



**So wird Content KI-fähig –
GEO von Antwort bis Zitat.**



Inhaltsverzeichnis

Einleitung	2
GEO / AEO	2
Definition der Begriffe	2
Differenzierung & Analogien zu SEO	3
Mögliche Ziele der Auffindbarkeit	4
Im Pre-Training	4
In der Websuche	5
Technische Grundlagen	7
Struktur und Markup	7
Semantisches HTML und Inhaltshierarchie	8
LLM-Interne Optimierungsprinzipien	8
Aktuelle Best Practices	9
Für Content-Verfasser	9
Für Marketer	9
Für technische Struktur / Plattform-Verantwortliche	10
Auswertung von GEO-Maßnahmen	10
Besonderheiten für Online-Publisher	10
Besonderheiten für Werbungtreibende	11
Analytics-Tools	11
Zukunftsblick	13
Maschinenlesbare Evidenz und Quellenketten	13
Unterschiedliche User Interfaces	13
Standardisierung von KI-spezifischen Protokollen und Labels	13
Multimodaler Content: verknüpfte Text-, Audio-, Bild- und Videoeinheiten	14
Glossar	14
Quellen	15
Impressum	16

Einleitung

Die wachsende Verbreitung von LLM-basierten Bots und AI-Suchfunktionen verändert die Rolle klassischer Medienangebote fundamental: Inhalte werden zunehmend nicht mehr direkt von Nutzer*innen konsumiert, sondern von KI-Systemen vorverarbeitet, zusammengefasst und in eigenen Interfaces ausgespielt. Publisher stehen damit vor einer neuen strategischen Herausforderung: Die eigenen Inhalte sollen für KI-Crawler und Bots auffindbar sein. Ziel ist es, bei relevanten Themen möglichst häufig in KI-generierten Antworten zitiert zu werden – idealerweise mit expliziter Quellenangabe und einer Verlinkung, die auf weiterführende Informationen im Original verweist.

In der Praxis kann diese Zielsetzung jedoch mit bestehenden oder geplanten Content-Lizenzierungsdeals mit LLM-Labs in Konflikt geraten, wenn dort Nutzungsrechte, Zugriffspfade oder Sichtbarkeitslogiken abweichend geregelt sind. Für Medienhäuser rückt damit eine neue Kernfrage in den Mittelpunkt: Wie müssen Inhalte strukturiert und technisch bereitgestellt werden, damit sie von LLMs überhaupt erfasst, richtig verstanden und in Antworten sichtbar gemacht werden – im Pre-Training wie in der laufenden Websuche und Inferenz. Der vorliegende Leitfaden führt dazu entlang der Konzepte GEO und AEO durch die Ziele der Auffindbarkeit, erläutert die technischen Grundlagen (Markup, semantische Struktur, LLM-interne Prinzipien), leitet konkrete Best Practices für Autorität, Content-Qualität, Maschinenlesbarkeit und Distribution ab und skizziert schließlich, welche zusätzlichen Anforderungen sich durch neue KI-User-Interfaces, Protokolle und multimodale Formate ergeben.

GEO / AEO

Mit Generative Engine Optimization (GEO) und Answer Engine Optimization (AEO) verschiebt sich der Schwerpunkt der Optimierung: weg von der Frage, wie Inhalte möglichst weit oben in Ergebnislisten erscheinen, hin zur Frage, wie sie in KI generierten Antworten vorkommen und dort als glaubwürdige Referenz gelten. Im Kern geht es also darum, Content so zu gestalten, dass generative und antwortbasierte Systeme ihn finden, verstehen, gliedern und zitieren können.

Definition der Begriffe

GEO

(Generative Engine Optimization) ist ein methodischer Ansatz zur Aufbereitung digitaler Inhalte, damit diese von generativen KI- und Assistenzsystemen erfasst, semantisch verstanden und als präferierte Quelle priorisiert werden. Ziel ist die direkte Integration der Informationen in synthetisierte Antworten oder autonome Prozesse. Damit ergänzt GEO die klassische Suchmaschinenoptimierung um neue Dimensionen, insbesondere Zitierfähigkeit, Datenvalidität, semantische Eindeutigkeit und Markenautorität innerhalb KI-gestützter Interaktionen.¹

In der Praxis werden GEO und AEO häufig zusammen mit Begriffen wie AI SEO, LLM Optimization (LLMO) oder Generative Search Optimization verwendet. Die Fachliteratur und aktuelle Branchenartikel² betonen, dass all diese Akronyme im Kern dasselbe Problem adressieren: Wie müssen Inhalte aufgebaut sein, damit sie in KI-Antworten präsent sind – unabhängig davon, ob der Einstieg über eine klassische Suchmaschine (Google, Bing) oder direkt über einen Chatbot erfolgt. GEO und AEO lassen sich als Spezialisierungen von Content- und Suchoptimierung verstehen, die explizit auf KI-Antwortformate ausgerichtet sind. Klassische SEO-Signale wie Relevanz, Autorität und technische Sauberkeit bleiben wichtig, werden aber ergänzt um Anforderungen wie Zitierfähigkeit, semantische Klarheit und KI-gerechte Strukturierung.

AEO

(Answer Engine Optimization) beschreibt ein ähnliches Konzept aus einer anderen Perspektive: AEO fokussiert auf „Answer Engines“ – Systeme wie AI Overviews, Sprachassistenten (Alexa, Siri, Google Assistant) oder Chatbots (Bing Copilot, ChatGPT), die direkte Antworten ausspielen, statt eine Trefferliste zu präsentieren. AEO umfasst Strategien, mit denen Inhalte so aufgebaut werden, dass sie von diesen Answer Engines bevorzugt ausgewählt, zitiert und in kompakten Antwortformaten ausgegeben werden.

¹ <https://www.bvdw.org/wp-content/uploads/2026/02/The-Answer-Economy.pdf>

² <https://backlinko.com/seo-vs-geo>, <https://neilpatel.com/blog/geo-vs-aeo/>

Differenzierung & Analogien zu SEO

1. Zielsystem:

Suchmaschine vs.
Generative / Answer Engine

Klassische SEO optimiert für Suchmaschinen, die eine Liste von Dokumenten zurückgeben. Erfolg bedeutet, für relevante Suchanfragen möglichst weit oben im Suchmaschinen-ranking zu erscheinen und Klicks auf die eigene Seite zu generieren. GEO und AEO optimieren dagegen für Systeme, die (vorrangig) Textantworten generieren. Hier besteht der Erfolg darin, dass das eigene Dokument als Quelle, Zitat oder Paraphrase in die Antwort einfließt, selbst wenn die Nutzer*innen die Website nie besuchen.

2. Ausgabeform:

Linkliste vs. synthetisierte Antwort

In der klassischen Suche sind Dokumente voneinander getrennt. Jede URL belegt eine Ranking-Position. AEO/GEO-Szenarien funktionieren anders. Die KI synthetisiert eine Antwort aus mehreren Quellen, extrahiert dabei einzelne Sätze, Fakten, Zahlen und fügt diese zusammen. Dadurch reduziert sich die Überlappung zwischen Top-Google-Rankings und den von Chatbots zitierten Quellen deutlich.³ Für Content-Strategien bedeutet das: Ein Platz 1 in Suchmaschinenrankings garantiert nicht, in KI-Antworten überhaupt aufzutauchen. Umgekehrt können auch Quellen, die organisch nur mittelmäßig ranken, in KI-Snapshots prominent erwähnt werden, wenn sie besonders zitierfähige oder einzigartige Informationen liefern.

3. Erfolgsmaßstab:

Klick vs. Zitat

SEO misst Erfolg primär über Traffic-Metriken, wie Impressionen, Klicks, durchschnittliche Position, CTR, organische Conversions. GEO/AEO verschiebt den Fokus auf „**Citation Authority**“: Wie oft wird eine Marke oder ein Dokument von Answer Engines genannt, zitiert oder verlinkt? Aktuelle GEO-Guides und Toolanbieter sprechen von Kennzahlen wie „AI Mentions“, „AI Citations“ oder einem „LLM Visibility Index“, die messen sollen, wie präsent eine Domain in KI-Antworten über verschiedene Systeme hinweg ist.⁴

In dieser Logik wird eine Seite, die in AI Overviews oder Chatbot-Antworten regelmäßig als Quelle erscheint, strategisch wertvoll, selbst wenn der direkte Traffic gering ist. Sichtbarkeit und Einfluss verschieben sich also von der Klick-Ebene zur Antwort-Ebene.

Trotz dieser Unterschiede gibt es deutliche **Analogien**:

- **Viele Ranking-Signale bleiben gleich. GEO/AEO baut auf denselben Fundamenten auf wie SEO:** hochwertige Inhalte, klare Struktur, gute User Experience, Backlinks und starke Marken-Signals.
- **Topical Authority und E-E-A-T⁵ bleiben wichtig.** Sowohl Suchmaschinen als auch KI-Modelle bewerten, ob eine Quelle zu einem Themengebiet glaubwürdig ist. Wer über längere Zeit konsistent hochwertige Inhalte zu einem Cluster liefert und externe Erwähnungen (Earned Media) aufbaut, wird als Autorität wahrgenommen – sowohl im Ranking als auch in generativen Antworten.
- **Strukturierte Antworten sind die neue Featured-Snippet-Logik.** AEO kann man als Weiterentwicklung der Featured-Snippet-Optimierung verstehen: Anstatt nur ein Snippet in der Google-SERP zu besetzen, optimiert man heute für AI Overviews, Chatbot- und Voice-Antworten. Das Prinzip ist jedoch ähnlich: Fragen antizipieren, Antworten klar und knapp formulieren, danach ins Detail gehen.

Aus Sicht großer Medienhäuser ist GEO/AEO daher eher eine Erweiterung als ein Ersatz für SEO. Der Fokus verschiebt sich jedoch strategisch: weg von der reinen Ranking-Position hin zur Frage, ob und wie die eigene Marke in KI-generierten Antworten vorkommt. Dabei geht es auch um die Narrative, die die KI transportiert.

³ <https://ahrefs.com/blog/ai-search-overlap/>

⁴ <https://www.ewrdigital.com/llm-visibility-glossary>

⁵ Experience (Erfahrung), Expertise (Fachwissen), Authoritativeness (Autorität) und Trustworthiness (Vertrauenswürdigkeit)
<https://services.google.com/fh/files/misc/hsw-sqrg.pdf>

Mögliche Ziele der Auffindbarkeit

Bevor konkrete Maßnahmen diskutiert werden, muss klar sein, wie Auffindbarkeit überhaupt optimiert wird. Im Kontext von LLMs geht es dabei vor allem um zwei Ebenen: zum einen die Verankerung im Modell während des Pre-Trainings, zum anderen die Auffindbarkeit zur Laufzeit über Websuche und damit der Inferenz der Modelle.



Das Ziel des Pre-Trainings besteht darin, die parametrische Wissensbasis eines Large Language Models zu beeinflussen. Dabei geht es um die Informationen, die das Modell „im Kopf“ hat, bevor es später bei der Nutzung eine Websuche initiiert⁶. In der Praxis kann diese Zielsetzung jedoch mit bestehenden oder geplanten Content-Lizenzierungsdeals mit LLM-Labs kollidieren, wenn Trainingsnutzungsrechte, Exklusivität oder Distributionswege die gewünschte Offenheit und Verbreitung einschränken.

Ein LLM wird in der Pre-Training-Phase auf sehr großen, überwiegend historischen Text-, Bild- und Tonkorpora trainiert. Dabei lernt es Sprachmuster und inhaltliche Zusammenhänge, indem es aus riesigen Mengen von Tokens kontinuierlich Embeddings und Gewichte optimiert.⁷ Für die Content-Strategie bedeutet das: Wer es schafft, mit seinen Inhalten in diese Trainingsdaten hineinzugelangen, beeinflusst die semantische Landschaft, in der das Modell später arbeitet.⁸

Die zentrale Zielsetzung im Pre-Training ist daher:

1. Thematische Autorität im Modell verankern

Das Ziel besteht darin, dass das Modell eine Marke, ein Medium oder ein Themencluster als „Autorität“ wahrnimmt. Dies geschieht nicht, weil es einen expliziten Rankingfaktor gäbe, sondern weil die Inhalte konsistent, dicht und hochwertig in den Trainingsdaten vorkommen. Für sogenannte „Contextual Authority“ sollen Inhalte intern konsistent sein, sich an autoritativen Konzepten ausrichten und so die eigene Position im vektorisierten Wissensgraphen stärken.

Wenn ein LLM aus vielen Beispielen lernt, dass bestimmte Begriffe regelmäßig gemeinsam mit einer Marke auftreten, etwa ein Wirtschaftstitel im Kontext von Kapitalmarktanalysen, dann verankert sich diese Verbindung dauerhaft im Modell.

2. Embedding-Relevanz und semantische Kohärenz erhöhen

Da LLMs Text in hochdimensionale Vektoren überführen, ist es für die Auffindbarkeit entscheidend, wie klar ein Inhalt im semantischen Raum positioniert ist. Daher ist die „Embedding Relevance“ ein wichtiges Optimierungsziel: Inhalte sollten thematisch fokussiert und in sich kohärent sein, damit ihr Embedding klar einem Themenfeld zugeordnet werden kann.

Je stärker diese Kohärenz, desto wahrscheinlicher wird der Content später Prompts zugeordnet, die sich auf dieses Thema beziehen.

3. Klarheit, Disambiguierung und kanonische Begriffe

Pre-Training ist grob gesprochen „blind“ gegenüber späteren User Prompts. Unklare oder mehrdeutige Bezeichnungen im Trainingsmaterial führen daher direkt zu semantischer Unschärfe im Modell. Deshalb empfiehlt die Literatur⁹ die Nutzung kanonischer Begriffe und eindeutiger Formulierungen („Canonical Clarity and Disambiguation“): klare Definitionen, Verzicht auf unnötigen Jargon, explizite Auflösung von Mehrdeutigkeiten.

Das reduziert Halluzinationen und Fehlzuordnungen und erhöht die Chance, dass das Modell die richtigen Konzepte miteinander verknüpft.

⁶ <https://developers.openai.com/api/docs/bots>

⁷ <https://commoncrawl.org/>

⁸ <https://cdn.openai.com/papers/gpt-4.pdf>

⁹ <https://arxiv.org/html/2601.22108v1/>

4. Token-Effizienz und inhaltliche Dichte

Im Pre-Training kommt es auf jeden Token an, da er die Fähigkeit des Modells, Muster zu erkennen, beeinflusst. „Token Efficiency“ bedeutet, Redundanz zu reduzieren, ohne dabei den Informationsgehalt zu verlieren. Inhalte sollten deshalb prägnant, strukturiert und ohne unnötige Füllsätze formuliert sein. Das erhöht die Wahrscheinlichkeit, dass relevante Aussagen als eigenständige Muster gelernt werden, statt im Rauschen zu verschwinden.

5. Langfristige Wissenslücken füllen (Original Research)

Besonders interessant im Pre-Training sind Inhalte, die neue Daten oder Perspektiven liefern, für die es bislang wenig öffentlich zugängliches Material gibt. Originalstudien, Marktanalysen oder neuartige Methoden füllen „temporäre Lücken“ in der LLM-Wissensbasis und werden daher überdurchschnittlich häufig zitiert werden können. Veröffentlichungen, die neue Fakten etablieren, wirken im Pre-Training also wie ein starker „Anker“ im Wissensraum des Modells.

Operativ lässt sich das Pre-Training-Ziel so zusammenfassen: Content-Strategien zielen weniger auf einzelne Keywords als auf eine breite, wiederholte und hochwertige Präsenz in den für Trainingsdaten relevanten Quellen ab, darunter offene Webinhalte, zitierte Reports, wissenschaftliche Veröffentlichungen und News-Archive. Wer dort mit klar strukturierten, fachlich starken und originären Inhalten vertreten ist, maximiert seine Chancen, im nächsten Pre-Training-Zyklus in der Modellbasis sichtbar zu sein.



In der Websuche

Während das Pre-Training die statische Wissensbasis eines Modells formt, geht es in der Websuche von KI-Anwendungen um die dynamische Nutzung externer Quellen. Diese werden etwa über Retrieval-Augmented Generation (RAG) oder AI Overviews ermittelt. Hier steht weniger die Frage im Vordergrund, ob ein Modell eine Marke kennt, sondern ob es sie in einer konkreten Antwort heranzieht, zitiert und gegebenenfalls verlinkt.

Moderne KI-Suchsysteme kombinieren LLMs mit Retrieval-Mechanismen: Sie zerlegen die Nutzeranfrage, suchen passende Dokumente in einem Index (z.B. Webindex, Unternehmens-Vektordatenbank) und generieren daraus eine Antwort. Für die Auffindbarkeit ergeben sich daraus drei Zielrichtungen:

1. Passage-Level-Auffindbarkeit im RAG-Workflow

RAG-Systeme arbeiten in der Regel mit relativ kurzen Textsegmenten – meist Abschnitte oder Paragraphen. Ziel ist daher, Inhalte so zu strukturieren, dass einzelne Abschnitte für sich genommen eine sinnvolle Antwort auf eine Teilfrage darstellen. In dieser „Answer-First Structure“ empfiehlt es sich, die Kernantwort in den ersten 40–60 Wörtern eines Abschnitts zu platzieren, ergänzt durch Listen und Tabellen als „extractable formats“. Auf dieser Basis können RAG-Systeme gezielt die passende Passage extrahieren und in die Antwort einbauen.¹⁰

2. Maschinenlesbarkeit und Schema für AI-Overviews & Agentic Browsing

Für die Sichtbarkeit in Google AI Overviews und anderen generativen Suchfunktionen ist eine klare technische Struktur der Website nach wie vor relevant. Entscheidend sind insbesondere die grundlegenden SEO-Anforderungen. Inhalte müssen von Suchmaschinen robots (Crawlern) erfasst (crawlbar) und indexiert werden können und für die Darstellung in der Google-Suche geeignet sein. Semantisches HTML, eine nachvollziehbare Seitenstruktur und strukturierte Daten können die maschinelle Verarbeitung unterstützen. Sie sind jedoch weder eine Garantie für die Aufnahme in AI Overviews noch spezielle Voraussetzungen dafür.

Schema.org-Markup ist ein essenzieller Bestandteil einer soliden SEO-Strategie und kann die Eignung für bestimmte Rich Results verbessern. Es wird für Artikel und andere zum jeweiligen Inhalt passende Typen eingesetzt. Für die generativen Suchfunktionen von Google ist strukturiertes Markup jedoch nicht zwingend erforderlich. Zudem existiert kein spezielles Schema.org-Markup für AI Overviews. Entscheidend ist, dass die ausgezeichneten Angaben mit dem sichtbaren Seiteninhalt übereinstimmen und einen tatsächlichen semantischen Mehrwert bieten.

¹⁰ <https://www.bvdw.org/wp-content/uploads/2025/12/Retrieval-Augmented-Generation.pdf>
<https://www.airops.com/blog/featured-snippet-optimization>

Die llms.txt ist eine noch nicht etablierte Konvention, mit der Websites eine maschinenlesbare Zusammenfassung sowie Verweise auf zentrale Inhalte für LLMs und KI-Agenten bereitstellen können. Google gibt an, dass die Datei für die klassische Google-Suche und deren generative Funktionen, einschließlich AI Overviews, ohne Einfluss ist. Google Search nutzt die Datei nicht für Sichtbarkeit oder Ranking. Außerhalb der Google-Suche kann die „llms.txt“ jedoch für agentenbasiertes Browsing relevant werden, indem sie KI-Agenten die Orientierung auf einer Website erleichtert. Ihr Einsatz ist daher getrennt von der Optimierung für AI Overviews zu bewerten und hängt davon ab, ob agentische Nutzungsszenarien zum Zielbild der Website gehören.¹¹

3. Sichtbarkeit als zitierfähige Quelle in KI-Suchergebnissen

Ein wesentlicher Unterschied zur Pre-Training-Optimierung liegt in der Rolle externer Autoritätsquellen. Klassische SEO fokussiert stark auf die eigene Domain Authority; GEO/AEO im Websuche-Kontext muss darüber hinaus den Bias generativer Suchdienste berücksichtigen: LLMs zeigen eine systematische Bevorzugung von Earned Media und Drittanbieter-Quellen (Reddit, Wikipedia, große Publikationen) gegenüber markeneigenem Content.¹²

Entsprechend verschiebt sich das Ziel, nur die eigene Webseite zu optimieren, hin zu einem Citation-Network-Ansatz: Marken- und Fachinhalte sollten auf maßgeblichen Drittplattformen präsent sein, damit LLMs dort zitierfähige Passagen finden. Der Aufbau von Zitat-Netzwerken, einer Präsenz auf bevorzugten Plattformen und Digital-PR sind neue Kernelemente der Content-Strategie.

Konkret lassen sich für die Inferenz folgende Unterziele formulieren:

a. In AI Overviews / AI Mode auftauchen:

Ziel ist, für relevante Suchanfragen als Quelle innerhalb der generativen Zusammenfassung genannt zu werden – idealerweise mit Markennamen und Link. Strukturierte Daten, saubere Überschriftenhierarchie und klar abgegrenzte Antwortblöcke erhöhen diese Wahrscheinlichkeit deutlich.

b. In Chatbot-Antworten als Referenz erscheinen:

Bei Chatbots wie ChatGPT, Perplexity oder Copilot, die eine Websuche oder dedizierte RAG-Pipelines nutzen, lautet das Ziel: als verlinkte oder zitierte Quelle in der Antwort aufzutauchen – auch dann, wenn die Nutzer*innen nicht klicken.

c. Korrekte und markenkonforme Darstellung sicherstellen:

Neben der reinen Sichtbarkeit geht es in der Inferenz um Inhaltssouveränität: Wie beschreibt die KI die Marke, das Produkt, die Dienstleistung? Durch konsistente Botschaften in eigenen und externen Quellen lässt sich steuern, welche Narrative sich in AI-Outputs durchsetzen.

4. Strukturelle Sauberkeit und Maschinenlesbarkeit als Mindeststandard

Während im Pre-Training die inhaltliche Tiefe und thematische Autorität im Vordergrund stehen, fokussiert die Optimierung für Websuche/Retrieval auf die strukturelle Sauberkeit und Maschinenlesbarkeit. Ohne solide technische Struktur können AI-Systeme selbst hochrangige Seiten übergehen, wenn der Inhalt nicht ausreichend zitierfähig ist.

Zusammengefasst:

- Im Pre-Training ist das Ziel, die semantische Identität und Autorität einer Marke in der internen Wissensbasis eines LLMs durch kohärente, klare, originäre Inhalte zu verankern.
- In Websuche und Inferenz geht es darum, in konkreten Antwortsituationen durch technisch saubere, fein segmentierte, Schema gestützte Inhalte und eine starke Präsenz in bevorzugten Drittquellen als Quelle ausgewählt und korrekt zitiert zu werden.

Beide Ebenen sind komplementär: Die Relevanz im Pre-Training entscheidet darüber, was ein Modell über eine Marke „weiß“, die Relevanz in der Inferenz darüber, wann und wie dieses Wissen in der Interaktion mit Nutzer*innern sichtbar wird.

¹¹ <https://llmstxt.org/>

<https://developer.chrome.com/docs/lighthouse/agentive-browsing/llms-txt>

<https://developers.google.com/search/docs/fundamentals/ai-optimization-guide>

¹² <https://ahrefs.com/blog/most-cited-domains-ai-mode/>

Technische Grundlagen

Bestimmte technische Grundlagen sind die Voraussetzung dafür, dass Inhalte überhaupt von LLMs erfasst, verstanden und in Antworten genutzt werden können. Für GEO und AEO genügt es nicht, gute Texte zu verfassen – entscheidend ist, wie diese im Markup verankert, semantisch strukturiert und für Maschinen lesbar gemacht werden. Erst wenn Struktur, HTML-Hierarchie und LLM-gerechte Aufbereitung stimmen, können inhaltliche Qualität und Autorität voll wirken.



Struktur und Markup

Technische Grundlage jeder GEO/AEO-Strategie ist, dass Inhalte maschinenlesbar und zuverlässig parsbar sind. Parser sind Computerprogramme, die Texte analysieren, verstehen und in ihre Einzelteile zerlegen. Suchmaschinen und LLMs arbeiten zwar mit unterschiedlichen Architekturen, benötigen aber beide eine klare, konsistente Struktur, um Text sinnvoll zu zerlegen, einzuordnen und später wieder abzurufen.

Zentral ist dabei der Einsatz von strukturierten Daten (Schema Markup). Über Schema.org-Typen wie Article, FAQPage, HowTo, Organization oder LocalBusiness lassen sich Inhalte in JSON-LD explizit annotieren: Titel, Autor, Veröffentlichungsdatum, Schrittfolgen, Fragen und Antworten, Unternehmensdaten. Schema-Markup ist dadurch zunehmend ein Game-Changer für AEO, weil es Struktur, Klarheit und Abrufbarkeit digitaler Inhalte verbessert und KI-Systemen hilft, Material präzise zu zitieren.

JSON-LD ist dabei das von Google bevorzugte Format, und spezifische Typen wie FAQPage oder HowTo werden explizit als essenzielle Bausteine für Q&A- und How-to-Inhalte hervorgehoben.¹³

Gleichzeitig bleibt klassisches HTML das primäre Trägermedium. Entscheidende Informationen sollten immer im sichtbaren Text vorkommen und nicht ausschließlich im Schema-Code versteckt sein. Image-basierte PDFs, Scans oder Bilder mit eingebettetem Text gelten als ungeeignet für die KI-Verarbeitung, weil sie nicht ohne Weiteres tokenisiert und in Embeddings überführt werden können.

Hinzu kommen KI-spezifische Konventionen. Aktuell wird insbesondere die Datei „llms.txt“ diskutiert, die im Domain-Root platziert wird und LLMs sowie KI-Agenten eine maschinenlesbare Übersicht zentraler Inhalte bereitstellen soll. Auf die klassische Google-Suche oder AI Overviews hat sie keinen Einfluss, sie könnte jedoch für agentenbasiertes Browsing relevant werden. In Kombination mit Sitemaps und Schema-Markup entsteht so eine mehrschichtige Signalisierung: Was gibt es auf der Seite, wo liegt es und wie darf es verwendet werden?

In der Summe definiert Struktur und Markup den technischen Rahmen, innerhalb dessen Media-Content für LLMs überhaupt adressierbar wird – sowohl im Pre-Training (Auswahl der Trainingsdaten) als auch in der Inferenz (Auswahl und Zitierung einzelner Passagen).

¹³ <https://developers.google.com/search/docs/appearance/structured-data/intro-structured-data>
<https://developers.google.com/search/blog/2019/05/new-in-structured-data-faq-and-how-to>

Semantisches HTML und Inhaltshierarchie

Über das Vorhandensein von Markup hinaus kommt es darauf an, wie Inhalte in HTML strukturiert sind. LLMs und AI-Suchsysteme nutzen die semantische Struktur einer Seite, um Abschnitte, Themen und Zusammenhänge zu erkennen.

Best Practice ist der konsequente Einsatz von semantischen HTML5-Elementen wie:

```
<header>
<nav>
<main>
<article>
<section>
<aside>
<footer>
```

Diese Elemente liefern Kontextsignale: Wo beginnt der eigentliche Artikel? Welche Bereiche sind Navigation, welche Zusatzinformationen? Das erleichtert sowohl die Chunking-Strategie für RAG-Systeme als auch die Identifikation relevanter Abschnitte für AI Overviews.

Ebenso wichtig ist eine saubere Überschriftenhierarchie (H1 → H2 → H3 ohne Sprünge). Konsistent aufgebaute Headline-Strukturen erhöhen die Wahrscheinlichkeit, von KI-Modellen zitiert zu werden deutlich, weil die Modelle Themenbäume und Unterpunkte klarer erkennen. In der Praxis bedeutet das: ein zentraler H1-Titel pro Seite, H2-Überschriften für Hauptabschnitte, H3-Unterpunkte für Detailebenen – keine willkürlichen Formatierungen via div oder span.

Für die tatsächliche Extrahierbarkeit sind zudem bestimmte Formate hilfreich. Definition Lists (<dl>, <dt>, <dd>) eignen sich für Glossare und Begriffserklärungen; Tabellen für Datenvergleiche, Feature-Matrizen oder KPI-Übersichten; Listen und nummerierte Schritte für How-to-Inhalte.

Diese Strukturen liefern „harte Kanten“ im Text, an denen LLMs Informationen als selbstständige Einheiten erkennen und herauslösen können.

Ein häufig unterschätzter Aspekt ist der Fehlerfall: Inhalte, die hinter Tabs, Akkordeons oder nur per JavaScript nachgeladen werden, sind aus KI-Perspektive riskant. Daher sollten zentrale Inhalte nicht in UI-Elementen versteckt werden, sodass diese sicher von Crawlern gerendert werden können.

Für GEO/AEO sollten wichtige Textpassagen immer auch in einer Form vorliegen, die ohne komplexes Rendering zugänglich ist. Daraus lässt sich eine einfache Heuristik ableiten: Wenn ein Mensch einen Artikel als logisch gegliedert, gut scanbar und klar strukturiert erlebt, steigt die Wahrscheinlichkeit, dass auch ein LLM denselben Text korrekt segmentiert, einordnet und später wieder nutzt.

LLM-Interne Optimierungsprinzipien

Die dritte technische Ebene betrifft nicht mehr primär HTML, sondern die interne Verarbeitung von Text im LLM. Dies umfasst sowohl Embeddings als auch Tokenisierung und Abruflogik.

Ein zentrales Prinzip ist die bereits angesprochene **Token Efficiency**. LLMs verarbeiten Text als Folge von Tokens; redundante oder stark aufgeblähte Formulierungen erzeugen Rauschen, ohne zusätzlichen Informationswert. Redundanzen sollten also reduziert werden, um Klarheit zu wahren, Präzision und Interpretation durch KI-Systeme zu verbessern. Praktisch bedeutet das: kurze, klare Sätze, Vermeidung von Füllphrasen, Konzentration auf inhaltlich relevante Aussagen.

Eng damit verbunden ist **Embedding Relevance**: Texte werden in hochdimensionale Vektoren eingebettet; je fokussierter und thematisch konsistenter ein Inhalt, desto klarer seine Position im semantischen Raum. Ziel ist es also die semantische Stärke und thematische Kohärenz von Inhalten zu verbessern, damit sie relevanten Prompts zuverlässiger zugeordnet werden. Artikel sollten klar einem Thema zugeordnet sein, anstatt viele unterschiedliche Themen oberflächlich zu mischen.

Ein weiteres Prinzip ist **Canonical Clarity und Disambiguation**. Mehrdeutige Begriffe oder kreative, aber unübliche Formulierungen erschweren es dem Modell, die richtige Entität oder das richtige Konzept zu treffen. Daher sollten kanonische Begriffe verwendet und Mehrdeutigkeiten aktiv aufgelöst werden. Dies kann etwa durch präzise Definitionen, Klammerzusätze oder Kontextformulierungen geschehen.

Schließlich spielt **Prompt-Kompatibilität** eine Rolle: Inhalte sollten so formuliert sein, dass sie typischen Nutzerfragen entsprechen. Das bedeutet, dass Überschriften und Einleitungen häufige Frageformen („Was ist...?“, „Wie funktioniert...?“) spiegeln und die Antwort direkt darunter liefern. Dadurch erhöht sich die Wahrscheinlichkeit, dass ein LLM genau diesen Abschnitt als passende Antwortpassage selektiert.

Diese internen Optimierungsprinzipien greifen in beiden Zielsystemen. Im Pre-Training tragen sie dazu bei, dass die richtigen Muster gelernt werden. In der Inferenz helfen sie, dass die passenden Passagen eines Dokuments tatsächlich abgerufen und verwendet werden.

Aktuelle Best Practices

Für Content-Verfasser

- ✓ **Expertise sichtbar machen:** Praxis-Erfahrung und Fachwissen durch Fallstudien, konkrete Projekte, Datenbeispiele und nachvollziehbare Methoden belegen, statt nur Meinungen zu formulieren.
- ✓ **Autor*innen klar ausweisen:** Vollständige Autorenprofile mit Rolle, Qualifikation und ggf. externen Profilen nutzen, um Autorität und Vertrauen erkennbar zu machen.
- ✓ **Primärquellen konsequent nutzen:** Zentrale Aussagen mit Originalstudien, Datensätzen oder Primärreports belegen und Sekundärquellen nur ergänzend einsetzen.
- ✓ **Answer-first schreiben:** Die Hauptfrage eines Abschnitts früh und klar beantworten, idealerweise in den ersten Sätzen, danach Details, Beispiele und Vertiefungen liefern.
- ✓ **„Meta Answers“ einbauen:** Knappe, in sich geschlossene Kernaussagen formulieren, die ohne weiteren Kontext verständlich sind und sich direkt als Zitat eignen.
- ✓ **Eigene Daten und Studien entwickeln:** In Originalforschung, exklusive Analysen und empirische Reihen investieren, um einzigartige, schwer substituierbare Inhalte zu schaffen.
- ✓ **Sprache an Nutzerfragen ausrichten:** Typische Frageformen in Überschriften und Einstiegen nutzen und die Antwort unmittelbar anschließend geben, in klarer, dialognaher Sprache.

Für Marketer

- ✓ **Präsenz auf bevorzugten Plattformen sichern:** Sichtbarkeit auf Quellen aufbauen, die LLMs häufig nutzen, zum Beispiel Wikipedia, relevante Foren, Fachportale und große Medienmarken.
- ✓ **Digital PR und Co-Citations stärken:** Gezielt Kooperationen, Gastbeiträge und gemeinsame Inhalte mit anderen Autoritäten planen, um Co-Citations und thematische Nähe im Wissensraum der Modelle zu erzeugen.
- ✓ **Kernbotschaften kanalübergreifend konsistent halten:** Zentrale Aussagen in eigenen Kanälen, Earned Media und Drittplattformen möglichst ähnlich formulieren, damit LLMs stabile Markenmuster erkennen und reproduzieren.

Für technische Struktur / Plattform-Verantwortliche



Semantisches HTML und klare Überschriftenstruktur nutzen: Inhalte mit Elementen wie `<article>`, `<section>`, `<header>`, `<footer>` und einer sauberen H1/H2/H3-Hierarchie strukturieren, damit LLMs Abschnitte als sinnvolle Chunks erkennen.



Schema-Markup gezielt einsetzen: JSON-LD für Typen wie Article, FAQPage und HowTo verwenden, um Struktur, Rollen und Inhaltstypen maschinenlesbar zu machen.



„Blinde“ Formate vermeiden: Wichtige Inhalte nicht ausschließlich in Bild-PDFs, Scans oder rein skriptbasierten Oberflächen ablegen, sondern immer eine zugängliche HTML- oder strukturierte Textversion bereitstellen.



KI-Crawler bewusst steuern: In robots.txt und verwandten Dateien KI-Bots explizit berücksichtigen und nicht pauschal blockieren, wenn Inhalte im Pre-Training und in AI-Overviews sichtbar sein sollen.¹⁴

Auswertung von GEO-Maßnahmen

Die Auswertung von GEO-Maßnahmen folgt einer anderen Logik als die klassische SEO-Messung. Während bei letzterem Rankings, Klicks und CTR im Mittelpunkt stehen, geht es bei GEO vor allem um Sichtbarkeit in KI-Antworten, Zitierfähigkeit und Empfehlungsstärke. Entscheidend ist also nicht nur, ob Inhalte gefunden werden, sondern auch, ob sie von generativen Systemen als vertrauenswürdige Quelle ausgewählt, korrekt wiedergegeben und im besten Fall aktiv empfohlen werden.

Zentrale Kennzahlen sind dabei die Brand Visibility, also der Anteil der KI-Antworten mit Marken- oder Quellennennung, die Citation Rate, also die Häufigkeit tatsächlicher Referenzen, sowie der Share of Voice im Vergleich zum Wettbewerb. Hinzu kommen das Sentiment und die Unterscheidung zwischen bloßer Erwähnung und aktiver Empfehlung. Gerade diese Differenz ist relevant, da Sichtbarkeit allein noch keinen strategischen Wert garantiert. Eine Marke kann häufig genannt werden, ohne im entscheidenden Moment als bevorzugte Lösung zu erscheinen. Ergänzend dazu gehört das Bot- und Referral-Monitoring zur Analyse, also die Beobachtung, welche KI-Crawler Inhalte auslesen und ob über KI-Interfaces überhaupt noch Traffic auf die eigene Website zurückfließt.

Methodisch ist die GEO-Analyse kein punktueller Report, sondern ein kontinuierlicher Beobachtungsprozess. Wichtig sind insbesondere Lückenanalysen: Für welche Fragen werden Wettbewerber zitiert, die eigene Quelle aber nicht? Ebenso relevant ist die Untersuchung realer Prompt- und Dialogverläufe, da sich Empfehlungen in Multi-Turn-Konversationen deutlich verändern können. Da Zitationen in generativen Oberflächen deutlich volatiler sind als klassische Rankings, ist ein engmaschigeres Monitoring als im traditionellen SEO sinnvoll.

Besonderheiten für Online-Publisher

Für Publisher ist die GEO-Analyse besonders anspruchsvoll, da Sichtbarkeit und Traffic zunehmend auseinanderfallen. Wenn KI-Systeme Inhalte direkt in der Oberfläche zusammenfassen, verliert der Klick auf die Ursprungsseite an Bedeutung. Deshalb müssen Publisher stärker messen, wie häufig ihre Inhalte als Quelle in KI-Antworten auftauchen, wie hoch ihr thematischer „Share of Voice“ ist und in welchem Umfang ihre Inhalte in Zero-Click-Szenarien genutzt werden.

Hinzu kommt, dass generative Systeme häufig nicht die Ursprungsquelle, sondern Drittquellen und Earned Media bevorzugen. Daher reicht es für Publisher nicht aus, nur die eigene Domain zu beobachten. Auch relevant ist, ob eigene Inhalte, Recherchen oder Marken auf externen Plattformen aufgegriffen und in KI-Antworten einbezogen werden.

¹⁴ <https://www.bvdw.org/wp-content/uploads/2025/12/Auswirkungen-von-KI-Suche-auf-Publisher-Geschäftsmodelle.pdf>

Besonderheiten für Werbungtreibende

Für Werbungtreibende verschiebt sich die Auswertung von der Frage nach der Sichtbarkeit zur Frage nach der Empfehlung. Im KI-Kontext genügt es nicht, in den Antworten aufzutauchen. Entscheidend ist, ob die Marke als bevorzugte Lösung genannt wird. Entsprechend wird das Verhältnis von Erwähnung zu Empfehlung zu einer zentralen Kennzahl. Gleichzeitig verkürzt sich der klassische Funnel, da KI-Assistenten den Vergleich und die Vorauswahl teilweise selbst übernehmen. Werbungtreibende müssen deshalb messen, wie präsent ihre Marke entlang der gesamten KI-gestützten Buying Journey ist – von der ersten Problemdefinition bis zur konkreten Kaufempfehlung.

Auch hier spielt der Bias zugunsten externer Quellen eine zentrale Rolle. KI-Systeme stützen sich häufig stärker auf Testberichte, Foren, Vergleichsportale oder Fachmedien als auf Unternehmenswebsites. Deshalb umfasst die GEO-Analyse für Werbungtreibende auch die Frage, wie die Marke auf Drittplattformen dargestellt wird, mit welcher Tonalität sie erscheint und ob Produkteigenschaften im richtigen Kontext wiedergegeben werden. Klassische Klick-Metriken verlieren damit an Aussagekraft. Wichtiger werden Share of Voice, Sentiment, kontextuelle Genauigkeit und maschinenlesbare Vertrauenssignale wie Bewertungen, Zertifikate oder strukturierte Produktinformationen. Für Werbungtreibende misst GEO-Erfolg somit nicht die Klicks, sondern die Qualität und Dominanz der eigenen Präsenz im KI-Dialog.

Analytics-Tools

In der Praxis wird GEO aktuell vor allem mithilfe spezialisierter Monitoring-Tools ausgewertet. Diese messen, ob und wie häufig Marken, Domains oder Inhalte in KI-Systemen wie ChatGPT, Perplexity, Gemini oder Google AI Overviews erscheinen. Typische Funktionen sind Sichtbarkeitstracking auf Prompt-Basis, Wettbewerbsvergleiche, Quellen- und Zitationsanalysen, Sentiment-Auswertungen sowie automatisierte Reports.

Inhaltlich lassen sich die Tools meist drei Analyseebenen zuordnen:

1

Die reine Sichtbarkeitsmessung

Kommt eine Marke in KI-Antworten vor?

2

Die Quellen- und Reputationsanalyse

Welche Domains, Zitate und Narrative werden von den Modellen aufgegriffen.

3

Die Gap-Analyse

Analyse für welche Prompts Wettbewerber zitiert werden, die eigene Marke aber nicht.

Wichtig ist allerdings, die Ergebnisse methodisch einzuordnen. Nicht alle Tools messen dieselben Plattformen, Oberflächen oder Prompt-Sets. Entsprechend unterscheiden sich auch die Ergebnisse. Für Publisher sind vor allem Zitations-, Quellen- und Crawler-Analysen relevant, für Advertiser zusätzlich Sentiment-, Empfehlungsstärke- und Wettbewerbsanalysen entlang der Buying Journey. GEO-Analytics ist somit kein Standardreport, sondern ein je nach Geschäftsmodell unterschiedlich gewichtetes Set komplementärer Auswertungen.

Die aktuell verfügbaren GEO-Tools unterscheiden sich vor allem dadurch, ob sie SEO-nah arbeiten, modellübergreifende Sichtbarkeit messen, Quellen- und Zitationsanalysen in den Vordergrund stellen oder komplette Konversationsverläufe auswerten. Für Publisher sind insbesondere Citation-, Quellen- und Crawler-Analysen relevant, während für Advertiser zusätzlich die Aspekte Sentiment, Empfehlungslogik, Funnel-Bezug und Conversion-Messung wichtiger werden.

Tool	Schwerpunkt	Erfasste Plattformen / Modelle	Zentrale Analysefunktionen
SE Ranking	Kombination aus klassischem SEO und GEO	Google AI Overviews, AI Mode, ChatGPT, Gemini, Perplexity	KI-Ergebnis-Tracking, Wettbewerbsanalyse, Visibility Score, Sentiment- und Prompt-Insights
Otterly.AI	Monitoring von AI-Suchprompts und Markensichtbarkeit	ChatGPT, Perplexity, Google	Prompt-Tracking, Erwähnungsanalyse, Keyword-Recherche für Conversational Search, Wettbewerbsanalyse, Reporting
Peec AI	Multi-LLM-Sichtbarkeit und Quellenanalyse	Je nach Quelle u. a. ChatGPT, Perplexity, Google AI Overviews, Claude, Gemini, DeepSeek, Grok	Multimodell- bzw. Multi-LLM-Tracking, Ländersteuerung, Wettbewerbs- und Quellenanalyse, automatisierte Reports
Rankscale AI	Reputations- und Citation-orientierte GEO-Analyse	Wichtige LLMs wie ChatGPT, Gemini etc.	Quellen-Mapping, Citation-Analyse, Sentiment, AI-Readiness-Audits
Superlines	Enterprise-orientiertes GEO-Tracking	Globale AI-Plattformen / Sprachmodelle	Lokalisierte Konversationsanalysen, Custom-Tracking, API- und Exportoptionen, Fokus auf DSGVO-Konformität
SISTRIX	SEO-nahe Auswertung von AI Overviews	Vor allem Google AI Overviews	Keyword-basierte AI-Overview-Analyse, Verbindung von SEO-KPIs und AI-Overview-Sichtbarkeit
Finseo	GEO-Monitoring mit Audit-, Crawler- und Conversion-Bezug	Google AI Overviews, ChatGPT, Perplexity	GEO-Audits, Messung von AI-Crawler-Aktivitäten, Unterscheidung zwischen Bot- und Human-Traffic, GA-Integration für Conversion- / ROI-Tracking
Metrix.ai	Dialog- und Intent-orientierte Analyse	KI-Konversationen / GEO- und SEO-Umfelder	Monitoring ganzer Konversationsverläufe, Zielgruppen-Mapping, Keyword2Prompt-Generator, semantische Trefferanalyse
Meikai	Enterprise-GEO / AI-Visibility- & Citation-Intelligence	u. a. ChatGPT, Google AI, Perplexity (weitere „und mehr“)	Citation Tracking, Competitor Benchmarking, priorisierte Actions / Opportunities für Optimierung
Minddex	GEO-/AI-Visibility-Messung & Aktivierung (Share-of-Answer/ Share-of-Voice)	ChatGPT, Gemini, Perplexity	GEO Score, AI Share of Voice, tatsächliche Zitate & Quellen je Query („Share-of-Answer“), Empfehlungen zur Sichtbarkeitssteigerung
Similarweb	Messung von AI-Chatbot-Referral-Traffic (Impact/Attribution statt „nur“ Erwähnungen)	Referral-Traffic aus AI-Chatbots (u. a. ChatGPT, Perplexity, Claude, Copilot, Gemini, DeepSeek)	Eigener Traffic-View für AI Chatbot Referrals, inkl. Top Landing Pages und Top Prompts, Trend- & Wettbewerbsvergleich
Adobe LLM Optimizer	Enterprise-Lösung für GEO/AI-Search & Discovery Optimization	Fokus auf Brand Presence, Empfehlungen & Optimierungen	Automatische Insights / Recommendations, Identifikation technischer Visibility-Gaps („what AI agents can't read“), One-click Deployment von Optimierungen, Analyse von Offsite-Quellen (Social/Community)

Zukunftsblick

Entscheidend für die kommenden Jahre ist die Frage, wie Inhalte so modelliert und strukturiert werden müssen, dass LLMs sie über unterschiedliche Kontexte und Interfaces hinweg zuverlässig verstehen, begründen und wiederverwenden können.

Maschinenlesbare Evidenz und Quellenketten

Mit wachsendem Druck auf Fact-Checking und Halluzinationskontrolle werden LLMs Quellen bevorzugen, bei denen Behauptung und Beleg sauber verbunden sind. Inhalte mit originellen Statistiken und klarer Attribution erzielen bereits heute eine deutlich höhere Zitierrate in KI-Antworten.¹⁵

Künftig braucht jede zentrale Aussage eine eigene, klar abgegrenzte Einheit: kurze Kernaussage, dazugehörige Zahl oder Studie und präzise Quellenangabe (DOI, Link, Datensatz). Tabellen, „Key Findings“-Boxen und Evidenz-Abschnitte werden damit zu bevorzugten Andockpunkten für LLMs.

Strukturell heißt das, dass Recherchen als Beweisketten modelliert werden. Statt „Quelle im Fließtext“ entsteht ein Content-Design, in dem Evidenzblöcke maschinell isolierbar sind und sich in Antworten samt Begründung übernehmen lassen.

Unterschiedliche User Interfaces

Der gleiche Inhalt muss künftig in AI Overviews, Chat-Interfaces, Voice-Ausgaben und Agent-Workflows funktionieren. Dafür braucht es mehrschichtige Inhalte: eine sehr kurze Kernaussage, eine kompakte, Answer-First formulierte Zusammenfassung und eine ausführliche Ebene mit Kontext und Differenzierung.

AI Overviews nutzen vor allem die mittlere Ebene; Chatbots und Voice-Systeme kombinieren Frageüberschriften und kurze Antwortblöcke; Agents greifen zudem auf klar strukturierte Handlungsbausteine („Voraussetzungen“, „Schritte“, „Optionen“) zu.

Kernprinzip: Jeder relevante Abschnitt muss auch für sich funktionieren – als eigenständiger Antwortbaustein, der sich ohne den gesamten Artikel in verschiedenen Interfaces verwenden lässt. Das verstärkt die heute schon empfohlene Answer-First-Struktur und die konsequente Segmentierung in logisch abgeschlossene Chunks.

Standardisierung von KI-spezifischen Protokollen und Labels

Neben dem „Was“ eines Inhalts wird das „Wie darf er genutzt werden?“ zum strategischen Faktor. Diskutierte Formate wie llms.txt ergänzen robots.txt: Statt nur Zugriffe zu verbieten, geben sie LLMs aktive Hinweise, welche Bereiche für Training oder Inferenz vorgesehen sind und wo „High-Value-Content“ liegt.

Auf Dokument- und Abschnittsebene werden maschinell auswertbare Labels wichtig: Nutzungsrechte (Training vs. nur Retrieval), Aktualität, Sensitivität (z.B. YMYL) und Gattung (News, Kommentar, Satire, Analyse, Faktencheck). Heute sind diese Informationen häufig nur visuell oder im CMS sichtbar; künftig müssen sie über strukturierte Daten oder Header explizit für LLMs zugänglich sein.

Damit werden Inhalte zu Datenobjekten mit Fachinhalt plus Policies. Nur wer klare, standardisierte Signale zu Rechten, Kontext und Aktualität sendet, kann erwarten, dass seine Inhalte von KI-Systemen systematisch und regelkonform genutzt werden.

¹⁵ <https://arxiv.org/html/2311.09735v3>



Multimodaler Content: verknüpfte Text-, Audio-, Bild- und Videoeinheiten

LLMs entwickeln sich schnell von Textsystemen zu multimodalen Modellen. Für Medienhäuser bedeutet das: Nicht der einzelne Artikel, Podcast oder Clip ist die Einheit, sondern das verknüpfte Wissensobjekt dahinter.

Zentral sind dabei drei Anforderungen:

- Vollständige, gut strukturierte Transkripte für Audio und Video, idealerweise mit Zeitstempeln und klaren Abschnitten, damit LLMs Aussagen präzise verorten und zitieren können.
- Konsequente Alt-Texte und Bildbeschreibungen, weil Inhalte, die ausschließlich im Bild ohne Beschreibung liegen, für KI faktisch unsichtbar bleiben.
- Explizite Verknüpfung der Formate über IDs und Metadaten (hasPart, isPartOf, about), sodass ein Modell erkennen kann, dass Artikel, Grafik, Video und Podcast dieselbe Story repräsentieren.

Damit verschiebt sich der Fokus von „Text plus Bewegtbild“ hin zu multimodalen Wissenspäckchen, die LLMs je nach Interface unterschiedlich anordnen können: als kurze Textzusammenfassung, als begründete Antwort mit Grafik oder als Audio-Snippet mit verlinktem Transkript. Je sauberer diese Struktur modelliert ist, desto optimaler kann Content in der nächsten LLM-Generation genutzt werden.

Glossar

Abkürzung	Begriff	Erklärung
LLM	Large Language Model	KI-Modell, das anhand großer Textmengen Sprachmuster lernt und Antworten generiert.
GEO	Generative Engine Optimization	Optimierung von Inhalten, um in Antworten generativer Systeme zitiert oder als Quelle genutzt zu werden.
AEO	Answer Engine Optimization	Optimierung von Inhalten für Systeme, die direkte Antworten statt Linklisten ausspielen.
SEO	Search Engine Optimization	Optimierung für klassische Suchmaschinen-Rankings und organischen Traffic.
LLMO	LLM Optimization	Maßnahmen, um Inhalte LLM-freundlich aufbereitet, auffindbar und zitierfähig zu machen.
E-E-A-T	Experience, Expertise, Authoritativeness, Trustworthiness	Qualitätsrahmen zur Beurteilung von Glaubwürdigkeit und Fachkompetenz.
AIO	Artificial Intelligence Optimization	Optimierung von Inhalten, um in KI-Systemen zitiert zu werden



Quellen

<https://www.bvdw.org/wp-content/uploads/2026/02/The-Answer-Economy.pdf>
<https://backlinko.com/seo-vs-geo>, <https://neilpatel.com/blog/geo-vs-aeo/>
<https://ahrefs.com/blog/ai-search-overlap/>
<https://www.ewrdigital.com/llm-visibility-glossary>
Experience (Erfahrung), Expertise (Fachwissen), Authoritativeness (Autorität) und Trustworthiness (Vertrauenswürdigkeit)
<https://services.google.com/fh/files/misc/hsw-sqrg.pdf>
<https://developers.openai.com/api/docs/bots>
<https://commoncrawl.org/>
<https://cdn.openai.com/papers/gpt-4.pdf>
<https://arxiv.org/html/2601.22108v1/>
<https://www.bvdw.org/wp-content/uploads/2025/12/Retrieval-Augmented-Generation.pdf>
<https://www.airops.com/blog/featured-snippet-optimization>
<https://llmstxt.org/>
<https://developer.chrome.com/docs/lighthouse/agentive-browsing/llms-txt>
<https://developers.google.com/search/docs/fundamentals/ai-optimization-guide>
<https://ahrefs.com/blog/most-cited-domains-ai-mode/>
<https://developers.google.com/search/docs/appearance/structured-data/intro-structured-data>
<https://developers.google.com/search/blog/2019/05/new-in-structured-data-faq-and-how-to>
<https://www.bvdw.org/wp-content/uploads/2025/12/Auswirkungen-von-KI-Suche-auf-Publisher-Geschaeftsmodelle.pdf>
<https://arxiv.org/html/2311.09735v3>
<https://www.averi.ai/blog/building-citation-worthy-content-making-your-brand-a-data-source-for-llms>
https://en.wikipedia.org/wiki/Artificial_intelligence_optimization
<https://www.searchenginejournal.com/microsoft-explains-how-to-optimize-content-for-ai-search-visibility/557893/>
https://www.youtube.com/watch?v=rvLD6WJ4_fc
<https://searchengineland.com/google-ai-overviews-hurt-click-through-rates-454428>
<https://backlinko.com/seo-vs-geo>
<https://shellypalmer.com/2025/06/the-aeo-tactical-playbook-machine-readable-websites-answer-engine-optimization/>
<https://lseo.com/generative-engine-optimization/the-role-of-e-e-a-t-in-generative-engine-optimization/>
<https://www.searchenginejournal.com/the-ethical-side-of-ai-in-seo/516171/>
<https://asia.hsm.ai.org/2025/10/13/what-is-geo-vs-seo-and-why-it-matters-now/>
<https://arxiv.org/abs/2509.08919>
<https://totheweb.com/wp-content/uploads/2025/03/2025-03-ToTheWeb-AI-Search-GEO-Checklist-for-content-discovery.pdf>
<https://seranking.com/>
<https://otterly.ai/>
<https://peec.ai/>
<https://rankscale.ai/>
<https://www.superlines.io/>
<https://www.sistrix.de/>
<https://www.finseo.ai/>
<https://metrix.ai/>
<https://www.meikai.ai/>
<https://minddex.ai/>
<https://www.similarweb.com/>
<https://business.adobe.com/products/llm-optimizer.html>



Bundesverband Digitale Wirtschaft (BVDW) e.V.

Der Bundesverband Digitale Wirtschaft (BVDW) e. V. ist die Interessenvertretung für Unternehmen, die digitale Geschäftsmodelle betreiben oder deren Wertschöpfung auf dem Einsatz digitaler Technologien beruht. Mit seinen Mitgliedern aus der gesamten Digitalen Wirtschaft gestaltet der BVDW bereits heute die Zukunft – durch kreative Lösungen und modernste Technologien. Als Impulsgeber, Wegweiser und Beschleuniger digitaler Geschäftsmodelle setzt der Verband auf faire und klare Regeln und tritt für innovationsfreundliche Rahmenbedingungen ein. Dabei hat der BVDW immer Wirtschaft, Gesellschaft und Umwelt im Blick. Neben der DMEXCO, der führenden Fachmesse für Digitales Marketing und Technologien, und dem Deutschen Digital Award richtet der BVDW auch den CDR-Award, die erste Preisverleihung im DACH-Raum für Digitale Nachhaltigkeit und Verantwortung sowie eine Vielzahl von Fachveranstaltungen aus.

OVK

Der Online-Vermarkterkreis (OVK) im BVDW ist die Interessenvertretung der Online-Display- und -Video-Vermarkter am deutschen Werbemarkt. Er setzt sich für die Stärkung des nationalen Online-Werbemarktes und die Erhaltung seiner Angebotsvielfalt ein. Gemeinsam mit den Marktpartnern entwickelt und fördert er Standards und Regelwerke. Als OVK liefert er mit seinen Mitgliedern Orientierung und stellt Markttransparenz her. Er agiert lösungsorientiert; Qualität, Nachhaltigkeit und Zukunftsfähigkeit stehen im Mittelpunkt der Arbeit.

Kontakt

Nicole Dreyer, Senior Programm Managerin, dreyer@bvdw.org

Bundesverband Digitale Wirtschaft (BVDW) e.V.

Obentrautstraße 55, 10963 Berlin

www.bvdw.org

