consistently corroborate the same expertise narrative.

PRIMARY OWNER

PR + partnerships

EXECUTIVE DECISION

Do not commission more undirected content.

Start with clinician-role clarity and the highest-intent lost-query clusters, then align structured data and external profiles before re-measuring.

ILLUSTRATIVE 90-DAY IMPLEMENTATION PLAN

The benchmark defines the backlog; implementation closes the highest-priority gaps.

Priorities below are fictional and would require clinic approval, medical review, access, and resourcing.

DAYS 0-30

Evidence foundation

- Confirm clinician roles and treatment ownership
- Create priority query briefs
- Repair procedure and aftercare evidence
- Define medical-review gates

DAYS 31-60

Entity consistency

- Align clinic, doctor, treatment, and location data
- Update key profiles and structured data
- Resolve conflicting public information
- Map missing corroborating sources

DAYS 61-90

Authority + QA

- Publish approved priority improvements
- Strengthen credible third-party profiles
- Complete implementation QA
- Re-run the frozen benchmark

MEASUREMENT GATE

The same query set, classification rules, market, language, and comparison scope are used for the re-benchmark. Observable change is reported without promising rankings, recommendations, patients, or revenue.

DIAGNOSIS FIRST. IMPLEMENTATION SECOND.

The benchmark tells the clinic what is happening and why. The monthly engagement does the work required to address the prioritized gaps.

01 / AI VISIBILITY BENCHMARK

Diagnosis

A defined, one-time research engagement.

- Freeze market, language, treatments, and competitors
- Collect and classify comparable observations
- Identify lost queries and source gaps
- Deliver executive diagnosis and prioritized roadmap

COMMERCIAL QUESTION

Decision: what should we fix first?

02 / MONTHLY IMPLEMENTATION

Execution

An ongoing, evidence-led improvement engagement.

- Build clinician and procedure evidence
- Align entities, profiles, and structured data
- Strengthen credible supporting sources
- Coordinate reviews, QA, and comparable re-measurement

COMMERCIAL QUESTION

Decision: who owns the work and how is progress measured?

TRANSPARENT BY DESIGN

A benchmark is a dated sample with explicit rules - not a permanent AI score.

A real report records the provider, model where available, date, market, language, query set, run count, classification logic, exclusions, and evidence archive.

01

Freeze scope

Clinic, market, language, treatments, competitors, platforms, and query families.

02

Collect independently

Repeat the same frozen queries without carrying conversation context between observations.

03

Classify evidence

Keep mentions, explicit recommendations, top positions, and supporting sources separate.

04

Review ambiguity

Manually review unknown entities, failures, partial answers, and unclear recommendation language.

05

Report denominators

Show counts, rates, exclusions, limitations, and preserved examples for auditability.

06

Re-measure comparably

Repeat the agreed scope after implementation to observe change under similar conditions.

LIMITATIONS

• AI outputs vary by provider, model, date, language, context, and availability.

• No ranking, recommendation, patient-volume, or revenue outcome is guaranteed.

• This service does not assess treatment quality, patient suitability, or clinical superiority.

• This sample contains no real clinic data and no live AI observations.

Start with evidence.

hello@medicailgrowth.com | medicailgrowth.com