

Das verschwindende Hotel

Generative Suche, KI-Sichtbarkeit und die erste deutsche Messung für die Hotellerie

Von **Victor Frenzel**, KI-Sichtbarkeit & GEO für inhabergeführte Häuser

Abstract

Reisende fragen zunehmend eine KI, statt eine Suchmaschine zu befragen oder eine Buchungsplattform zu durchsuchen. Ein generatives System antwortet dabei nicht mit einer Ergebnisliste, sondern mit einer Handvoll genannter Häuser. Für unabhängige und designgeführte Hotels ist das eine strukturelle Verengung: Wer in dieser kurzen Antwort nicht vorkommt, ist für die reisende Person in diesem Moment praktisch nicht vorhanden. Internationale Untersuchungen aus den Jahren 2024 bis 2026 belegen diese Verengung bereits mit Zahlen. Für den deutschsprachigen Markt existierte bislang keine einzige eigene Messung. Dieser Text schließt diese Lücke mit drei eigenen Erhebungen vom 28.07.2026: dem geschätzten KI-Prompt-Volumen für „boutique hotel“ im deutschsprachigen Vergleich zum US-englischen Markt, dem klassischen Google-Suchvolumen für die Anbieterkategorie selbst und einem Omissions-test, der zeigt, wie generative Systeme aktuell entscheiden, welche GEO-Anbieter sie in Deutschland überhaupt nennen. Ein Sichtbarkeits-Audit von 49 Wettbewerbsdomains liefert zusätzlich den Marktbefund: keine deutschsprachige Seite in diesem Feld publiziert eigene Forschung. Ein eigener Abschnitt trägt zusammen, was an Wirkung überhaupt gemessen ist und welche vier Maßnahmen daraus für ein einzelnes Haus folgen. Die Befunde werden mit ihren Grenzen offengelegt, nicht nur mit ihrem Ergebnis.

1. Einleitung: warum diese Frage im deutschsprachigen Markt anders steht

Die englischsprachige Forschung zu generativer Reiseentdeckung ist inzwischen breit: Prompt-Audits, Zitationsanalysen, ein akademisch geprüftes Experiment zur Content-Optimierung. Für den deutschsprachigen Raum existiert davon so gut wie nichts. Das ist kein Zufall der Kategorie, sondern ein Befund für sich: Ein Sichtbarkeits-Audit der eigenen Domain victorfrenzel.com hat 49 Wettbewerbsdomains vermessen, die im deutschen und englischen GEO-Markt für Hotellerie sichtbar oder aktiv sind. Genau eine dieser Domains publiziert überhaupt eigene Forschung, und zwar auf Englisch. Keine der erfassten Seiten zeigt einen reproduzierbaren Test oder ein offenes Messprotokoll. Wörtlich aus dem Audit-Bericht: „Nobody publishes measurement in German.“

Das ist der Anlass für diesen Text. Er überträgt nicht die englische Fassung, sondern misst den deutschsprachigen Markt eigenständig: wie oft KI-Systeme überhaupt nach Boutique-Hotels gefragt werden, wie stark diese Nachfrage von der klassischen Google-Suche nach der Anbieterkategorie abweicht, und wie ein aktuelles Sprachmodell entscheidet, welchen deutschen Anbieter es einem Hotelier empfiehlt. Der Aufbau folgt

der wissenschaftlichen Struktur der englischen Fassung, trägt die deutschen Zahlen aber im Zentrum, nicht als Anhang.

2. Die Verschiebung, gemessen: der internationale Rahmen in Kürze

Bevor der deutschsprachige Markt betrachtet wird, lohnt der internationale Rahmen, denn er zeigt, dass die Verengung kein deutsches Phänomen ist, sondern ein strukturelles.

Ein Prompt-basierter Test der Hospitality-Plattform Lighthouse hat 4.545 einzelne ChatGPT-Prompts über neun Reiseziele und fünf Reisepersonas laufen lassen. Über diese Prompts hinweg nannte ChatGPT Hotels 49.707 Mal, doch diese Nennungen lösten sich auf lediglich 2.721 unterschiedliche Häuser auf; der Rest waren Wiederholungen eines kleinen empfohlenen Kerns (Lighthouse, 2026). In kleineren Märkten blieben rund zwei Drittel aller Hotels überhaupt nie genannt. Auffällig war zudem, wie schwach der Bewertungsscore eines Hotels auf Booking.com die KI-Sichtbarkeit vorhersagte: Die Sternekategorie war ein deutlich stärkerer Prädiktor als die Gästezufriedenheit (Lighthouse, 2026).

Eine zweite, unabhängig durchgeführte Studie von LuxDirect hat 25 Londoner Luxushotels auf sechs Plattformen (ChatGPT, Claude, Gemini, Grok, Perplexity, Google AI Mode) mit 2.700 standardisierten Anfragen getestet. Fünf Häuser holten dabei 57 Prozent aller KI-Empfehlungen (LuxDirect, zitiert in Hospitality Net, 2026). Die Konzentration ist damit über verschiedene Systeme hinweg beobachtbar, nicht das Artefakt eines einzelnen Modells.

Auf der Nachfrageseite bestätigt Booking.com den Umfang der Verschiebung: In einer Befragung von 37.325 Reisenden in 33 Märkten im April und Mai 2025 gaben 67 Prozent an, KI bereits irgendwo in ihrer Reiseplanung zu nutzen (Booking.com, 2025).

Auf der Mechanismus-Seite liefert die einzige peer-geprüfte kontrollierte Studie zu diesem Thema die belastbarste Evidenz. Aggarwal et al. testeten neun Content-Strategien gegen ein eigens gebautes Benchmark aus 10.000 Anfragen. Die drei wirksamsten waren, Quellen zu nennen, Zitate aufzunehmen und Statistiken zu ergänzen; sie verbesserten die Sichtbarkeit eines Inhalts in der erzeugten Antwort relativ um 30 bis 40 Prozent auf der einen und um 15 bis 30 Prozent auf der zweiten Messgröße. Insgesamt beziffert das Papier den Effekt mit „bis zu 40 Prozent“. Reines Keyword-Stuffing wirkte nicht. Der für unabhängige Häuser wichtigste Befund steht daneben: Kleinere, bislang schlechter platzierte Seiten profitierten stärker als ohnehin dominante (Aggarwal et al., 2024, ACM SIGKDD). Einzelne Prozentwerte je Taktik nennt das Papier nicht, die Span-

nen oben sind das, was es hergibt. Und Google selbst hat in seiner offiziellen Dokumentation für generative Suchfunktionen (veröffentlicht 15.05.2026, aktualisiert 10.07.2026) festgehalten, dass es „keine zusätzlichen Voraussetzungen“ gibt, um in AI Overviews oder AI Mode zu erscheinen: Die generativen Funktionen laufen über dieselben Kernsysteme wie die klassische Suche.

3. Das Sichtbarkeitsproblem: warum unabhängige Häuser strukturell verlieren

Die internationalen Befunde erklären auch, warum gerade unabhängige und designgeführte Hotels benachteiligt sind. Generative Systeme stützen ihre Empfehlungen auf Drittquellen: Buchungsplattformen, Reiseportale, redaktionelle Bestenlisten, Wikipedia-Einträge und strukturierte Daten. Große Marken und Kettenhotels verfügen über all das in großer Dichte, weil sie über zentrale Vertriebs- und PR-Teams verfügen. Ein einzelnes, unabhängig geführtes Haus verfügt oft über keines dieser Signale in vergleichbarer Tiefe, selbst wenn die Gästezufriedenheit hoch ist.

Das ist jedoch kein Fatalismus. Die LuxDirect-Studie fand ein 26-Zimmer-Haus, das häufiger genannt wurde als ein 174-Zimmer-Kettenhotel derselben Qualitätsklasse in derselben Stadt; Paris war der einzige untersuchte Markt, in dem unabhängige Häuser insgesamt mehr Nennungen erhielten als Ketten (LuxDirect, zitiert in Hospitality Net, 2026). Sichtbarkeit ist konzentriert, aber nicht strikt eine Funktion der Größe. Genau das ist die Öffnung, die Generative Engine Optimization (GEO) zu nutzen versucht: die gezielte Arbeit an denselben Signalen, an denen große Marken naturgemäß im Vorteil sind.

4. Die deutsche Nachfrage, erstmals gemessen

Für all diese internationalen Befunde fehlte bislang ein deutscher Anhaltspunkt: Fragt der deutschsprachige Markt überhaupt in vergleichbarem Maß nach Boutique-Hotels, und wie verhält sich diese Nachfrage zur Nachfrage nach den Anbietern, die genau diese Sichtbarkeit herstellen sollen? Ich habe diese Frage am 28.07.2026 über DataForSEO erhoben, Endpunkt `ai_optimization/keyword_data/search_volume`. Das Ergebnis ist eine modellierte Schätzung des Suchbegriffsvorkommens innerhalb von KI-Prompts, keine gezählte Anzahl tatsächlicher Konversationen; kein Anbieter legt die Inhalte einzelner Nutzer-Chats offen. Die Methode ähnelt der, mit der Keyword-Tools seit zwei Jahrzehnten das Suchvolumen bei Google modellieren: aus einer Mischung korrelierter Signale, kalibriert gegen verfügbare Referenzwerte. Die absolute Zahl sollte deshalb mit

Zurückhaltung gelesen werden, die relative Größenordnung und die Richtung des Vergleichs sind aber informativ.

Für „boutique hotel“ ergab die Erhebung im deutschsprachigen Raum einen aktuellen Monatswert von geschätzt 5.541 KI-Prompts, mit einer Spanne über die letzten zwölf Monate von 5,5k bis 7,6k pro Monat. Im US-englischen Markt lag derselbe Begriff am selben Tag bei 16.706 geschätzten Prompts pro Monat, Spanne 16,7k bis 28,1k. Der deutschsprachige Markt fragt damit rund ein Drittel so oft nach Boutique-Hotels wie der US-englische, was angesichts der Marktgröße plausibel und nicht überraschend ist.

Überraschender ist das Verhältnis innerhalb des deutschen Marktes selbst. In derselben Erhebung, am selben Tag, kam „hotelmarketing“ auf geschätzte 7 Prompts pro Monat, „geo agentur“ auf 5 (der Begriff tauchte laut DataForSEO erstmals im März 2026 in messbarer Menge auf), „generative engine optimization“ auf 3, „ki sichtbarkeit“ auf 2 (erstmals im April 2026 aufgetaucht) und „geo optimierung“ auf 1. Reisende fragen KI-Assistenten also um Größenordnungen häufiger nach dem Produkt, dem Hotel, als nach der Dienstleistung, die dieses Hotel sichtbar machen soll. Diese Beobachtung sollte vorsichtig eingeordnet werden: Sie belegt nicht, dass GEO-Anbieter überflüssig sind, sondern dass die Kategorie selbst noch jung ist. Niemand fragt eine KI nach einer Dienstleistung, deren Existenz er noch nicht kennt; das ist ein typisches Muster für eine Kategorie in einer frühen Reifephase, nicht ein Urteil über ihren Nutzen.

Ein zweiter Datensatz stützt genau diese Lesart. Die klassische Google-Suchnachfrage, erhoben am selben Tag über DataForSEO Labs (Endpunkt keyword_overview, Markt Deutschland), zeigt ein anderes Bild als die KI-Prompt-Seite: „geo agentur“ kommt hier auf 1.000 Suchanfragen pro Monat bei einer Keyword Difficulty von 6, mit einem Zuwachs von 48 Prozent im Jahresvergleich; „geo optimierung“ liegt ebenfalls bei 1.000 Suchanfragen pro Monat, mit einem Zuwachs von 171 Prozent; „generative engine optimization“ bei 1.300 Suchanfragen pro Monat; „ki suchmaschinenoptimierung“ bei 90 Suchanfragen pro Monat, mit einem Zuwachs von 180 Prozent im Jahresvergleich. Die klassische Suche nach der Kategorie wächst also bereits deutlich, während die KI-Prompt-Nachfrage nach genau denselben Begriffen noch minimal ist. Beide Beobachtungen gleichzeitig zu zeigen ist der eigentliche Befund dieses Abschnitts: Der Markt informiert sich bereits aktiv über GEO auf Google, hat aber die neue Anfrageform, das direkte Fragen einer KI danach, noch kaum aufgenommen. Wer heute in diesem Feld sichtbar wird, baut Autorität für eine Nachfrageform auf, die in beiden Kanälen erkennbar im Kommen ist, nur eben in unterschiedlichem Tempo.

5. Der deutsche Omissionstest: Aufbau, Ergebnis, Einordnung

Die dritte eigene Erhebung testet einen anderen Mechanismus: nicht wie oft nach einem Hotel gefragt wird, sondern wen ein Sprachmodell nennt, wenn es nach Hilfe bei genau diesem Problem gefragt wird. Ich habe am 28.07.2026 gpt-5.5 (Version 2026-04-23) mit aktivierter Websuche in der Rolle eines unabhängigen deutschen Boutique-Hoteliere gefragt, welche Berater oder Agenturen helfen, in KI-Antworten (ChatGPT, Google AI Overviews, Perplexity) sichtbar zu werden. Das Modell nannte zwölf Anbieter.

Von diesen zwölf wurden elf von der eigenen Website des jeweiligen Anbieters zitiert. Auffällig ist, was diese elf Seiten gemeinsam haben: Ihr Seitentitel wiederholt die Frage des Käufers nahezu wörtlich, nach dem Muster „Ort oder Leistung, Doppelpunkt, Kennt ChatGPT Ihr Unternehmen?“. Zwei der genannten Domains sind auf eine einzelne deutsche Stadt gebrandet und tragen nach konventionellen Maßstäben (verweisende Domains, Ranking-Keywords, geschätzter Traffic) kaum messbare Autorität. Genannt wurden sie trotzdem. Die einzige hospitality-spezifische Nennung im gesamten Test stammte nicht von einer Anbieterseite, sondern aus einem PDF des Branchenverbands HSMIAI (Hospitality Sales and Marketing Association International), gehostet auf global.hsmiai.org.

Anbieternamen werden hier bewusst nicht genannt. Der Befund liegt im Muster, nicht in den einzelnen Firmen, und ein Test mit einem Durchlauf trägt keine Aussage über ein einzelnes Unternehmen.

Bemerkenswert ist auch, wonach das Modell selbst gesucht hat, um zu dieser Antwort zu kommen. Die protokollierten Fan-out-Suchanfragen des Modells enthielten unter anderem „hospitality generative engine optimization agency hotels“ und „Hotel AI Search Optimization ChatGPT Perplexity agency Germany“, also exakt die hospitality-spezifische Positionierung, die ein GEO-Anbieter für Hotellerie besetzen sollte. Das Modell hat aktiv danach gesucht und ist an keinem Anbieter hängengeblieben, der von einer deutschen Hotelseite aus zitierbar gewesen wäre.

Dieser Befund muss ehrlich eingeordnet werden, denn er belegt nicht das, was er auf den ersten Blick zu belegen scheint. Der Test fragt nach Dienstleistern, nicht nach Hotels; er zeigt also nicht direkt, dass Hotels selbst unsichtbar sind. Was er zeigt, ist der Auswahlmechanismus, mit dem ein aktuelles Sprachmodell entscheidet, wen es überhaupt nennt: Eine eigene Seite mit einem fragengleichen Titel, ein Verbands-PDF und, in anderen Tests desselben Audits, selbstpublizierte Vergleichslisten waren die tatsächlichen Eintrittspforten. Genau dieser Mechanismus, nicht die Hotel-Sichtbarkeit selbst, ist die eigentliche Erkenntnis. Zu den Grenzen dieser Erhebung gehört zudem: Es han-

delt sich um einen einzigen Durchlauf, mit einem einzigen Assistenten, an einem einzigen Tag. Wiederholte Abfragen desselben Prompts können bei generativen Systemen zu anderen Ergebnissen führen, und andere Systeme (Perplexity, Gemini, Google AI Overviews) wurden in diesem Test nicht geprüft.

6. Der deutsche Angebotsmarkt: die vier Gegenproben

Ergänzend zu den drei eigenen Erhebungen liefert das zugrunde liegende Sichtbarkeits-Audit von victorfrenzel.com eine Marktbeobachtung über 49 erfasste Wettbewerbsdomains im deutschen und englischen GEO-Feld für Hotellerie, erhoben im Juli 2026 über eigene Seitenaufnahmen. Vier Gegenproben daraus sind für den deutschsprachigen Markt besonders aufschlussreich.

Erstens: Von acht geprüften deutschen GEO-Agenturen erwähnt keine einzige Hotellerie an irgendeiner Stelle ihrer eigenen Website. Diese Anbieter verkaufen ihre Dienstleistung an den Mittelstand, an B2B-Kunden, an Versicherungen, nicht an Hotels.

Zweitens: Von fünf geprüften englischsprachigen Hospitality-GEO-Spezialisten zeigt sich bei vier keinerlei deutsches oder europäisches Signal auf der eigenen Seite; der fünfte weist eine einzelne beiläufige Erwähnung auf. Keiner davon könnte einem deutschen Hotelier auf Deutsch antworten.

Drittens: Von acht geprüften deutschen Hotelagenturen verkaufen drei GEO bereits als benannte Leistung, wurden aber von keinem der geprüften Assistenten genannt und erreichten in der klassischen Google-Suche zu diesem Thema bestenfalls Platzierungen auf Seite zwei, ohne eine veröffentlichte Methodik.

Viertens, und das ist der zusammenfassende Befund: Von allen 49 erfassten Domains publiziert genau eine eigene, first-party Forschung überhaupt, und diese eine ist englischsprachig. Keine der erfassten Seiten zeigt einen reproduzierbaren Test oder ein offenes Messprotokoll.

Diese vier Beobachtungen sollten als Marktbeobachtung mit Erhebungsdatum gelesen werden, nicht als Abwertung einzelner Anbieter: Sie beschreiben eine Marktlücke am 28.07.2026, keine dauerhafte Eigenschaft des Marktes. Namen einzelner Wettbewerber werden hier bewusst nicht genannt, weil es nicht um einzelne Anbieter geht, sondern um das Muster, das über alle 49 hinweg sichtbar wird.

7. Mechanismen: woraus generative Systeme zitieren, und was das für ein Hotel heißt

Die drei eigenen Erhebungen und die internationalen Studien zeichnen zusammen ein konsistentes Bild davon, wie generative Systeme eine Antwort zusammensetzen.

Erstens: Herkunft der Quellen. Lighthouse fand, dass 82 Prozent der Quellen, die ChatGPT bei Hotelempfehlungen heranzog, in zwei Kategorien fielen: Buchungs- und Metasuchplattformen sowie redaktionelle Reisemedien (Lighthouse, 2026). Die eigene Zitationsmessung für „boutique hotel“ im US-englischen Markt bestätigt dieses Muster auf Domänebene: Unter den am häufigsten zitierten Quellen liegen Expedia, Tripadvisor, die englische Wikipedia, ein Wörterbuch, zwei Fachmedien der Hotelbranche, Reddit, Booking.com und zwei Anbieter von Property-Management-Software. Keine Hotelmarketing-Agentur und kein einzelnes Hotel erscheint unter den zehn meistzitierten Domains.

Zweitens: Inhaltliche Merkmale. Wie in Abschnitt 2 dargestellt, belohnt das kontrollierte Experiment von Aggarwal et al. (2024) überprüfbare Statistiken, glaubwürdige Zitate und Quellenangaben mit den größten Sichtbarkeitsgewinnen, während reine Keyword-Häufung wirkungslos oder sogar schädlich war. Das erklärt auch, warum der deutsche Omissionstest gerade jene Seiten belohnte, deren Titel die Nutzerfrage direkt aufgriff: eine erkennbare, direkte Antwort auf eine konkrete Frage ist genau die Art von klarem, belegbarem Inhalt, den diese Systeme bevorzugen.

Drittens: Strukturierte und Entitätssignale. Weil generative Systeme strukturierte Daten über eine Entität abgleichen, verringert eine konsistente, maschinenlesbare Darstellung eines Hotels über die eigene Website, Buchungsplattformen und Branchenverzeichnisse hinweg die Mehrdeutigkeit darüber, was dieses Hotel ist. Lighthouse rät ausdrücklich dazu, Buchungs- und Metasuchlistings zu prüfen, weil KI-Systeme genau diese Listings nutzen, „um zu verifizieren, was sie bereits zu wissen glauben“ (Lighthouse, 2026).

Für ein einzelnes Hotel folgt daraus eine konkrete Konsequenz: Die eigene Website ist nicht die primäre Zitationsfläche für die Frage, die Reisende einer KI stellen. Die Zitationsfläche ist der Buchungsplattform-Eintrag, der Wikipedia-Artikel, falls einer existiert, die Erwähnung in der Fachpresse und der Diskussionsbeitrag in einem Forum. Wer ausschließlich die eigene Website mit Inhalten bespielt, optimiert eine Fläche, die generative Systeme kaum zitieren.

8. Was gemessen wirkt, und was ein Haus damit anfangen kann

Bis hierher beschreibt dieser Text ein Problem. Die Frage, ob sich daran überhaupt etwas ändern lässt, ist eine eigene, und sie ist die am schlechtesten belegte im ganzen Feld. Was es dazu gibt, steht hier mit seinen Grenzen.

Die belastbarste Grundlage bleibt das Experiment von Aggarwal et al. aus Abschnitt 2: Quellen nennen, Zitate aufnehmen, Statistiken ergänzen, und Keyword-Häufung sein lassen. Ein späterer Versuch, diese Rangfolge außerhalb des ursprünglichen Benchmarks zu reproduzieren, kam auf plus 40 Prozent für das Nennen von Quellen, plus 37 Prozent für ergänzte Statistiken und plus 22 Prozent für aufgenommene Zitate, bei plus 3 Prozent und ohne statistische Signifikanz für Keyword-Häufung (Oruesagasti, 2026, Preprint). Dieser Bericht verdient dieselbe Vorsicht wie die Anbieterstudien weiter oben: Er stammt aus der Forschungsabteilung eines kommerziellen Anbieters, ist nicht peer-geprüft, und er misst britische Glücksspielmarken, keine Hotels. Sein Wert liegt darin, dass die Reihenfolge der Taktiken in einem anderen Feld dieselbe blieb, nicht darin, dass er die Effektgröße unabhängig bestätigt.

Eine Primärquelle gehört daneben gestellt, weil sie von der Plattform selbst kommt. Google hält in seiner Dokumentation zu generativen Suchfunktionen fest, dass es „keine zusätzlichen Voraussetzungen“ gibt, um in AI Overviews oder im AI Mode zu erscheinen: Diese Funktionen laufen über dieselben Kernsysteme wie die klassische Suche (Google Search Central, veröffentlicht 15.05.2026, aktualisiert 10.07.2026). Das nimmt eine Abkürzung weg und bestätigt eine Richtung. Es gibt keinen eigenen Schalter, und wessen Seiten diese Systeme nicht sauber einlesen können, hat eine Vorbedingung zu klären, bevor irgendeine der Taktiken oben überhaupt greifen kann.

Für ein einzelnes Haus laufen die Befunde dieses Textes auf vier Maßnahmen zu, die auch dann vertretbar bleiben, wenn einzelne Taktiken in einem Jahr nicht mehr wirken. Erstens: die eigenen Angaben über alle Kanäle hinweg gleich halten, also Website, Buchungsplattformen und Verzeichnisse, weil generative Systeme genau diese Einträge abgleichen, um zu prüfen, was sie zu wissen glauben (Lighthouse, 2026). Zweitens: redaktionelle Erwähnungen und Aufnahme in Bestenlisten als Vertriebsweg behandeln, nicht als Imagepflege, denn diese Seiten sind die Fläche, aus der zitiert wird, nicht die eigene. Drittens: die eigenen Texte belegt und klar schreiben, mit Zahlen und benannten Quellen, weil das die einzige Eigenschaft ist, für die ein kontrolliertes Experiment einen Effekt gemessen hat. Viertens: den Ausgangswert messen, bevor irgendetwas verändert wird, sonst lässt sich später nicht sagen, ob eine Maßnahme etwas bewirkt hat.

Alle vier verbessern Signale, auf die vermutlich auch ein künftiges System zurückgreift. Was ein Haus nach dieser Quellenlage nicht tun sollte, ist die 40 Prozent aus einem

Benchmark von 2024 in einem fachfremden Feld als zugesicherte Wirkung für die eigene Hotellerie zu lesen.

9. Was daraus folgt, und was ausdrücklich nicht folgt

Aus den fünf Erhebungen dieses Textes, drei eigenen und den zitierten internationalen Studien, ergeben sich einige vorsichtig formulierte Schlüsse.

Es folgt: Die deutschsprachige Nachfrage nach Boutique-Hotels innerhalb von KI-Prompts existiert bereits in messbarer Größenordnung, deutlich kleiner als im US-englischen Markt, aber nicht klein. Es folgt: Der Markt für GEO-Beratung selbst befindet sich im deutschsprachigen Raum in einer sehr frühen Phase, erkennbar an geringem KI-Prompt-Volumen bei gleichzeitig stark wachsendem klassischem Suchvolumen für dieselben Begriffe. Es folgt: Der Auswahlmechanismus, mit dem ein Sprachmodell entscheidet, wen es nennt, begünstigt aktuell fragengleiche Seitentitel, Verbandsdokumente und Drittquellen gegenüber reiner Domainautorität. Es folgt: Im deutschsprachigen Markt existiert derzeit keine Seite, die Hotellerie-Spezialisierung mit veröffentlichter, nachprüfbarer Messung verbindet.

Es folgt ausdrücklich nicht: dass einzelne deutsche Hotels heute schon nachweislich unsichtbar sind, denn der Omissionstest dieses Textes prüft Dienstleister, nicht Hotels. Es folgt nicht, dass GEO-Anbieter in Deutschland überflüssig sind, nur weil die direkte KI-Prompt-Nachfrage nach ihnen noch gering ist; das widerspräche der parallel gemessenen, stark wachsenden klassischen Suchnachfrage. Es folgt auch nicht, dass die hier beschriebenen Mechanismen dauerhaft gültig bleiben: Generative Systeme ändern sich häufig und teils unangekündigt, und ein Muster, das an einem Tag beobachtet wurde, kann sich in wenigen Monaten verschieben.

10. Grenzen, offene Fragen, Methodenkritik

Jede der drei eigenen Erhebungen dieses Textes ist eine Momentaufnahme, keine wiederholte Stichprobe. Das KI-Prompt-Volumen ist eine modellierte Schätzung, keine gezählte Menge tatsächlicher Konversationen, und sollte mit derselben Vorsicht gelesen werden, die für jedes Keyword-Volumen-Tool gilt. Der Omissionstest wurde einmal pro Sprache und mit nur einem Assistenten (ChatGPT) durchgeführt; der deutsche Durchlauf lief auf gpt-5.5-2026-04-23, der englische auf gpt-4o-2024-08-06, weshalb Unterschiede zwischen beiden Antworten mit dem Modellunterschied vermengt sind und nicht als Sprach- oder Markteffekt gelesen werden dürfen. Perplexity, Gemini und Google AI Overviews wurden nicht geprüft, und eine erneute Abfrage desselben Prompts hätte möglicherweise ein anderes Ergebnis geliefert, weil generative Systeme

nicht deterministisch antworten. Alle drei eigenen Messungen wurden am selben Tag, dem 28.07.2026, erhoben; generative Systeme ändern sich bekanntermaßen häufig, weshalb diese Zahlen einen Tag beschreiben, keinen dauerhaften Zustand.

Auch die zugrunde liegenden internationalen Studien haben Grenzen, die offen benannt gehören. Lighthouse und LuxDirect sind Anbieterstudien, keine peer-geprüfte Forschung, und ihre Methodik ist nicht vollständig offengelegt; ihre Autoren verkaufen verwandte Produkte. Ihre Befunde stützen sich jedoch gegenseitig und decken sich mit dem einzigen kontrollierten akademischen Experiment (Aggarwal et al., 2024), was ihnen Glaubwürdigkeit verleiht, ohne sie zu unabhängiger Forschung zu machen.

Offene Fragen bleiben: Wie stabil sind die hier beschriebenen Muster, wenn sich die zugrunde liegenden Modelle weiterentwickeln? Übersetzt sich eine gemessene KI-Sichtbarkeit tatsächlich in Buchungen, gegeben die dokumentierte Vertrauenslücke gegenüber KI an der eigentlichen Buchungsstelle? Und wie wird sich das Verhältnis zwischen klassischer Suchnachfrage und KI-Prompt-Nachfrage im deutschsprachigen Markt in den kommenden zwölf Monaten entwickeln, wenn beide, wie in Abschnitt 4 gezeigt, aktuell mit sehr unterschiedlicher Geschwindigkeit wachsen? Dieser Text beantwortet diese Fragen nicht, er dokumentiert den Ausgangspunkt, von dem aus sie beantwortbar werden.

Zitierfähige Einzelsätze (gesammelt)

1. „Boutique hotel“ erzeugte im deutschsprachigen Markt am 28.07.2026 geschätzte 5.541 KI-Prompts pro Monat, mit einer Zwölfmonatsspanne von 5.500 bis 7.600 (eigene Erhebung, DataForSEO, Endpunkt `ai_optimization/keyword_data/search_volume`, 28.07.2026).
2. Im US-englischen Markt lag „boutique hotel“ am selben Tag bei geschätzten 16.706 KI-Prompts pro Monat, Zwölfmonatsspanne 16.700 bis 28.100 (eigene Erhebung, DataForSEO, 28.07.2026).
3. Am selben Tag, in derselben deutschen Erhebung, kam „geo agentur“ auf geschätzte 5 KI-Prompts pro Monat, „generative engine optimization“ auf 3 und „geo optimierung“ auf 1, gegenüber 5.541 für „boutique hotel“ (eigene Erhebung, DataForSEO, 28.07.2026).
4. Die klassische Google-Suche nach „geo optimierung“ wuchs im Jahresvergleich um 171 Prozent auf 1.000 Suchanfragen pro Monat, während die KI-Prompt-Nachfrage nach demselben Begriff bei geschätzt 1 pro Monat lag (eigene Erhebung, DataForSEO, 28.07.2026).
5. Am 28.07.2026 nannte gpt-5.5 in der Rolle eines unabhängigen deutschen Boutique-Hoteliere zwölf GEO-Anbieter, elf davon zitiert von der jeweils eigenen Anbieterseite mit einem Titel, der die gestellte Frage nahezu wörtlich wiederholte (eigene Erhebung, 28.07.2026).
6. Die einzige hospitality-spezifische Nennung im deutschen Omissionstest vom 28.07.2026 stammte nicht von einer Anbieterseite, sondern aus einem PDF des Branchenverbands HSMAI (eigene Erhebung, 28.07.2026).

7. Von 49 im Rahmen des zugrunde liegenden Sichtbarkeits-Audits erfassten Wettbewerbsdomains publiziert genau eine eigene Forschung, und diese ausschließlich auf Englisch (Sichtbarkeits-Audit victorfrenzel.com, erhoben Juli 2026).
8. Ein Prompt-basierter Test von Lighthouse fand unter 4.545 ChatGPT-Prompts nur 2.721 unterschiedlich genannte Hotels bei insgesamt 49.707 Nennungen (Lighthouse, 2026).
9. Im einzigen kontrollierten, peer-geprüften Experiment zu diesem Thema verbesserten Quellenangaben, Zitate und Statistiken die Sichtbarkeit eines Inhalts relativ um 30 bis 40 Prozent, während Keyword-Häufung wirkungslos blieb; schlechter platzierte Seiten profitierten stärker als dominante (Aggarwal et al., 2024, ACM SIGKDD).
10. Google hält für seine generativen Suchfunktionen fest, dass es „keine zusätzlichen Voraussetzungen“ gibt, um in AI Overviews oder im AI Mode zu erscheinen (Google Search Central, veröffentlicht 15.05.2026, aktualisiert 10.07.2026).

Diesen Text zitieren

Frenzel, V. (2026). *Das verschwindende Hotel: KI-Sichtbarkeit im DACH-Markt* (Fassung 2, August 2026). Abgerufen von <https://victorfrenzel.com/de/research/das-verschwindende-hotel/>

Die englischsprachige Fassung erscheint als *The Disappearing Hotel* unter victorfrenzel.com/research.html und ist zusätzlich auf [ResearchGate](#) archiviert (Stand Juni 2026). Das [ResearchGate-Profil des Autors](#) führt die archivierten Arbeiten.

Quellen

Aggarwal, P., Murahari, V., Rajpurohit, T., Kalyan, A., Narasimhan, K. & Deshpande, A. (2024). GEO: Generative Engine Optimization. In Proceedings of the 30th ACM SIGKDD Conference on Knowledge Discovery and Data Mining (KDD '24) (S. 5 bis 16). Association for Computing Machinery. <https://doi.org/10.1145/3637528.3671900> (Preprint: arXiv:2311.09735, veröffentlicht am 16.11.2023.)

Booking.com. (2025, 23. Juli). Booking.com releases the Global AI Sentiment Report. Abgerufen am 28.07.2026 von <https://news.booking.com/bookingcom-releases-the-global-ai-sentiment-report/>

Frenzel, V. (2026, 28. Juli). Eigene Erhebung: KI-Prompt-Volumen für „boutique hotel“ (US/en und DE/de) [Datenerhebung]. Erhoben über DataForSEO, Endpunkt [ai_optimization/keyword_data/search_volume](#). Rohdaten beim Autor.

Frenzel, V. (2026, 28. Juli). Eigene Erhebung: deutsches Suchvolumen für GEO-Begriffe [Datenerhebung]. Erhoben über DataForSEO Labs, Endpunkt [dataforseo_labs/google/keyword_overview](#), Markt Deutschland. Rohdaten beim Autor.

- Frenzel, V. (2026, 28. Juli). Eigene Erhebung: Omissionstest deutscher GEO-Anbieter [Datenerhebung]. Erhoben über eine ChatGPT-Live-Sitzung (gpt-5.5-2026-04-23, Websuche aktiviert). Rohdaten und Antwortprotokoll beim Autor.
- Google. (2026, 15. Mai, aktualisiert 10. Juli). Optimizing your website for generative AI features on Google Search. Google Search Central. Abgerufen am 28.07.2026 von <https://developers.google.com/search/docs/fundamentals/ai-optimization-guide>
- Lighthouse. (2026, 12. Juni). The AI recommendation is the new battleground for hotel distribution [Forschung präsentiert bei Luminate 2026 von B. Reiter]. Hospitality Net. Abgerufen am 28.07.2026 von <https://www.hospitalitynet.org/news/4132964/>
- LuxDirect. (2026). CS7 London study: Five London hotels capture 57% of AI recommendations. Berichtet in Hospitality Net (2026, 16. März). Abgerufen am 28.07.2026 von <https://www.hospitalitynet.org/opinion/4131422/>
- Oruesagasti, J. (2026, 5. März). Algorithmic trust and compliance: Benchmarking brand notability for UK iGaming entities in generative search engines [Preprint, technischer Bericht der Interamplify Research Division, UK]. arXiv:2603.12282. Abgerufen am 26.08.2026 von <https://arxiv.org/pdf/2603.12282>
- Sichtbarkeits-Audit victorfrenzel.com. (2026, Juli). Interner Bericht: 49 Wettbewerbsdomains im deutschen und englischen GEO-Markt für Hotellerie, Kapitel 4, 8 und 9. Unveröffentlichter interner Bericht des Autors.

Victor Frenzel arbeitet als GEO-Berater für Hotellerie und unterstützt unabhängige, designgeführte Häuser dabei, in KI-Antworten benannt zu werden. Mehr dazu auf victorfrenzel.com.