

crawlen google

Category: Online-Marketing

geschrieben von Tobias Hager | 30. Januar 2026



Crawlen Google: So tickt der Suchmaschinen-Bot wirklich

Du hast Content geballert, SEO-Plugins installiert und jeden zweiten Blogartikel zum Thema „Keyword-Recherche“ gelesen – und trotzdem rankt deine Seite nicht? Dann solltest du dich mit dem Typen beschäftigen, der bei Google die Drecksarbeit macht: dem Googlebot. Denn was der nicht versteht, sieht auch niemand. Und was der nicht crawlt, existiert für Google schlichtweg nicht. Willkommen in der Welt des Crawlings – brutal technisch, gnadenlos ehrlich und absolut entscheidend für deinen Erfolg.

- Wie der Googlebot wirklich arbeitet – und warum Crawling kein Mythos ist
- Unterschied zwischen Crawling, Indexierung und Rendering (Spoiler: viele verwechseln das immer noch)
- Welche Parameter der Googlebot beim Crawlen deiner Website berücksichtigt

- Warum Crawl-Budget kein Marketing-Gag ist, sondern harte Realität für große Sites
- Wie du Crawling-Probleme mit Tools wie Logfile-Analyse, Screaming Frog und GSC erkennst
- Was Google priorisiert – und wie du deine Seite für den Crawl optimierst
- Die größten Crawling-Killer: JavaScript, Endlosschleifen, kaputte Links und Wildwuchs
- Warum technische Hygiene entscheidend ist, um überhaupt gecrawlt zu werden
- Eine Schritt-für-Schritt-Anleitung für eine crawlingfreundliche Seitenarchitektur
- Wie du den Googlebot lieben lernst – oder ihn zumindest nicht verärgerst

Was ist Crawling? Der technische Kern von Googles Suchmaschinerie

Crawling ist der erste Schritt im SEO-Spiel – und gleichzeitig der am wenigsten verstandene. Viele werfen „crawlen“, „indexieren“ und „rendern“ gedankenlos in einen Topf. Der Unterschied ist aber gewaltig: Beim Crawling durchsucht Google deine Website nach neuen oder aktualisierten Inhalten. Beim Indexieren werden diese Inhalte in den Google-Index aufgenommen. Und beim Rendern entscheidet sich, ob Google überhaupt versteht, was du da gebaut hast – insbesondere bei JavaScript-lastigen Seiten.

Der Googlebot ist der Crawler von Google. Er arbeitet wie ein hochautomatisierter Webscraper, der sich durch deine Seitenstruktur hangelt, Links folgt, HTTP-Header auswertet, robots.txt-Dateien respektiert und anhand von internen Signalen entscheidet, wie tief er in deine Seite eintaucht. Dabei kommt ein komplexes Zusammenspiel aus Infrastruktur, Priorisierung und Budgetierung zum Einsatz, das den meisten Website-Betreibern völlig entgeht.

Crawling ist keine Wohltätigkeit. Google hat begrenzte Ressourcen und crawlt selektiv. Wenn deine Seite technisch schlecht aufgestellt ist, schlechte Signale sendet oder unnötig viele URLs produziert, wirst du schlichtweg weniger gecrawlt. Und was nicht gecrawlt wird, kann auch nicht indexiert – und damit nicht gerankt – werden. Punkt.

Besonders bei großen Websites mit tausenden von URLs entscheidet das Crawling über Leben und Tod im organischen Ranking. Wenn du denkst, das betrifft dich nicht, weil du nur ein paar Seiten hast – falsch gedacht. Auch kleine Seiten können durch technische Fehler komplett aus dem Crawl fallen. Und dann stehst du da mit deinem SEO-optimierten Content, den niemand sieht.

So funktioniert der Googlebot in der Praxis – und worauf er achtet

Der Googlebot arbeitet nicht wie ein Mensch, sondern wie ein Bot. Klingt logisch, ist aber entscheidend. Er sieht keine Bilder, klickt nicht auf Buttons und hat kein Interesse an deinem Design. Was er will, ist Klarheit. Struktur. Effizienz. Und Ressourcen, die er interpretieren kann – vorzugsweise in HTML, nicht in JavaScript-Murks, der erst beim zweiten Rendering auftaucht.

Der Crawl beginnt mit der Analyse deiner robots.txt-Datei. Was hier blockiert ist, wird nicht gecrawlt – Punkt. Danach prüft Google die XML-Sitemaps, Linkstrukturen und interne Verlinkung, um zu erkennen, welche Seiten Priorität haben. Auch HTTP-Statuscodes, Canonical-Tags und Redirects spielen eine Rolle – denn sie beeinflussen, wie der Bot deine Inhalte einordnet.

Wichtige technische Signale, auf die der Googlebot achtet:

- HTTP-Statuscodes: 200 = alles gut, 301/302 = Weiterleitung (aber wie viele?), 404 = Fail
- Crawl-Delay und Server-Antwortzeiten: Langsame Server? Dann reduziert Google dein Crawl-Rate automatisch
- robots.txt und Meta Robots: Blockst du versehentlich wichtige Seiten? Dann sieht Google sie nicht
- Canonical-Tags: Falsch gesetzt? Dann wird die falsche Seite indexiert (oder gar keine)
- XML-Sitemap: Veraltet oder fehlerhaft? Dann crawlt Google ins Leere

Der Googlebot nutzt außerdem sogenannte Caffeine-Systeme im Backend. Diese priorisieren neue und oft aktualisierte Inhalte – Stichwort „Freshness“. Wenn deine Seite veraltet oder statisch wirkt, sinkt die Crawl-Frequenz. Umgekehrt gilt: Wer regelmäßig aktualisiert (und dabei nicht ins Duplicate-Content-Loch fällt), bekommt mehr Aufmerksamkeit vom Crawler.

Crawl-Budget verstehen und effektiv nutzen

Das Crawl-Budget ist kein Marketing-Buzzword, sondern harte Realität. Es beschreibt die Menge an URLs, die Google innerhalb eines bestimmten Zeitraums auf deiner Domain crawlt. Dieses Budget ist endlich – und wird durch viele Faktoren beeinflusst: Servergeschwindigkeit, Domainautorität, Anzahl der internen Links, Seitenstruktur, Crawling-Fehler, Duplicate Content und mehr.

Besonders bei großen Websites oder E-Commerce-Shops mit Tausenden von

Produktseiten ist das Crawl-Budget kritisch. Wenn du Google mit tausenden irrelevanten oder fehlerhaften URLs zumüllst – etwa durch Filterkombinationen oder Session-IDs – verschwendest du Ressourcen. Die Folge: Wichtige Seiten werden seltener oder gar nicht gecrawlt.

Um dein Crawl-Budget effizient zu nutzen, beachte folgende Punkte:

- Vermeide Duplicate Content und parameterbasierte URL-Explosionen
- Setze Canonical-Tags korrekt und konsistent
- Blockiere unwichtige Seiten gezielt über robots.txt oder Meta Robots
- Halte deine XML-Sitemap schlank und aktuell
- Behebe Crawling-Fehler regelmäßig (404, 500 etc.)
- Nutze Pagination (rel=prev/next) oder strukturierte interne Verlinkung

Google selbst sagt, Crawl-Budget sei für die meisten Seiten kein Problem. Das ist halb wahr. Für Blogs mit 20 Seiten vielleicht nicht. Für skalierende Projekte oder Shops mit 50.000 URLs ist es existenziell. Du willst indexiert werden? Dann verschwende Googles Ressourcen nicht – sonst wirst du einfach ignoriert.

Tools für Crawling-Analyse: Was wirklich hilft – und was nur blendet

Du kannst nur optimieren, was du verstehst. Und verstehen kannst du nur, was du misst. Deshalb sind technische SEO-Tools kein Luxus, sondern Pflicht. Die Google Search Console ist dein Einstieg – sie zeigt dir gecrawlte Seiten, Crawling-Fehler, Indexierungsstatus, Mobilfreundlichkeit und mehr. Aber sie ist limitiert.

Für tiefere Crawling-Analysen brauchst du Tools wie:

- Screaming Frog: Simuliert den Googlebot und zeigt dir Statuscodes, Meta-Daten, Canonicals, Redirects, Linkstruktur und mehr
- Sitebulb: Visuelle Crawling-Analyse mit UX-Fokus und technischer Tiefe
- Logfile-Analyse (z. B. mit Screaming Frog Log Analyzer): Zeigt dir, welche Seiten Google wirklich besucht – nicht nur theoretisch
- Ahrefs/Semrush: Zeigen Crawling-Fehler, Broken Links, Duplicate Content etc.

Besonders wertvoll ist die Logfile-Analyse. Sie zeigt dir schwarz auf weiß, welche Seiten der Googlebot wie oft besucht, welche er ignoriert, und wo Serverfehler auftreten. Das ist die ungeschönte Wahrheit – keine Schätzung, kein Tool-Fake, sondern echte Daten aus deinem Server-Backend.

Die häufigsten Crawling-Killer – und wie du sie abstellst

Viele Websites scheitern am Crawling – nicht aus Böswilligkeit, sondern aus Ignoranz. Die häufigsten Fehler sind:

- JavaScript-Overkill: Inhalte, die nur per JavaScript nachgeladen werden, werden oft nicht gecrawlt oder verstanden
- Endlosschleifen: Kalender, Filter oder Paginierung ohne Limit erzeugen Millionen nutzlose URLs
- Kaputte interne Links: 404-Seiten oder falsche Weiterleitungen bremsen den Crawl
- robots.txt-Fehler: Wichtige Ressourcen (z. B. CSS/JS) werden geblockt
- Unsaubere Canonicals: Google weiß nicht, welche Seite die „echte“ ist

Der Weg zur Crawling-Effizienz ist brutal einfach – aber erfordert Disziplin:

1. Analysiere deinen Status mit Screaming Frog & Logfile-Tools
2. Bereinige kaputte Links, überflüssige Seiten, Dubletten
3. Strukturiere deine Website flach und sauber
4. Setze Canonicals, robots.txt und XML-Sitemap korrekt ein
5. Teste regelmäßig mit Search Console, Lighthouse & Co.

Fazit: Wer den Crawler versteht, gewinnt das SEO-Spiel

Der Googlebot ist kein Mythos, keine Blackbox und kein Gegner – er ist ein maschineller Besucher, der klare Regeln versteht. Wenn du ihm den Weg ebnest, indexiert er deine Seite. Wenn du ihn verwirrst, blockierst oder mit JavaScript zuschüttest, ignoriert er dich. So einfach – und so brutal – ist das.

Technisches SEO beginnt beim Crawling. Wer hier patzt, braucht sich über Rankings nicht zu wundern. Du willst Sichtbarkeit? Dann hör auf, nur an Content zu denken – und fang endlich an, deinen Code, deine Struktur und deinen Crawl-Flow ernst zu nehmen. Denn Content ohne Crawl ist wie ein Buch im Safe: Niemand wird es je lesen.