



Christina Elmer ist Professorin für Digital Journalism/Data Journalism an der TU Dortmund. Zuvor war sie Mitglied der Chefredaktion von „Spiegel online“ und Leiterin des Ressorts Datenjournalismus.

Assistenten im Digitalen: Wo lernende Algorithmen schon heute Redaktionen unterstützen.

Von Christina Elmer

Sie erschließen Archive und transkribieren Interviews, übersetzen Texte und selektieren Empfehlungen. In redaktionellen Workflows finden sich bereits heute an vielen Stellen lernende Algorithmen – von der Recherche über die Produktion bis zur Distribution digitaler Inhalte. Sie werden eingesetzt, wo sich Abläufe einerseits automatisieren lassen und zudem Systeme benötigt werden, die aus der Analyse großer Datenmengen Regeln ableiten und auf neue Inputs anwenden können. Streng genommen reicht das Feld der künstlichen Intelligenz über lernende Algorithmen hinaus und umfasst ebenso die Automatisierung einfacherer Aufgaben, für die eigentlich menschliche Intelligenz notwendig wäre. Dennoch sollen in diesem Beitrag lernende Systeme im Fokus stehen. Denn verglichen mit der Automatisierung klar definierter Prozesse bergen sie ungleich größere Potentiale wie auch Herausforderungen: Für den Einsatz im Journalismus müssen sie angepasst und idealerweise spezifisch trainiert werden. Ihre Verwendung erfordert ein tiefes Verständnis der Technologie und ihrer Grenzen. Und schließlich gilt es, diese Komplexität der Leserschaft gegenüber transparent zu machen, ohne das Vertrauen in die Redaktion aufs Spiel zu setzen.

Automatisierung und Verantwortung

Denn noch immer steht der über Leichen gehende Terminator einer förderlichen Kommunikation zu künstlicher Intelligenz im Weg, stellvertretend für die häufige Assoziation mit humanoiden, autark agierenden Robotern. So ist es nicht verwunderlich, dass sich Redaktionen in Projekten mitunter klar von lernenden Algorithmen abgrenzen: „Mit künstlicher Intelligenz hat das übrigens nichts zu tun: Die Texte werden regelbasiert erzeugt“, erklärte „Der Spiegel“ in einem Hintergrund über automatisiert generierte Artikel zu mehreren Landtagswahlen (vgl. Pauly 2021). Basierend auf den offiziellen Wahlergebnissen beschreiben die Artikel jeweils die Resultate für alle Wahlkreise und visualisieren diese in ebenfalls datenbasiert produzierten Infografiken. In diesem Kernbereich der politischen Berichterstattung Systeme einzusetzen, die derartige Berichte frei formulieren, erschien zum aktuellen Entwicklungsstand

unverantwortlich – zu groß das Risiko einer publizistisch und presserechtlich schädlichen Fehlinterpretation der Daten.

Diese Abwägung ist auch in den folgenden Anwendungsbereichen grundlegend. Und doch zeigen die Beispiele eine breite Varianz sehr spezifischer Systeme, die Journalist:innen nicht ersetzen, aber doch wirkungsvoll unterstützen können.

Recherche: Quellen, Investigation, Verifikation

Terrabyte an Daten, ob archiviert oder in Echtzeit: Ohne Support lassen sich die heute verfügbaren digitalen Informationsquellen mit vertretbarem Aufwand nicht in der Recherche bewältigen. Das gilt besonders für Datenspuren von Ereignissen, die sich über Soziale Medien, Foren und Blogs erstmals bemerkbar machen. Um derartige Signale nicht zu verpassen, arbeiten einige Redaktionen mit Systemen, die Onlinequellen kontinuierlich scannen und nach Mustern durchsuchen (bspw. „Data-minr“ und „NewsWhip“). Auch die eigenen Archive werden für Redaktionen mittels lernender Algorithmen zunehmend besser nutzbar. Sei es über die automatisierte Verschlagwortung von Textinhalten, die sich somit gezielter durchsuchen lassen, oder die Digitalisierung audiovisueller Inhalte. So modernisiert der „Westdeutsche Rundfunk“ (WDR) bereits seit 2015 sein Archivsystem gemeinsam mit dem Fraunhofer-Institut für Intelligente Analyse- und Informationssysteme IAIS (vgl. Maroni et al. 2020). Unter anderem werden dabei Video- und Audiodateien automatisiert transkribiert, verschlagwortet und in Segmente zerlegt. Inhalte können infolgedessen von Redakteur:innen und Dokumentar:innen gezielter recherchiert und einfacher wiederverwendet werden.

Datenleaks als Quelle und Herausforderung

Während Archive gemeinhin bereits strukturierte Informationen enthalten, sind investigative Journalist:innen regelmäßig mit umfangreichen Leaks konfrontiert, die sich aus unterschiedlich gut erschließbaren Quellen diverser Formate zusammensetzen. So besteht das bislang umfangreichste Datenleak des Pandora-Papers-Projekts aus rund 11,9 Millionen Dateien, darunter vor allem Dokumente, Bilder und E-Mails. Für die Analyse derart komplexer Quellensammlungen kommen in den beteiligten Redaktionen Systeme zum Einsatz, die auf lernenden Algorithmen basieren. Die vom internationalen Recherchenetzwerk ICIJ genutzte Software „Neo4j“ identifiziert unter anderem Entitäten wie Unternehmen und Personen in den

Daten, um diese in einer grafischen Oberfläche als Netzwerk darstellen zu können (vgl. Hunger 2021). Eine optische Exploration von Beziehungsnetzwerken wird somit möglich und erlaubt gezielte vertiefende Recherchen. Auch in investigativen Projekten kleineren Umfangs helfen lernende Algorithmen dabei, verräterische Muster zu identifizieren – etwa in Satellitendaten, als Hinweise auf Bernsteinminen (vgl. Herasymenko et al. 2018) oder Solaranlagen (vgl. de-Lima-Santos/Salaverría 2021), oder als Belege für geheimdienstliche Überwachung in Positionsdaten von Flugzeugen (vgl. Aldhous 2017).

Verifikation ist zu einer zentralen Aufgabe geworden

Neben der Analyse von Datenleaks hat auch die Verifikation von Aussagen und Quellen in den vergangenen Jahren an Relevanz gewonnen. Angesichts strategisch gesteuerter Desinformationskampagnen und einer starken Verbreitung von Verschwörungserzählungen wappnen sich Redaktionen mit Werkzeugen, die den Prozess des Faktencheckens unterstützen. Systeme wie „ClaimBuster“ verwenden lernende Algorithmen, um Aussagen in digitalen Quellen zu identifizieren, zu klassifizieren und automatisiert mit bereits überprüften Fakten und falsifizierten Behauptungen abzugleichen. Das Ergebnis ist kein eindeutiges Urteil, sondern vielmehr eine Entscheidungshilfe. Schließlich zeichnen sich irreführende Informationen auch dadurch aus, dass Fakten in einen falschen Kontext gesetzt, satirisch, sarkastisch oder anderweitig umgedeutet werden. Für lernende Algorithmen stellt Desinformation daher eine besondere Herausforderung dar – inhaltlich wie auch technologisch. So werden aktuell mit teilweise beachtlichen Fördermitteln Systeme entwickelt, um Deepfakes zu erkennen oder zumindest Indizien für sie zu finden. Bereits Ende 2020 präsentierte Microsoft eine Anwendung zur Authentifizierung von Fotos und Videos, die anhand bekannter Manipulationen trainiert wurde und gemeinsam mit Medienpartnern weiterentwickelt wird (vgl. Burt 2020).

Produktion: Transkription, Translation, Synthetische Medien

Aus gesprochenem Wort wird Text: Was lernende Algorithmen bei der Transkription von Audiodateien leisten können, hat in vielen Redaktionen für große Entlastung gesorgt. Mussten Interviews bisher mühsam von Hand übertragen werden, können

das inzwischen Werkzeuge übernehmen, deren Ergebnisse nur noch sporadisch verbessert werden müssen (bspw. „Trint“). Ähnliche Systeme sind im Einsatz, wenn Videos automatisiert untertitelt und für Podcasts oder Videobeiträge maschinell erstellte Transkripte angeboten werden. Kaum ein lernendes Assistenzsystemen ist wohl schneller in Redaktionen akzeptiert worden, wird das Verschriftlichen doch als besonders zeitraubende und stupide Tätigkeit empfunden. Und auch hinsichtlich der Barrierefreiheit kann ein Transkript mit kleineren Fehlern noch als Fortschritt wahrgenommen werden.

Ebenfalls weit verbreitet sind lernende Systeme in Newsrooms, wenn fremdsprachige Inhalte übersetzt werden müssen. Einerseits verwenden viele Redakteur:innen in ihren Recherchen inzwischen Tools wie das in Köln entwickelte „DeepL“, dessen Algorithmen mit einer großen Menge an bilingualen Online-Texten trainiert wurden. Zusätzlich helfen derartige Werkzeuge bei der Produktion mehrsprachiger Inhalte. So demonstriert das BBC News Lab im Projekt „Frank“ (abgeleitet von Lingua franca), wie automatisierte Übersetzungen den Austausch von Beiträgen über Sprachgrenzen hinweg ermöglichen. Angesichts der 43 Sprachen, in denen das Medienhaus veröffentlicht, erscheint dieser Ansatz besonders vielversprechend – für eine stärkere Sichtbarkeit der Inhalte und eine effizientere Nutzung interner Ressourcen. Zumal die übersetzten Inhalte vor der Veröffentlichung noch eine Redigatur durchlaufen (vgl. BBC News Lab 2021).

Kann KI auch moderieren?

Auch die Suchmaschinenoptimierung (SEO) gehört in Digitalmedien zur Qualitätssicherung. Dabei geht es unter anderem um Verlinkungen oder zusätzliche Überschriften, die eine Auffindbarkeit der Inhalte über Suchmaschinen verbessern sollen. Lernende Algorithmen unterstützen Redaktionen etwa darin, diese SEO-Zeilen wirkungsvoll zu formulieren. Der Axel-Springer-Verlag entwickelte gemeinsam mit Amazon ein System, welches dafür direkt im Content Management System einen Vorschlag generiert und das Suchvolumen prognostiziert – basierend auf dem von Google erstellten BERT-Modell (vgl. Anastasiu 2020).

Weniger gut erschlossen ist aktuell das Feld synthetisch erstellter Medieninhalte. Während die chinesische Nachrichtenagentur Xinhua bereits 2018 einen künstlich generierten Moderator die TV-Nachrichten verlesen ließ, wird die Technologie in Europa bislang vor allem in Machbarkeitsstudien getestet. Ursächlich ist sicher auch die Sorge vor einem Vertrauensverlust

beim Publikum, wenn Journalist:innen ihrem Publikum gegenüber nicht mehr persönlich auftreten. Dabei ließen sich gerade Audio- und Videoinhalte über eine synthetisch erzeugte und somit personalisierbare Präsentation womöglich deutlich weiter verbreiten. Um dieses Potential zu erforschen, entwickelte das WDR-Innovationsteam eine künstliche Nachbildung der Stimme einer bekannten Moderatorin (vgl. WDR Innovation Hub 2021). Für eine Bewertung ist zudem entscheidend, ob nur die Moderationen, oder auch die Inhalte selbst von lernenden Algorithmen (etwa GPT-3) generiert werden. Letzteres würde fundamentale Fragen auf – nach den medienethischen Grundsätzen der Redaktion, presserechtlichen Pflichten und einer verantwortungsvollen Kommunikation.

Distribution: Personalisierung, Scoring, Empfehlungen

Geht es hingegen nicht um die Erstellung journalistischer Inhalte, sondern primär um deren Verbreitung und Monetarisierung, lassen sich viele Medienhäuser schon heute von lernenden Algorithmen helfen. Mehrere Verlage verwenden automatisierte Systeme zur Aussteuerung ihrer Paywall, zur personalisierten Kundenansprache oder um Abonnent:innen mit einem hohen Kündigungsrisiko zu identifizieren (vgl. Karle/Wiegand 2022).

Medienethisch herausfordernder ist die Berechnung von Metriken zu zentralen Qualitätsmerkmalen mithilfe lernender Algorithmen. So lassen sich einige Redaktionen das Abonnement-Potential eines Artikels bestimmen, die Häufigkeit weiblicher und männlicher Protagonisten oder die Konstruktivität der Berichterstattung (vgl. ebd.). Derartige Metriken können bei redaktionellen Entscheidungen helfen, diese jedoch auch unverhältnismäßig stark beeinflussen – schließlich lassen sich andere, womöglich wichtigere Qualitätskriterien weniger gut in Metriken übersetzen. Zudem bergen derartige Methoden die Gefahr, dass Erfolgsmuster aus der Vergangenheit reproduziert werden und die eigentliche Strategie keine Berücksichtigung findet. Denn das würde eine Operationalisierung voraussetzen – eine Übersetzung der redaktionellen Ziele in ein Format, das von Algorithmen prozessiert werden kann. Für die Steigerung von Abonnements mag das noch einfach sein, für eine facettenreiche, ausgewogene oder konstruktive Berichterstattung ist es eine große Herausforderung.

Diese Schwierigkeit zeigt sich auch bei lernenden Systemen, die Inhalte empfehlen oder deren Auswahl und Darstellung be-

einflussen sollen. Solche Algorithmen müssten in der Lage sein zu verstehen, was redaktionelle Relevanz bedeutet – ein Konstrukt, das ohnehin unklar definiert ist und von den wenigsten Redaktionen in ein klares Framework übersetzt wurde. Eine Ausnahme bildet das öffentlich-rechtliche Radio Schweden mit einem System, das die Positionierung von Inhalten auf Playlists und Websites automatisch vornimmt, basierend auf der Einschätzung der Redakteur:innen (vgl. Beckett 2020). Diese geben direkt im Redaktionssystem ein, wie bedeutsam und aktuell ein Beitrag ist, und inwiefern er auf klar definierte „public service values“ einzahlt, etwa die regionale Nähe zum Publikum. Derart eindeutige Kriterien sind es letztlich, die den meisten Empfehlungsalgorithmen fehlen, zumal wenn sie in Bereichen außerhalb des Journalismus trainiert wurden.

Der Ansatz aus Schweden zeigt: Der Kernbereich redaktioneller Arbeit wird für Algorithmen nur in einem hybriden System greifbar, in dem auch menschliche Expertise kontinuierlich wirksam ist. Geht es jedoch um einzelne Arbeitsschritte, die sich klar abgrenzen lassen, sind lernende Systeme in Redaktionen schon heute an vielen Stellen eingebunden. Wie viel sie Journalist:innen letztlich abnehmen können, hängt nicht nur von der technischen Entwicklung ab. Entscheidend wird sein, ob Redaktionen die damit zusammenhängenden Möglichkeiten erkennen, offen explorieren und kompetent mit ihren Qualitätskriterien zusammenführen. Der Überblick über Applikationen in Recherche, Produktion und Distribution zeigt jedenfalls eine große Vielfalt und ein mindestens ebenso großes Potential.

Literatur

- Aldhous, Peter (2017): *We Trained A Computer To Search For Hidden Spy Planes. This Is What It Found.* <https://www.buzzfeednews.com/article/peteraldhous/hidden-spy-planes>.
- Anastasiu, Cristian et al. (2021): *DeepTitle – Leveraging BERT to generate Search Engine Optimized Headlines.* <https://arxiv.org/pdf/2107.10935.pdf>.
- BBC News Labs (2021): *Frank – machine translation of BBC articles.* <https://bbcnewslabs.co.uk/projects/Frank/>.
- Beckett, Charlie (2020): *An algorithm for empowering public service news.* <https://blogs.lse.ac.uk/polis/2020/09/28/this-swedish-radio-algorithm-gets-reporters-out-in-society/>.
- Burt, Tom (2020): *New Steps to Combat Disinformation.* <https://blogs.microsoft.com/on-the-issues/2020/09/01/disinformation-deepfakes-newsguard-video-authenticator/>.

- de-Lima-Santos, Mathias Felipe/Salaverria, Ramón (2021): *From Data Journalism to Artificial Intelligence: Challenges Faced by La Nación in Implementing Computer Vision in News Reporting*. In: *Palabra Clave*, 24. Jg., H. 3, DOI: 10.5294/pacla.2021.24.3.7.
- Herasymenko, Vlad et al. (2018): *Leprosy of the land*. https://texty.org.ua/d/2018/amber_eng/.
- Hunger, Michael (2021): *Exploring the Pandora Papers with Neo4j*. <https://neo4j.com/developer-blog/exploring-the-pandora-papers-with-neo4j/>.
- Karle, Roland/Wiegand, Markus (2022): *Medienbranche und KI*. In: *Kress Pro*, H. 2.
- Maroni, Dirk et al. (2020): *Die ARD Mining-Plattform – Künstliche Intelligenz im produktiven Einsatz für die automatisierte Erschließung im multimediale Produktionsprozess*. In: *FKT – Fernseh- und Kinotechnik*, H. 5, S. 22-27.
- Pauly, Marcel (2021): *So lassen wir Dutzende Wahlkreisanalysen automatisiert schreiben*. <https://www.spiegel.de/backstage/ltw-in-baden-wuerttemberg-und-rheinland-pfalz-2021-automatisierte-wahlkreis-analysen-a-701c2ffb-coa9-4ffb-8575-ac27c2184a0e>.
- WDR Innovation Hub (2021): *Synthetische Medien. Zukünfte der Medienproduktion mit Künstlicher Intelligenz*. https://zukunft.wdr.de/assets/pdf/WDR-Zukunftsreport_SynthetischeMedien.pdf.

Alle Internetquellen zuletzt aufgerufen am 14.7.2022.