

KI-CONTENT- STUDIE TEIL II

KI-TEXTE UND GOOGLE-
RANKINGS: WIE GUT
RANKEN KI-INHALTE?

GPT-40





KI STUDIE TEIL II

KI-TEXTE UND GOOGLE-RANKINGS:
WIE GUT RANKEN KI-INHALTE?

ZIEL



Vorrangiges Ziel ist die Evaluierung der Rankingfähigkeit von KI-Texten.

AUSGANGSSITUATION

Im Testzeitraum von **01. Mai bis 31. August 2024** hat die eo KI-Task-Force **118 Webseitexte** mit den Chatbots ChatGPT-4o (OpenAI) und Claude 3.5 Sonnet (Anthropic) generiert. Der wichtigste Punkt der eo KI-Content-Studie Teil II war zu prüfen, ob KI-Texte **Google's E-E-A-T-Modell** erfüllen können.

STUDIENOBJEKTE
Claude 3.5 Sonnet
& ChatGPT-4o



SHORT FACTS

- E-E-A-T-Check bei GPT-4o & Claude 3.5 Sonnet
- E-E-A-T-Check anhand SEO Qualitätsstandards
- Manuelle, empirische Bewertung der Textqualität
- Bewertung des Erfüllungsgrades der Kriterien in Prozent
- Menschliche eology SEO-Texte als Benchmark
- Testdomain die-familie.org zur Prüfung, wie Google mit KI-Inhalten umgeht
- 150 Texte und passende Bilder für die Domain
- Betextung der thematischen Verzeichnisse durch KI (GPT-4o, Claude 3.5 Sonnet) und ausgebildete eo SEO-Texter
- Analyse der Rankingentwicklung von KI-Texten vs. menschliche SEO-Texte

ERGEBNIS

“Die empirische Untersuchung von E-E-A-T bei KI-Texten durch die eology Online-Marketing-Experten, die Fachexperten und die SEO Performance der Domain die-familie.org verdeutlicht, dass **KI-Inhalte noch nicht alle Kriterien des SEO Qualitätsstandards** gleichermaßen und konsistent **erfüllen**.”

LEARNINGS

SEO-TEXTERSTELLUNG

Blogbeitrag 1.000 Wörter

Zeitersparnis

53 %

Unterkategorieseite

Einleitungstext (300 W.) &
vier Teasertexte (à 100 Wörter)

Zeitersparnis

46 %

Hauptkategorieseite

Einleitungstext (300 W.) &
vier Teasertexte (à 100 Wörter)

Zeitersparnis

37 %

INHALT

Vorwort	4
Short Facts zur eo KI-Content-Studie Teil II	6
Studienbeschreibung: was, wann, wie, wer? Vorgehen und mehr	7
Bewertung der Texte durch eo KI-Task-Force	14
Bewertung der Texte durch Experten	22
Bewertung durch Google – SEO-Performance	30
In 5 Schritten zu rankingfähigen Texten	35
Realitätscheck: Wie viel Zeit spart KI wirklich?	37
Fazit & Ausblick	41

VORWORT

In einem Jahr seit der Veröffentlichung des ersten Teils unserer KI-Content-Studie ist viel passiert. Die KI-Systeme haben sich rasant weiterentwickelt, es sind neue Modelle auf den Markt getreten mit neuen Anwendungsmöglichkeiten bis hin zu KI-Suchmaschinen. Die rasante Entwicklung hat die Content-Erstellung revolutioniert, die Auswirkungen sind zu spüren: Chatbots zu nutzen, wird im Online-Marketing immer selbstverständlicher. Die Konferenzen sind voll mit Vorträgen zu den Einsatzmöglichkeiten von KI.

Der „Hype“ hat Gründe: Viele Unternehmen setzen große Hoffnungen in die KI-Nutzung. Eine Umfrage von Forbes Advisor zeigt:

- 64 % der Unternehmen glauben, dass KI ihre Produktivität steigert.
- 59 % erwarten Kosteneinsparungen.
- 42 % sehen Chancen zur Prozessoptimierung.

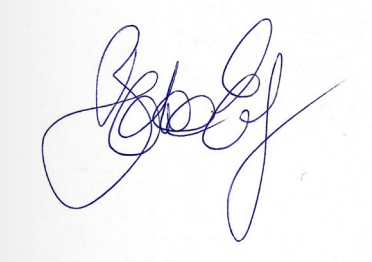
Diese Zahlen verdeutlichen den hohen Stellenwert von KI im Content-Marketing. Unternehmen erhoffen sich nicht nur Effizienzsteigerungen, sondern auch die Möglichkeit, kreative Prozesse zu automatisieren und Ressourcen zu sparen. Doch wie realistisch sind diese Erwartungen? Wie viel Zeit lässt sich beispielsweise bei der Texterstellung tatsächlich einsparen?

KI-Systeme wie ChatGPT-4o (OpenAI) oder Claude 3.5 Sonnet (Anthropic) können mittlerweile Texte erzeugen, die sich auf den ersten Blick kaum von menschlichen unterscheiden. Doch wie leistungsfähig sind diese Inhalte wirklich, wenn es um ihre Sichtbarkeit in Suchmaschinen wie Google geht? Und können sie die gleichen Standards erfüllen wie menschlich erstellte Inhalte?

Dieses Whitepaper beleuchtet das Potenzial und die Grenzen von KI-gestütztem Content im Hinblick auf SEO. Ihr erfahrt, wie KI-Inhalte auf Google ranken, welche Herausforderungen sie mit sich bringen und welche Schritte erforderlich sind, um sie rankingfähig zu machen. Dabei gehen wir auch auf spezifische technische, strategische und kreative Aspekte ein, die für die Erstellung hochwertiger Inhalte entscheidend sind.

Darüber hinaus werden wir die Grenzen der Automatisierung erörtern und aufzeigen, wie menschliche Expertise eine zentrale Rolle spielt. Ihr bekommt praktische Handlungsempfehlungen, wie ihr das Beste aus KI-Tools herausholt und dennoch die hohen Anforderungen an qualitativ hochwertige Inhalte erfüllen könnt.

Wir wünschen sehr viel lehrreiche Erkenntnisse beim Lesen dieses Whitepapers



Dr. Beatrice Eiring

Head of Content Creation eology GmbH, Leitung eo KI-Content-Studie

SHORT FACTS ZUR EO KI-CONTENT-STUDIE TEIL II

- E-E-A-T-Check bei GPT-4o & Claude 3.5 Sonnet
- E-E-A-T-Check anhand des SEO Qualitätsstandards
- Manuelle, empirische Bewertung der Textqualität
- Bewertung des Erfüllungsgrades der Kriterien in Prozent
- Menschliche eology SEO-Texte als Benchmark
- Testdomain die-familie.org zur Prüfung, wie Google mit KI-Inhalten umgeht
- 150 Texte und passende Bilder für die Domain
- Betextung der thematischen Verzeichnisse durch KI (GPT-4o, Claude 3.5 Sonnet) und ausgebildete eo SEO-Texter
- Analyse der Rankingentwicklung von KI-Texten vs. menschliche SEO-Texte

STUDIENBESCHREIBUNG: WAS, WANN, WIE, WER? VORGEHEN UND MEHR

Im Testzeitraum von 01. Mai bis 31. August 2024 hat die eo KI-Task-Force 118 Websitetexte mit den Chatbots ChatGPT-4o (OpenAI) und Claude 3.5 Sonnet (Anthropic) generiert. Die von den Chatbots erstellten Texte wurden evaluiert, be- und eingewertet und mit Inhalten verglichen, die von erfahrenen Textern der eology GmbH verfasst wurden:

Hinweis

Bei der Bezeichnung „menschlich erstellt“ bedarf es einer Definition, die Eindeutigkeit schafft. Denn nicht jeder Mensch besitzt die Fähigkeit (Kenntnisse der deutschen Sprache, redaktionelle Ausbildung etc.) und das Talent (Affinität und Begabung für das Schreiben), um hochwertige Texte zu erstellen. Im Fall der eo KI-Content-Studie kommt der Aspekt der SEO-Kenntnisse hinzu.

Aus diesem Grund werden die KI Texte in der eo KI Studie mit Texten ausgebildeter SEO Redakteure verglichen und anhand des Qualitätsstandards bewertet, dem die eology Content Creation für hochwertige Mehrwert-Inhalte mit hohem Ranking-Potential folgt.

Vorrangiges Ziel der eo KI-Content-Studie Teil II ist die **Evaluierung der Rankingfähigkeit von KI-Texten**. Google bewertet Inhalte nach verschiedenen Kriterien, um sicherzustellen, dass Nutzer die relevantesten und qualitativ hochwertigsten Ergebnisse erhalten. Ein **zentraler Maßstab ist das sogenannte E-E-A-T-Modell** (Experience, Expertise, Authoritativeness, Trustworthiness). Diese Kriterien dienen der Bewertung der Qualität von Inhalten und der Vertrauenswürdigkeit von Quellen. Demnach müssen Webseiteninhalte Expertise, Erfahrung, Autorität und Vertrauenswürdigkeit vermitteln. Ist dies nicht gegeben, ranken die Webseiten schlechter oder werden sogar abgestraft. Dies kann die gesamte Domain und nicht nur einzelne Seiten mit schlechten Inhalten treffen.

Bedeutung von E-E-A-T

Das E-E-A-T-Modell stellt hohe Anforderungen an die Qualität von Inhalten. Jedes Kriterium hat dabei seine eigene Gewichtung und Relevanz:

Experience (Erfahrung): Inhalte sollten zeigen, dass der Autor oder die Quelle praktische Erfahrung in dem behandelten Thema hat. Dies kann durch Beispiele, Fallstudien oder Praxisberichte unterstrichen werden.

Expertise (Fachwissen): Der Autor sollte Fachkenntnisse besitzen, die über allgemeines Wissen hinausgehen. Google erkennt Expertise durch die Tiefe und Genauigkeit der bereitgestellten Informationen.

Authoritativeness (Autorität): Die Quelle muss als autoritativ wahrgenommen werden. Dies geschieht durch Verweise, Erwähnungen und Verlinkungen von vertrauenswürdigen Seiten.

Trustworthiness (Vertrauenswürdigkeit): Inhalte sollten frei von Fehlern sein und sich auf vertrauenswürdige Quellen stützen. Transparenz, wie etwa die Angabe der Autoren, stärkt das Vertrauen.

Besonders kritisch ist es, wenn E-E-A-T **bei sogenannten Your-Money-Your-Life-Inhalten** nicht erfüllt ist. Seiten, die sich mit den Themen Gesundheit, Finanzen und rechtliche Informationen befassen. Die Themen, für die Texte mithilfe der KI für die Studie erstellt wurden, reichen von soften Themen wie „Kindergeburtstag feiern“ bis hin zu medizinischen Themen wie die Krankheit „Scharlach“ oder rechtliche Themen wie „Kündigungsschutz in der Schwangerschaft“.

Der wichtigste Punkt der eo KI-Content-Studie Teil II war zu prüfen, ob KI-Texte Google's E-E-A-T-Modell erfüllen können. **Diese Prüfung erfolgte in 3 Schritten:**

1. Evaluierung des Erfüllungsgrades der Zielgruppenorientierungskriterien des SEO Qualitätsstandards durch die Experten und Expertinnen der eo KI-Task-Force
2. Evaluierung von E-E-A-T durch Experten und Expertinnen aus den Fachbereichen Jura, Medizin, Finanzen
3. Evaluierung der Rankingfähigkeit über die Testdomain die-familie.org

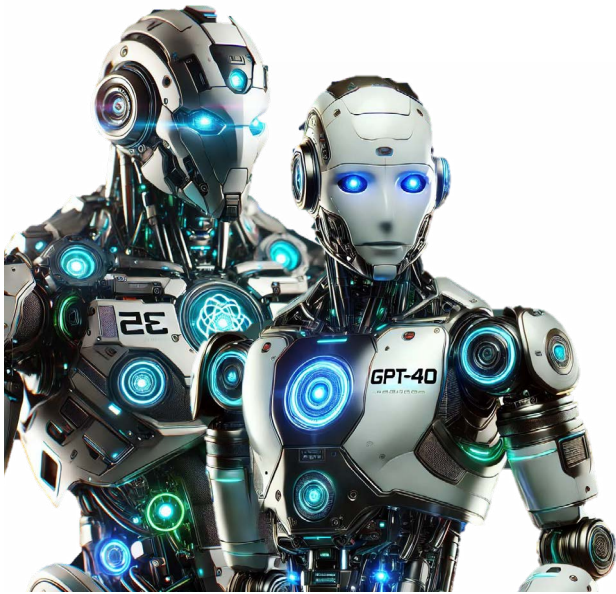
Hinweis:

Da es sich um eine empirische Studie handelt, der ein begrenztes Text-Korpus zugrunde liegt mit dem Ziel, das SEO Performance Potential zu beurteilen, können andere Methoden sowie andere Testszenarien und Use Cases zu abweichenden Ergebnissen gelangen. Die eo KI-Task-Force freut sich über konstruktives Feedback, fachlichen Austausch und Diskussionen.

DIE EO KI-TASK-FORCE

- **Dr. Beatrice Eiring, Sprachwissenschaftlerin, Head of Content Creation eology GmbH:** „Ich bin davon überzeugt, dass AI die SEO Texterstellung langfristig und nachhaltig verändert. Denn sie sorgt für eine Produktivitätssteigerung, die aus unternehmerischer Sicht immer erstrebenswert ist. Doch sollten dabei Aspekte, die ich (und andere) seit Jahren predige und die Suchmaschinen (Google!) anstreben, wie Mehrwert, Einzigartigkeit und Nutzerfokussierung nicht vernachlässigt werden! Ob die nicht mehr aufzuhaltenden Veränderungen insgesamt zum Guten sind, darauf haben am Ende nur wir Menschen Einfluss...“
- **Dr. Bettina Zehner, Sprachwissenschaftlerin, Content Consultant eology GmbH:** „Ich war früher schon richtig schnell bei der Contenterstellung. Jetzt mit ChatGPT bin ich noch schneller. Doch eine sorgfältige Prüfung des Outputs relativiert die neugewonnene Schnelligkeit. Denn ohne den Text auf seine Richtigkeit hinsichtlich Inhalt, Aufbau und Rechtschreibung zu prüfen, würde ich keinen AI Content veröffentlichen oder an Kunden weitergeben.“
- **Antonios Smyrniaios, Kultur- und Literaturwissenschaftler, KI & Prompt Expert eology GmbH:** „Ein KI-Tool ist am Ende des Tages eben ein Werkzeug. Die Qualität der Arbeit hängt davon ab, wie man es einsetzt. In den Händen einer Person, die weiß, was die KI kann und was (noch) nicht, zeigt sie auch ihr größtes Potential. Somit ist KI eine Stütze und schafft ein solides Fundament, auf dem der Mensch qualitativen Output schaffen kann.“
- **Elisa Greubel, Kommunikationsexpertin, KI & Prompt Expert eology GmbH:** „Die Prompt-Erstellung für ChatGPT ist ein iterativer Prozess. Es gibt nicht DEN perfekten Prompt – vielmehr muss ein Prompt stetig an die jeweilige Anforderung angepasst werden, um so möglichst guten Output zu erhalten.“
- **Larissa Köberlein, Sprachwissenschaftlerin, Teamlead Content Creation eology GmbH:** „KI kann als Sparringspartner bei vielen Aufgaben rund um die Content Creation unterstützen. Die Ergebnisse sind allerdings nicht 1:1 verwendbar – sie können vielmehr als Grundlage dienen, aus der gemeinsam mit menschlicher Kreativität rundum stimmiger Content mit echtem Mehrwert entsteht. Eine gründliche Qualitätssicherung des KI-Outputs ist dabei ein Muss.“
- **Christoph Herbig-Fleißner, Sprachwissenschaftler, Teamlead Content Creation eology GmbH:** „Bereits heute kann KI bei der Texterstellung unterstützen. Allerdings ist es wichtig zu wissen, wo die KI an ihre Grenzen stößt. Auch wenn die Qualität der Texte in vielen Bereichen hoch ist, kann die KI einem erfahrenen Texter nicht das Wasser reichen.“
- **Theresa Brieg, Sprach- und Literaturwissenschaftler, Social Media Consultant eology GmbH:** „KI als Hilfsmittel für Social Media Content öffnet Türen zu grenzenlosen kreativen Möglichkeiten, sei es in textueller oder visueller Form. Sie ist ein kraftvolles Werkzeug, um Content-Ideen zu generieren und zu verfeinern. Doch der Schlüssel zum Erfolg liegt im richtigen Umgang und der Überprüfung und Feinjustierung des Outputs durch einen Menschen. Denn unsere Emotionen und einzigartigen Charakterzüge sind das, was wir der KI voraus haben und ihr mitgeben müssen, um unigen Content für Brands zu kreieren, der begeistert und berührt.“





DIE STUDIENOBJEKTE:
GPT & CLAUDE

Warum ChatGPT?

- **“Lebensalter“:** Am 30.11.24 wurde der OpenAI Chatbot 2 Jahre alt.
- **Userzahlen:** Innerhalb von 2 Monaten nach dem Launch erreichte der Chatbot 100 Millionen Nutzer. Weltweit verzeichnet das KI-System circa 1,5 Milliarden Besuche pro Monat.
- **Multimodales Modell:** ChatGPT könnt ihr für verschiedene Aufgaben entlang der Content Creation einsetzen:
 - Textoutput erstellen
 - Mit Bing auf das Internet zugreifen
 - Mit Advanced Data Analysis Dokumente wie PDFs, Excel-Dateien und Co. lesen und auswerten
 - Über DALL E-Schnittstelle Bilder generieren
 - Über Plugins weitere Funktionen nutzen wie Canva für Grafiken oder Slide Maker für Präsentationen
- ChatGPT-4o gibt es seit Mai 2024. Das „o“ in GPT-4o steht für „omni“, was die Fähigkeit des Modells unterstreicht, verschiedene Eingabeformen wie Text, Bilder und Audio zu verarbeiten und darauf zu reagieren.
- **Verbesserungen** im Vergleich zu seinen Vorgängern: höhere Geschwindigkeit und Effizienz, indem es Texte doppelt so schnell generiert und dabei 50 % kostengünstiger ist.
- Seit Dezember 2024 gibt es o1 und o1-mini. Die Unterschiede seht ihr in der Tabelle. Da diese beiden Versionen nach der Testphase gelauncht wurden, sind sie nicht in Teil II der Studie eingeflossen.

Unterschiede zwischen den GPT-Modellen

GPT-3.5	Dieses Modell war hauptsächlich für textbasierte Aufgaben konzipiert und konnte Texte generieren sowie einfache Anweisungen verarbeiten.	Erweiterte die Möglichkeiten durch die Einführung von multimodalen Eingaben, einschließlich Text und Bild, und verbesserte das Verständnis komplexer Anweisungen.
GPT-4	Einschränkungen: Begrenzte Kontextverarbeitung und fehlende multimodale Fähigkeiten.	Höhere Rechenanforderungen und Kosten im Vergleich zu GPT-3.5.
GPT-4o	Bietet eine verbesserte Leistung mit der Fähigkeit, Text, Bilder und Audio in Echtzeit zu verarbeiten. Schneller und kosteneffizienter als GPT-4 und erzielt herausragende Ergebnisse in mehrsprachigen und visuellen Benchmarks.	Obwohl das Modell vielseitiger ist, kann es in bestimmten spezialisierten Aufgaben den neueren Modellen unterlegen sein.
GPT-o1	Eingeführt im Dezember 2024, ist o1 darauf ausgelegt, komplexe Probleme zu lösen, indem es mehr Zeit für die „Überlegung“ vor der Antwort verwendet. Es übertrifft GPT-4o in Bereichen wie Programmierung, Mathematik und wissenschaftlichem Denken.	Kann aufgrund der tiefergehenden Analyse längere Antwortzeiten haben.
GPT-o1-mini	Eine kompaktere und schnellere Version von o1, optimiert für weniger komplexe Aufgaben, bei denen Geschwindigkeit und Effizienz im Vordergrund stehen.	Weniger leistungsfähig bei der Lösung hochkomplexer Probleme im Vergleich zum vollständigen o1-Modell.

Tabelle 1: Unterschiede zwischen den GPT-Modellen

Zusammenfassung: Die Entwicklung der GPT-Modelle zeigt einen klaren Trend hin zu größerer Vielseitigkeit, Geschwindigkeit und Effizienz. Während GPT-3.5 und GPT-4 den Grundstein legten, brachte GPT-4o erhebliche Verbesserungen in der Verarbeitung multimodaler Eingaben. Die o1-Modelle konzentrieren sich auf tiefere Analysefähigkeiten und erweitern die Einsatzmöglichkeiten der KI in spezialisierten Bereichen. Mehr über die Evolution von GPT-3.5 bis GPT-4o lest ihr im Kapitel [„ChatGPT-Evolution“](#).

Warum Claude?

- Claude 3.5 Sonnet ist ein fortschrittliches KI-Sprachmodell des Unternehmens Anthropic
- **Anthropic:** 2021 von ehemaligen OpenAI-Mitarbeitern gegründet; ihre KI ist nach Claude Shannon benannt, dem „Vater der Informationstheorie“
- **Hauptziel:** KI-Modelle zu entwickeln, die sicher, zuverlässig und ethisch sind. Ein zentrales Anliegen von Anthropic ist die Förderung von KI-Systemen, die in Übereinstimmung mit menschlichen Werten agieren und potenzielle Risiken minimieren. Anthropic positioniert sich mit Claude als führender Anbieter von sicheren und leistungsstarken KI-Lösungen, die vielseitig einsetzbar und gleichzeitig auf den Schutz von Nutzerinteressen ausgerichtet sind.
- **Kennzeichen:** hohe Verarbeitungsgeschwindigkeit und ein erweitertes Kontextfenster von 200.000 Tokens ermöglicht die Verarbeitung umfangreicher Texte
- **Schwerpunkt auf Sicherheit:** Claude wurde entwickelt, um Missbrauch und schädliche Interaktionen so weit wie möglich zu minimieren. Das Modell basiert auf „Constitutional AI“, einem Trainingsansatz, der klare Regeln und Werte integriert, damit die KI konsistent und sicher agiert.
- **Einsatzmöglichkeiten:** Automatisierte Kundeninteraktion, Unterstützung bei kreativen Prozessen wie Schreiben und Brainstorming, Datenanalyse und -Zusammenfassung, Programmierhilfe
- **Verbesserte Dialogfähigkeiten:** Die Modelle sind darauf trainiert, natürliche und nuancierte Konversationen zu führen.
- In verschiedenen Disziplinen erzielt Claude 3.5 Sonnet bessere Ergebnisse als GPT-4.

DAS VORGEHEN

Der E-E-A-T-Check der KI-Inhalte erfolgte in 3 Schritten:

1. Schritt: Evaluierung der Zielgruppenorientierungskriterien

Zunächst hat die eo KI-Task-Force die Texte von ChatGPT-4o und Claude 3.5 Sonnet gemäß des SEO-Qualitätsstandards eingewertet. **E-E-A-T** hat sie **anhand der Zielgruppenorientierungskriterien** des Standards bewertet. Denn diese beschreiben detaillierter die Begriffe von Google's Modell. Detail-Informationen zum SEO Qualitätsstandards findet ihr im Whitepaper zum ersten Teil der Studie [hier](#).

Wie die beiden neuen Chatbots im Vergleich zu den älteren abgeschnitten haben und wie gut die eo KI-Task-Force E-E-A-T bewertet hat, lest ihr im nächsten Kapitel „Bewertung der Texte durch eo KI-Task-Force“.

2. Schritt: Bewertung durch Experten

Im 2. Schritt hat die eo KI-Task-Force Texte aus den Fachbereichen Jura, Medizin und Finanzen **Experten und Expertinnen zur Prüfung** vorgelegt. Die Experten bewerteten die fachliche Tiefe der Texte und die inhaltliche Korrektheit. Die Bewertung der Experten lest ihr im Kapitel [„Bewertung durch Google – SEO-Performance“](#).

3. Schritt: Beobachtung der Ranking-Entwicklung der Test-Domain

Die KI-Texte wurden auf die **Testdomain die-familie.org** hochgeladen. Die Domain ist nach Hauptkategorien mit mindestens 2 Unterverzeichnissen strukturiert. In einem Unterverzeichnis liegen immer nur Texte von einer Quelle. Quelle ist neben ChatGPT und Claude auch die eo KI-Task-Force. Die **menschlichen Texte dienen als Vergleichsgruppe**. Durch die Unterverzeichnisse soll gewährleistet werden, dass sich Sichtbarkeitsentwicklungen durch Keyword-Rankings nach den Quellen einfach unterscheiden lassen.

Nach dem Launch der Domain am 02.09.2024 wurde ein Monitoring aufgesetzt und ein regelmäßiges Reporting erstellt mit wichtigen Kennzahlen:

- Sichtbarkeitsentwicklung der Domain
- Sichtbarkeitsentwicklung der Unterverzeichnisse
- Beste Keyword-Rankings
- Top-30-Keyword-Rankings nach Quelle

Detail-Informationen zur SEO Performance findet ihr im Kapitel [„In 5 Schritten zu rankingfähigen Texten“](#).

BEWERTUNG DER TEXTE DURCH EO KI-TASK-FORCE

Die beiden neuen Testobjekte der Studie, GPT-4o und Claude 3.5 Sonnet, hat die eo KI-Task-Force genau wie die älteren Modelle nach dem SEO-Qualitätsstandard bewertet. Die folgenden Diagramme **vergleichen alle Chatbots in den 3 Kriterienbereichen** miteinander. Die Studie zeigt, dass KI-Inhalte in vielen Bereichen solide Ergebnisse liefern, jedoch **in einigen Schlüsselkategorien hinter den menschlichen Texten** zurückbleiben. Dies trifft vor allem auf die älteren Chatbots zu, die Neueren sind leistungsfähiger und konnten einige Schwächen ihrer Vorgänger überwinden, wie ihr im Kapitel „**Vergleich der KI-Systeme: GPT-4 und Claude 3.5 Sonnet**“ nachlesen könnt.

1. Sprachliche Kriterien:

- KI-Texte überzeugen oft durch ihre Flüssigkeit und allgemeine Verständlichkeit.
- Dennoch gibt es Schwächen:
 - Nominalisierungen: KI nutzt überproportional häufig lange Substantivkonstruktionen.
 - Sprachliche Varianz: Wiederholungen und fehlende Kreativität mindern die Textqualität.
 - Fehlende Zielgruppensensibilität: Texte sind oft generisch und greifen die spezifischen Bedürfnisse der Zielgruppe nicht auf.

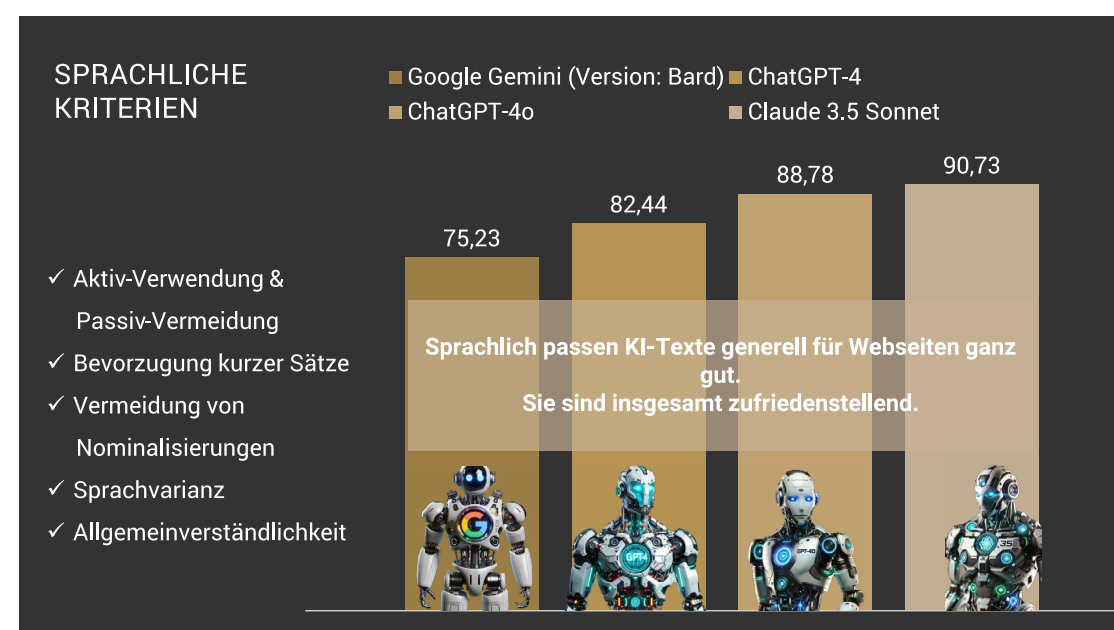


Abbildung 1: Sprachliche Kriterien

2. SEO-Kriterien

- SEO bleibt eine der größten Schwächen von KI-generierten Texten.
- Trotz Fortschritten in der Technologie haben KI-Texte oft Probleme bei der richtigen Keyword-Verwendung:
 - **Keywordspam:** Keywords werden manchmal unnatürlich oft und ohne grammatikalische Anpassung eingebaut.
 - **Fehlender Fokus:** Viele KI-Modelle verfehlen einen klaren Keyword-Fokus, was zu schlechterer SEO-Performance führt.
 - **Unzuverlässigkeit:** Essenzielle Keywords und W-Fragen fehlen immer wieder, insbesondere bei Systemen wie Claude 3.5 Sonnet.
- Gefahr: Diese Schwächen können dazu führen, dass Webseiten mit KI-generierten Texten ohne menschliche Nachbearbeitung Rankings verlieren oder gar nicht erst Sichtbarkeit erlangen.

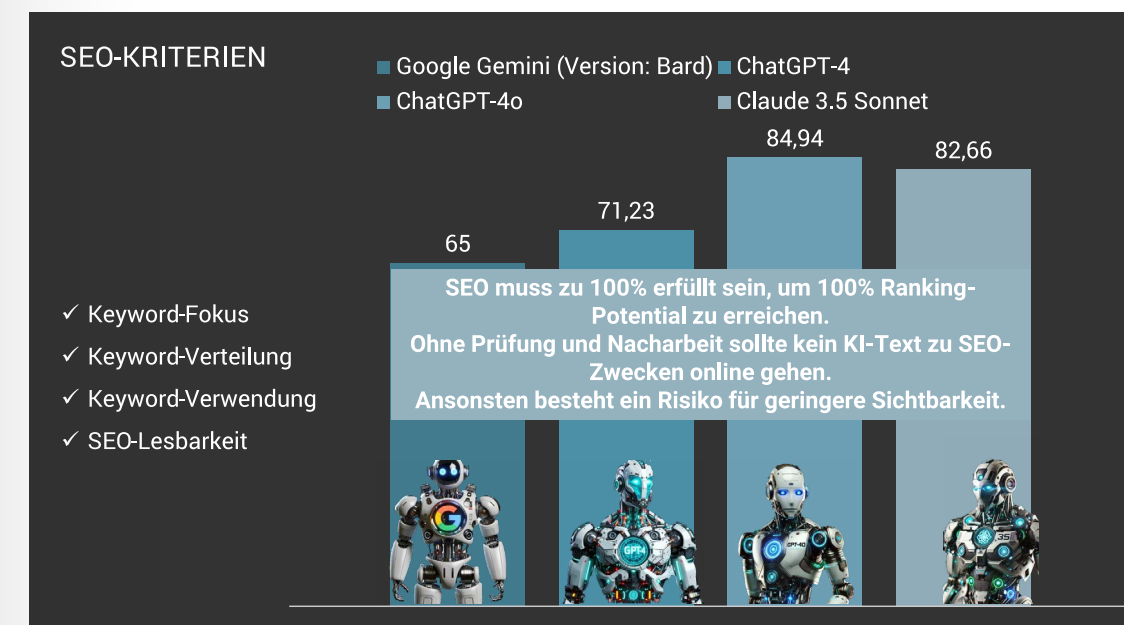


Abbildung 2: SEO-Kriterien

3. Zielgruppenorientierung

- KI liefert nicht immer passgenaue Inhalte.
- Besonders auffällig ist das Fehlen von:
 - **Korrektheit von Zahlen, Daten und Fakten:** Diese werden oft ungenau oder unvollständig dargestellt. Zahlen und Aussagen waren in einigen Fällen falsch oder nicht belegt.
 - **Mehrwert:** Konkrete Beispiele, die den Leser direkt ansprechen, fehlen häufig, genauso wie wichtige Detail. Das hat den Gesamteindruck beeinträchtigt.
 - **Kreativität:** KI-Texte waren oft weniger originell und boten kaum innovative Ansätze.
 - **Passendes Wording und Zielgruppenorientierung:** Die Sprache war oft zu generisch und wenig zielgruppenspezifisch.
 - **Leseransprache:** Emotionaler Bezug und Ansprache auf Augenhöhe fehlten.

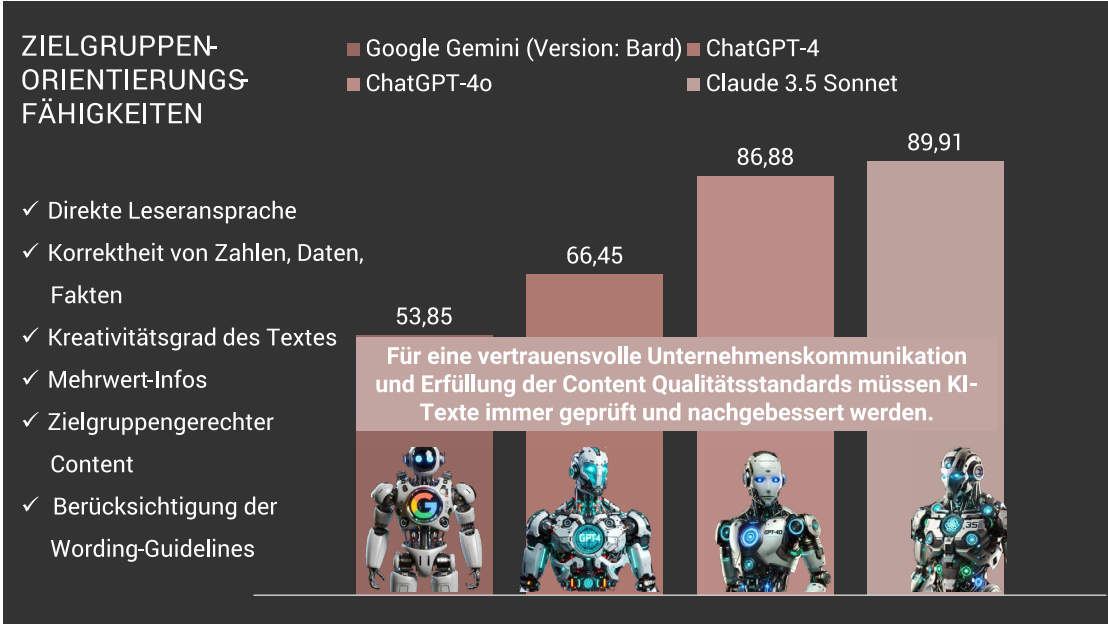


Abbildung 3: Zielgruppenorientierungsfähigkeiten

Die folgende Tabelle stellt die Erfüllungsgrade der Einzelkriterien aus allen 3 Bereichen für die getesteten Chatbots im Detail dar:

Kriterien-bereich	Kriterien	Google Gemini (Version: Bard)	ChatGPT 3.5 & 4	ChatGPT 4o	Claude 3.5 Sonnet
sprachlich	Aktiv-Verwendung	78,84 %	90 %	87,17 %	90,91 %
	Kurze Satzlängen	87,69 %	81,2 %	92,83 %	91,82 %
	Vermeidung Nominalisierung	71,92 %	80,4 %	88,91 %	88,64 %
	Sprachvarianz	56,92 %	70,2 %	82,39 %	87,73 %
	Allgemein-verständlichkeit	80,77 %	90,4 %	92,61 %	94,55 %
	Mittelwert	75,23 %	82,44 %	88,78 %	90,73 %
SEO	Keyword-Fokus	49,23 %	70,4 %	86,52 %	71,36 %
	Keyword-Verteilung	57,69 %	73,6 %	86,09 %	81,14 %
	Keyword-Verwendung	90 %	69,79 %	78,03 %	78,41 %
	SEO-Lesbarkeit	63,08 %	71 %	87,17 %	92,5 %
	Mittelwert	65 %	71,23 %	84,94 %	82,66 %
Ziel-gruppen-orientierung	Direkte Leseransprache	41,54 %	61,2 %	82,61 %	86,59 %
	Kreativitätsgrad	42,31 %	59 %	82,39 %	88,18 %
	Mehrwertinfos	52,31 %	60,8 %	83,26 %	92,95 %
	Zielgruppen-gerecht	56,92 %	70,2 %	88,04 %	92,27 %
	Berücksichtigung Wording	63,85 %	75,8 %	91,52 %	88,86 %
	Korrekte Zahlen, Daten, Fakten	87,5 %	93 %	89,52 %	87,27 %
	Mittelwert	53,85 %	66,45 %	86,88 %	89,91 %

Abbildung 4: Erfüllungsgrad der Einzelkriterien

CHATGPT-EVOLUTION

Die KI Task Force hat die älteren GPT-Modelle mit dem neueren GPT-4o verglichen. Es hat definitiv eine Evolution der KI stattgefunden, sie zeigt klare Fortschritte. Die folgende Tabelle stellt diese überblicksartig dar.

GPT-3.5	GPT-4	GPT-4o
Keywordspam	Texte lesen sich flüssiger	Empathischer
Forsche/provokative Zielgruppenansprache	Weniger Wortwiederholungen	Die Texte sind zielgruppen-gerechter, was Sprache und Inhalt anbelangt
Pseudo-kreativ: „Ausfallschritt-Charme“	Versteht komplexe Aufgaben besser & ist kreativer als GPT-3.5	Zuverlässiger bei SEO
Wenig Fingerspitzengefühl	Reine Wechsel zwischen Überschriften & Aufzählungen	Schafft nun längere Texte, allerdings meist nur nach erneuter Aufforderung mit inhaltlichen Redundanzen
Häufig Falschinfos	Gut klingende Aussagen, die aber nicht plausibel sind	Nach wie vor kein perfekter Text ohne Nacharbeitsaufwand
Textlänge maximal 500 – 800 Wörter	Verweigert gerne mal die Arbeit	
Fauler: schreibt tendenziell noch weniger Wörter als GPT-4		

Tabelle 2: ChatGPT-Vergleich

VERGLEICH DER KI-SYSTEME:
GPT-4 UND CLAUDE 3.5 SONNET

Die beiden neueren KI-Systeme hat die eo KI-Task-Force direkt miteinander verglichen. Das KI-Modell von Anthropic hat die Task Force komplett neu aufgenommen, daher war ein Vergleich mit dem OpenAI-Modell eine aufschlussreiche Analyse. **Beide Modelle sind von ihrer Leistungsfähigkeit ungefähr auf gleichem Niveau.** Sie haben jedoch individuelle Stärken. Gesondert zu erwähnen ist bei Claude: Die Texte des Anthropic Modells wirken öfters **natürlicher, menschlicher und emotionaler** als die von ChatGPT. Stilistisch haben sie oft besser zu den Anforderungen und Instruktionen gepasst. Dafür war **Claude bei SEO unzuverlässiger**, Halluzinationen und falsche Zahlen, Daten und Fakten findet man bei beiden Chatbots.

Stärken von GPT-4:

- Zuverlässigere SEO-Ergebnisse
- Höhere Textlängen möglich (nach erneuter Aufforderung)
- Verbesserte Sprachflüssigkeit und komplexe Aufgabenbewältigung

Stärken von Claude 3.5 Sonnet:

- Emotionalisierende Texte und zielgruppengerechte Inhalte
- Kreative Überschriften
- Berücksichtigung von SEO-Elementen wie Infoboxen oder Tabellen

E-E-A-T-CHECK BEI GPT-4O UND CLAUDE 3.5 SONNET

Den E-E-A-T-Check hat die eo KI-Task-Force bei Texten von ChatGPT-4o und Claude 3.5 Sonnet durchgeführt. Dabei hat sie die **Zielgruppenorientierungsfähigkeiten** der Chatbots als **Bewertungsgrundlage** herangezogen. Diese Kriterien beschreiben im Detail die 4 Begriffe des E-E-A-T-Modells. Die Grafik stellt dar, wie GPT-4o und Claude 3.5 Sonnet in den einzelnen Kriterien sowie im gesamten Durchschnitt abgeschnitten haben.

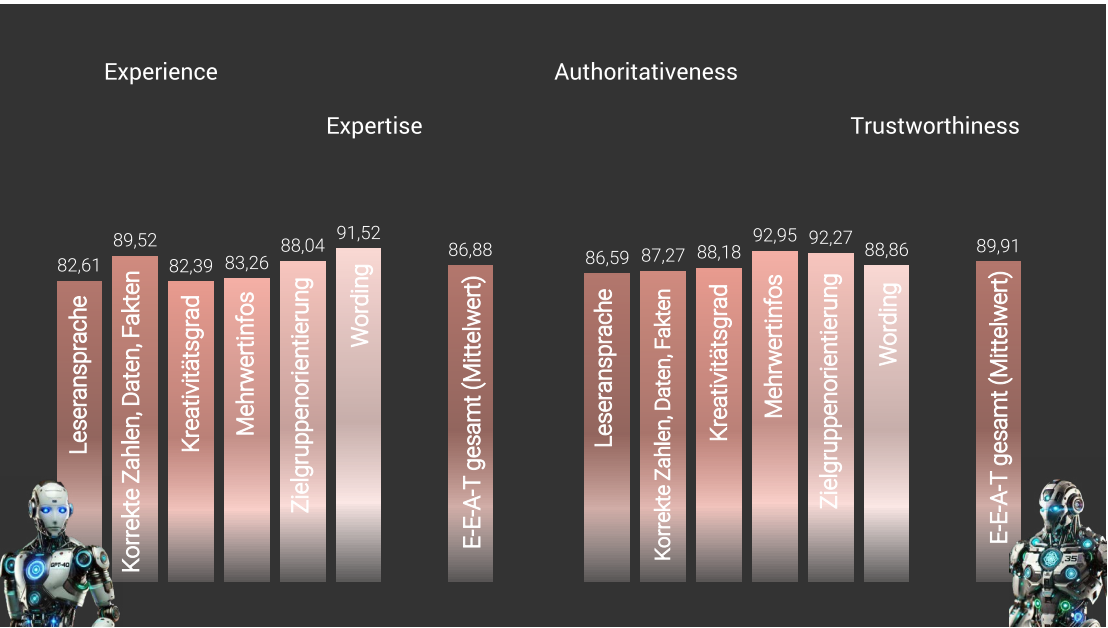


Abbildung 5: E-E-A-T-Check

Die Werte stellen den **Durchschnitt über alle Themen** dar. Darunter waren soft Themen wie „Kindergeburtstag feiern“, aber auch fachliche Themen wie „Kündigungsschutz während der Schwangerschaft“. Die Werte liegen alle bei über 80 % Erfüllungsgrad, teilweise über 90 %. Das sind keine schlechten Werte, es sind aber auch keine 100 %. Nun fällt es bei soften Themen nicht so ins Gewicht, wenn Kriterien nicht zu 100 % erfüllt sind. Anders sieht es jedoch bei **Your-Money-Your-Life-Themen** aus: Diese werden von Google besonders genau auf E-E-A-T geprüft. Bei schlechten Ergebnissen wird der Inhalt oder **im schlimmsten Fall die gesamte Domain** abgestraft. Bei YMYL-Themen gibt es tatsächlich einige Besonderheiten bzw. Qualitätsmängel.

YMYL-Inhalte stellen eine besondere Herausforderung dar. Die KI-Texte enthielten hier oft:

- **Falsche oder unscharfe Fakten:** Zahlen und Statistiken sind nicht immer korrekt. Beispielsweise wurden falsche Einkommensgrenzen für Versicherungen angegeben.
- **Fehlende Keywords = fehlende Informationen:** KI „vergisst“ gerne Keywords und mit den Keywords fehlen dann Informationen. Für holistische Texte ist das Keywordset aber entscheidend, um alle Zielgruppenfragen zu beantworten. Beispielsweise fehlten die Keywords „Mund“ und „Kinn“ im Text zu Scharlach und damit fehlte ein wichtiges Symptom der Krankheit, der sogenannte Milchbart, eine weiße Fläche um Mund und Kinn.
- **Redundanzen:** Inhalte wurden häufig wiederholt, ohne Mehrwert zu schaffen. Die Texte sind lang, bleiben aber sehr oberflächlich.
- **Eingeschränkte Kreativität:** KI folgt den Prompts, ohne originelle Ansätze einzubringen.
- **Halluzinationen:** Immer wieder halluziniert KI oder bringt weltfremde Inhalte in die Texte. Beispielsweise schlägt die KI vor zur Kindersicherung am Gartenteich mit Mosaik oder Bodenmalerei sichere Bereiche zu markieren oder mit einem Spielgerät die Kinder vom Wasser abzulenken. Ohne weitere Maßnahmen sind diese Empfehlungen keinesfalls kindersicher.

BEWERTUNG DER TEXTE DURCH EXPERTEN

Um das „Bauchgefühl“ bezüglich der Mängel von fachlichen KI-Texten bei E-E-A-T, das aus der Online-Marketing-Expertise der eo KI-Task-Force resultierte, zu überprüfen, wurden die KI-Texte aus dem YMYL-Bereich Experten und Expertinnen vorgelegt.

Experte: Rechtsanwalt Florian Kaiser	
Bewertung & Feedback	Grundsätzlich sind die rechtlichen Texte von der KI ok für den Online-Marketing-Zweck. Manche Formulierungen klingen komisch. Es fehlen teilweise Infos und ebenso sind Fehler enthalten.
Beispiele	– Fehlt: Zitate der erwähnten Urteile, damit man sie nochmals nachlesen kann, z.B. ArbG Würzburg, Urteil vom xx.xx.xx., Az xy. – Falsch: Der Kündigungsschutz gilt vom ersten Tag der Schwangerschaft bis 4 Monate nach der Entbindung – und nicht bis zur Entbindung wie im ersten Satz von der KI beschrieben. Erst weiter unten im Text wird es dann richtig gestellt, was aber für Verwirrung sorgen könnte.
Experte: Vermögens- und Unternehmerberater Max Hausmann	
Bewertung & Feedback	Die KI-Texte sind nicht per se schlecht. Sie bieten einen Grobüberblick, mit dem man sich als ersten Schritt informieren kann. ChatGPT sagt viel, wodurch der Eindruck entsteht, man wäre umfassend informiert. Jedoch ist vieles zu unspezifisch. Die KI pauschalisiert und erzeugt dadurch Missverständnisse.
Beispiele	– Bsp. „Die Familienhaftpflicht ist somit eine Erweiterung der normalen Haftpflichtversicherung und bietet speziell auf Familien zugeschnittenen Schutz.“ Experte: unspezifisch, was ist “normal”, es gibt Single-, Paar- und Familienversicherung – Bsp. „Liegt das Einkommen des privatversicherten Elternteils über der Versicherungspflichtgrenze, müssen die Kinder privat versichert werden.“ Experte: unspezifisch, hier verwechseln Laien oft die Beitragsbemessungsgrenze und die Jahresentgeltgrenze, das müsste hier eineindeutig erklärt werden – Bsp. „Beispielrechnung für eine private Krankenversicherung: Ehepartner (30 Jahre): 350 Euro/Monat, Kind 1 (3 Jahre): 150 Euro/Monat, Kind 2 (5 Jahre): 150 Euro/Monat = Gesamtkosten pro Monat: 650 Euro. Diese Beispielrechnung zeigt, dass die private Krankenversicherung für Familien erheblich teurer sein kann als die gesetzliche Familienversicherung.“ Experte: für die GKV macht die KI keine Beispielrechnung; zudem: nicht repräsentativ, teurer muss nicht „teurer“ bedeuten, weil man die Leistungen vergleichen muss.

Experte: Leo Störche, Dienstbeleghebamme am Leopoldina Krankenhaus	
Bewertung & Feedback	Fachlich sind die Texte grundsätzlich korrekt, Praxisbezug und Anwendungstipps sind vorhanden. Jedoch sind die Aussagen teilweise doppelt beschrieben, teilweise zu kurz geraten, in dem Sinne, dass nicht alle (emotionalen) Aspekte umfassend abgedeckt sind. Manche Infos waren per Definition nicht korrekt.
Beispiele	– Bsp: Vorteile des Wochenbetts: Bindung: ‚Die intensive Zeit, die Du mit Deinem Baby verbringst, ist entscheidend für die emotionale Entwicklung. Das ist grundsätzlich korrekt. Aus Expertensicht kommt das jedoch emotional zu kurz. Es gibt viele Fälle, in denen diese anfängliche Bindung gestört sein kann, z.B. Frühchen auf der Intensivstation, Mutter nach der Geburt auf der Intensivstation, Wochenbettdepressionen, Mütter/Väter, die keine bzw. zu kurz Elternzeit nach der Geburt haben, etc. Wenn ich das als Mutter lese und ein Frühchen bekomme, denke ich mir: F**k, jetzt ist das mit der Bindung alles dahin! – Bsp: Babyschlaf im 1. Monat: ‚Neugeborene schlafen durchschnittlich 16-18 Stunden pro Tag, allerdings in kurzen Intervallen von 2-3 Stunden.‘ Das ist zu allgemein und einfach formuliert. Es gibt Kinder, die nur 10 Minuten schlafen und danach fit sind, etc. So ist es zu arg verallgemeinert. – Bsp: ‚...wie lange Du im Wochenbett liegen solltest...‘. Wochenbett bedeutet nicht grundsätzlich „liegen“, sondern einfach „Schonung“. Außerdem spricht man im sog. „Frühwochenbett“ nicht von „zwei Wochen“, sondern von 10 Tagen. Aber auch hier bedeutet es eben nicht, dass man die ganze Zeit liegen muss. – Bsp: Geburtsverletzungen: ‚Diese Verletzungen benötigen spezielle Pflege und Geduld‘. Geduld ja, aber Pflege braucht ein Dammriss grundsätzlich erstmal nicht. Klar kann man z.B. Sitzbäder, etc. machen, nötig ist es aber nicht.“

Experte: Gynäkologe, Prof. Dr. med. Michael Weigel, Chefarzt Gynäkologie Leopoldina Krankenhaus Schweinfurt	
Bewertung & Feedback	Die Texte sind erstaunlich gut für eine KI. Allerdings wirkt er nicht wie von einem Experten, da Fachbegriffe falsch verwendet sind. Was ich da nicht ganz verstehe, ist diese erhebliche Redundanz. Viele Textabschnitte, die sich mit ihren Informationen immer wieder wiederholen und somit keine tiefergehenden Informationen liefern, sondern oberflächlich bleiben. Die emotionale Intelligenz fehlt. Schwangerschaft und Geburt sind hoch emotionale Themen, bei denen gerade erstmalige Eltern viele Fragen haben und sich durch Fachtexte rückversichern wollen. Die KI schafft durch viele Formulierungen keine Sicherheit, sondern schürt teilweise Ängste und Unsicherheiten. Außerdem sind fachlich falsche Informationen enthalten. Die Texte wirken nicht so, als wären sie von Fachexperten erstellt worden. Daher kann auch nicht das bei solchen Themen benötigte Vertrauen nicht aufkommen. Es wirkt eher so, „Zitat: als hätte die Pressestelle eines Krankenhauses die Texte veröffentlicht“. Auch der Aufbau der Texte ist teils nicht logisch. Es wird viel thematisch hin und her gesprungen und vieles kommt eben doppelt und dreifach vor, sodass es redundant wird. Gerade bei solchen sensiblen Themen, die sehr fachlich sind, einfach und verständnisvoll erklärt werden sollten, ist es wichtig, dass der Nutzer sich abgeholt fühlt. Aber die Texte sind aufgrund ihrer falschen und fehlenden Informationen sowie der fehlenden emotionalen Intelligenz wenig vertrauenswürdig, geben den Nutzern nicht das Gefühl von Wissen, Sicherheit und Expertise.

Experte: Gynäkologe, Prof. Dr. med. Michael Weigel, Chefarzt Gynäkologie Leopoldina Krankenhaus Schweinfurt	
Beispiele	– Bsp: „Der Epiduralraum, der Bereich um das Rückenmark...“ Arzt: Sobald Nutzer „Rückenmark“ lesen, haben sie Angst, dass da irgendwas verletzt wird. Das würde man den Leuten daher nicht so sagen. Dort, wo man die Anästhesie für die PDA setzt, da ist streng gar kein Rückenmark mehr. Sollte also besser ausformuliert und erklärt werden. – Bsp: “Sobald die PDA wirkt, sind die Beine oft taub und eine aufrechte Position ist nicht mehr möglich. Das bedeutet, dass Du während der Wirkung der PDA normalerweise nicht aufstehen oder herumlaufen kannst.“ Arzt: Ist falsch, mit dem heutigen Stand der PDA kann man meistens laufen. – Bsp: „Für geplante oder ungeplante Kaiserschnitte ist die PDA eine gängige Anästhesieform.“ Arzt: falsch, typisch ist da die spinale Anästhesie und bei geplanten Kaiserschnitten macht man keine PDA. – Bsp: “Einige Frauen verspüren während der fruchtbaren Tage einen erhöhten Sexualtrieb.“ Arzt: Hierfür muss ich erstmal den Grenzwert kennen, um zu definieren, das etwas erhöht ist. Ist unglücklich, unsensibel und unschön formuliert, man spricht ja nicht vom Trieb, wenn man über Frauen redet. Besser: sexuelle Lust. Einerseits falsche oder fehlende Informationen und andererseits kommt wieder die fehlende emotionale Intelligenz der KI zum Vorschein. – Bsp: „Temperaturmethode: Ein Temperaturanstieg von etwa 0,2 bis 0,5 Grad Celsius signalisiert, dass der Eisprung stattgefunden hat. Diese Methode erfordert regelmäßiges und genaues Messen, um verlässliche Ergebnisse zu liefern.“ Arzt: Ungenau beschrieben. Du misst zwar täglich Deine Temperatur, aber das musst Du morgens vor dem Aufstehen machen. , sobald Du mal aktiver geworden bist, sind diese 0,2 Grad locker erreicht. – Bsp: „Wenn die Geburt sehr schnell voranschreitet und die Mutter bereits in der Übergangsphase ist, bleibt keine Zeit mehr für die Anlage einer PDA.“ Arzt: Welche Übergangsphase soll die Mutter sein? Es gibt keine Übergangsphase, das ist falsch und Unsinn. Wenn die Geburt zu schnell voranschreitet, dann ist das kein Zeitmangel.

Expertin: Kinderärztin, anonym	
Bewertung & Feedback	Sprachlich sind die Texte einfach geschrieben, verständlich als schnelle Info für Familien. Das Wichtigste ist so weit gut erfasst. Aber es ist nicht alles richtig und manches ist Unnötigerweise doppelt beschrieben.
Beispiele	– Bsp: "Notiere die Dauer des Krampfes. Wenn er länger als fünf Minuten dauert, rufe den Notarzt." Ärztin: nicht warten! Sofort Notarzt. – Bsp: „Ein Fieberthermometer ist ein unverzichtbares Hilfsmittel, um die Körpertemperatur bei Babys genau zu messen. Die rektale Messung gilt als die genaueste Methode.“ Ärztin: unspezifisch, hier sollte man beschreiben, in welchem Alter man welche Methode anwendet. – Bsp: „Scharlach Symptome: Kopfschmerzen und Übelkeit: Häufige Begleitscheinungen des Fiebers.“ Ärztin: Bauchschmerzen und Erbrechen sind auch häufige Symptome. Sie fehlen. – Bsp: „Himbeerzunge: Eine auffällige, rote Zunge mit einer erdbeerartigen Oberfläche, die nach einigen Tagen weißlich belegt sein kann.“ Ärztin: Falsch. Erst ist sie weißlich belegt, dann sieht sie aus wie eine Himbeere. – Bsp: "Rötung im Rachen: Der Rachen kann stark gerötet und entzündet sein." Ärztin: Falsch, ist immer hochrot. – Bsp: „Besonders in den ersten 24 Stunden nach Beginn der Antibiotikatherapie ist die Ansteckungsgefahr hoch.“ Ärztin: Richtig ist: 24 – 48h nach Beginn mit dem Antibiotikum ist man nicht mehr ansteckend.

Expertin: Lehrerin, Alessia Merletti, Lehramtsassessorin	
Bewertung & Feedback	Auf den ersten Blick erscheint der Artikel wie von einem Experten geschrieben, bei genauerem Lesen fallen einer Fachkraft jedoch einige Ungenauigkeiten auf. Die Beispiele sind häufig nicht gut gewählt, da sie nicht repräsentativ, nicht korrekt und wenig anschaulich für den Leser sind. Auf einer Skala von 1 – 10 würde ich den Text hinsichtlich der Vertrauenswürdigkeit bei 7 einordnen, da jemand, der sich mit LRS noch nie auseinandergesetzt hat, einen guten ersten Einblick in die Materie bekommt. Jedoch wäre eine Überarbeitung durch eine Fachkraft wünschenswert, um den Praxisbezug und die Anwendungstipps durchweg korrekt und authentisch darzulegen.
Beispiele	– Bsp: Praxisbezug & Anwendungstipps: Ab dem zweiten Teil „Was kann ich gegen LRS tun?“ sind die genannten Maßnahmen redundant. An dieser Stelle würde ich mir konkrete Maßnahmen zur individuellen Unterstützung wünschen, die so tatsächlich im Unterrichtsalldag umgesetzt werden. Bei Lesestörung beispielsweise geeignete Schriftarten bei Arbeitsblättern, größere Zeilenabstände, farbige Kennzeichnung von Silben, differenzierte Hausaufgaben. Bei Rechtschreibstörung Nutzung spezieller Stifte, Nutzung eines Computers, differenzierte Hausaufgaben. – Bsp: Was Experten noch ergänzen würden: Häufig gehört auch eine unleserliche Handschrift zu den Anzeichen für LRS. – Bsp: LRS wird häufig erst spät erkannt, da stellenweise keine elterliche Zustimmung für die Überprüfung vorliegt, die Grundschule versäumt den entsprechenden Test durchzuführen, LRS in der Muttersprache diagnostiziert werden muss und es einigen Schülern gelingt, in der Grundschulzeit Fehler durch andere Strategien clever auszugleichen – Bsp: "Lernbegleitung: Einsatz von Schulbegleitern oder Nachhilfelehrern zur individuellen Förderung." Lehrerin: LRS allein ist kein Grund für einen Schulbegleiter, diese werden nur in wenigen Fällen genehmigt, wenn weitere Problemfelder hinzukommen. – Bsp: "Im Gegensatz zu einer allgemeinen Leseschwäche oder Schreibschwäche hat Legasthenie nichts mit mangelnder Intelligenz oder unzureichender schulischer Förderung im Unterricht zu tun." Lehrerin: Würde ein Pädagoge so nicht ausdrücken. – Bsp: „Einfügen von Buchstaben: Zusätzliche Buchstaben werden eingefügt, wie „Füßball“ statt „Fußball“. Lehrerin: Ich würde „ü“ statt „u“ nicht als zusätzlichen Buchstaben beschreiben. – Bsp: „Die Schwierigkeiten beim Lesen und Schreiben stehen in keinem Zusammenhang mit ihrem allgemeinen kognitiven Potential.“ Lehrerin: Legasthenie ist eine Teilleistungsstörung. Den Begriff Teilleistungsstörung fände ich an dieser Stelle wichtig. – Bsp: „Die Diagnose von Legasthenie erfolgt meist durch spezialisierte Psychologen oder pädagogische Fachkräfte.“ Lehrerin: Schulpsychologen oder Kinder- und Jugendpsychiater diagnostizieren LRS, nicht jede pädagogische Fachkraft hat die Befähigung dazu.

Experte: Erzieher, Frank Salwiczek-Merletti, staatlich anerkannter Erzieher	
Bewertung & Feedback	Der KI-Text zeigt ein modellbasiertes Expertenwissen und verweist regelmäßig auf die Individualität jedes Kindes in der Eingewöhnungsphase. Ebenso ist es aber wichtig die unterschiedlichen Gegebenheiten jeder Einrichtung zu berücksichtigen. Sei es das Konzept, die Personalsituation oder die räumlichen Gegebenheiten der Kita. Für viele Einrichtung zeigt der KI-Text einen guten Praxisbezug und Anwendungsprozess der Eingewöhnung. Die Handhabung der Eingewöhnung in einer Einrichtung, hängt häufig von jahrelangen Erfahrungen des Einrichtungsteams und auch stellenweise von den Vorgaben der Leitung oder des Einrichtungsträger ab. Der Text ist gut und ich würde auf einer Skala 8 - 9 geben."
Beispiele	– Bsp: Fehlende Fachtermini. Zu verallgemeinerte Fachkraftbezeichnungen. Es gibt im Arbeitsbereich einer Kita viele unterschiedliche Fachkräfte, von Erzieher/innen über Kinderpfleger/innen, Heilerziehungspfleger/innen und noch weiteres Personal. – Bsp: "Eine gute Eingewöhnung zeichnet sich durch Flexibilität, Einfühlungsvermögen und eine enge Zusammenarbeit zwischen Eltern und Erziehern aus." Erzieher: Fachpersonal; Nicht nur Erzieher/innen übernehmen Eingewöhnungen. – Bsp: "Ablauf einer Eingewöhnung." Erzieher: zu allgemein geschrieben und teilweise falsch, da jedes Kind ein anderes Tempo hat. Bspw. erkunden viele Kinder weitere Räumlichkeiten erst nach längerer Zeit. Kontaktaufnahme zu einer Bezugsperson sollte vor der Erkundung der Räumlichkeiten geschehen. – Bsp: "Bezugserzieherin und allgemein Erzieher." Erzieher: KI schreibt immer nur von „Erziehern“, die Aufgaben kann auch anderes Fachpersonal übernehmen. Falsch beschrieben: Kinder suchen sich selbst eine Bezugsperson aus, nicht umgekehrt. Bei offenen Konzepten ist das sogar gruppenübergreifend möglich. – Bsp: "Erstkontakt zur Erzieherin." Erzieher: Erster Kontakt ab Tag 1 und nicht erst ab Tag 3. – Bsp: "Verabschiedung." Erzieher: Eltern und Kind sollten sich gemeinsam verabschieden, damit das Kind weiß: Wenn Mama oder Papa kommen, werden ich abgeholt. Ein kurzer Austausch über den ersten Tag kann beim Anziehen stattfinden.

Tabelle 3: Bewertung von E-E-A-T von Experten aus den Bereichen Recht, Finanzen, Medizin und Erziehung

Es gibt einige Gemeinsamkeiten im Feedback der Experten:

- **Ordentlicher Basistext:** Alle Experten finden die KI-Texte grundsätzlich in Ordnung
- **Mangel an Details und Genauigkeit:** die Texte sind oberflächlich und oft einseitig in der Beschreibung von Sachverhalten; beispielsweise inkorrekte Verwendung von Fachtermini
- **Falsche Zahlen, Daten, Fakten:** In allen Fachbereichen finden sich falsche Informationen; Falschinformationen im Finanz- und Medizinbereich können dramatische Auswirkungen haben auf die Vermögens- und Gesundheitssituation, wenn Entscheidungen auf den Informationen fußen oder Empfehlungen der KI Folge geleistet werden. Beispiel: Fieberkrampf im Text zu „Scharlach“: Hier empfiehlt die KI die Zeit zu stoppen, wie lange ein Fieberkrampf beim Kind anhält. Wenn er länger als 5 Minuten andauert, soll man den Notarzt rufen. Laut der Kinderärztin soll man direkt den Notarzt anrufen.
- **Mangel an Empathie:** Texte wirken oft kühl und unpersönlich.
- **Mangelnde emotionale Intelligenz:** Wichtige Nuancen bei sensiblen Themen wie Schwangerschaft oder Gesundheit fehlten. Die Aussagen der KI können bestehende Unsicherheiten und Ängste sogar fördern.

Den Eindruck aus dem E-E-A-T-Check der eo KI-Task-Force bestätigen die Fachexperten. GPT-4o und Claude 3.5 Sonnet generieren **keine schlechten Texte**, aber **Expertenstatus bekommen sie nicht hin**. Eine der Expertinnen wünscht sich ganz nachdrücklich, dass die KI-Texte immer auf fachliche Richtigkeit und emotionale Aspekte geprüft werden, um Falschinformationen zu verhindern und keine Ängste zu schüren.

E-E-A-T in YMYL-Texten von Chatbots

1. Die Rolle von Fachwissen und menschlicher Expertise

Gerade bei komplexen Themen wie rechtlichen oder medizinischen Inhalten reicht die KI nicht aus. Menschliche Expertise ist notwendig, um sicherzustellen, dass Inhalte korrekt, verständlich und vollständig sind.

2. Probleme bei sensiblen Themen

Fehlende emotionale Intelligenz kann Vertrauen zerstören. Besonders bei emotionalen oder persönlichen Themen wie Schwangerschaft oder Familienplanung ist dies problematisch.

BEWERTUNG DURCH GOOGLE – SEO-PERFORMANCE

Ob KI-Texte ranken oder nicht, entscheidet am Ende aber immer noch Google. Daher war der 3. Schritt der E-E-A-T-Prüfung ein Realitycheck, **wie sich die SEO-Performance der Test-Domain entwickelt.**

Die Test-Domain die-familie.org ist ein informationales Blog-Portal für Familien. Die **wichtigsten Themen für Familien** sind abgebildet in Form der Hauptkategorien, darunter liegen mindestens zwei Unterverzeichnisse. Jedes Unterverzeichnis ist mit Texten aus einer Quelle befüllt: ausgebildeter SEO Redakteur, ChatGPT-4o oder Claude 3.5 Sonnet. Die Unterverzeichnis-Struktur soll es ermöglichen, **Aussagen über die Performance** der Texte verschiedener Quelle zu treffen.

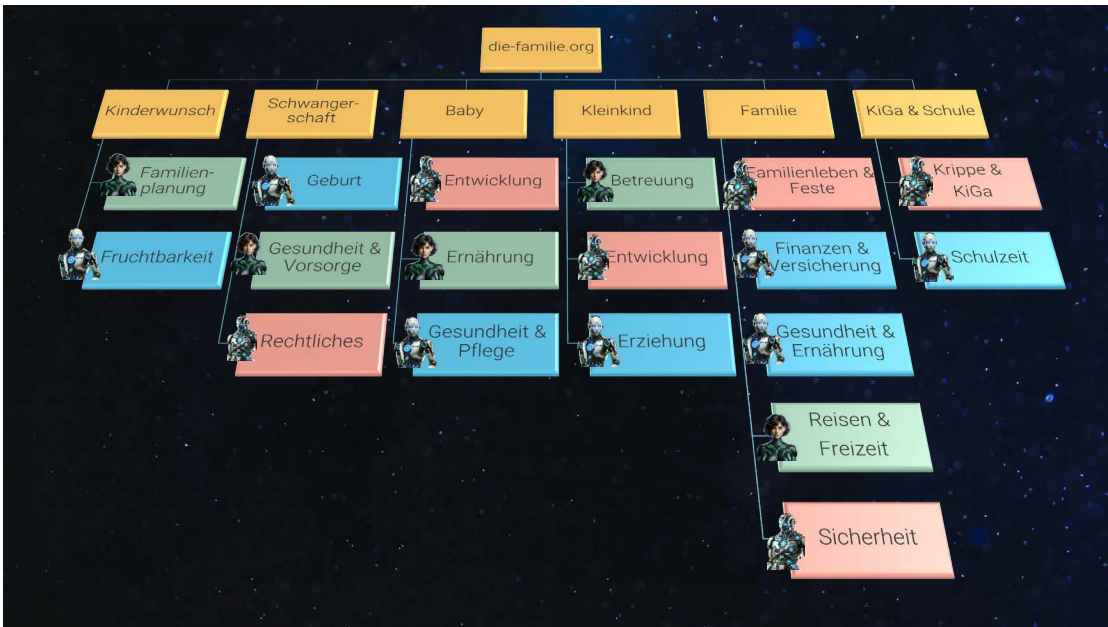


Abbildung 6: Aufbau der Seite die-familie.org

Als Hypothese für die Entwicklung der Testdomain wurde definiert:

Die Unterverzeichnisse mit den menschlichen Texten haben eine höhere Sichtbarkeit sowie mehr und bessere Rankings in den Google SERPs. Google erkennt, dass menschliche Inhalte hochwertiger sind und E-E-A-T besser erfüllen.

Es haben sich sehr schnell nach dem Livegang am 02.09.2024 erste Keywordrankings für alle 3 Quellen ergeben. Jedoch ist bis heute **das typische Google-Dancing** zu beobachten, die Sichtbarkeit der Domain geht rauf und runter. Die Suchmaschine testet die Inhalte am Anfang, so dass es oft sehr lange dauert, bis sich stabile Rankings einstellen. Aktuell (29.01.2025) haben **3 Unterverzeichnisse einen eigenen Sichtbarkeitsverlauf**, 2 davon gehören zur menschlichen Quelle, eines zu Claude 3.5 Sonnet.

Die Tabelle stellt die Sichtbarkeitsentwicklung der Unterverzeichnisse dar (Stand: 29.01.2025).

Unterverzeichnis	Quelle	Sichtbarkeitsindex Wert
https://www.die-familie.org/kinderwunsch/fruchtbarkeit/	ChatGPT 4o	0,000
https://www.die-familie.org/kinderwunsch/familienplanung/	Mensch	0,000
https://www.die-familie.org/schwangerschaft/gesundheit-und-vorsorge	Mensch	0,000
https://www.die-familie.org/schwangerschaft/rechtliches/	Claude	0,000
https://www.die-familie.org/schwangerschaft/geburt/	ChatGPT 4o	0,000
https://www.die-familie.org/baby/ernaehrung/	Mensch	0,001
https://www.die-familie.org/baby/entwicklung/	Claude	0,001
https://www.die-familie.org/baby/gesundheit-pflege/	ChatGPT 4o	0,000
https://www.die-familie.org/kleinkind/betreuung/	Mensch	0,002
https://www.die-familie.org/kleinkind/entwicklung/	Claude	0,000
https://www.die-familie.org/kleinkind/erziehung/	ChatGPT 4o	0,000
https://www.die-familie.org/familie/familienleben-feste/	Claude	0,000
https://www.die-familie.org/familie/finanzen/	ChatGPT 4o	0,000
https://www.die-familie.org/familie/sicherheit/	Claude	0,000
https://www.die-familie.org/familie/gesundheit-ernaehrung/	ChatGPT 4o	0,000
https://www.die-familie.org/familie/reisen-freizeit/	Mensch	0,000
https://www.die-familie.org/kindergarten-und-schule/krippe-und-kindergarten/	Claude	0,000
https://www.die-familie.org/kindergarten-und-schule/schulzeit/	ChatGPT 4o	0,000

Tabelle 4: Sichtbarkeitsentwicklung der Unterverzeichnisse

Nun könntet ihr vielleicht direkt denken: Menschliche Texte performen besser. Betrachten wir allerdings die **Keywordrankings mit den besten Positionen**, dann dominieren Texte ganz deutlich von einer Quelle: **Claude hat die besten Keywordrankings** (Stand: 29.01.2025).

Tabelle 27 stellt die Top-Keywordrankings mit ihrer Quelle dar.

Keyword	Position	Veränderung	URL	Quelle
Lernentwicklungs- gespräch Ablauf	10	+ 1	https://www.die-familie.org/kindergarten- und-schule/schulzeit/lernentwicklungsge- sprach	ChatGPT 4o
Vorbereitung Entwicklungs- gespräch Kita	11	Neu	https://www.die-familie.org/kindergarten- und-schule/krippe-und-kindergarten/ent- wicklungsgespraech-kita	Claude
Kindergarten Entwicklungs- gespräch	13	Neu	https://www.die-familie.org/kindergarten- und-schule/krippe-und-kindergarten/ent- wicklungsgespraech-kita	Claude
Baby und Familie Entwicklungs- kalender zum Ausdrucken	13	Neu	https://www.die-familie.org/baby/entwick- lung/entwicklungskalender-babys	Claude
Kinderentwicklung mit 3 Jahren	14	Neu	https://www.die-familie.org/kleinkind/ent- wicklung/kinderentwicklung-mit-3-jahren	Claude
Teich Kindersiche- rung Abdeckung	20	Neu	https://www.die-familie.org/familie/si- cherheit/teich-pool-kindersicher	Claude
Ab wann krabbeln Babys frühestens	20	Neu	https://www.die-familie.org/baby/entwick- lung/ab-wann-krabbeln-babys	Claude
Kinderentwicklung mit 2 Jahren	23	Neu	https://www.die-familie.org/kleinkind/ent- wicklung/kinderentwicklung-mit-2-jahren	Claude
Poolumrandung kindersicher	23	Neu	https://www.die-familie.org/familie/si- cherheit/teich-pool-kindersicher	Claude
Bester Schulranzen für Erstklässler	26	Neu	https://www.die-familie.org//kindergarten- und-schule/schulzeit/bester-schulranzen	ChatGPT 4o
Wann sitzen Babys im Durchschnitt	28	+ 7	www.die-familie.org/baby/entwicklung/ ab-wann-sitzen-babys	Claude
Wie warm darf Babymilch sein	30	Neu	https://www.die-familie.org/baby/ernaeh- rung/babymilch-zubereiten	Mensch

Tabelle 5: Top-Keywordrankings

Wenn man sich die Keywordphrasen genau anschaut, fällt auf: 3 Keywordrankings hat der Text zur Kindersicherung an Pool und Gartenteich, den die eo KI-Task-Force sehr stark in Bezug auf weltfremde Halluzinationen bemängelt hat (siehe Kapitel 4.3). Google sieht hier also keine Qualitätsmängel, im Gegensatz zu menschlichen „Qualitätsprüfern“. Es liegt also **in unserer Verantwortung**, die, die wir Webseiteninhalte online stellen, die **KI-Texte genau zu prüfen** und Qualitätsmängel auszugleichen.

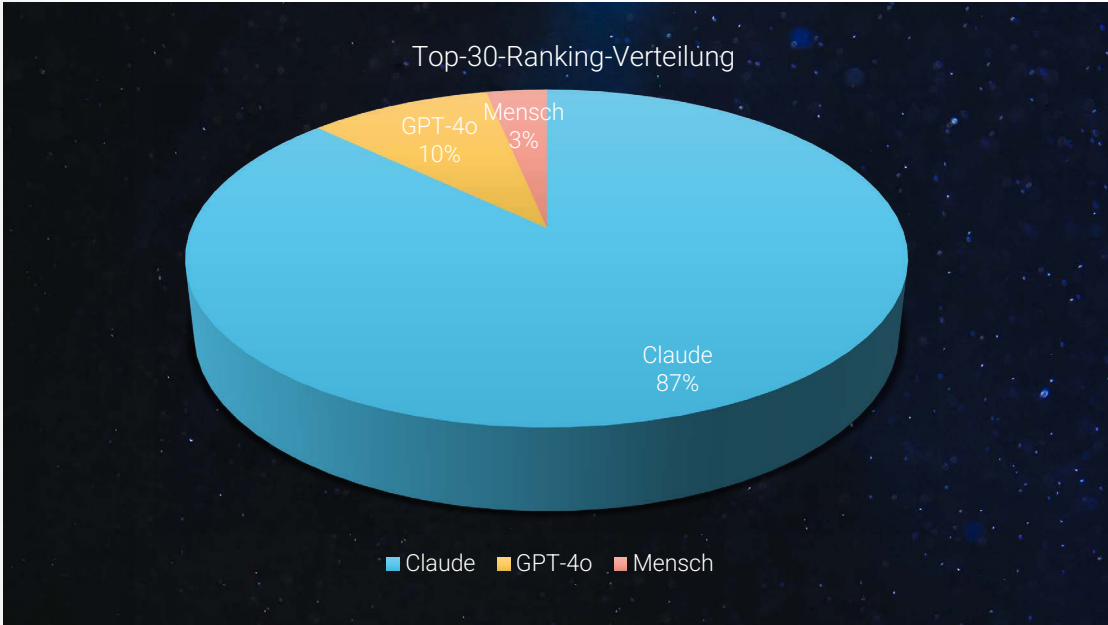


Abbildung 7: Top-30-Ranking-Verteilung

Im Kuchendiagramm sind die aktuellen (Stand: 29.01.2025) **Top-30-Keywordrankings** in der Verteilung nach der Quelle dargestellt. Bei den aktuellen Rankings, die die-familie.org in den Top-30 erzielt hat, **liegt Claude mit 87 % deutlich vorne**. Auf menschliche Texte entfallen gerade mal 3 % der Top-30-Rankings.

Es zeigt sich nach momentanem Stand: E-E-A-T scheint bei KI-Texten, vor allem bei Inhalten von Claude 3.5 Sonnet, auszureichen. **Google rankt KI-Inhalte nicht schlechter**, sondern eher besser als menschliche Texte. Platzierungen in den Top-10 konnten die Chatbots bislang jedoch nicht erreichen. Dies kann verschiedenste Ursachen haben.

Die Domain ist seit dem 02.09.2024 online. Rankings entwickeln sich langsam und langfristig, daher sind die Ergebnisse noch nicht absolut aussagekräftig. Über die Zeit erst wird sich zeigen, **ob sich die Hypothese doch noch bestätigt**. Das Netz verändert sich beständig und weitere Einflussfaktoren kommen auf die Rankings zu. Zudem testet Google oft Inhalte, lässt sie zunächst gut ranken und **strafft sie später ab**. Abstrafungen können zudem mit späteren Updates ebenfalls noch eintreffen.

Für Informationen zur weiteren Entwicklung der Domain schaut doch am besten immer mal wieder auf unserer Website oder auf dem LinkedIn-Profil unserer Dr. Bea vorbei:
<https://www.linkedin.com/in/dr-beatrice-eiring-digitalkommunikationsexpertin/>

Anmerkung SEO-Maßnahmen:



Die suchmaschinenoptimierten Texte von den Chatbots und dem eo Content Team wurden auf die Domain hochgeladen, zudem wurde die technische Fehlerfreiheit der Domain sichergestellt. Weitere, ggf. manipulative SEO-Maßnahmen, wurden nicht getätigt.

IN 5 SCHRITTEN ZU RANKINGFÄHIGEN TEXTEN

Was zweifelsfrei erwiesen ist: **Nach dem SEO-Qualitätsstandard erstellte KI-Texte sind rankingfähig** und können die Sichtbarkeit einer Domain genauso unterstützen wie menschliche Inhalte. Allerdings ist die Erstellung solcher Texte nicht damit getan, einfach einen Prompt einzugeben und das Ergebnis online zu stellen. **Es bedarf menschlicher Zuarbeit und Kontrolle.**

Die eo KI-Task-Force hat aus der Studie heraus Best Practices ermittelt, um zu rankingfähigen Texten zu kommen. **In 5 Schritten könnt ihr das auch:**

1. Schritt: Chatbot briefen: Bei ChatGPT einen CustomGPT anlegen mit allen notwendigen Projektinstruktionen. Auf diese greift die KI immer zurück, ihr müsst sie daher nicht immer wieder neu eingeben. Das spart Zeit und sorgt langfristig für besseren Output. Bei Claude geht das nicht. Hier kopiert ihr die Projektinstruktionen einfach in den Chat. Die KI meldet dann, dass sie die Aufgabe verstanden hat und fasst diese noch einmal zusammen.

Folgende Instruktionen sind wichtig:

- Keywordset
- Infos zur Verwendung der Keywords (Verteilung, Häufigkeit, Positionierung im Text)
- Textlänge
- Zielgruppe und Lesersprache
- Stil, Tonalität und weitere Infos zum Wording
- Formaler Textaufbau (Überschriften, Aufzählungen, Tabellen)
- Inhaltlicher Textaufbau, ggf. Gliederung

- 2. Schritt: Chatbot testen:** Prompt erstellen mit allen für den einzelnen Text relevanten Informationen (Thema, Keywordset). Da ihr alle Projektinstruktionen bereits an den Chatbot übermittelt habt, braucht ihr in den Prompt nur noch die Infos für den einzelnen Text reinschreiben. Er lautet dann: „Hallo GPT, erstellst Du mir bitte einen Text zum Thema „XY“ mit diesen Keywords: Hauptkeyword=Thema: [Hauptkeyword] Nebenkeywords: [Nebenkeywords] W-Fragen: [W-Fragen]. Danke.“
- 3. Schritt: Output prüfen:** Output prüfen auf alle Projektanforderungen: SEO, Keywords, Wording, Rechtschreibung sowie Inhalt
- 4. Schritt: Anpassungen vornehmen:** Wenn der Output nicht passt, gebt der KI Anweisungen für Änderungen, passt die Projektinstruktionen an oder optimiert den Prompt. Ladet Beispieltex te hoch mit eurem gewünschten Stil, lasst die KI den Stil beschreiben und packt die Stilbeschreibung in die Projektinstruktionen. Wenn der Chatbot wiederholt zu wenig Text liefert, kann es Sinn machen, dass ihr diese Information doch in den Prompt eingeben müsst, selbst wenn diese in den Instruktionen steht. Insbesondere ChatGPT neigt dazu, zu wenig Text zu liefern und muss immer wieder daran erinnert werden.
- 5. Schritt: Bilder generieren:** Instruktionen für die Bilderstellung in den Projektinstruktionen hinterlegen. Bei einem CustomGPT hinterlegt ihr auch den Stil für die Bilder, denn dann habt ihr eine Chance auf einen einheitlichen Stil. Allerdings ist ChatGPT mit der Schnittstelle zu DALL-E nicht die beste Bild-KI. Daher funktioniert nicht alles reibungslos. Für eine hochwertige und professionelle Bildgenerierung nutzt besser eine spezielle Bild-KI wie Midjourney.

REALITÄTSHECK: WIE VIEL ZEIT SPART KI WIRKLICH?

SEO-Texterstellung mit KI-Support spart Zeit. **Es sind aber nicht 100 % Zeitersparnis.** Um hochwertige, rankingfähige Texte zu bekommen, muss der Mensch mitarbeiten. Im Sparring – **einer den Fähigkeiten entsprechenden Aufgabenverteilung** – kommt ihr zu guten Ergebnissen.

Die folgende Tabelle stellt die von der eo KI-Task-Force empfohlene Aufgabenverteilung dar.

Menschlicher Anteil	KI-Anteil
Briefing erstellen	Inhalte recherchieren
Keywords recherchieren	Inhalte recherchieren
Prompting Texterstellung	Bilderstellung
Lektorat: Rechtschreibung, SEO- & Inhaltsprüfung	
Nachbessern: fehlende Keywords & inhaltliche Fehler	
Prompting Bilderstellung	
Interne Links setzen	
Meta-Daten erstellen	
Text online stellen	

Tabelle 6: Aufgabenverteilung von Mensch und KI

Vielleicht erstaunt es, dass in der Liste des menschlichen Anteils auch die Keywordrecherche steht. Diese lässt sich von KI ebenfalls abbilden oder mit Schnittstellen automatisieren. Für uns ist die **Keywordrecherche ein enorm wichtiger Teil** der Texterstellung. Aus der Keywordrecherche ergibt sich die Gliederung für den Text, dieser soll möglichst holistisch sein und alle Nutzerinteressen abbilden. Daher überlässt die eo KI-Task-Force der KI nicht die Auswahl der Keywords, sondern recherchiert diese **manuell mit Tools** wie dem Google Ads Keyword Planner, answerthepublic, quaro und termlabs und strickt daraus eine sinnvolle Gliederung für den KI-Text.

Auf der Basis dieser Aufgabenverteilung hat die eo KI-Task-Force die Zeitersparnis beispielhaft für 3 Seiten- bzw. Content-Typen ermittelt, die bei der Testdomain erstellt wurden. Die Ergebnisse sind in der Tabelle dargestellt.

Seitentyp	Menschlicher Anteil	KI Anteil	Zeitersparnis
Blogbeitrag 1.000 Wörter	• Briefing: 5 min. • Keywordrecherche: 10 min. • Prompting Claude 3.5 Sonnet Text: 2 min • Keyword- & Inhaltsprüfung: 25 min • Einfügen fehlender Keywords: 15 min • Prompting GPT-4o Bildgenerierung: 1 min • Interne Links setzen: 15 min • Meta-Daten erstellen: 15 min Gesamt: 88 min.	• Generierung der Texte durch Claude 3.5 Sonnet: 5 min • Generierung und Anpassung der Bilder durch GPT-4o: 2 min Gesamt: 7 min.	Ohne KI: 200 min Mit KI: 95 min Zeitersparnis: 53 %
Unterkategorie-seite Einleitungstext 300 Wörter Teasertexte für Blogbeiträge 4 Stück à 100 Wörter	• Briefing: 5 min. • Keywordrecherche: 15 min. • Prompting gesamt: 7 min. • Keyword- & Inhaltsprüfung: 10 min • Einfügen fehlender Keywords: 3 min • Interne Links setzen: 15 min • Meta-Daten erstellen: 15 min Gesamt: 70 min.	• Generierung der Texte durch GPT-4o: 2 min • Generierung und Anpassung der Bilder durch GPT-4o: 9 min Gesamt: 11 min.	Ohne KI: 150 min Mit KI: 81 min Zeitersparnis: 46 %

Seitentyp	Menschlicher Anteil	KI Anteil	Zeitersparnis
Hauptkategorie-seite Einleitungstext 300 Wörter Teasertexte für Unterkategorien 4 Stück à 100 Wörter	• 13 Briefing: 5 min. • Keywordrecherche: 15 min. • Prompting gesamt: 12 min. (v.a. für die Anpassung der Bilder) • Keyword- & Inhaltsprüfung: 14 min • Einfügen fehlender Keywords: 5 min • Interne Links setzen: 15 min • Meta-Daten erstellen: 15 min Gesamt: 81 min.	• Generierung der Texte: 2 min • Generierung und Anpassung der Bilder: 8 min Gesamt: 10 min.	Ohne KI: 125 min Mit KI: 91 min Zeitersparnis: 37 %

Tabelle 7: Zeitersparnis beim Einsatz von KI

1. Beispiel Blogbeitrag: Der **große zeitliche Vorteil von KI** besteht in der **Inhaltsrecherche und der Texterstellung**. Hier brauchen Menschen einfach deutlich länger. Jedoch könnt ihr euch keineswegs auf **Fakten** verlassen, die die Chatbots liefern. Diese müsst ihr **immer kontrollieren**. Bei soften Themen wie „Kindergeburtstag feiern“, der Text, den ihr als Beispiel in der Tabelle findet, geht das recht schnell. Bei Fachthemen aus Jura, Finanzen und Medizin, die deutlich **faktenlastiger** sind, **dauert das Lektorat länger**. Dadurch relativiert sich wiederum die Zeitersparnis. Dies trifft analog zu, wenn ihr viele und ggf. komplexe Projektanforderungen habt.

2. Beispiel Unterkategorie-Seite: Bei der **Unterkategorie-Seite** ergibt sich bei den kurzen Texten eine **immer noch große Zeitersparnis**. Diese entsteht durch die Art der Texte – Teasertexte, die jeden Blogbeitrag zusammenfassen. Um diese Texte zu erstellen, müsste ein Mensch jeden Text noch einmal lesen. Die eo KI-Task-Force hat Word-Dokumente mit den Texten hochgeladen, die KI hat den Inhalt zuverlässig ausgelesen und die Teasertexte erstellt.

3. Beispiel Hauptkategorie-Seite: Die **Hauptkategorie hat die geringste Zeitersparnis**, da der Mensch hier selbst nicht lange für die Teasertexte gebraucht hätte, die die Unterverzeichnisse thematisch zusammenfassen.

An den 3 Beispielen zeigt sich, dass es nicht eine generelle **Zeitersparnis** gibt, sondern diese immer **von verschiedenen Faktoren abhängt**:

- Textlänge
- Thema
- Anforderungen
- Output der KI
- Tagesform der KI: ja, sie hat sie ebenfalls.

Eine völlig konsistente Qualität über einen längeren Zeitraum könnt ihr nicht erwarten. Je schlechter oder unpassender der Output ist, desto höher ist euer Zeitaufwand.

FAZIT & AUSBLICK

Die empirische Untersuchung von E-E-A-T bei KI-Texten durch die eology Online-Marketing-Experten, die Fachexperten und die SEO Performance der Domain die-familie.org verdeutlicht, dass **KI-Inhalte noch nicht alle Kriterien des SEO Qualitätsstandards gleichermaßen und konsistent erfüllen**. Während bei soften Themen E-E-A-T gut bis sehr gut ausfällt, zeigen sich bei YMYL-Themen Schwächen in der **Datenkorrektheit und der inhaltlichen Tiefe**.

Nichtsdestotrotz: KI revolutioniert die Texterstellung und bietet enorme Potenziale für Effizienzsteigerung. Es erfordert aber **immer menschliche Kontrolle und Nacharbeit**, um qualitative Standards zu erfüllen und um vor allem zu vermeiden, dass sich Falschinformationen verbreiten oder das Netz mit minderwertigen Inhalten geflutet wird. Langfristig wird der **Mensch die entscheidende Rolle** spielen, um sicherzustellen, dass KI-Inhalte einzigartig, nutzerorientiert und mehrwertstiftend bleiben. Der Einsatz von KI-Tools sollte daher immer als **hybride Lösung** betrachtet werden, bei der Technologie und menschliche Expertise Hand in Hand gehen. Die Zukunft liegt in der **Zusammenarbeit von Mensch und Maschine**. Ihr solltet KI als Werkzeug betrachten, das eure Arbeit ergänzt, aber nicht ersetzt. Mit einem durchdachten Ansatz könnt ihr das Beste aus beiden Welten nutzen.

Es grüßt herzlich

Die eo KI-Task-Force





UNSERE MISSION:

WIR ERSCHAFFEN SICHTBARKEIT IN EINER DIGITALEN WELT UND REALISIEREN VISIONEN DURCH SMARTE UND GANZHEITLICHE KONZEPTE.

SEARCH ENGINE OPTIMIZATION

Wir entwickeln eine durchdachte SEO-Strategie, beraten kompetent bei der Umsetzung, übernehmen das Monitoring der SEO-Erfolge und erstellen aussagekräftige Reports. So schaffen wir organische Reichweite.

BRAND & PERFORMANCE MARKETING

Mit Brand & Performance Marketing erreichen wir Marketingziele zielsicher. Wir wählen die richtigen Kanäle (Search, Social, Display Ads) aus und steuern die Anzeigen kreativ und datengesteuert.

CONTENT CREATION

Wir helfen dabei, die richtige Content-Strategie für Websites, Onlineshops oder Content-Marketing-Kampagnen zu definieren. Unsere erfahrenen Redakteure erstellen Inhalte genau nach den Kundenbedürfnissen.

CONTENT OUTREACH

Wir sind Spezialisten für die Verbreitung von Inhalten und nutzen Online-PR, Content-Seeding-Kampagnen und Linkmarketing, um sicherzustellen, dass die Website, Marke oder Kampagne unserer Kunden im Internet wahrgenommen wird.

SOCIAL MEDIA MARKETING

Wir entwickeln eine klare Social Media Strategie, erstellen Social Content und schalten Anzeigen auf den verschiedenen Plattformen. So nutzen wir die Synergien von organischen und bezahlten Inhalten perfekt und erreichen die Zielgruppe da, wo sie sich aufhält.

TRACKING & ANALYTICS

Daten sind die Grundlage für unsere Arbeit in allen Bereichen. Deshalb ist ein sauberer Aufbau extrem wichtig. Denn nur so können wir die Auswirkungen unserer gemeinsamen Arbeit richtig nachvollziehen. Wir unterstützen unsere Kunden kompetent.

DEINE ANSPRECHPARTNERIN BEI EOLOGY



DR. BEATRICE EIRING

Head of Content Creation
b.eiring@eology.de
 09381 / 58 290 - 15

AUTOREN

Dr. Beatrice Eiring
 Dr. Bettina Zehner
 Antonios Smyrniaios
 Elisa Greubel
 Larissa Köberlein
 Christoph Herbig-Fleißner
 Theresa Brieg

BILDQUELENNACHWEISE

eology GmbH
 ChatGPT
 Claude

KONTAKT

eology GmbH

Spitalstraße 23
 97332 Volkach
 09381 / 58290 00

info@eology.de
www.eology.de

AUSKUNFT ÜBER DIE GESELLSCHAFT

Registergericht: Amtsgericht Würzburg

HRB-Nr: 10610
 UST-Nr: 257/125/70116
 USD-ID: DE-270186142

GESCHÄFTSFÜHRER

Daniel Unger
 Axel Scheuering



DISCLAIMER: Die Inhalte dieses Whitepapers wurden mit größter Sorgfalt erstellt. Die eology GmbH übernimmt jedoch keine Gewähr für die Richtigkeit und Aktualität der bereitgestellten Informationen. Die Nutzung der Inhalte erfolgt auf eigene Gefahr des Nutzers. Alle Inhalte unterliegen dem deutschen Urheber- und Leistungsschutzrecht. Jede Verwendung bedarf der vorherigen schriftlichen Zustimmung durch die eology GmbH. Die unerlaubte Vervielfältigung oder Weitergabe einzelner oder kompletter Inhalte ist nicht gestattet. Änderungen, auch ohne vorherige Bekanntgabe, vorbehalten. Die Wiedergabe von Gebrauchsnamen, Handelsnamen, Warenbezeichnungen etc. in diesem Werk berechtigt auch ohne besondere Kennzeichnung nicht zur Annahme, dass solche Namen im Sinne der Warenzeichen- und Markenschutz-Gesetzgebung als frei zu betrachten wären und daher von jedermann benutzt werden dürfen.