

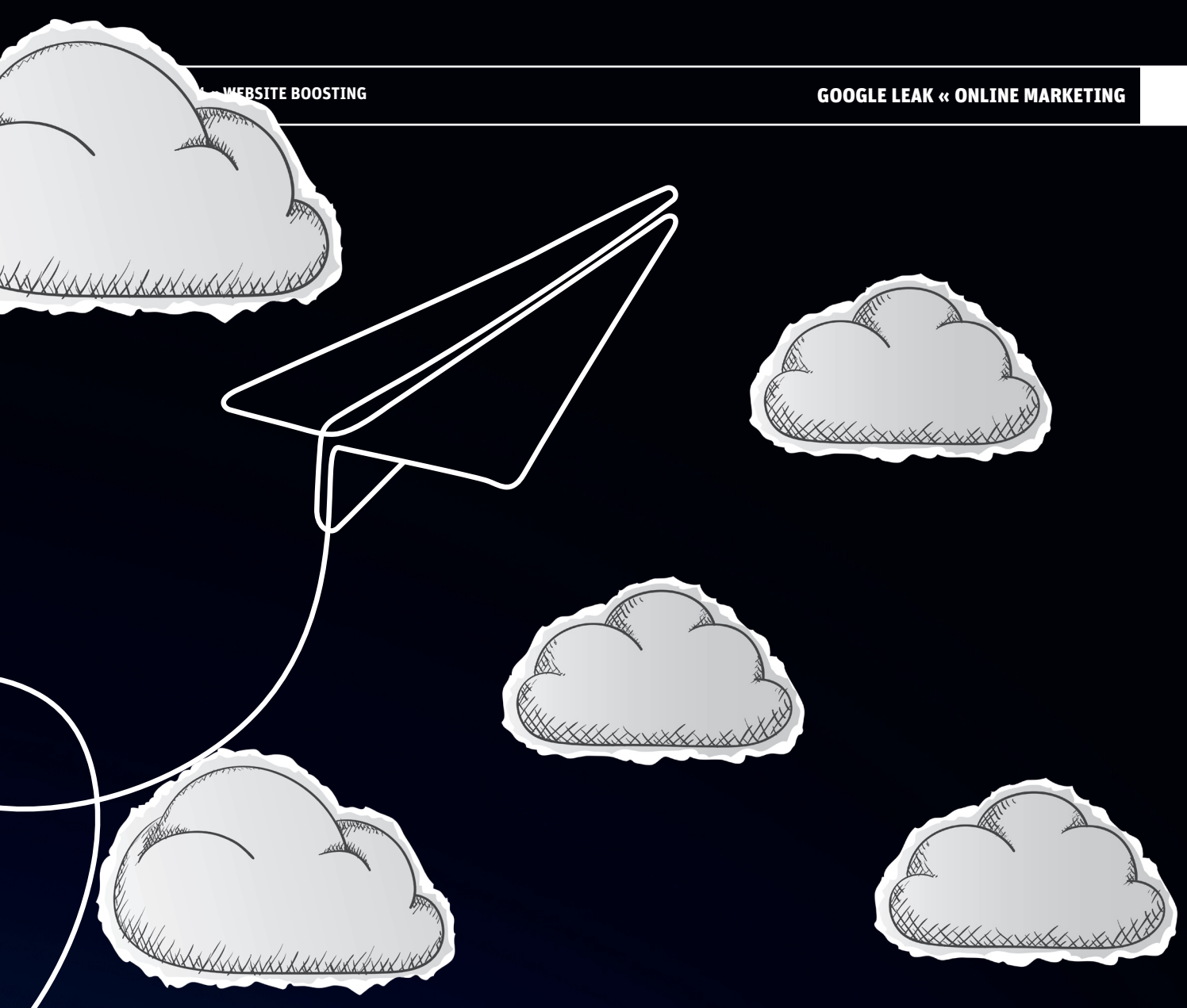
WIE FUNKTIONIERT RANKING BEI GOOGLE? EIN DOKUMENT MACHT SICH AUF DIE REISE IN DIE TOP TEN

Mario Fischer

Dass wir aus den Google-Leaks oder den öffentlichen FTC-Dokumenten aus Anhörungen nicht wirklich die genaue Funktionsweise des Rankings herauslesen können, dürfte wohl jedem klar sein. Der Aufbau der organischen Suchergebnisse ist mittlerweile – nicht zuletzt durch den Einsatz von Maschine Learning – derart komplex, dass selbst die Google-Mitarbeiter sagen, die an den Ranking-Algorithmen arbeiten, sie könnten auch nicht mehr erklären, warum ein Treffer auf eins oder zwei stünde.

Wir kennen keine Gewichtungen der vielen Faktoren und natürlich auch nicht das genaue Zusammenspiel. Trotzdem ist es wichtig, sich mit dem Aufbau der Suchmaschine vertraut(er) zu machen, um verstehen zu können, warum gut optimierte Seiten nicht ranken oder umgekehrt scheinbar kurze und nicht optimierte Ergebnisse manchmal ganz oben in den Rankings auftauchen. Der wichtigste Aspekt ist sicherlich, dass man den Blick deutlich für das erweitern muss, was wirklich wichtig ist. DAS kann man aus den ganzen vorliegenden Informationen schon sehr gut herauslesen.

Jeder, der sich auch nur am Rande mit Ranking beschäftigt, sollte diese Erkenntnisse ins eigene Mindset aufnehmen. Sie werden Ihre Websites mit völlig anderen Augen sehen und weitere Metriken in Ihre Analysen, Planungen und Entscheidungen einfließen lassen! Am besten legen Sie sich die große Beilage aus dieser Ausgabe mit der Übersicht beim Lesen daneben, das erleichtert das Verständnis der Zusammenhänge.



Im Beitrag „Hey, Google, du hast da was leaken lassen!“ von Johan von Hülsen in dieser Ausgabe wird erklärt, welche Kerninformationen das Leak enthält, wie die Module prinzipiell aufgebaut sind und welche Learnings sich daraus ergeben. Dieser Beitrag zeigt, wie die einzelnen Systeme der Suchmaschine zusammenhängen und wie sie arbeiten.

Es ist extrem schwierig, ein wirklich valides Bild von der Struktur der Systeme zu zeichnen. Die Informationen im Web sind durchaus in der Interpretation verschieden und unterscheiden sich teilweise in den Begrifflichkeiten, obwohl das Gleiche gemeint ist. Ein Beispiel: Das System, das zum Aufbau einer SERP (Suchergebnisseite) für die optimale Platznutzung zuständig ist, heißt Tangram. In einzelnen Google-Dokumenten wird es aber auch als

Tetris bezeichnet, was wohl der Anlehnung an das bekannte Spiel geschuldet ist. Wir haben nahezu an die 100 Dokumente gesichtet, analysiert, strukturiert, verworfen, neu strukturiert, und das in wochenlanger Kleinstarbeit viele Male hintereinander. Der folgende Beitrag kann und will daher weder den Anspruch auf Vollständigkeit noch auf

TIPP

Zum besseren Überblick finden Sie in dieser Ausgabe eine herausnehmbare Beilage, in der die (vermutete) Struktur zumindest grob dargestellt wird. Dieser Überblick dürfte hinsichtlich des Detaillierungs- und Informationsgehalts weltweit vermutlich einzigartig sein.

DocID #73923328

https://www.domain.de/bleistift_A (kanonisch)

https://www.domain.de/sonderangebote/bleistift_A

https://www.domain.de/en/pencil_A

https://www.domain.de/fr/crayon_A

<https://www.bleistiftwelt.de/seite32>

...

Abb. 1: Alexandria sammelt URLs zu einem Dokument.

formale Korrektheit erheben. In die Waagschale können wir nur Leiden-schaft, die Floskel „mit bestem Wissen und Gewissen“ und eine gehörige Portion Columbo werfen. Das hier ist dabei herausgekommen.

Am besten lässt sich das Ranking verstehen, wenn man ein Dokument im Weg auf seine Reise durch die einzelnen Systeme begleitet.

Ein Dokument wartet auf den Besuch des Google-Bots

Publiziert man eine neue Website, wird sie nicht sofort indexiert. Google muss erst einmal Kenntnis von der URL bekommen. Das geschieht in der Regel entweder über eine aktualisierte Sitemap oder über einen Link, der von einer bereits bekannten URL dort hingestellt wurde. Seiten, die häufiger besucht werden wie zum Beispiel die Startseite, bringen diese Linkinformation Google natürlich schneller zur Kenntnis.

Google nutzt sieben verschiedene Arten von PageRank.

Das Trawler-System sorgt dafür, dass der neue Inhalt abgerufen wird, und merkt sich vor, wann diese URL erneut besucht wird, um sie wegen möglicher Änderungen erneut zu bewerten. Diese Aufgabe übernimmt der sogenannte Scheduler. Im Store-server wird entschieden, ob die URL weitergereicht wird oder ob sie in die Sandbox gestellt wird. Die Existenz dieser Box hat Google über die Jahre abgestritten, die Leaks legen jedoch nahe, dass dort vor allem (vermutete) Spam-Seiten und Seiten mit geringem Wert eingestellt werden. Am Rand sei erwähnt, dass Google offenbar einen Teil des Spams durchleitet, wahrscheinlich zur weiteren Analyse zum Training der Algorithmen. Unser fiktives Dokument passiert diese Schranke. Von

TIPP

Geben Sie Ihren Dokumenten plausible und stimmige Datumswerte mit auf den Weg. Unter anderem werden `bylineDate` (Datum im Quellcode), `syntaticDate` (extrahiertes Datum aus URL und/oder Title) und `semanticDate` (gezogen aus dem lesbaren Content) verwendet. Aktualität per Datumsänderung zu faken, kann durchaus zum Downranking (Demotion) führen. Im Attribut `lastSignificantUpdate` wird festgehalten, wann die letzte bedeutende Änderung an einem Dokument gemacht wurde. Kleinigkeiten oder Tippfehler auszubessern, beeinflusst diesen Zähler nicht.

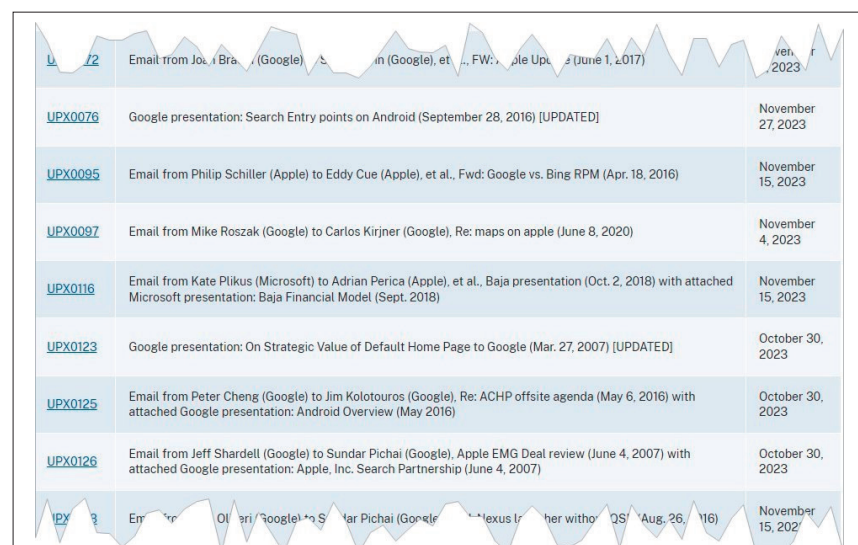
unserem Dokument abgehende Links werden extrahiert und nach intern oder extern abgehend sortiert. Diese Informationen werden von anderen Systemen vor allem zur Linkanalyse und zur Berechnung des PageRanks verwendet. Dazu später mehr. Gefundene Links zu Bildern werden an den ImageBot übergeben, der sie teils mit deutlichem Zeitverzug aufruft und (zusammen mit gleichen oder ähnlichen Bildern) in einen Bildcontainer einstellt. Trawler verwendet offenbar einen eigenen PageRank zur Justierung der Crawl-Fre-

quenz. Hat eine Website mehr Traffic, steigt diese Crawl-Frequenz an (`ClientTrafficFraction`).

„There is no sandbox.“ – Google

Alexandria – die große Bibliothek

Das Indexingsystem von Google heißt Alexandria. Dort wird eine eindeutige DocID vergeben. Ist der Inhalt schon bekannt, zum Beispiel durch Dubletten, wird keine neue ID erzeugt, sondern die URL zu einer bereits bestehenden DocID hinzugefügt. Wichtig: Google unterscheidet zwischen URL und Dokument! Ein Dokument kann aus einem Set von URLs bestehen, die mehr oder weniger den gleichen Content beinhalten. Auch unterschiedliche Sprachversionen, sofern sie sauber gekennzeichnet sind, werden hier eingestellt. Auch URLs von anderen Domains werden hier einsortiert. Das bedeutet, dass die Signale aller hier verwendeten URLs über die gemeinsame DocID gelten. Google bestimmt bei Dubletten eine kanonische Variante und diese wird später im Ranking



DocID	Document Title	Date
UPX0076	Google presentation: Search Entry points on Android (September 28, 2016) [UPDATED]	November 27, 2023
UPX0095	Email from Philip Schiller (Apple) to Eddy Cue (Apple), et al., Fwd: Google vs. Bing RPM (Apr. 18, 2016)	November 15, 2023
UPX0097	Email from Mike Roszak (Google) to Carlos Kirjner (Google), Re: maps on apple (June 8, 2020)	November 4, 2023
UPX0116	Email from Kate Plikus (Microsoft) to Adrian Perica (Apple), et al., Baja presentation (Oct. 2, 2018) with attached Microsoft presentation: Baja Financial Model (Sept. 2018)	November 15, 2023
UPX0123	Google presentation: On Strategic Value of Default Home Page to Google (Mar. 27, 2007) [UPDATED]	October 30, 2023
UPX0125	Email from Peter Cheng (Google) to Jim Kolotouros (Google), Re: ACHP offsite agenda (May 6, 2016) with attached Google presentation: Android Overview (May 2016)	October 30, 2023
UPX0126	Email from Jeff Shaddell (Google) to Sundar Pichai (Google), Apple EMG Deal review (June 4, 2007) with attached Google presentation: Apple, Inc. Search Partnership (June 4, 2007)	October 30, 2023
UPX0127	Email from Oliver Schaefer (Google) to Sundar Pichai (Google), Lexus lawsuit (Aug. 26, 2016)	November 15, 2023

Abb. 2: Neben den Leaks sind die Beweisdokumente aus Anhörungen und Prozessen der US-Justiz gegen Google eine nützliche Quelle für Recherchen. Dort finden sich sogar interne E-Mails! (siehe einfach.st/justicegov)

ausgegeben. Das erklärt vermutlich auch, warum manchmal andere URLs an nahezu gleicher Stelle ranken, weil sich die Bestimmung der als „original“ (kanonisch) erkannten URL durchaus im Lauf der Zeit ändern kann.

Da es für unser Dokument nur diese eine Variante im Netz gibt, bekommt es eine eigene DocID.

Einzelne Segmente unserer Seite werden nach relevanten Keyword-Phrasen durchsucht und in den Suchindex geschoben. Dort wird die „Hitlist“ (alle wichtigen Wörter der Seite) zunächst in den direkten Index geschickt, der die mehrfach auftretenden Keywords pro Seite zusammenfasst. Jetzt passiert ein wichtiger Schritt. Die einzelnen Keyword-Phrasen werden in den Wortkatalog des invertierten Index (Wortindex) eingegliedert. Dort stehen bereits das Wort Bleistift und alle wichtigen Dokumente, die dieses Wort enthalten. Da unser Dokument das Wort Bleistift prominent und mehrfach enthält, steht es fortan im Wortindex mit seiner DocID beim Eintrag „Bleistift“. Das alles ist natürlich nur eine vereinfachte Beschreibung. Die DocID bekommt für Bleistift einen algorithmisch berechneten IR-Score (IR = Information Retrieval) zugewiesen, der später für die Aufnahme in die Posting List verwendet wird. In unserem Dokument wurde das Wort Bleistift beispielsweise im Text mit Fettschrift ausgezeichnet und ist in der H1 enthalten (gespeichert in „AvrTermWeight“). Solche und andere Signale erhöhen den IR-Score.

Als wichtig erachtete Dokumente schiebt Google in den sogenannten HiveMind, also den Hauptspeicher. Daneben gibt es noch die schnellen SSDs und herkömmliche HDDs (Tera-Google), in denen Informationen langfristig abgelegt werden, die nicht mit hoher Antwortgeschwindigkeit benötigt werden. Google hält Dokumente und Signale dafür im Hauptspeicher. Dazu muss man wissen, dass Fachleute

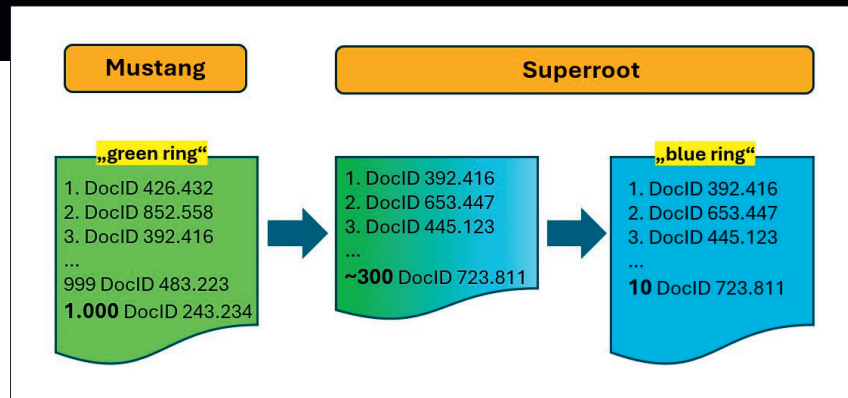


Abb. 3: Mustang erzeugt 1.000 potenzielle Ergebnisse und Superroot filtert diese auf zehn Ergebnisse.

HINWEIS

Das Google-Caffein-Hardware-Update, das den sogenannten Google Dance abgelöst hat, gibt es in dieser Form nicht mehr. Nur der Name ist geblieben. Google arbeitet mittlerweile mit unzähligen Micro-Services, die miteinander kommunizieren und Attribute für Dokumente erzeugen, die von den unterschiedlichsten Ranking- und Re-Ranking-Systemen als Signale verwendet werden und mit denen die neuronalen Netze zur Vorhersage trainiert werden.

schätzen, dass zumindest vor dem KI-Boom der letzten Zeit etwa die Hälfte der weltweit eingesetzten (Web-)Server bei Google stehen. Eine gigantische Armada, die im Clusterverbund arbeitet, womit auch die vielen Millionen Hauptspeichereinheiten zusammenschaltbar sind. Ein Google-Engineer hat einmal auf einer Konferenz angemerkt, dass es theoretisch möglich wäre, das gesamte Web (!) bei Google im Hauptspeicher abzulegen! Interessant ist, dass zum Beispiel Links (auch Backlinks), die im HiveMind gespeichert sind, offenbar deutlich mehr Gewicht bekommen. Learning nebenbei: Links von wichtigen Dokumenten zählen daher ungleich mehr, Links von URLs aus TeraGoogle (HDD) weniger oder möglicherweise gar nicht.

Im Repository werden weitere Informationen und Signale für jede DocID dynamisch gespeichert (PerDocData).

Darauf greifen später viele Systeme zu, wenn es um die Feinarbeit der Relevanz geht. Es ist nützlich, zu wissen, dass dort die letzten 20 Versionen eines Dokuments abgelegt sind (via Crawler-ChangerateURLHistorie). Google kann also auch Änderungen über die Zeit hinweg aus- und bewerten. Will man ein Dokument inhaltlich beziehungsweise thematisch verändern, muss man theoretisch 20 (Zwischen-)Versionen erzeugen, damit man die bisherigen Signale als Altlast loswird. Hier liegt auch der Grund, warum das Neubefüllen einer sogenannten Expired Domain (eine Domain, die bereits vorher im Web aktiv war und aufgegeben oder verkauft wurde, zum Beispiel wegen einer Insolvenz) keinen Ranking-Vorteil bringt. Ändert sich der Admin-C einer Domain UND gleichzeitig der thematische Inhalt, lässt sich das an dieser Stelle maschinell leicht erkennen. Google stellt dann alle Signale auf null und die vermeintlich wertvolle alte Domain bietet keinerlei Vorteile mehr gegenüber einer komplett neu registrierten Domain.

QBST: Jemand sucht nach „Bleistift“

Gibt jemand im Suchschlitz von Google Bleistift als Suchwort ein, fängt QBST mit der Arbeit an. Die Suchphrase wird analysiert und bei mehreren Wörtern werden die als relevant erkannten Wörter zur Abfrage an den Wortindex vorgesehen. Die Termgewichtung ist

dabei relativ komplex. Hier kommen unter anderem das bekannte Rank-Brain, DeepRank (ehemals BERT) und RankEmbeddedBert zum Einsatz. Die relevanten Begriffe (hier nur einfach „Bleistift“) werden an den sogenannten Ascorer übergeben.

Ascorer: Der „green ring“ entsteht

Der Ascorer extrahiert aus dem invertierten Index die ersten 1.000 DocIDs für „Bleistift“, absteigend nach dem IR-Score. Gemäß internen Dokumenten wird diese Liste als „green ring“ bezeichnet, in der Branche ist sie als Posting List bekannt. Der Ascorer ist Teil des als Mustang bezeichneten Ranking-Systems. Dort finden noch weitere Filterungen statt, zum Beispiel durch Deduplizierung über SimHash (eine Art Fingerprint für ein Dokument), das Passagen-System, Systeme zur Erkennung von Originalinhalten, hilfreichen Content etc. Ziel ist es, durch verschiedene Filter die Liste der 1.000 Kandidaten am Ende auf die berühmten „ten blue links“, den sogenannten „blue ring“, einzugrenzen. Unser Bleistift-Dokument hat es in die Posting List geschafft und steht fiktiv an Stelle 132. Gäbe es nicht noch weitere Systeme, wäre es für uns hier zu Ende.

Superroot: Aus 1.000 mach zehn!

Das als Superroot bezeichnete System ist für das Re-Ranking zuständig und erledigt somit die Feinarbeit, den „green ring“ (1.000 DocIDs) auf den „blue ring“ mit nur noch zehn Ergebnissen zu reduzieren: Diese Aufgabe wird von den sogenannten Twiddlern und NavBoost erledigt. Wahrscheinlich sind hier noch andere Systeme im Einsatz, aber die genaue Struktur und Zusammensetzung ist wegen der hier teils sehr vagen Informationen nicht erkennbar.

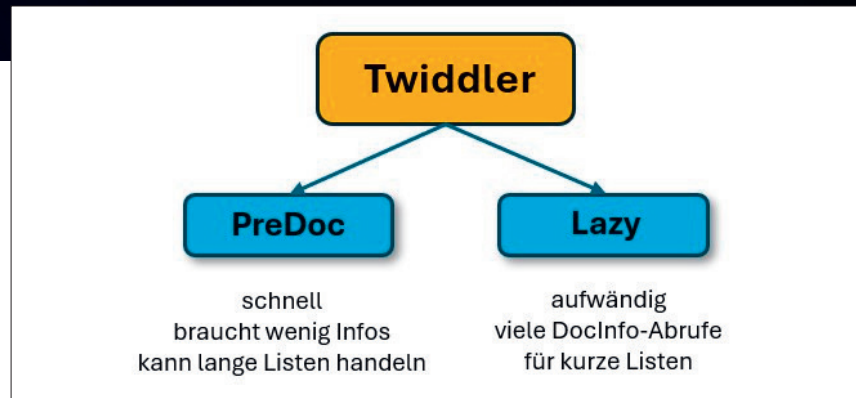


Abb. 4: Über 100 unterschiedliche Twiddler reduzieren die potenziellen Suchergebnisse und sortieren sie um.

Filter über Filter: die Twiddler

Aus verschiedenen Dokumenten geht hervor, dass mehrere Hundert Twiddler-Systeme im Einsatz sind. Man kann sich einen Twiddler als eine Art Plug-in wie bei Wordpress vorstellen. Jeder Twiddler hat ein eigenes Filterziel. Der Grund dafür ist, dass sie relativ einfach zu erstellen sind und keine der großen Ranking-Algorithmen verändert werden müssen, wie sie im Ascorer zum Einsatz kommen. Letztere sind sehr komplex und selbst kleine Änderungen würden wegen der möglichen Nebenwirkungen einen großen Planungs- und Programmieraufwand bedeuten. Das wäre für diverse Tests deutlich zu schwierig. Die Twiddler hingegen laufen parallel oder nacheinander ab und wissen nicht, was die anderen Twiddler tun.

Es gibt prinzipiell zwei Typen von Twiddlern. PreDoc-Twiddler können mit dem gesamten Set von mehreren Hundert DocIDs arbeiten, weil sie keine weiteren oder nur wenige Informationen benötigen. Im Gegensatz dazu brauchen die Twiddler vom Typ „Lazy“ mehr Informationen, zum Beispiel aus der PerDocData-Datenbank. Das dauert entsprechend länger und ist aufwendiger. Daher schmelzen die PreDocs die Posting List zunächst auf deutlich weniger Einträge ab und dann startet man mit den langsameren Filtern. Das spart enorm Rechenkapazität und Zeit.

Die einen Twiddler verändern den IR-Score positiv oder negativ, andere passen die Ranking-Position an. Da unser Dokument neu im Index ist, könnte ein Twiddler, der dafür sorgen soll, dass auch aktuelle Dokumente eine Chance auf Ranking haben, den IR-Score beispielsweise mit dem Faktor 1,7 multiplizieren. Wir rutschen von Platz 132 entsprechend weiter nach oben auf Platz 81. Ein anderer Twiddler sorgt für mehr Vielfalt (strideCategory) in den SERPs und wertet inhaltlich ähnliche Dokumente ab. Über den hinterlegten Demotion-Score und die erzeugte Abwertung verlieren mehrere Dokumente über uns ihre Position und wandern hinter uns. Unser Bleistift-Dokument macht zwölf Plätze gut und steht nun auf 69. Ein weiterer Twiddler

TIPP

Wer mehr über die Arbeit der Quality Rater wissen möchte, findet grundlegende Erklärungen von Olaf Kopp dazu bereits in Ausgabe 28 und online unter einfach.st/sqrbeitrag. Die aktuelle Version der Guidelines ist direkt bei Google als PDF unter einfach.st/goguide abrufbar.



hat die Regel, (bei bestimmten Suchanfragen) nur maximal drei Seiten von Blogs zuzulassen. Wir steigen auf 61.

Unsere Seite wurde beim Attribut „CommercialScore“ mit einer Null (für „Ja“) versehen, da die Systeme im Mustang-System bei der Analyse eine Verkaufsabsicht erkannt haben. Nehmen wir an, Google weiß, dass der Sucheingabe von „Bleistift“ sehr häufig eine sogenannte „refined search“ mit „Bleistift kaufen“ folgt. Der Sucheingabe von Bleistift wird also zumindest teilweise eine kommerzielle beziehungsweise transaktionale Absicht zugeordnet. Ein weiterer Twiddler ist dafür zuständig, Ergebnisse für den vermuteten Such-Intent beizumischen, und gibt uns einen Boostfaktor um 20 Positionen. Wir rutschen auf Position 41 weiter nach vorne.

Erneut schlägt ein weiterer Twiddler zu, der dafür sorgt, dass Seiten mit Verdacht auf Spam maximal auf Seite drei, sprich Position 31, ranken dürfen (die berühmte „Seite-drei-Strafe“). Unter BadURL-demonteindex wird eine beste Position festgelegt, über die das Dokument beim Ranking nicht hinauskommen darf. Dafür gibt es die Attribute DemoteForContent, DemoteForForwardlinks und DemoteForBacklinks. Das trifft für drei Dokumente über uns zu und wir springen eben durch dieses Freiwerden auf Position 38. Selbstverständlich könnte unser Dokument genauso gut auch eine Abwertung erfahren, aber um es nicht zu kompliziert zu machen, bleibt unsere Reise davon verschont. Lassen wir noch einen letzten beispielhaften Twiddler ins Spiel, der auswertet, wie stark unsere einzelne Bleistift-Seite von dem durch Embeddings berechneten Thema unserer Domain entfernt ist. Da sich unsere Website ausschließlich mit Schreibgeräten beschäftigt, ist das gut für uns und schlecht für 24 andere Dokumente.

Stellen wir uns zur Erklärung eine Preisvergleichsseite vor, die natürlich

thematisch vielfältig aufgestellt ist, aber auch eine „gute“ Seite zum Thema Bleistift hat. Das das Thema dieser Seite durch die sonstige Vielfalt stark abweicht, würde sie von diesem Twiddler eine Abwertung bekommen. Hierfür können Attribute wie „siteFocusScore“ oder „siteRadius“ herangezogen werden, die eben diese thematische Entfernung wiedergeben. Unser IR-Score wird noch einmal multipliziert und andere Ergebnisse werden verschlechtert. Wir landen auf Position 14.

Wie erwähnt gibt es Twiddler für völlig unterschiedliche Zwecke. Die Entwickler sind praktisch frei, mit neuen Filtern, Multiplikatoren oder bestimmen Positionseinschränkungen zu experimentieren. Es ist sogar möglich, ein Ergebnis gezielt nur hinter oder vor einem anderen Ergebnis ranken zu lassen. In einem der geleakten internen Dokumente von Google wird sogar davor gewarnt, bestimmte Arten von Twiddler-Funktionen nur zu nutzen, wenn man wirklich weiß, was man tut, und nur nach Rücksprache mit dem Search-Kernteam. „If you think you understand how they work [bezogen auf „merge cluster categories“, Anm. der Red.], trust us: you don't. We're not sure that we do either.“ Das ist im „Twiddler Quick Start Guide“ zu Superroot zu lesen, der der Redaktion vorliegt.

Ebenso gibt es Twiddler, die nur Annotations erstellen und diese der DocID mit auf den Weg in die SERP geben. Dort erscheint dann beispielsweise ein Bild im Snippet oder der Titel und/oder die Description werden später dynamisch umgeschrieben.

Übrigens: Wer sich in der Pandemie gewundert hat, warum zu Suchen rund um Corona plötzlich das Bundesgesundheitsministerium praktisch überall auf Platz eins zu finden war: Das war dem Einsatz eines Twiddlers zu verdanken, der via „queriesForWhichOfficial“ nach Sprache und Land offizielle Ressourcen boosten kann.

Natürlich hat man wenig bis keinen Einfluss, ob und wie Twiddler das eigene Ergebnis umsortieren. Aber es ist wichtig, deren Arbeitsweise zu verstehen, um zum Beispiel Ranking-Schwankungen oder „unerklärliche Rankings“ besser deuten zu können. Als Learning kann man festhalten, dass es sich auf jeden Fall lohnt, häufiger Blicke in die SERPs zu werfen und die Arten von Ergebnissen dort zu beachten. Erhält man auch bei einer Variation einer Suchphrase immer nur eine bestimmte Anzahl zum Beispiel an Foren- oder Blogbeiträgen? Wie viele Ergebnisse sind transaktional, informational oder navigational? Tauchen immer wieder die gleichen Domains auf oder variiert das mit einer bereits leichten Veränderung der Suchphrase? Wenn man feststellt, dass jeweils maximal drei Online-Shops im Ergebnis enthalten sind, ist es gegebenenfalls wenig sinnvoll, selbst mit einer solchen Seite ranken zu wollen. Vielleicht hilft das Ausweichen auf einen eher informationsorientierten Content? Voreilige Schlüsse sollte man allerdings jetzt noch nicht ziehen, es kommt später noch ein weiterer Mitspieler dazu: das System NavBoost.

Exkurs: Quality Rater und RankLab – hier werkeln Menschen

Weltweit arbeiten mehrere Tausend sogenannte Quality Rater für Google, um bestimmte Suchergebnisse zu bewerten und neue Algorithmen und/oder Filter zu testen, bevor sie „live“ gehen. Google selbst erklärt dazu: „Ihre Bewertungen haben keinen direkten Einfluss auf das Ranking“ (*einfach.st/grater8*). Das stimmt im Kern, aber wie wir gleich sehen werden, haben diese Votings durchaus indirekt einen gewichtigen Einfluss beim Ranking. Das Prinzip ist wie folgt: Die Rater bekommen vom System URLs oder Suchphrasen (Suchergebnisse) und

beantworten vorgegebene Fragen, die offenbar ausschließlich auf Mobilgeräten beurteilt werden sollen. Eine beispielhafte Frage wäre zum Beispiel „Ist klar, wer diesen Inhalt verfasst hat und wann? Ist erkennbar, ob die Person eine fachliche Expertise zu diesem Thema hat?“. Die Antworten auf diese und viele weitere Fragen werden gespeichert und zum Training per Machine Learning verwendet. Die Maschine versucht, herauszufinden, welche Merkmale gute und vertrauenswürdige Seiten haben und wie trennscharf diese gegenüber schlechten Seiten verwendbar sind. Somit ist man nicht mehr darauf angewiesen, dass sich Menschen im Search-Team bei Google „ausdenken“, welche Kriterien man für besser oder schlechter ansetzen könnte, sondern Algorithmen finden via Deep Learning Muster anhand des Trainings der menschlichen Bewerter. Ein Gedankenexperiment soll das verdeutlichen. Angenommen, Menschen antworten auf die Frage, ob ein Inhalt vertrauenswürdig erscheint, intuitiv und ohne dass es ihnen vielleicht bewusst ist, mit Ja, wenn ein Autorenbild und ein voller Name zu sehen ist. Gegebenenfalls gibt es noch einen Link zur Biografie auf LinkedIn. Seiten, die diese Merkmale nicht haben, werden als weniger vertrauenswürdig angesehen. Trainiert man nun ein neuronales Netz mit allen Seitenmerkmalen und den Bewertungen „Ja“ oder „Nein“, findet dieses diese trennende Eigenschaft heraus und verwendet sie nach einigen positiven Testläufen als Ranking-Signal. Solche Tests laufen in der Regel mindestens über 30 Tage. Dann könnte es sein, dass die Verwendung eines Autorenbilds mit Vor- und Nachname plus verlinktem LinkedIn-Profil Seiten (vielleicht via Twiddler?) pusht oder Seiten ohne dieses Merkmal abwertet. Dazu könnte passen, dass Google offiziell abstreitet, dass man auf Autoren achtet. Aus den Leaks weiß man jedoch, dass es Attribute wie „isAuthor“ gibt, und auch

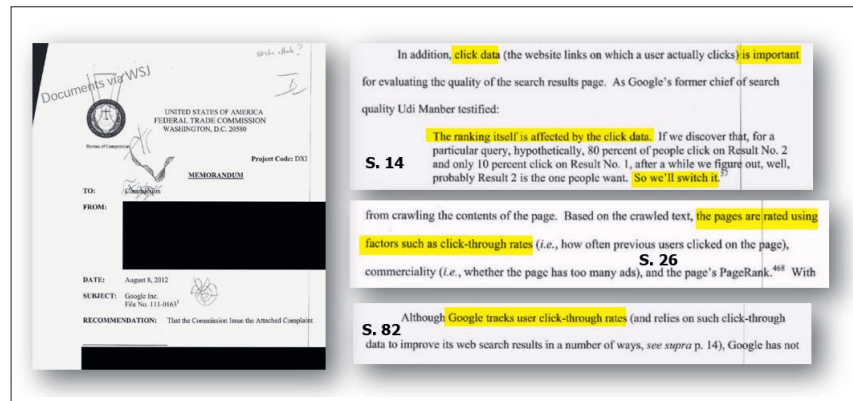


Abb. 5: Seit August 2012 (!) war offiziell klar, dass Klickdaten das Ranking verändern.

der Begriff des Autoren-Fingerprintings taucht via Attribut „AuthorVectors“ auf, das den Ideolekt (die individuelle Verwendung von Begriffen und Formulierungen) eines Autors – erneut über Embeddings – unterscheid- beziehungsweise identifizierbar macht.

Die von den Ratern vergebenen Bewertungen werden in einem sogenannten Information Satisfaction Score (IS-Score) zusammengefasst. Trotz der hohen Zahl an menschlichen Bewertern liegt natürlich nur für einen geringen Bruchteil an URLs ein IS-Score vor, der von null bis 100 geht. Daher wird dieser wie oben beschrieben auf andere Seiten mit ähnlichen Mustern hochgerechnet und für das Ranking verwendet. Google sagt dazu: „A lot of documents have no clicks but can be important.“ Wird für ein Dokument vom System eine Bewertung benötigt, weil Hochrechnungen aus bestimmten Gründen nicht machbar sind, wird es automatisch an die Rater geschickt, die dann einen Scorewert erzeugen.

Im Zusammenhang mit den Quality Ratern taucht an einigen Stellen das Attribut „golden“ auf. Es gibt also möglicherweise eine Art Goldstandard für Dokumente beziehungsweise Dokumentarten. Man kann also davon ausgehen, dass die Erfüllung der Erwartungen der menschlichen Tester das eigene Dokument in Richtung eines solchen Goldstandards treibt. Und es ist nicht unwahrscheinlich, dass es auch einen oder gar mehrere Twiddler gibt, die

DocIDs mit „golden“ eine ordentliche Multiplikatorwirkung mit auf den Weg geben oder sie gar fix und direkt in die Top Ten schieben.

Google sortiert mehr als 200 Milliarden Spam-Seiten aus – täglich!

Während die Quality Rater in der Regel keine Vollzeitangestellten von Google sind oder in Fremdfirmen auf Rechnung arbeiten, sitzen im sogenannten RankLab Google-eigene Experten. Dort führt man Experimente durch, setzt neue Twiddler auf und prüft, ob diese oder nur feinjustierte Twiddler die Ergebnisqualität erhöhen oder auch einfach nur mehr Spam ausfiltern. Twiddler, die sich bewährt und bewiesen haben und deren Einsatz bei allen Ergebnissen sinnvoll ist, werden gegebenenfalls in das vorgelagerte Mustang-System überführt. Zur Erinnerung: Dort laufen die komplexen, rechenintensiven und ineinandergreifenden Algorithmen ab.

Aber was wollen die User? NavBoost kann das richten!

Unser Bleistift-Dokument hat es noch nicht ganz geschafft. Im System Superroot gibt es ein weiteres Kernsystem, das einen offenbar ganz wesentlichen Einfluss auf die Reihenfolge von Ergebnissen hat: NavBoost. Dieses System nutzt sogenannte Slices und hält

TIPP

Einige Metriken, die gespeichert werden, sind unter anderem badClicks und goodClicks. Die Zeitdauer, die ein Suchender auf der Zielseite bleibt, und die Information, wie viele weitere Seiten er dort in welcher Zeit betrachtet (Chrome-Daten), fließen höchstwahrscheinlich in diese Bewertung ein. Ein nur kurzer Abstecher zu einem Suchergebnis und eine schnelle Rückkehr zu den Suchergebnissen und weitere Klicks auf andere Ergebnisse können die Zahl der badClicks erhöhen. Das Suchergebnis, das in einer Suchsitzung den letzten „guten“ Klick hatte, wird als lastLongestClick festgehalten. Dabei werden die Daten gesquasht, also verdichtet, damit sie statistisch normalisiert werden und weniger anfällig für Manipulationen sind. Hat eine Seite, ein Cluster von Seiten oder die Startseite einer Domain generell gute Besuchermetriken (Chrome-Daten), wirkt sich das über NavBoost positiv aus.

Über die Analyse von Bewegungsmustern innerhalb einer Domain oder über Domaingrenzen hinweg kann man sogar feststellen, wie gut die Nutzerführung über die Navigation ist. Da Google nach jeweils ganzen Such-Sessions misst, kann es theoretisch im Extremfall sogar passieren, dass erkannt wird, dass ein ganz anderes Dokument für eine Suchanfrage als passend erachtet wird. Verlässt ein Suchender innerhalb einer Suche die Domain, die er im Suchergebnis angeklickt hat, und geht zu einer anderen Domain (weil vielleicht sogar von dort aus verlinkt wurde) und bleibt dort als erkennbares Ende der Suche, könnte dieses „End“-Dokument gezielt über NavBoost künftig nach vorne gespült werden, sofern es im Set des Auswahlrings vorhanden ist. Dazu wäre allerdings ein starkes statistisch relevantes Signal von vielen Suchenden nötig.

für Mobile und Desktop sowie für lokale Bezüge unterschiedliche Datensets.

Bisher hat Google immer offiziell abgestritten, dass User-Klicks für das Ranking herangezogen werden. In den von der FTC veröffentlichten Dokumenten ist auch eine interne E-Mail-Anweisung zu finden, in der darauf hingewiesen wurde, dass der Umgang mit Klickdaten nicht an die Öffentlichkeit getragen werden darf. Das darf man Google nicht zum Vorwurf machen, denn diese Unwahrheit hat zwei wichtige Dimensionen. Zum einen gäbe es einen von Medien erzeugten Aufschrei über den „Datenkraken“, der uns jetzt auch noch beim Surfen „überwacht“. Ziel der Verwendung von User-Klicks ist jedoch nicht, einem User hinterherzuschneffeln, sondern Zugang zu statistisch relevanten Metriken zu bekommen. Der einzelne Nutzer ist dafür völlig uninteressant. Hier vertreten Datenschützer ganz sicher eine völlig andere Meinung, aber es macht zumindest verständlich, warum dies abgestritten wurde. In den FTC-Dokumenten heißt es übrigens an mehreren Stellen, dass die Klickdaten für das Ranking verwendet würden. Auch das dafür zuständige System NavBoost wird 54-mal erwähnt (in der Anhörung vom 18.04.2023). Bereits im Jahr 2012 kam bei einer offiziellen Anhörung heraus, dass Klickdaten das Ranking beeinflussen.

Seither ist also klar, dass sowohl das Klickverhalten auf Suchergebnisse als auch der Traffic auf einer Website oder Webseite für das Ranking herangezogen werden. Die Auswertung des Suchverhaltens kann Google recht einfach vornehmen, da Suche, Klicks, erneute Suche, erneute Klicks etc. direkt in den SERPs im Google-System anfallen. Immer wieder wurde öffentlich laut vermutet, Google würde die Bewegungsdaten einer Domain aus Google Analytics ableiten können, weshalb man dieses System nicht ver-

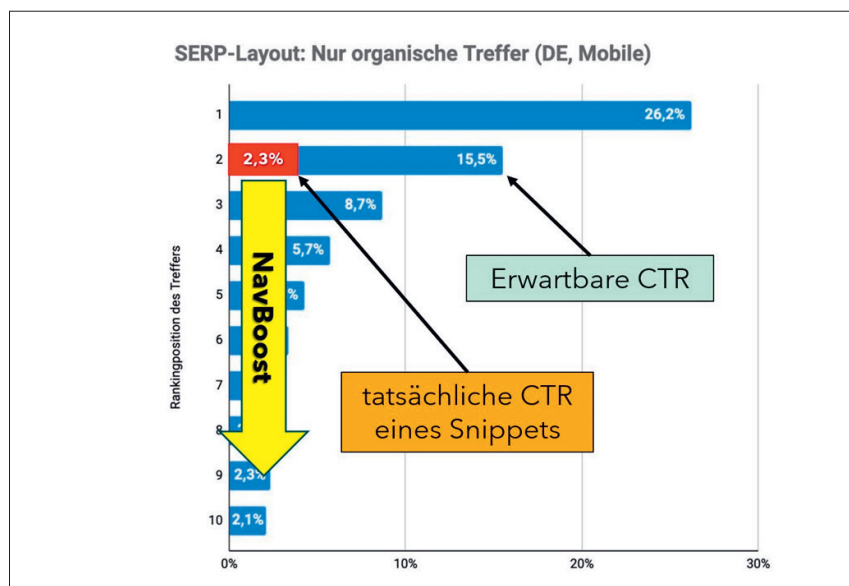


Abb. 6: Weicht die „expected_CTR“ signifikant vom tatsächlichen Wert ab, werden die Rankings entsprechend angepasst (Quelle: Johannes Beus, mit Overlays der Redaktion).

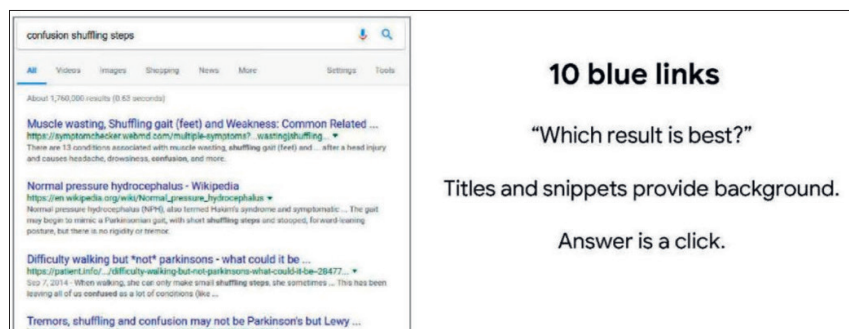


Abb. 7: Folie aus einer Google-Präsentation (Quelle: Trial Exhibit - UPX0228: U.S. and Plaintiff States v. Google LLC; www.justice.gov/d9/2023-09/416665.pdf)

wenden solle. Es gibt mehrere Gründe, warum das zu kurz greift. Zum einen würde man damit ja nicht an alle Bewegungsdaten einer Domain kommen. Zum anderen, und das ist eher der gewichtigere Beweis, nutzen je nach Studie über 60 % aller Menschen den Google-Chrome-Browser, weit über drei Milliarden User sind es also. Und da ein Browser alle aufgerufenen Websites nach Hause funkt, hat Google statistisch gesehen eine Stichprobe von über der Hälfte (n = 60 %) aller Bewegungen im Web zur Auswertung zur Verfügung. Im Puzzle für gute Rankings wäre Chrome sogar ein Schlüsselement, heißt es in den Anhörungen. Übrigens werden sogar ganz offiziell die Core-Web-Vital-Signale über die Chrome-Datenbank gesammelt und im Wert „chromeInTotal“ aggregiert.

„We don't use anything from Chrome for ranking.“

– Google

Es passiert also. Die erwähnte negative Publicity einer „Überwachung“ ist der eine Grund, der andere ist sicherlich darin zu sehen, dass man die Auswertung von Klick- und Bewegungsdaten als Aufforderung an Spammer und Trickser verstanden wissen möchte, noch mehr Traffic über Botsys-

TIPP

SEO-Experten und Datenanalysten berichten seit Längerem, dass sie bei einer umfassenden Überwachung der eigenen Klickraten das folgende Phänomen feststellen können: Erscheint ein Dokument für eine Suchanfrage neu in den Top Ten und bleibt die CTR deutlich hinter den Erwartungen zurück, kann man einen Rückgang des Rankings innerhalb weniger Tage (je nach Suchvolumen) beobachten. Umgekehrt steigt das Ranking oft, wenn die CTR bezogen auf den Rank deutlich höher liegt. Man hat nur wenig Zeit, zu reagieren, das Snippet

bei schlechter CTR anzupassen (in der Regel über eine Optimierung des Titles und der Description), damit mehr Klicks eingesammelt werden. Ansonsten verschlechtert sich die Position und ist anschließend nicht so leicht zurückzuerobieren. Man vermutet Tests hinter diesem Phänomen. Bewährt sich ein Dokument, darf es bleiben. Mögen die Suchenden es nicht, verschwindet es wieder. Ob das tatsächlich mit NavBoost zusammenhängt, ist natürlich weder klar noch abschließend beweisbar.

teme zu faken, um bessere Rankings zu erreichen. Man mag das Abstreiten nicht gut finden, die Gründe dahinter sind jedoch durchaus zumindest verständlich.

„We don't use clicks for ranking.“ – Google

Wenden wir uns zuerst den Klicks im Suchergebnis zu. Für jede Ranking-Position in den SERPs gibt es eine mittlere, zu erwartende Klickrate (CTR). Dies kann man sich als eine Art Leistungserwartungsschwelle vorstellen. Johannes Beus hat auf der diesjährigen CAMPiXX in Berlin (siehe Bericht in dieser Ausgabe) eine Analyse vorgestellt. So entfallen statistisch gesehen nach

seinen Zahlen auf Platz eins 26,2 % der Klicks. Platz zwei bekommt noch 15,5 % aller Klicks ab. Fällt die tatsächliche CTR eines Snippets deutlich unter die zu erwartende Klickrate, wird dies registriert und das System NavBoost korrigiert die letzten noch übrig gebliebenen DocIDs entsprechend. Hat ein Ergebnis also in der Vergangenheit statistisch signifikant deutlich mehr oder deutlich weniger Suchende zu einem Klick animiert, bewegt NavBoost das Dokument entsprechend nach oben oder unten in der Reihenfolge (Abbildung 6). Das ist nur konsequent, denn die Klicks sind ja nichts anderes als eine Abstimmung der Suchenden, wie gut sie ein Ergebnis anhand des Titles und der Description und auch bezüglich der Domain für ihre Suche empfinden. In einem offiziellen Dokument kann man dies sogar nachlesen, wie Abbildung 7 zeigt.

Für unser Bleistift-Dokument liegen noch keine oder nicht genügend CTR-Werte vor, da es ja noch neu ist. Ob eine Korrektur via CTR-Abweichung für Dokumente, für die keine Daten vorliegen, einfach ignoriert wird, ist nicht bekannt. Wahrscheinlich ist es allerdings, denn es geht ja darum, die Stimme der Suchenden einfließen zu lassen. Natürlich könnte es auch möglich sein, dass die CTR ähnlich wie bei Google Ads der Qualitätsfaktor auf der Basis von anderen Werten zunächst

TIPP

Google hat bei einer Anhörung selbst ein gutes Beispiel für Freshness gegeben. Der Becher der Firma Stanley ist wahrscheinlich auch hierzulande ein Begriff. Sucht man nach dem Stanley Cup, erscheint der berühmte Becher. Ist allerdings zeitgleich der Stanley Cup, ein Pokalspiel im Eishockey, sorgt NavBoost durch das geänderte Suchbeziehungswise Klickverhalten dafür, dass Informationen rund um die Spiele ganz oben auftauchen. Freshness bezieht sich dabei eben nicht auf neue, also „frische“ Dokumente, sondern auf geändertes Such-

verhalten. Nach Angaben von Google gibt es täglich über eine Milliarde (das ist kein Schreibfehler) neue Verhaltensweisen in den SERPs! Jede Suche und jeder Klick trägt also zum Learning von Google bei. Die Vermutung, dass Google alles über Saisonalitäten wüsste, ist so wahrscheinlich nicht richtig. Google erkennt feingranular Änderungen in den Suchabsichten und passt das System laufend an – was bei uns die Illusion erzeugt, dass Google tatsächlich „verstehen“ würde, was Suchende wollen.

hochgerechnet wird.

Durch die Auswertungen der Leaks liegt die begründete Vermutung nahe, dass Google für neue und noch unbekannte Seiten umfassende Informationen aus dem „Umfeld“ dieser Seite für eine Hochrechnung von Signalen verwendet. Über „NearestSeedversion“ wird vermutlich der PageRank der Startseite (HomePageRank_NS) auf neue Seiten übertragen, solange denen noch kein eigener Wert zugewiesen werden konnte. Und über „pnavClicks“ wird offenbar sogar die Wahrscheinlichkeit für Klicks via Navigation berechnet und zugewiesen. Die Berechnung beziehungsweise Aktualisierung des PageRanks ist sehr aufwendig und rechenintensiv. Daher wird offenbar mittlerweile der Nachfolger PageRank_NS als Metrik verwendet. NS steht für „nearest seed“ und bedeutet salopp erklärt, dass ein Set von Seiten als zusammengehörig betrachtet wird und der PageRank-Wert der bekannten Seiten auf die neue Seite (vorläufig oder dauerhaft, das ist unklar) übertragen wird. Wahrscheinlich werden auch für andere wichtige Signale Werte von „Nachbarn“ herangezogen, damit neue Seiten überhaupt eine Chance haben, im Ranking nach oben zu kommen. Sie haben ja bisher weder nennenswert Besucher noch Backlinks. Und auch die Zurechnung vieler Signale

TIPP

Anders als vermutet liefert Google wenig tatsächlich personalisierte Suchergebnisse aus. Tests haben wohl ergeben, dass die Modellierung des User-Verhaltens und deren Änderungen bessere Ergebnisse liefern als die Auswertung persönlicher Vorlieben einzelner User. Das ist bemerkenswert. Die Vorhersage über neuronale Netze ist mittlerweile besser für uns passend als die eigene Surf- und Klickhistorie. Einzelne Präferenzen wie zum Beispiel eine Vorliebe für Videocontent fließen allerdings nach wie vor in die persönlichen Ergebnisse ein.

250 Jahre Faber-Castell

★★★★☆ [37]

Geschichte

Der Bleistift im Wandel der Zeit

Wo und wann der Bleistift zur Welt kam, ist wissenschaftlich nicht eindeutig nachgewiesen. Um seine Erfindung, die weit zurück geht, ranken sich aber viele Geschichten. Fest steht: Die Karriere des Bleistifts fängt mit einem falschen Namen an.

Abb. 8: Im Fernsehen läuft aktuell ein Beitrag über die Begriffsentstehung des Worts „Bleistift“ zur Feier von 250 Jahren Faber-Castell.

passiert nicht in Echtzeit, sondern mit teilweise deutlichem Zeitverzug.

Die Klickmetriken für Dokumente werden offenbar nach neuesten Erkenntnissen über einen Zeitraum von 13 Monaten (ein Monat Überschneidung im Jahr für Vorjahresvergleiche) aufbewahrt und ausgewertet.

„We do not understand documents. We fake it. So we watch how people react to documents.“

– Google

Da unsere fiktive Domain gute Besuchermetriken aufweist und unter anderem auch durch die Werbung als bekannte Marke genügend sogenannten Direct-Traffic bekommt (das ist ein wichtiges und gutes Signal!), erbt unser neues Bleistift-Dokument von den älteren und erfolgreichen Seiten entsprechend positive Signale. Und es passiert. NavBoost gibt uns einen Boost von Platz 14 auf fünf. Wir landen in der Short List, dem „blue ring“, den Top Ten. Diese Liste wird jetzt zusammen mit neun anderen organischen Ergebnissen an das sogenannte GWS weitergereicht.

Das GWS – wo alles ein Ende und einen neuen Anfang findet

Der Google Webserver (GWS) ist für den Zusammenbau und die Auslieferung der Suchergebnisseite (SERP)

verantwortlich. Diese besteht ja nicht nur aus den zehn blauen Links, sondern aus Ads, Bildern, Google-Maps-Ansichten, „Weiteren Fragen“, „andere suchten auch“ oder anderen Elementen. Das alles muss austariert, angepasst und am Ende auch vom Raumbedarf her organisiert werden. Für die geometrische Flächenoptimierung ist Tangram zuständig. Es berechnet, wie viel Platz Elemente brauchen und wie viele Ergebnisse in einzelne „Boxen“ passen würden. Das System Glue „klebt“ dann alles zusammen und an die richtigen Stellen. Ein Teil innerhalb der organischen Ergebnisse ist auf Platz fünf unser Bleistift-Dokument.

Sofern nicht das sogenannte Cookbook dazwischenfunkt! Es enthält Systeme namens FreshnessNode, InstantGlue (reagiert jeweils in Zeiträumen von 24 Stunden mit etwa zehn Minuten Zeitverzug) und InstantNavBoost. Sie können „zur Laufzeit“ weitere Signale erzeugen, die vor allem mit Aktualität zu tun haben. Sie greifen in den sprichwörtlich letzten Hundertstelsekunden ein und können das Ranking abermals verändern.

Gerade läuft im Fernsehen ein Beitrag über 250 Jahre Faber-Castell und die Mythen, die sich um die Entstehung des Worts „Bleistift“ ranken. Tausende Zuschauer greifen innerhalb weniger Minuten zum Smartphone oder Tablet und fangen an, zu googeln. Das ist durchaus nicht ungewöhnlich. FreshnessNode registriert einen starken Anstieg der Suche nach „Bleistift“ und anhand der User-Klicks, dass

dahinter kein Kaufinteresse steht, sondern dass Informationen gesucht und gewollt werden. Genau in diese Ausnahmesituation platzt unser beispielhaft beobachtetes Dokument mit seinem Ranking. InstantNavBoost sorgt dafür, dass alle transaktionalen Ergebnisse eliminiert und in Echtzeit durch informationale Ergebnisse ersetzt werden. InstantGlue verändert den „blue ring“ und unser Dokument, das ja eindeutig als verkaufsorientiert klassifiziert wurde, verschwindet und wird ersetzt.

Pech. Durch dieses herbeifantasierte Ende unserer Ranking-Reise soll vor allem deutlich werden, dass es nicht unbedingt am Dokument, an richtigen oder falschen SEO-Maßnahmen mit wirklich gutem und nützlichem Content liegen muss, wenn es mit dem Ranking nicht klappt. Es gibt viele Umstände, verändertes Suchverhalten (kurzfristig oder langsam auf Dauer) oder neue Signale für andere Dokumente, die das Ranking beeinflussen. Die bisherige Sicht „Ich habe doch ein tolles Dokument und gute Arbeit gemacht“ muss dahingehend erweitert werden.

Die Zusammenstellung der Suchergebnisse ist ein sehr komplexer Prozess mit Tausenden von Einflussmöglichkeiten. Mit Zehntausenden von Tests, die im Livebetrieb vom SearchLab über Twiddler verursacht werden. Es kann schon reichen, dass Dokumente, von denen man Backlinks bekommt, vom HiveMind auf die unwichtigere Ebene der SSDs oder gar in TeraGoogle verschoben werden, die Signale dadurch abgeschwächt werden oder ab jetzt fehlen, sodass sich das feine Zünglein an der Ranking-Waage verändert. Am eigenen Dokument muss sich rein gar nichts geändert haben. Gerade John Müller von Google hat immer wieder nach den Updates der letzten Monate darauf hingewiesen, dass man bei

einem Ranking-Rückgang oft gar nichts falsch gemacht hat. Allein verändertes User-Verhalten oder andere Umstände können dafür sorgen, dass es nicht mehr so klappt wie früher. Mögen Suchende zu einem Thema aktuell ausführlichere Informationen und tendieren im Lauf der Zeit dazu, lieber kürzere Texte zu lesen, reagiert NavBoost ganz automatisch und wertet auf oder ab. Der IR-Score im Alexandria-System oder im Ascorer hat sich deswegen um kein Jota verändert!

Eines der wichtigsten Learnings ist wohl, dass SEO sehr viel breiter betrachtet werden muss. Passen Dokument und Such-Intent nicht wirklich zusammen, nützen weder eine Title- noch eine Content-Optimierung via WDF/IDF.

Die Multiplikatorhebel in Twiddlern und NavBoost können zum Teil einen viel größeren Impact auf das Ranking haben als eine On-Page-, On-Site oder Off-Site-Optimierung. Bremsen diese beiden Systeme, nützen ein paar PS mehr auf der Seite rein gar nichts.

Wir wollen aber mit einem guten Ende aus unserer Reise gehen. Denn die Auswirkungen der Sendung über Bleistifte im Fernsehen sind nur kurzfristiger Natur. Geht der Such-Boom wieder zurück, lässt uns Freshness-Node beim nächsten Suchvorgang ungeschoren und wir tauchen letztlich dann doch auf Platz fünf auf.

Und hier nimmt alles auch wieder einen Anfang durch das Sammeln der Klickdaten. Auf Position fünf wäre eine CTR von etwas über 4 % erwartbar (sofern die Daten aus den diversen Studien einigermaßen richtig sind, hier stammen sie von Johannes Beus von SISTRIX). Können wir die dauerhaft erreichen, dürfen wir jetzt auch mit einem Verbleib in den Top Ten rechnen. Alles ist gut.

Ihre Take-aways

- ☑ Sorgen Sie dafür, dass Sie auch genügend Traffic aus anderen Kanälen und nicht nur aus der Suche bekommen. Auch Traffic aus den vermeintlich unsichtbaren Bereichen wie Social-Media-Plattformen hilft. Über den Chrome-Browser beziehungsweise die URL weiß Google, von woher wie viele zu Ihnen kommen, auch wenn der Crawler die Seiten dieser Plattformen nicht aufrufen kann.
- ☑ Stärken Sie unter allen Umständen Ihre Brand beziehungsweise die Bekanntheit Ihres Domain-Namens. Durch die Wiedererkennung in den Suchergebnissen steigt Ihre Klickrate (sofern man Sie mag). Auch viele Treffer im sogenannten Long Tail (leicht zu ranken) machen Ihre Domain bei Suchenden bekannter. Vieles in den Leaks deutet darauf hin, dass es so etwas wie eine „Site Authority“ als Ranking-Faktor gibt.
- ☑ Versuchen Sie, die Suchabsichten Ihrer Besucher und zumindest in Teilen deren Such-Journey besser zu verstehen. Nutzen Sie Tools wie zum Beispiel Semrush oder SimilarWeb, um zu sehen, von woher Ihre Besucher kommen, vor allem aber wohin sie nach Ihrem Besuch bei Ihnen gehen. Analysieren Sie diese Domains. Haben diese Informationen, die Ihnen (noch) auf Ihren Landingpages fehlen? Ergänzen Sie diese nach und nach, um zur „Endstation“ der Suchreise zu werden. Denken Sie daran, Google speichert jeweils alle zusammenhängenden Such-Sessions und weiß daher genau, was die Suchenden brauchen beziehungsweise wo sie es aufgesucht haben!
- ☑ Kontrollieren Sie Ihre CTR bei den Suchtreffern und optimieren Sie Title und Description zu mehr Klickaffinität.

Angeblich wirken sich bei wenigen wichtigen Wörtern Großbuchstaben positiv auf die CTR aus, weil das optisch hervorsticht. Testen Sie das bei sich.

- ✓ Der Title ist ein wesentliches Entscheidungskriterium, ob Ihr Dokument für eine Suchphrase in den „green ring“ aufgenommen wird. Title-Optimierung ist Chefsache!
- ✓ Wenn Sie sogenannte Akkordeons verwenden, in denen Sie wichtigen Content „verstecken“, der erst aufgeklickt werden muss, prüfen Sie, ob bei diesen Seiten die Bounce-Rate höher ist als der Durchschnitt. Sieht ein Suchender nicht sofort, dass er bei Ihnen richtig ist und muss viel klicken, steigt die Wahrscheinlichkeit für badClick-Signale.
- ✓ Seiten, die niemand aufruft (Webanalytics) oder die über längere Zeiträume kein gutes Ranking erzielen, sollten Sie gegebenenfalls entfernen. Auch schlechte Signale werden auf angrenzende Seiten vererbt! Publizieren Sie ein neues Dokument in einem „schlechten“ Seitencluster, hat die neue Seite wenige Chancen. „deltaPageQuality“ misst offenbar tatsächlich die qualitative Differenz zwischen einzelnen Dokumenten einer Domain oder eines Clusters.
- ✓ Ein klarer Seitenaufbau, eine klare Nutzerführung und eine gute First Impression sind nicht nur „nice to have“, sondern durch NavBoost nicht selten sogar auch überhaupt erst entscheidend für Top-Rankings.
- ✓ Je länger sich Besucher auf Ihrer Site aufhalten, desto besser werden Ihre Domain-Signale, die auf alle Unterseiten abstrahlen. Versuchen Sie, zum Endziel, zu „Destination“ zu werden. Alle Informationen fin-

det man bei Ihnen, noch wo anders weitersuchen zu müssen, sollte nicht passieren.

- ✓ Bauen Sie besser bestehende Dokumente inhaltlich immer weiter aus, aktualisieren Sie diese und machen Sie sie dadurch besser, als ständig neue Dokumente zu erzeugen. Der „ContentEffortScore“ versucht, zu ermitteln, wie viel Aufwand in die Erstellung eines Dokuments eingeflossen ist. Gute Bilder, Videos, Tools, einzigartiger Content zählen auf dieses wichtige Signal ein.
- ✓ Richten Sie (Zwischen-)Überschriften passend auf die folgenden Textblöcke aus. Bei der thematischen Vermessung durch sogenannte Embeddings (Vektorisierung des Texts) werden Übereinstimmungen oder eben falsch verwendete Überschriften sehr viel besser erkannt als bei einem rein lexikalischen Vergleich.
- ✓ Versuchen Sie, über Ihr Webanalytics-Tool wie Google Analytics das Engagement der Besucher sauber zu tracken. Das hilft bei der Beurteilung und Lückensuche. Prüfen Sie die Bounce-Rate Ihrer Landingpages (Ranking-Treffer). Ist sie zu hoch, suchen Sie nach plausiblen Gründen und steuern Sie gegen. Google kennt alle diese Daten durch den Chrome Browser.
- ✓ Backlinks von frischen oder hoch bewerteten Seiten mit viel Traffic, die im Arbeitsspeicher (HiveMind) hinterlegt sind, bringen deutlich mehr gute Signale, Backlinks von Seiten, die niemand besucht oder klickt, wirken dagegen nicht (mehr). Links aus dem gleichen Land helfen mehr. Die thematische Nähe von der den Link gebenden und der den Link empfangenden Seite ist besonders

wichtig für eine gute Linkbewertung. Und offenbar gibt es doch die umstrittenen „toxischen“ Backlinks in dem Sinn, dass sie Ihren Score verringern.

- ✓ Sie können sich auch zunächst auf ein gutes Ranking bei weniger umkämpften Keywords auf ein gutes Ranking konzentrieren und somit einfacher positive Usersignale aufbauen.
- ✓ Wörter vor und nach einem Link werden ebenfalls für das Ranking erfasst, nicht nur der Ankertext. Sorgen Sie für passend „umfließende“ Texte. „Hier klicken“ ist vor über 20 Jahren schon keine gute Idee gewesen.
- ✓ Das Disavow-File zum Entwerten schlechter Links wird im gesamten Leak nicht einmal erwähnt. Offenbar wird es von den Algorithmen nicht beachtet und hat rein dokumentativen Charakter für die Spam-Fighter.
- ✓ Wenn Sie Autorenhinweise einsetzen, sollten diese auch auf anderen Websites auffindbar sein und dort eine fachlich passende Expertise haben. Weniger (aber gute) Autoren sind offenbar besser als möglichst viele. Google kann wohl laut einem Patent Content vom Experten bis hin zum Laien geschrieben klassifizieren.
- ✓ Und zu guter Letzt: Erzeugen Sie exklusiven, hilfreichen, umfassenden und gut strukturieren Content, zumindest bei den wichtigen Seiten. Zeigen Sie, dass Sie wirklich Ahnung haben über das, was Sie schreiben. Im Idealfall können Sie das sogar nachweisen? Einfach Text von jemandem verfassen lassen, damit etwas da steht – kann man machen. Aber dann sollte man die Ranking-Erwartungen nicht allzu hoch setzen!