
METHODOLOGIE · AUSGABE 2026

Wie aus einer *Messung* Sichtbarkeit wird.

Die GEO-Methodik von Bavaria AI — wie wir Marken in ChatGPT, Perplexity, Google AI Overviews und Gemini auffindbar machen. Erst diagnostizieren, dann das Fundament legen, dann wachsen.

Marken sterben nicht durch ~~Klicks~~, sondern durch *Stille*.

Die Suche verändert sich von der Liste blauer Links zur fertigen Antwort. Wer heute wissen will, welche Software, welche Kanzlei oder welchen Dienstleister er wählen soll, fragt zunehmend ChatGPT, Perplexity oder Google AI Overviews — und bekommt eine Empfehlung, keine zehn Treffer. In dieser Antwort kommt eine Handvoll Marken vor. Der Rest existiert für den Käufer in diesem Moment nicht.

Das ist die eigentliche Verschiebung: **Sichtbarkeit heißt nicht mehr ranken, sondern zitiert werden.** Ein Unternehmen kann bekannt sein, Pressecoverage haben und trotzdem in der kaufentscheidenden KI-Antwort fehlen — weil die Modelle es an dieser Stelle nicht als relevante Option führen. Klassisches SEO optimiert für die Position in einer Trefferliste. *Generative Engine Optimization* (GEO) optimiert dafür, in der generierten Antwort selbst vorzukommen.

Diese Bewegung ist keine ferne Prognose, sondern messbar im Gange. Die Zahlen darunter ordnen ein, warum Marken jetzt handeln — und warum GEO Zeit braucht, kein Schalter ist.

51%

der B2B-Software-Käufer
starten ihre Recherche heute mit
KI-Chatbots

G2 · „THE ANSWER ECONOMY“ ·
2026

+91%

höhere Paid-Click-Rate bei
Marken, die in *Google AI
Overviews* zitiert werden

SEER INTERACTIVE ·
SEPTEMBER 2025

4–12

Monate, bis *Suchmaschinen-
Arbeit* messbar wirkt —
Autorität entsteht über Zeit

GOOGLE SEARCH CENTRAL

Alle externen Kennzahlen sind als Quelle und Stichtag ausgewiesen. Sie beschreiben den Markt, nicht ein zugesichertes Ergebnis für ein einzelnes Mandat. KI-Antworten schwanken und ändern sich laufend.

Erst *messen*. Dann das Fundament legen. Dann wachsen.

Jedes Projekt folgt demselben Weg. Wir beginnen mit einer belastbaren Diagnose, bauen darauf ein maschinenlesbares Fundament und heben im laufenden Mandat kontinuierlich die Autorität. Alle Leistungen sind als Tätigkeit beschrieben — kein bestimmtes Ergebnis wird garantiert.

I

DIAGNOSE

Messung über *vier Engines*

GEO-Score über vier KI-Engines und vier Dimensionen, mit Gate-Status und priorisierter Roadmap. Der Ausgangspunkt in Zahlen — an einem dokumentierten Stichtag.

II

FUNDAMENT

Entität und *Substanz*

Schema und sameAs, Wikidata und Knowledge Panel, Kennzahlen als extrahierbarer Text. Die Grundlage, auf der KI-Systeme eine Marke überhaupt zitieren können.

III

WACHSTUM

Laufende *Umsetzung*

Content, Earned Media, Review-Profile und Entitäts-Pflege — plus immer-an-Messung über alle vier Engines. Monat für Monat, in kalibrierter Frequenz.



RESULTAT

Belegbare *KI-Präsenz*

Von der ersten Nennung zur konsistenten Zitation — transparent gemessen, nachvollziehbar dokumentiert, ehrlich in der Roadmap.

Der rote Faden ist eine Haltung, keine Tool-Liste: **Wir behaupten nicht, wir messen.** Jede Zahl gehört zu einem Stichtag. Jeder Hebel ist als Tätigkeit definiert. Und wir sagen offen, was wir bewusst weglassen. Die folgenden Kapitel gehen die drei Phasen der Reihe nach durch.



Wir messen, *bevor* wir bauen.

Am Anfang steht kein Maßnahmenpaket, sondern eine Messung. Wir erheben den **GEO-Score** über vier KI-Engines — *ChatGPT*, *Perplexity*, *Google AI Overviews* und *Gemini* — und über vier Dimensionen. Dazu läuft ein eingefrorenes, dokumentiertes Prompt-Set live gegen die Engines, das die tatsächliche Sichtbarkeit gegen die kaufentscheidenden Fragen der Kategorie prüft. Ergänzt um einen Share-of-Model-Snapshot gegen zwei bis drei Wettbewerber entsteht ein Bild in Zahlen: Wo steht die Marke heute, an diesem Stichtag, in KI-Antworten?

Der Score ist kein Selbstzweck. Er übersetzt sich in eine priorisierte Roadmap über drei Horizonte — Sofortmaßnahmen, 30 Tage, 90 Tage —, damit klar ist, welcher Hebel zuerst zählt und welcher warten kann.

LIEFERGEGENSTÄNDE · GEO-AUDIT

Stichtags-Snapshot

- › **Audit-Bericht** mit Score je Dimension, Gate-Status und Wettbewerbs-Snapshot
- › **SoV-Baseline** über die vier Engines mit dokumentiertem Prompt-Set
- › **Entitäts- und Autoritäts-Befund** (Wikidata, Knowledge Panel, Review-Präsenz)
- › **Priorisierte Roadmap** über drei Horizonte — inklusive „was wir bewusst weglassen“
- › **Strategiegespräch** zur Einordnung, Founder-geführt

01

Content

ANSWER-FIRST ·
EXTRAHIERBARE CLAIMS

02

Technik

SCHEMA · CRAWLBARKEIT
· LLMS.TXT

03

Autorität

ENTITÄT · EARNED
MEDIA · REVIEWS

04

Sichtbarkeit

CITATION · SHARE OF
MODEL

Ein Score, der *nicht schönfärbt*.

Nicht jede Nennung in einer KI-Antwort ist gleich viel wert. Manche Zitationen stehen auf belegbarem Fundament, andere auf einer Quelle, die morgen verschwunden sein kann. Unser Score trennt beides — sichtbar.

Wir klassifizieren jede gefundene Citation nach ihrer Belastbarkeit in **Trust-Tiers**. Daraus ergeben sich zwei Werte statt einer geschönten Gesamtzahl: der *Total-Score* zählt alle Nennungen, der *Durable-Score* nur die reproduzierbar belegbaren. Die Lücke zwischen beiden ist die ehrlichste Kennzahl im ganzen Audit — sie zeigt, wie viel der aktuellen Sichtbarkeit auf Sand gebaut ist.

TIER 1 · DURABLE

Reproduzierbar belegbar. Zitationen, die auf stabilen, autoritativen Quellen stehen und über Messläufe hinweg wiederkehren. Nur diese zählen in den *Durable-Score*.

TIER 2 & 3 · VOLATIL

Nennungen auf schwächeren oder wechselnden Quellen. Zählen in den Total-Score, aber nicht als belastbare Autorität — Kandidaten, die zu Tier 1 gehärtet werden müssen.

CLIFF-EDGE

Vorübergehende Quellen, die Engines heute noch zitieren, aber bald entfernen. Sichtbarkeit auf Abruf — wir markieren sie als Risiko, statt sie mitzuzählen.

*Dazu kommen **harte Gate-Deckel**: Fehlt das Fundament in einer Dimension, deckelt das den effektiven Score — unabhängig davon, wie gut die anderen Dimensionen aussehen. So verhindern wir, dass ein hoher Substanz-Wert eine strukturelle Lücke überstrahlt.*

Das Ergebnis ist ein Score, dem man trauen kann, weil er die eigene Ausgangslage eher unter- als überschätzt. Er ist reproduzierbar: Dasselbe Prompt-Set, dieselbe Methode, jederzeit wiederholbar. Und er ist ehrlich datiert — ein Stichtag, kein Trend.

II

PHASE II · FUNDAMENT

Substanz, bevor Reichweite.

KI-Systeme zitieren, was sie verstehen. Bevor eine Marke Reichweite aufbaut, muss sie für Maschinen eindeutig lesbar sein: als klar definierte **Entität** und mit Inhalten, aus denen sich Aussagen sauber extrahieren lassen. Das ist die Arbeit der zweiten Phase — im Implementation Sprint gebündelt, einbau-fertig geliefert.

Wir konsolidieren die Entitäts-Signale (konsistente *sameAs*-Ziele über die zentralen Profile), bereiten den Wikidata-Eintrag vor und richten die relevanten Review-Profile ein. Parallel überführen wir die entscheidenden Kennzahlen aus Grafiken in extrahierbaren Text — denn was nur im Bild steht, ist für ein Sprachmodell nicht vorhanden. Darauf setzt eine Content-Hub-Architektur mit ersten Answer-First-Seiten auf.

LIEFERGEGENSTÄNDE · IMPLEMENTATION SPRINT

einbau-fertig

- › **Konsolidierungs-Konzept** für Domains und Marken
- › **Wikidata-Antrag** plus dokumentierte *sameAs*-Liste
- › **Review-Profile** aufgesetzt und übergeben (z.B. OMR, Capterra, G2)
- › **Schema-Snippets und llms.txt**, validiert, plus Quality-Fix-Liste
- › **Content-Hub-Struktur** mit Answer-First-Seiten

Ehrliche Einordnung. Schema und llms.txt sind enthaltene Fundament-*Hygiene* — notwendige Grundlage, aber kein magischer Hebel und kein alleiniger Treiber der KI-Sichtbarkeit. Wir liefern Content und Code einbau-fertig plus Anleitung; das Deployment im CMS macht der Kunde (oder wir gegen gesonderte Vereinbarung). Und die Aufnahme in Wikidata entscheidet die Community, nicht wir.

III

PHASE III · WACHSTUM

Autorität entsteht über Zeit.

Sichtbarkeit in generativen Engines ist kein einmaliges Projekt, sondern eine Kadenz. Im laufenden Mandat arbeiten wir an vier Hebeln — flankiert von einer immer-an-Messung, die den Fortschritt gegen die Baseline sichtbar hält. Wie viele Hebel parallel laufen und in welcher Frequenz, hängt von der Stufe ab.

Hebel 01

Content

Answer-First-Produktion, einbau-fertig, konsequent auf Zitierbarkeit ausgelegt.

Hebel 02

Earned Media & Digital PR

Client-fronted vorbereitet: Story, Daten-Hook, Presselisten, Pitches. Der Kunde tritt als Quelle auf.

Hebel 03

Reviews

Aktivierung und Pflege der Review-Profile — Volumen anstoßen, Aktualität halten.

Hebel 04

Entität & Off-Page

sameAs-Konsistenz, Knowledge-Panel-Propagation, Verzeichnisse und Off-Page-Präsenz.

Dazu als Baseline, nicht als Hebel: die **immer-an-Messung** — SoV-Tracking gegen die dokumentierte Baseline plus regelmäßiger Bericht. Drei Stufen skalieren Umfang und Frequenz:

Foundation

Ein Hebel im Fokus, Baseline-Messung, monatlicher Call.

Momentum

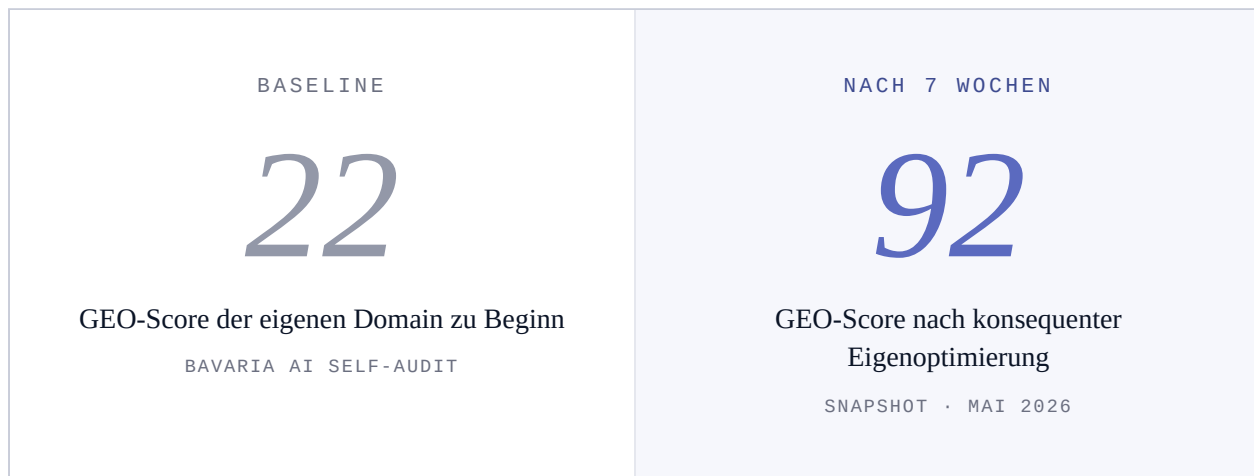
Mehrere Hebel parallel, höhere Frequenz, laufende Share-of-Model-Analyse.

Authority

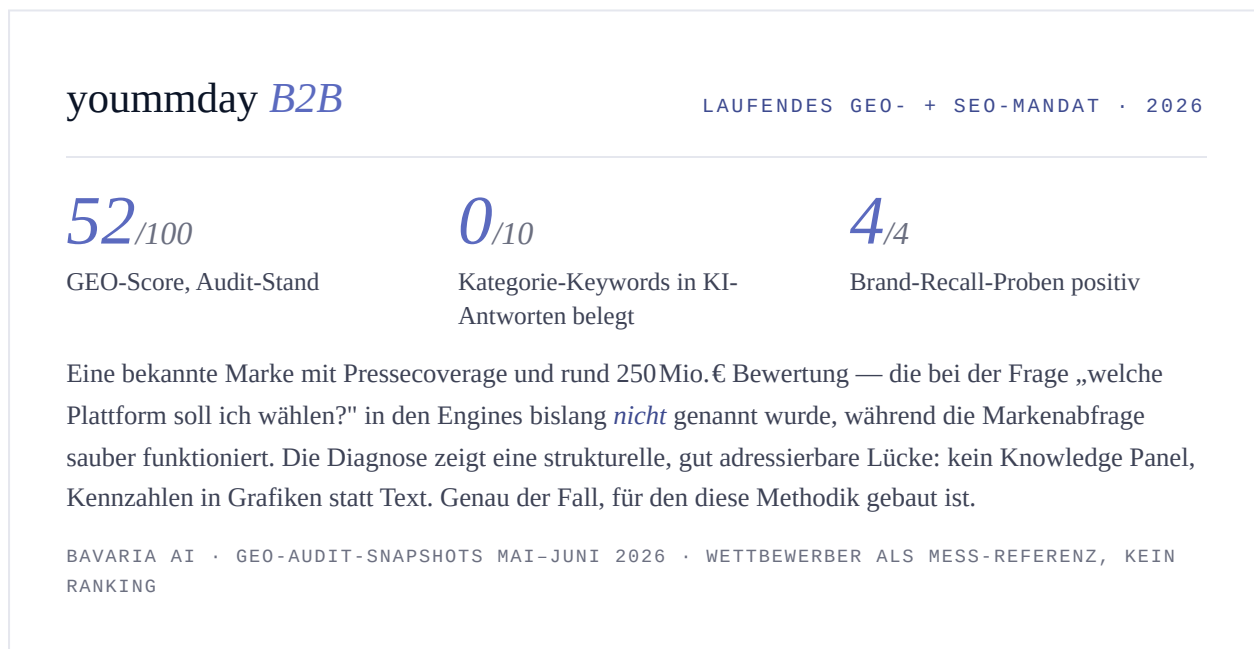
Alle Hebel, höchste Frequenz, tiefste Analyse, fester Strategie.

Zahlen mit *Datum*. Ein Stichtag, kein Trend.

Die Methodik zeigt sich am eigenen Haus und an Mandaten — jede Zahl mit dokumentiertem Stichtag, als erbrachte Tätigkeit, nicht als garantiertes Ergebnis.



Der belastbarste Beweis ist die eigene Marke: Bavaria AI hat die eigene Methodik zuerst auf sich selbst angewandt und den GEO-Score in sieben Wochen von **22 auf 92** gehoben — der einzige Case mit vollständiger, öffentlich nachvollziehbarer Before/After-Datenlage.



Wettbewerber dienen ausschließlich als objektive Mess-Referenz, nicht wertend. Alle GEO-Scores sind zeitpunktbezogene Momentaufnahmen über vier KI-Engines — KI-Antworten schwanken laufend.

Tätigkeit, kein Erfolg. Und wir sagen, was wir weglassen.

GEO arbeitet gegen Black-Box-Systeme. Wer hier Rankings garantiert, verkauft eine Illusion. Wir schulden die fachgerechte, sorgfältige Bemühung — und sind transparent an ihren Grenzen.

WAS WIR BEWUSST NICHT TUN

- × Keine garantierten Rankings oder Citations — LLMs sind Black Boxes
- × Keine gekauften Backlinks, kein Massen-Outreach in Ihrem Namen
- × Keine Wikipedia-Auftrags-Einträge ohne echte Notability
- × Keine Phantom-Standort-Einträge ohne physische Präsenz
- × Kein Schönfärben — volatile Sichtbarkeit zählen wir nicht als Autorität

WORAUF SIE SICH VERLASSEN KÖNNEN

- ✓ Eine **dokumentierte, reproduzierbare** Methodik statt Bauchgefühl
- ✓ Jede Zahl mit **Snapshot-Datum** — ein Stichtag, kein Trend
- ✓ Deliverables **einbau-fertig** plus Anleitung, in Ihrer Hand
- ✓ Eine ehrliche Roadmap — inklusive „*was wir weglassen*“
- ✓ Founder-geführt, ohne Account-Manager-Schicht dazwischen

Für Earned Media und Digital PR gilt: Wir recherchieren, entwickeln den Daten-Hook, schreiben die Pitches und bereiten Sie als Quelle vor — **den Versand und den Auftritt als Quelle übernehmen aber Sie**, nicht wir. Diese Grenze ist kein Detail, sondern Teil der Methode: Autorität, die zählt, muss echt sein.

— IHRE AUSGANGSLAGE

Wo stehen Sie *heute* in KI-Antworten?

Ein GEO-Audit über vier Engines zeigt Ihre aktuelle Sichtbarkeit, den Abstand zur Kategorie und die größten Hebel. Transparent, belegt, ohne Erfolgsversprechen.

KOSTENLOSES ERSTGESPRÄCH • 30 MINUTEN

calendly.com/bavaria-ai-info/30min

Ein 30-minütiges Erstgespräch mit kostenlosem Mini-Audit. Mehr zur vollständigen Methodik und den Leistungspaketen auf bavaria-ai.com — oder in unseren [Case Studies](#).

CHATGPT • PERPLEXITY • GOOGLE AI OVERVIEWS • GEMINI • *Mia san KI.*