



Algorithmen und Meinungsbildung

Eine grundlegende Einführung

Katharina A. Zweig · Oliver Deussen
Tobias D. Krafft

Einleitung

Wir stehen unmittelbar vor einer Bundestagswahl und viele befürchten, dass das Internet und die allgegenwärtigen Algorithmen der Suchmaschinen und sozialen Medien auch diese Wahl beeinflussen könnten. Die Sorge ist nicht unbegründet, denn der letzte Präsidentschaftswahlkampf der USA im Herbst 2016 war in vielerlei Hinsicht wegweisend, insbesondere bei der Nutzung der sozialen Medien zur Beeinflussung der Wählermeinung. Dieser neue Weg der Wählergewinnung war schon in den social-media-basierten Kampagnen von Obama angelegt: Es war daher auch für 2016 von vornherein offensichtlich, dass ein großer Teil der Wähler im Netz gewonnen oder verloren werden würde. Nach dem für viele unerwarteten Ausgang der Wahl ist die Sorge nun groß, dass bei der deutschen Bundestagswahl im Herbst 2017 Algorithmen missbraucht werden könnten, um Wählerstimmen zu fangen. In diesem Artikel soll daher eine Grundlage dafür gelegt werden, wie Algorithmen dafür genutzt werden können, um mögliche Wähler über soziale Medien oder Suchmaschinen gezielt anzusprechen. Die hinter diesen Diensten stehenden Algorithmen und ihre prinzipielle Funktionsweise werden deswegen zuerst behandelt. Richtig mächtig werden Empfehlungssysteme aber erst durch die Personalisierung, also die Einbindung von Algorithmen der künstlichen Intelligenz, die bestimmte Verhaltensweisen der Nutzer kategorisieren und danach ihre Empfehlungen aussprechen.

Es gibt nun mehrere Möglichkeiten, wie diese Algorithmen genutzt werden könnten, um die öffentliche Meinung zu manipulieren: (1) Die dahinterstehenden Entwicklerinnen und Entwickler

bei den großen Intermediären nutzen sie, um die Bevölkerung zu manipulieren. (2) Die Mischung aus menschlichen Nutzern und algorithmischen Systemen führt zur Abschottung von Gruppen mit bestimmten politischen Meinungen und (3) die Algorithmen werden von außen manipuliert. Insbesondere der zweite Punkt wird in diesem Artikel behandelt, nämlich die Bildung sogenannter Filterblasen und Echokammern, während die Frage nach der Möglichkeit interner und externer Manipulation durch den Artikel von Krafft und Zweig („Ein Faktencheck – Ließ ein Algorithmus Trump triumphieren?“, in diesem Heft) diskutiert wird. Es zeigt sich, dass das „Agenda Setting“, das in den Fokus Rücken von Themen in die öffentliche Diskussion, heute von mehr Akteuren betrieben werden kann als noch vor einigen Jahren. Dies kann positive und negative Auswirkungen haben, da es auf der einen Seite mehr Meinungsvielfalt zulässt und auf der anderen Seite auch mehr Beeinflussung erlaubt.

In jedem Fall scheint es nötig zu sein, dass demokratische Gesellschaften weiterhin darüber wachen, wer die öffentliche Meinung in welcher Form beeinflussen kann. Insgesamt kann damit festgestellt werden, dass Algorithmen zwar der Motor für viele Veränderungen gesellschaftlicher

DOI 10.1007/s00287-017-1050-5
© Springer-Verlag Berlin Heidelberg 2017

Katharina A. Zweig · Tobias D. Krafft
TU Kaiserslautern, FB Informatik,
Algorithm Accountability Lab,
Gottlieb-Daimler-Str. 48, 67663 Kaiserslautern
E-Mail: zweig@cs.uni-kl.de, krafft@cs.uni-kl.de

Oliver Deussen
Universität Konstanz, Visual Computing,
Gebäude Z, Raum 708, Postfach 698
E-Mail: oliver.deussen@uni-konstanz.de

Zusammenfassung

In diesem Artikel geben wir eine grundlegende Einführung in die algorithmischen Empfehlungssysteme und wie sie – unter Umständen – Filterblasen und Echokammern in sozialen Medien erzeugen könnten. Der Term Filterblase beschreibt dabei das Phänomen, dass wir von Algorithmen hauptsächlich solche Themen wieder vorgeschlagen bekommen, die wir schon mögen. Als Echokammern bezeichnet man Freundesgruppen, die hauptsächlich aus Leuten mit ähnlicher Meinung bestehen, in denen also jede Aussage widerhallt. Auch wenn es noch keine Studien gibt, die wirklich nachweisen, dass Menschen heutzutage durch die Wirkung von Algorithmen tatsächlich in dichteren Filterblasen leben oder in der Bildung von Echokammern bestärkt werden, ist doch klar, dass mit Hilfe dieser Algorithmen unsere Meinungsbildung manipuliert werden könnte. Daher sprechen wir uns für eine sinnvolle Überwachung von solchen Algorithmen aus, um eine solche Manipulation überhaupt entdecken zu können.

Prozesse sind, aber der Algorithmus selbst, als nur ein Teil und zudem der am besten verstandene Teil des algorithmischen Entscheidungssystems, nicht das alleinige Ziel einer Überprüfung durch die Ge-

sellschaft sein kann. Der Artikel schließt daher mit einer Diskussion des Begriffes „Algorithm Accountability“ und der Frage, ob und wie algorithmische Entscheidungssysteme kontrolliert werden sollten.

Algorithmische Empfehlungssysteme

Algorithmische Empfehlungssysteme treffen aus einer großen Menge von Informationen, z. B. in Form von Dokumenten oder Webseiten, eine grundlegende Auswahl und sortieren diese Auswahl dann nach geeigneten Kriterien so, dass möglichst relevante Informationen zuerst angezeigt werden (zur grundlegenden Definition eines Algorithmus s. Abb. 1, zu der des algorithmischen Empfehlungssystems Abb. 2).

Um ein solches Empfehlungssystem aufzubauen, gibt es mehrere Möglichkeiten. Man kann zum Beispiel eine Reihe von Experten befragen, wann welches Dokument in welcher Reihenfolge auszugeben sei. Die von den Experten angegebenen Regeln bilden dann ein sogenanntes Expertensystem [13]. Solche Expertensysteme wurden insbesondere in der Medizin in den Achtziger- und Neunzigerjahren des letzten Jahrhunderts aufgebaut. Für die Verwaltung von Dokumenten – auch den digitalen – wurde lange versucht, sie mit menschlicher Hilfe so zu kategorisieren, wie man es vorher auch mit analogen Dokumenten getan hatte. Das war auch die Strategie von Yahoo!, die damit Mitte der Neunziger als börsennotiertes Unternehmen sehr erfolgreich waren.

Ein **Algorithmus** ist eine definierte Handlungsvorschrift, die für jede mögliche **Eingabe** von Informationen eine **Ausgabe** generiert, die bestimmte Eigenschaften hat.

Beispiel Navigation

Eingabe: Startpunkt und Zielort, Straßenkarten, aktuelle Verkehrsdaten

Ausgabe: eine mögliche Fahrstrecke mit der kürzesten erwarteten Dauer unter Einbeziehung der vorhandenen Daten.

Der sogenannte **Dijkstra-Algorithmus** garantiert die Berechnung einer optimalen Strecke basierend auf den Eingaben.



+ Startpunkt
+ Ziel
+ aktuelle
Verkehrsdaten

=

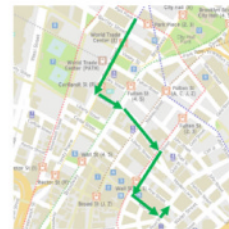


Abb. 1 Definition des Begriffes Algorithmus. Für Informatiker muss die Definition noch ergänzt werden um den Hinweis, dass eine definierte Handlungsanweisung nur dann ein Algorithmus ist, wenn ihre Ausführung nach endlicher Zeit zu einer Lösung führt.

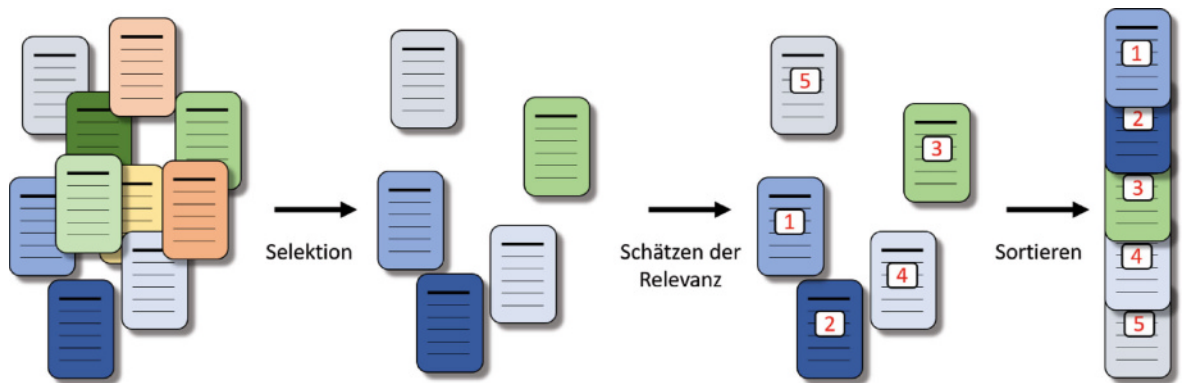


Abb. 2 Algorithmische Entscheidungssysteme entnehmen zuerst eine Teilmenge von Dokumenten, Informationen oder Produkten aus einer großen Datenbank (Selektion). Dann schätzen sie die allgemeine oder individuelle Relevanz der Teilmenge an Dokumenten (oder Informationen oder Produkten). Danach werden die Elemente dann sortiert und dem Nutzer präsentiert. Durch die Selektion und Sortierung werden wichtige Akzente gesetzt, die die Nutzer beeinflussen können.

Google PageRank: Die geniale Idee der Google-Gründer Sergei Brin und Larry Page aber machte Suchmaschinen erst wirklich erfolgreich [1]: die Idee, dass Webseiten von ihren Designern nicht willkürlich miteinander verlinkt werden, sondern meistens eine Webseite auf eine ähnliche verweist, die der Designer für sinnvoll hält als Ergänzung zu seiner eigenen. Ihre Idee war daher, diese Information zu nutzen (Abb. 3). Sie stellten sich einen zufälligen Webnutzer vor (Random Surfer Modell), der sich von Webseite zu Webseite über die Links hangelt und von Zeit zu Zeit auf eine völlig zufällig gewählte Webseite springt. Für jede Webseite maßen

sie, wie lange sich dieser zufällige Webnutzer hier finden würde, und nahmen diesen Wert als Qualitätsmaß an. Glücklicherweise muss man für die Berechnung dieses Maßes nicht wirklich einen Zufallsnutzer simulieren – es gibt eine mathematische Gleichung, mit der man bestimmen kann, wie lange der Nutzer wo sein würde.

An diesem Algorithmus sieht man übrigens auch schon, dass Algorithmen trotz ihrer vermeintlichen Objektivität durchaus subjektive Annahmen beinhalten können. In diesem Fall nahmen Brin und Page an, dass eine Verlinkung von einer Webseite auf die andere ein Qualitätsurteil der Webdesignerin

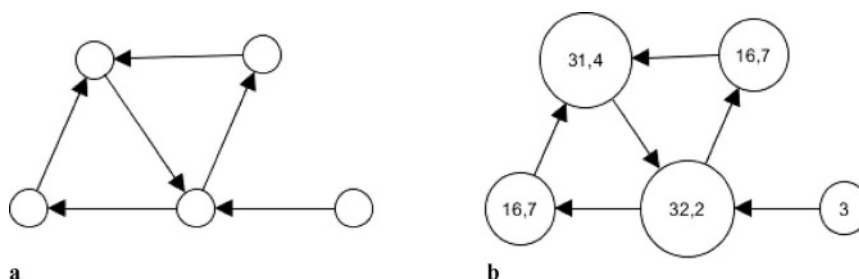


Abb. 3 Berechnung des PageRanks: (a) Ein kleiner, künstlicher Webgraph mit fünf „Webseiten“, die aufeinander verlinken. Der Zufallssurfer springt mit einer Wahrscheinlichkeit von 15 % auf irgendeine der fünf „Webseiten“ – jede einzelne wird also in jedem Schritt mit einer Wahrscheinlichkeit von 3 % angesprochen. Mit 85 % springt der Surfer von der Webseite, auf der er gerade ist, zu einer der verlinkten Webseiten, jeweils mit derselben Wahrscheinlichkeit. Die Größen in (b) sind proportional zum Anteil der Zeit, die sich der Surfer auf der jeweiligen Webseite aufhält, bzw. der Wahrscheinlichkeit, dass er gerade dort ist. Die ganz rechts unten stehende Webseite kann nur durch Zufall erreicht werden, also mit einer Wahrscheinlichkeit von 3 %. Die restlichen Prozentzahlen setzen sich jeweils aus der Wahrscheinlichkeit zusammen, dass der Surfer hier zufällig hingespungen ist (also 3 %), plus der Summe der Wahrscheinlichkeiten, dass der Surfer vorher auf einer benachbarten Webseite war, von der aus er die Seite erreichen kann. Die obere linke Webseite z. B. kann von den zwei Webseiten links unten und rechts oben mit einer Wahrscheinlichkeit von $2 \times 85\% \times 16,7\%$ erreicht werden und direkt angesprochen werden mit einer Wahrscheinlichkeit von 3 %. Das ergibt eine Gesamtwahrscheinlichkeit von 31,4 %. Die anderen Zahlen können entsprechend errechnet werden.

beinhaltet und dass die Links voneinander unabhängig auf andere Webseiten gesetzt werden. Diese Annahmen mögen in der Mitte der 1990er-Jahre noch gegolten haben – nach Veröffentlichung des Algorithmus dagegen galten sie schnell nicht mehr: Es entwickelte sich der Beruf des Suchmaschinenoptimierers („search engine optimization“ oder kurz SEO). Und davon gibt es gutartige und böartige Varianten („white hat“ und „black hat“). Die böartige Variante ließ sich sogenannte Linkfarmen einfallen, das sind Webseiten, die keinen anderen Daseinsgrund haben, als möglichst viele Links auf Webseiten zu enthalten, die dafür bezahlen. Andere überredeten besonders vertrauenswürdige Quellen, wie beispielsweise Webseiten von amerikanischen Universitäten, Links auf ihre weniger vertrauenswürdigen Webseiten zu verlinken. Google hat eine ganze Hilfeseite, was sie alles zu den „black hat“-Varianten der Suchmaschinenoptimierung zählen [3]. Webseiten, die solche Tricks enthalten, werden von Google abgestraft und bekommen sehr schlechte Positionen oder werden ganz aus der Ergebnisanzeige entfernt. Dabei sind es natürlich gerade die ersten zehn bis zwanzig Suchmaschinenergebnisse, die besonders wichtig sind (vgl. [2]).

Das Beispiel von Googles PageRank zeigt auch, dass die Offenlegung der Wirkweise eines algorithmischen Empfehlungssystems nicht immer sinnvoll ist, und zwar weder für die Firma noch für die Gesellschaft als Ganzes. Bei Suchmaschinen und auch bei anderen algorithmischen Empfehlungssystemen, wie z. B. Produktempfehlungen oder Facebooks „Newsfeed“, wird jede Form von Information immer auch zu gezielten Manipulationen führen. Bis heute ist Google trotzdem relativ transparent gegenüber den Suchmaschinenoptimierern (siehe die Webmaster Guidelines von Google [3]) und spielt daher ein Hase-und-Igel-Spiel mit den böartigen Manipulierern der Systeme. Daher ist der Page-Rank einer Webseite aber heute nur noch eines von ca. 200 Eigenschaften, nach denen die möglichen Suchergebnisse sortiert werden [4].

Insbesondere können heute ganz andere Informationen verarbeitet werden, die Hinweise darauf geben, welche Inhalte für uns als Nutzerinnen und Nutzer interessant sind: Dazu beobachtet man die menschlichen Nutzer und lernt mithilfe von Algorithmen, welche Art von Informationen sie suchen. Es wird gespeichert, auf welche Links einer Suchergebnisanzeige sie letztendlich klicken, wie lange

sie auf dieser Seite bleiben, ob danach andere Ergebnisse angeklickt werden oder weitere Suchen stattfinden. Auf der anderen Seite kann man auch Informationen über Inhalte effizient nutzen: Dazu gehören zum Beispiel der Autor oder die Autorin einer Nachricht, die Textlänge, wann sie erstellt wurde oder wie viele andere Nutzer damit interagiert haben. Algorithmen des sogenannten maschinellen Lernens können daraus Entscheidungsstrukturen aufbauen, die verschiedene Inhalte in Gruppen bündeln. Eine Methode dafür ist das sogenannte Clustering. Hier werden die Daten daraufhin untersucht, ob es Dokumente gibt, die in ähnlicher Weise, zum Beispiel bei der Suche nach wortverwandten Themen, gemeinsam angeklickt werden. Wenn dann ein neuer Nutzer mit einer ähnlichen Anfrage kommt, werden Dokumente dieser Kategorie bevorzugt angezeigt. Eine Anwendung dieser Art von Algorithmen findet sich in der Analyse von Papakyriakopoulos et al. zum Thema „Social Media und Microtargeting in Deutschland“ in diesem Heft.

Eine heute sehr gebräuchliche Entscheidungsstruktur ist der sogenannte Entscheidungsbaum. Entscheidungsbäume wurden auch in den zuerst erwähnten Expertensystemen benutzt, um Regeln miteinander zu verbinden. Die Algorithmen maschinellen Lernens können aber Entscheidungsbäume aus Daten mit einer zu kategorisierenden Eigenschaft selbstständig lernen. In unserem Beispiel können sie lernen, welche Webseiten tendenziell eher angeklickt werden und welche nicht (Abb. 4a). Dazu nehmen sie ebenfalls die schon beschriebene Datenmenge und durchsuchen sie nach derjenigen Eigenschaft, die die Webseiten am besten in angeklickte und nichtangeklickte Webseiten unterteilt. In der vereinfachten Abb. 4 sind Datenpunkte grün, die eine angeklickte Nachricht repräsentieren, und die anderen rot. Jede Nachricht wird mit nur zwei Eigenschaften beschrieben: der Länge in Worten und der Anzahl der Stunden seit ihrem Erscheinen. Der Algorithmus testet nun beide Eigenschaften darauf, welche besser in rote und grüne Datenpunkte unterteilt, und findet heraus, dass eine gute Trennung erfolgt, wenn man zuerst danach fragt, ob die Nachricht weniger als 3,5 Stunden alt ist (Abb. 4b). Aber auch unter den etwas älteren Nachrichten werden immer noch diejenigen angeklickt, die nicht älter als 6 Stunden sind und weniger als 200 Worte enthalten.

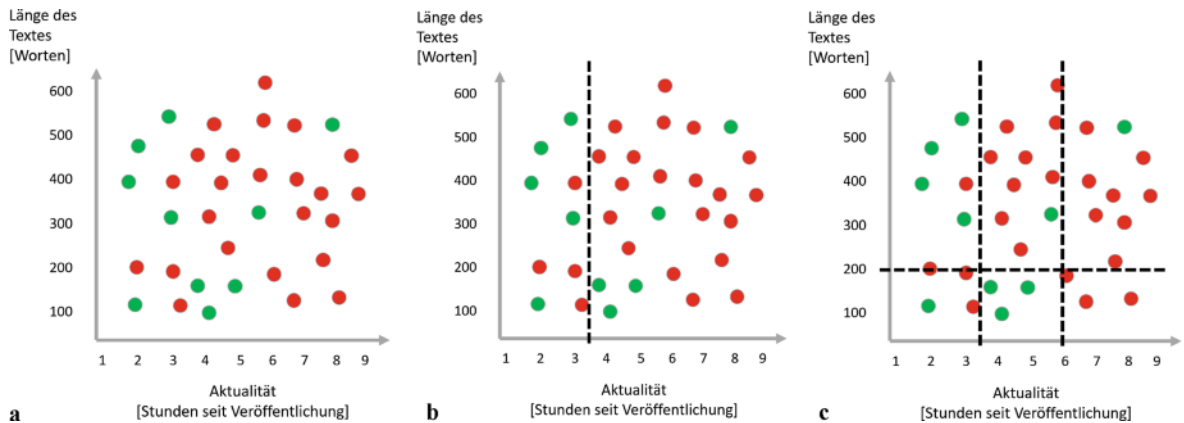


Abb. 4 Erstellen eines Entscheidungsbaumes: (a) Die Punkte stellen Nachrichten dar, die durch nur zwei Eigenschaften beschrieben werden (Stunden seit Veröffentlichung und Länge des Textes in Worten). (b) Ein Algorithmus kann automatisch feststellen, dass eine erste Trennung am besten bei einem Schwellwert von ca. 3,5 Stunden seit Veröffentlichung liegt. (c) Wenn die Veröffentlichung nicht länger als 6 Stunden her ist und der Text nicht mehr als 200 Worte enthält, ist die Wahrscheinlichkeit ebenfalls wieder hoch, dass die Nachricht angeklickt wird. Auch diese Schwellwerte können automatisch bestimmt werden.

Der Algorithmus funktioniert auch mit vielen Dimensionen, also vielen Eigenschaften, die den jeweiligen Datenpunkt beschreiben. In einem Entscheidungsbaum werden nach der ersten Entscheidung immer weitere Eigenschaften gesucht, die die Inhalte jeweils immer besser in angeklickte und nichtangeklickte Inhalte unterscheiden. Ein fiktives Resultat basierend auf Abb. 4 wird in Abb. 5 gezeigt.

Basierend auf einem solchen Entscheidungsbaum können nun neue Dokumente einsortiert werden, ob sie aufgrund der bisher gefundenen Regelmäßigkeit eher zu den angeklickten oder den nichtangeklickten Inhalten zählen werden. Ein solcher Entscheidungsbaum kann auch ganz personalisiert für individuelle Nutzer aufgebaut werden. Beispielsweise könnte ein soziales Medium lernen, dass manche Nutzer lieber Nachrichten von Freunden lesen, während andere sich gerne über das aktuelle Tagesgeschehen informieren lassen, oder individuelle Tagesgewohnheiten lernen. Bei Suchmaschinen könnten Informationen darüber gesammelt werden, welche Publikationsquellen besonders oft angeklickt werden, ob jemand eher längere Texte mag oder eher kürzere Texte oder gar gleich am liebsten wissenschaftliche Texte liest.

Am Beispiel wird auch klar, dass auch Algorithmen des maschinellen Lernens in den meisten Fällen keine perfekten Entscheidungsregeln finden können. In Abb. 5 ist zu erkennen, dass die Trefferquote, also der Anteil der korrekten Resultate

innerhalb der Gruppe, nur einmal einen Wert von 1 aufweist (zweite Gruppe von links), was einer perfekten Differenzierung entspricht. Würde man nun lediglich diese Gruppe als interessant ansehen, so würde man jedoch 6 der 10 grünen Datenpunkte übersehen, weshalb auch eine Betrachtung der Genauigkeit notwendig ist. Die Genauigkeit einer Gruppe („precision“) ist definiert als der Anteil der grünen Datenpunkte, die in dieser Gruppe enthalten sind. Im vorliegenden Beispiel wurden für beide Gütekriterien Schwellenwerte festgelegt, ab welchen angenommen wird, dass die Entscheidungen im Baum mit einer akzeptablen Wahrscheinlichkeit zu einer Klassifizierung eines grünen Elementes führen. Zum einen sollte eine Trefferquote von mehr als 50 % vorliegen, also dass Elemente, welche dem Pfad im Entscheidungsbaum entsprechend eingeteilt wurden, zu mehr als der Hälfte interessant sind. Zum anderen bedarf es einer Genauigkeit von mehr als 25 %, um so zumindest ein Viertel aller interessanten Artikel auf einmal zu erfassen. Sonst könnte das Ergebnis gegebenenfalls durch sehr kleine Klassen verschlechtert werden. Somit kategorisiert der Entscheidungsbaum die linken beiden Gruppen als interessant und die rechten beiden als uninteressant, obgleich ersichtlich ist, dass die Unterteilung nicht perfekt ist.

Die Bestimmung der perfekten Schwellwerte ist mit zunehmender Anzahl der Dimensionen sehr komplex, jedoch ist es in den meisten Fällen möglich, eine Balance zwischen Genauigkeit und Treffer-

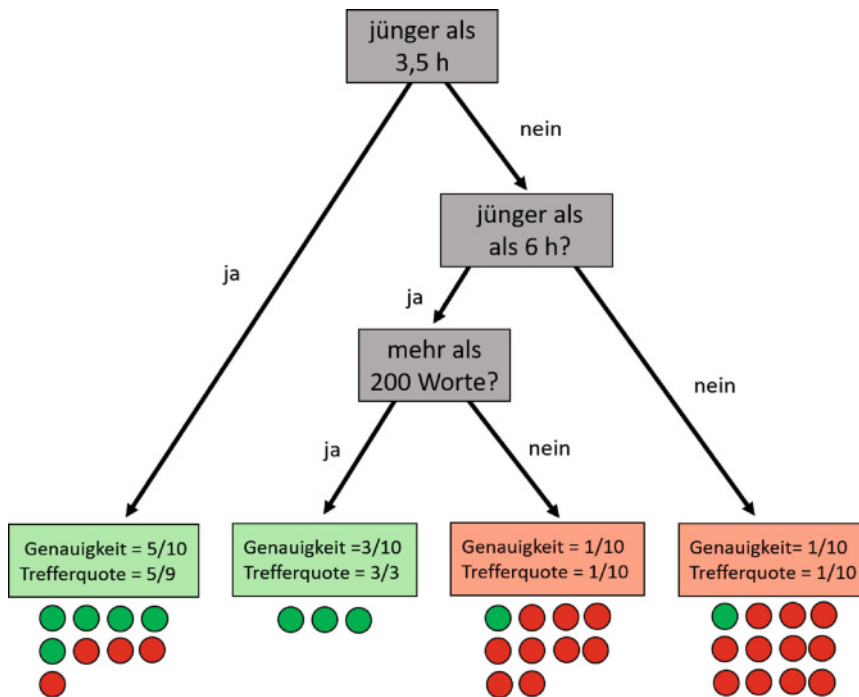


Abb. 5 Entscheidungsbaum resultierend aus den Schritten in Abb. 4. Für alle 4 resultierenden Gruppen sind die Genauigkeit (Precision) und Trefferquote (Recall) notiert. Diese aus dem Information Retrieval stammenden Qualitätsmaße (vgl. [7, S. 269]) setzen den Anteil der gesuchten Elemente, in diesem Fall die grünen, also interessanten Artikel, zu verschiedenen Größen ins Verhältnis. Im Beispiel unten links werden die 5 grünen Datenpunkte zu der Anzahl aller Elemente, auf welche die Eigenschaft „jünger als 3,5 h“ zutrifft (9), in Relation gesetzt und man erhält die Genauigkeit 5/9. Wenn in derselben Gruppe die grünen Punkte, also die Anzahl interessanter Artikel (5) im Vergleich zur Anzahl aller vorhandenen grünen Datenpunkte (10), betrachtet werden, ergibt sich die sogenannte Trefferquote mit 5/10.

quote festzulegen. Besonders wichtig ist es noch zu bemerken, dass Algorithmen des maschinellen Lernens oftmals schon vielfach implementiert wurden und meistens auf relativ einfachen statistischen Methoden beruhen. Man kann davon ausgehen, dass sie an sich nicht (mehr) fehlerhaft und für sich genommen auch objektiv sind. Trotzdem gibt es neben der gerade genannten Frage nach der Genauigkeit und Trefferquote viele weitere Modellierungsentscheidungen, die Designerinnen und Designer eines algorithmischen Entscheidungssystems treffen müssen: beispielsweise die genaue Auswahl des Kriteriums, nach dem kategorisiert werden soll. Als Menschen müssen wir uns schon fragen, ob die „Relevanz“ einer Nachricht wirklich am besten daran gemessen wird, wie oft sie angeklickt wird! Gerade im zurückliegenden Wahlkampf wurde immer wieder festgestellt, dass die Nachrichten mit der schlechtesten Qualität sehr viel mehr Aufmerksamkeit bekamen als diejenigen mit hoher Qualität [11, 12].

Filterblasen und Echokammern

Gerade diese Konzentration der Empfehlungssysteme auf die angeklickten Inhalte, zusammen mit der Personalisierung der Systeme und unserer Neigung zu homogenen sozialen Freundeskreisen, ist es, die manchen Menschen Sorgen bereitet. Was passiert mit einem sozialen Medium für einen Nutzer X, wenn dieser sich immer wieder Nachrichten eines bestimmten Typs ansieht? Nehmen wir dazu an, dass er Fan von mehreren Clubs und Vereinen unterschiedlicher Sportarten ist, verschiedene Sportmagazine liest und selbst auf der Onlineversion seiner Tageszeitung die Sportberichte bevorzugt. Es ist zu vermuten, dass bei einem starken Grad an Personalisierung dieser Nutzer im Laufe der Zeit nur noch Nachrichten bekommt, die vom Sport handeln. In einem klassischen Medium wie dem Fernsehen, dem Radio oder der Tageszeitung wäre er auch anderen Inhalten ausgesetzt, auch wenn er seine Aufmerksamkeit hauptsächlich den Sportinhalten widmen würde. Übertragen auf politische Meinungen ist

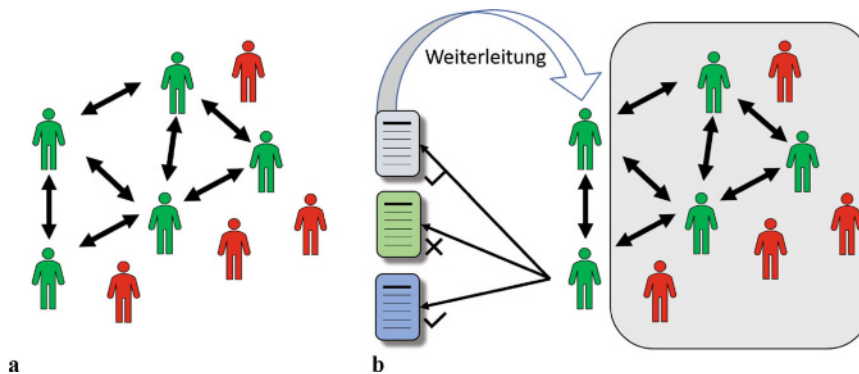


Abb. 6 Echokammern und ihre Gefahren: a) Echokammer, in der sich Mitglieder (z. B. Freunde, Vereinsmitglieder) gegenseitig in ihren Meinungen stützen und bestätigen. Farben rot und grün repräsentieren hier unterschiedliche Haltungen. Im Digitalen können Algorithmen lernen, welche Eigenschaften die Freunde einer Person aufweisen, um personalisiert vorgefiltert Personen vorzuschlagen, die diesen Kriterien entsprechen. Somit transferiert sich dieser Effekt noch verstärkt in die digitale Welt. b) Empfehlungssysteme lernen über Interaktionsverhalten der Freunde mit Webseiten, welche Artikel einer Person gefallen können. Die Algorithmen zum Vorschlagsystem für Freunde im Netzwerk sowie die Eigenschaft der Empfehlungssysteme, einer Person ähnliche Artikel vorzuschlagen, die seinem Freundeskreis gefallen, kann zu Filterblasen führen, in der soziale wie auch politische Extreme überhandnehmen können.

somit die Sorge, dass jemand, der sich mit einer extremen politischen Meinung lange beschäftigt, auch nur noch solche Inhalte bekommt, die diese extreme Meinung unterstützen. Die Sorge ist also, dass seine Meinungsbildung durch den starken Algorithmus im Filter nicht mehr so frei ist wie in klassischen Medien – das ist die klassische Filterblase, wie sie von Eli Pariser erstmalig beschrieben wurde [10].

Dazu kommt auf sozialen Medien ein weiterer Effekt (s. Abb. 6): Die Soziologen wissen schon lange, dass wir Menschen uns vor allen Dingen mit solchen Menschen wohlfühlen, die uns ähnlich sind. Normalerweise suchen wir uns also Personen, die uns im Alter, Bildungs- und Familienstand ähnlich sind und möglichst auch eine ähnliche politische Meinung haben. Da die Empfehlungssysteme der sozialen Medien lernen, mit welchen Personen wir uns gerne umgeben, versuchen sie uns weitere ähnliche Bekannte und Freunde vorzuschlagen. Über diesen zweiten Aspekt bekommen wir also verstärkt Nachrichten und Informationen aus demselben Spektrum von Meinungen, die wir schon unterstützen oder die sich mit unseren gut ergänzen.

Indem die Empfehlungssysteme insofern dem menschlichen Bestreben, sich mit Gleichgesinnten zu umgeben und auszutauschen, entsprechen, kommt es zu einer einseitigen Verlagerung der Realität. Auch wenn man sich im realen Leben, also in Clubs, Vereinen oder am Stammtisch, bevorzugt mit Menschen umgibt, die ähnliche Standpunkte

vertreten wie man selbst, lässt es sich an diesen realen Orten nicht immer vermeiden, auch mit Andersdenkenden in Kontakt zu kommen. Insbesondere bei sehr extremen politischen Standpunkten ist eine Konfrontation mit Andersdenkenden wahrscheinlich.

Im Gegensatz dazu bewegt man sich im Internet in geschlossenen Echokammern von Personen, die unsere Meinung teilen, sodass es zu einer verzerrten Realitätssicht kommen kann: Durch die permanente Bestätigung der eigenen Meinung entsteht der Eindruck, man gehöre zu einer Mehrheit, was im Falle extremer Meinungen bedenkliche Konsequenzen haben kann.

Diese beiden Phänomene, die Filterblase und die Echokammern, sind immer noch nicht hinreichend untersucht. Es gibt einige Studien, die eher darauf hinweisen, dass soziale Medien einen polarisierenden Effekt haben [5]. Davon unabhängig sind Studien, die fragen, welchen Anteil die Internetnutzung überhaupt an unserer Meinungsbildung hat. Dazu finden Sie auch in diesem Heft einen Artikel von Jan-Hendrik Schmidt, der zeigt, dass der Medienmix in Deutschland immer noch recht stabil ist und dass der Anteil der Nutzer, die ihre Nachrichten über soziale Medien bekommen, sehr gering ist. In Deutschland werden wir also weiterhin mit einer hohen Meinungsvielfalt konfrontiert, wenn wir öffentliches Radio hören, öffentliches Fernsehen sehen oder eine Tageszeitung lesen.

Überwachung von Algorithmen und algorithmischen Entscheidungssystemen

Trotzdem ist nicht zu leugnen, dass algorithmische Empfehlungssysteme von sozialen Netzwerken oder Suchmaschinen uns ganz gezielt manipulieren könnten. Es ist rein technisch denkbar, dass

ein Suchmaschinenanbieter seine eigene politische Präferenz in ein solches System einfließen lässt. Dieser Vorwurf wurde beispielsweise Google im letzten Präsidentschaftswahlkampf gemacht, als SourceFed behauptete, dass die Eingabe von „Hillary Clinton“ keine negativen Suchvervollständigungen erzeugen würde. Der Beweis lag darin, dass andere Suchmaschinen wie Bing oder Yahoo beispielsweise mit „indictment“ ergänzten („Anklage“). Es wäre natürlich technisch machbar, ganz explizit dem Algorithmus zu sagen, dass, wann immer jemand das Wort „Clinton“ eingibt, keinerlei Vervollständigung angegeben wird, die eine Liste von negativen Wörtern enthält. Den Faktencheck, ob eine solche interne Manipulation von Algorithmen im letzten Präsidentschaftswahlkampf wahrscheinlich ist, finden Sie im Artikel von Krafft und Zweig in diesem Heft. Die Autoren kommen zu dem Schluss, dass eine solche Manipulation durch die Intermediäre selbst nicht sehr wahrscheinlich ist. Da solche Manipulationen aber ohne Frage denkbar sind und in einer anderen politischen Situation vielleicht sogar gewinnbringend wären, ist es unausweichlich, dass die Gesellschaft beständig darauf prüft, ob diese Algorithmen versuchen, unsere Meinungsbildung zu manipulieren. Nick Diakopoulos hat dafür den Begriff der *Algorithm Accountability* geprägt – eine deutsche Zusammenfassung seines grundlegenden Artikels finden Sie ebenfalls in diesem Heft. Dabei geht es auch um die Frage, inwieweit Algorithmen von Dritten genutzt werden können, um die Gesellschaft zu manipulieren. Zu den strittigen Fragen gehört zum Beispiel die Nutzung des Microtargetings, also die algorithmische Bestimmung von Adressatengruppen für gezielte, feingranulare politische Werbung. Dass dies auch in Deutschland mit den öffentlich verfügbaren Daten grundsätzlich möglich ist, zeigen Papakriakopoulos et al. in ihrem Artikel „Social Media und Microtargeting in Deutschland“ in diesem Heft.

Dieser Begriff beschreibt, dass algorithmische Entscheidungssysteme so überwachbar gemacht werden müssen, dass ihre Nebeneffekte sichtbar

werden und auch einer natürlichen oder juristischen Person zurechenbar werden. Während wir für klassische Medien wie den Rundfunk oder das Fernsehen zahlreiche Kontrollen und Regulation eingeführt haben, ist dies für die sogenannten Intermediäre – also die Suchmaschinen, Videoplattformen wie YouTube oder soziale Medien – bisher nicht der Fall. Warum das so ist, erklären Wolfgang Schulz und Kevin Dankert in ihrem Artikel in diesem Heft. Und während solche algorithmischen Entscheidungssysteme, die unsere Nachrichten filtern und sie für uns vorsortieren, einen wichtigen Dreh- und Angelpunkt politischer Meinungsbildung darstellen könnten, gibt es noch ganz andere algorithmische Entscheidungs- und Empfehlungssysteme, die in gesellschaftlich hochrelevante Prozesse eingreifen. Obwohl wir in Deutschland davon noch nicht viele in der Anwendung sehen, können wir uns insbesondere in den USA davon überzeugen, wie viele menschliche Entscheidungen solche algorithmischen Empfehlungssysteme in naher Zukunft übernehmen könnten. **Da gibt es Algorithmen, die Lehrer automatisch in ihrer Leistung bewerten [8], die Terroristen identifizieren sollen [14], oder solche, die die Rückfallwahrscheinlichkeit von Kriminellen vorhersagen sollen. Auch in Deutschland sehen wir Algorithmen, die die Kreditwürdigkeit einer Person bestimmen sollen, und die ersten Firmen setzen Algorithmen der künstlichen Intelligenz ein, um Bewerbungen vorzuselektieren: Nur die erfolgreichen Kandidaten werden dann zu einem Gespräch eingeladen. In China soll es bald für alle Bürger den sogenannten Citizen Score geben [6]. Dieser bewertet, wie regierungstreu ein Bürger ist. Dazu zählt auch die Regierungstreue seiner Freunde und Verwandten auf Weibo, dem sozialen Medium in China. Ein solcher Algorithmus wird große Auswirkungen auf das soziale Geflecht haben, denn wenn die eigene wirtschaftliche Zukunft davon abhängt, welcher Meinung meine Freunde auf dem digital gut überwachten sozialen Netzwerk Weibo haben, dann wird es wohl nur wenige geben, die diese Freundschaft nicht offiziell kündigen werden.** Diese Melange aus privatrechtlichen und staatlichen algorithmischen Entscheidungssystemen erzeugt bei vielen den Wunsch nach einer unabhängigen Institution, die solche Entscheidungssysteme kontrolliert. Viktor Mayer-Schönberger und Kenneth Cukier haben in ihrem Buch *Big Data* dazu den sogenann-

ten Algorithmen-TÜV vorgeschlagen [9]. Für eine mögliche Kontrolle und Regulierung dieser (und anderer) Algorithmen setzen sich verschiedene Institutionen ein, die in einem journalistischen Beitrag von Anna Loll („Akteure im Bereich Informatik und Gesellschaft“) vorgestellt werden. Eine davon, die Landesmedienanstalt des Saarlands (LMS) hat uns ein Interview mit ihrem Direktor, Uwe Conradt, zur Verfügung gestellt, in der er seine Sicht auf die Notwendigkeit und Möglichkeit der Kontrolle von Algorithmen darstellt. Abschließend behandelt ein Artikel von Agata Królikowski und Jens-Martin Loebel aus der GI-Fachgruppe „Informatik und Gesellschaft“ die besondere Bedeutung von Fake-News.

In jedem Fall ist es wichtig, dass sich die Gesellschaft darüber klar wird, welche Chancen algorithmische Entscheidungssysteme bieten. In den USA wird vor allen Dingen darauf abgehoben, dass sie konsistent in ihren Entscheidungen sind und so gebaut werden können, dass sie objektive Entscheidungen treffen, die sich weder vom Geschlecht noch von der Hautfarbe beeinflussen lassen. Auf der anderen Seite könnten algorithmische Entscheidungssysteme den Eindruck machen, deutlich fehlerfreier zu arbeiten, als sie es in Wirklichkeit tun. **Denn bei vielen wichtigen gesellschaftlichen Fragen gibt es einfach keinen klaren Regelsatz, der alle Menschen in klare Kategorien sortieren würde. Da sich aber der Computer nicht verrechnet, könnte eine Einstufung eines solchen Systems sehr viel ernster genommen werden, als es die Datenlage eigentlich zulässt. Es handelt sich hierbei also um eine Verwechslung der Fehlerfreiheit eines Systems in Bezug auf seine mathematischen Komponenten und der Interpretierbarkeit seiner Ergebnisse.** Auch in diesem Sinne ist eine Überwachung der Qualität von algorithmischen Entscheidungssystemen

notwendig und gleichzeitig eine Schulung der Gesellschaft darüber, was diese Systeme leisten können und was sie nicht leisten können. Über ihre Rolle in der Meinungsbildung, insbesondere der politischen Meinungsbildung, informiert Sie dieses Heft. Wir freuen uns über Ihre Kommentare, Anregungen und Wünsche über weitere algorithmische Entscheidungssysteme und eine Analyse ihrer Qualität, um dieses Thema weiter mit Ihnen zu diskutieren.

Literatur

1. Brin S, Page L (1998) The anatomy of a large-scale hypertextual web search engine. *Computer networks and ISDN systems* 30:107–117
2. Epstein R, Robertson RE (2015) The search engine manipulation effect (SEME) and its possible impact on the outcomes of elections. *P Natl Acad Sci* 33:E4512–E4521
3. Google (2017a) Webspam content violations. <https://support.google.com/webmasters/answer/35769?hl=en>, letzter Zugriff: 1.5.2017
4. Google (2017b) Alles über die Suche – Algorithmen. <https://www.google.com/insidesearch/howsearchworks/algorithms.html?hl=de>, letzter Zugriff: 1.5.2017
5. Hagen L, In der Au A, Wieland M (2017) Polarisierung im Social Web und der intervenierende Effekt von Bildung. Eine Untersuchung zu den Folgen algorithmischer Medien am Beispiel der Zustimmung zu Merkels „Wir schaffen das!“. Sonderausgabe „Algorithmen, Kommunikation und Gesellschaft“ In: kommunikation @ gesellschaft 18. 20 pages.
6. Hanfeld M (10.10.2015) Punkte für gefälliges Verhalten. FAZ, <http://www.faz.net/video/medien/punktrichter-citizen-score-ueberwachung-in-china-13848403.html>, letzter Zugriff: 29.5.2017
7. Manning CD, Schütze H et al. (1999) *Foundations of statistical natural language processing*. MIT Press, London
8. O'Neil C (2016) *Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy*. Crown Publishing Group, New York
9. Mayer-Schönberger V, Cukier K (2013) *Big Data: Die Revolution, die unser Leben verändern wird*. Redline Verlag, München
10. Pariser E (2012) *Filter Bubble: Wie wir im Internet entmündigt werden*. Carl Hanser Verlag, München
11. Silverman C (2016) This analysis shows how viral fake election news stories outperformed real news on Facebook. Buzzfeed, 16.11.2016, <https://www.buzzfeed.com/craigsilverman/viral-fake-election-news-outperformed-real-news-on-facebook>, letzter Zugriff: 1.5.2017
12. Silverman C, Strapagiel L, Shaban H, Hall E, Singer-Vine J (2016) Hyperpartisan Facebook pages are publishing false and misleading information at an alarming rate. Buzzfeed, 20.10.2016, <https://www.buzzfeed.com/craigsilverman/partisan-fb-pages-analysis>, letzter Zugriff: 1.5.2017
13. Spreckelsen C, Spitzer K (2008) *Wissensbasen und Expertensysteme in der Medizin*. Vieweg + Teubner, GWV Fachverlage GmbH, Wiesbaden
14. The Intercept (2015) SKYNET: courier detection via machine learning. <https://theintercept.com/document/2015/05/08/skynet-courier/>, letzter Zugriff: 3.5.2017