

KI-Modelle in Medienunternehmen: empirische Befunde und ethische Reflexionen für die Regulierung

Michael Litschka

Zusammenfassung

Der Beitrag untersucht die Auswirkungen und Herausforderungen des Einsatzes von KI-Modellen (unter anderem auch Sprachmodellen) in Medienunternehmen aus empirischer und ethischer Perspektive. Basierend auf einer für eine Regulierungsbehörde durchgeführten Studie, die unter anderem Expert:inneninterviews und eine SWOT-Analyse umfasste, werden zentrale Chancen und Risiken für ökonomische Wertschöpfung und ethische Verantwortung identifiziert. Die Ergebnisse zeigen einerseits mögliche Effizienzgewinne und Innovationspotenziale, andererseits erhebliche ethische Problemfelder wie algorithmische Verzerrung, fehlender Datenschutz, mangelnde Transparenz und die Gefahr des Verlusts journalistischer Qualität. Für die wahrgenommenen Risiken wünschen sich die an der Studie teilnehmenden Medienunternehmen klarere und stärkere Regulierungsvorschriften. In der auf diese empirischen Befunde folgenden ethischen Diskussion beschreibt der Beitrag einige aktuelle medien- und KI-ethische Ansätze, wobei er insbesondere die Gerechtigkeitstheorien nach Rawls (hinsichtlich substanzieller Chancengleichheit und öffentlicher Legitimation) und Sen (hinsichtlich „vergleichender“ Gerechtigkeitsanalysen und Pluralismus) systematisch gegenübergestellt und für die Entwicklung pluralistisch ausgerichteter, gesellschaftlich akzeptierter Regulierungsinstrumente fruchtbar macht. Beide Ansätze zeigen (und stützen somit Regulierungsbefürworter:innen), dass nicht Regulierung per se unternehmerische Freiheiten einschränkt, sondern der ethisch problematische Einsatz von KI Werte wie Autonomie und Freiheit gefährdet. Der Beitrag schließt mit der Empfehlung, sowohl substanzielle Chancengleichheit als auch globale Pluralität und partizipative Stakeholder-Dialoge als Grundlagen für eine verantwortungsvolle KI-Regulierung im Mediensektor zu verankern.

1. Einleitung

Der vorliegende Beitrag beschreibt *einerseits* (siehe Kapitel 2) eine empirische Untersuchung, die im Auftrag einer österreichischen Regulierungsbehörde durchgeführt wurde und eine umfassende Analyse des Einsatzes von KI-Anwendungen, insbesondere Sprachmodellen, in österreichischen und internationalen Medienunternehmen sowie Plattformen beinhaltet (siehe dazu Belinskaya et al. 2024 und Pinzolits et al. 2025). Ziel war es, sowohl ökonomische als auch gesellschaftliche Dimensionen zu beleuchten, wobei eine SWOT-Analyse entlang der Medien-Wertschöpfungskette im Mittelpunkt stand. Der Fokus lag insbesondere auf der Selbstsicht der Akteur:innen, um deren subjektive Wahrnehmungen und strategische Einschätzungen differenziert herauszuarbeiten.¹ Da diese Erhebung neben den technologischen Chancen auch viele gesellschaftliche und ethische Risiken der KI-Nutzung offenbart und der Wunsch nach einer stärkeren Regulierung aufkam, erfolgt nach der Beschreibung einiger empirischer Erkenntnisse aus der Studie eine genauere Analyse der Frage, wie eine solche Regulierung auch ethisch fundiert werden kann.

Der Beitrag versucht somit *andererseits* (siehe Kapitel 3 und 4), die aufkommenden Fragen in einen medien- und KI-ethischen Rahmen zu stellen, der in der zugrundeliegenden empirischen Studie nur angedeutet wurde. Dazu werden insbesondere die Gerechtigkeitstheorien nach Rawls und Sen herangezogen, systematisch gegenübergestellt und für die Entwicklung pluralistisch ausgerichteter, gesellschaftlich akzeptierter Regulierungsinstrumente fruchtbar gemacht. Beide Ansätze argumentieren (und stützen somit Regulierungsbefürworter:innen), dass nicht Regulierung *per se* unternehmerische Freiheiten einschränkt, sondern der ethisch problematische Einsatz von KI-Werte wie Autonomie und Freiheit gefährden kann. Ebenso zeigen beide Ansätze, dass eine öffentliche Rechtfertigung zum Beispiel einer KI-Strategie eines Unternehmens oder einer KI-Regulierung durch Behörden private Rechtfertigungen gewinnorientierten Handelns (durch Medienunternehmen und Plattformen) aussticht.

1 Eine Folgestudie für denselben Auftraggeber mit dem Ziel einer repräsentativen Erhebung der Einschätzung der Bevölkerung zum Einsatz von KI in Medien wurde im Oktober 2025 präsentiert (Pinzolits et al. 2025).

2. SWOT-Analyse zur Selbsteinschätzung der befragten Medienunternehmen

Die dem Beitrag zugrundeliegende Studie (teilweise veröffentlicht in Belinskaya et al. 2024 und Pinzolit et al. 2025) folgte einem multi-methodischen Design:

- Literaturstudie zum *status quo* von KI-Anwendungen im Mediensektor
- Sechzehn leitfadengestützte Expert:inneninterviews zur praktischen Integration und Bewertung von KI-Technologien im Medienworkflow (Personen im Management von Printmedien, Radio, öffentlich-rechtlicher Rundfunk, Universität, HAW, Presseagentur, Unterhaltungsmedien, Kreativagentur, KI-Consulting, KI-Softwareentwicklung)
- Eine daraus folgende Analyse der Selbsteinschätzung dieser Expert:innen zu Chancen und Risiken von KI-gestützten Medienprozessen (SWOT-Analyse) entlang der relevanten Wertschöpfungsstufen (siehe auch Abbildung 1).
- Eine systematische Übersicht aktueller Regulierungsrichtlinien
- Überlegungen zu KI-ethischen Inputs für künftige Regulierungsaktivitäten.

Im Folgenden sollen einige Ergebnisse der Interviews und der daraus hervorgegangenen SWOT-Analyse dargestellt werden. Ziel ist es, die anhand von zwei Wertschöpfungsstufen (*creation* und *editing*) beispielhaft herausgefilterten Chancen und Risiken aus Sicht der Akteur:in darzustellen und für jene Medienunternehmen, die im Bereich Journalismus und Unterhaltung tätig sind, möglichst zu verallgemeinern. Für diese verallgemeinerten Problemfelder gilt es dann, medien- und KI-ethisch reflektierte Regulierungszugänge zu diskutieren. Die Literaturstudie und die Übersicht zu aktuellen Regulierungsrichtlinien spielen für die folgende Argumentation keine Rolle; die für diesen Beitrag neu detaillierten KI-ethischen Überlegungen folgen in Kapitel 3 und 4.

KI-Sprachmodelle werden bereits in vielen internationalen Medien genutzt: In Schweden nutzt beispielsweise *Aftonbladet* KI, um Texte in Rap-Songs umzuwandeln; *Reuters* und *Synthesia* haben einen vollautomatisierten Nachrichtenzusammenfassungsdienst mit virtuellem Moderator; *Associated Press* verwendet Automatisierungstechnologien zur Berichterstattung über *Minor League* Baseball; Bloomberg nutzt ein Sprachmodell, um automatisierte Antworten auf Kundenfragen zu liefern; in der Radiobranche wird mit *voice-cloning* und Sprachsynthese-Technologien experimentiert; in der Forschung werden *LLMs* (*Large Language Models*) für Textgene-

rierung und -analyse eingesetzt. Wenn wir eine (von mehreren möglichen) typische Wertschöpfungskette eines Medienunternehmens heranziehen (siehe Abbildung 1), ergeben sich Themenfelder für eine so genannte SWOT-Analyse, wie sie in der Betriebswirtschaftslehre gerne verwendet wird.

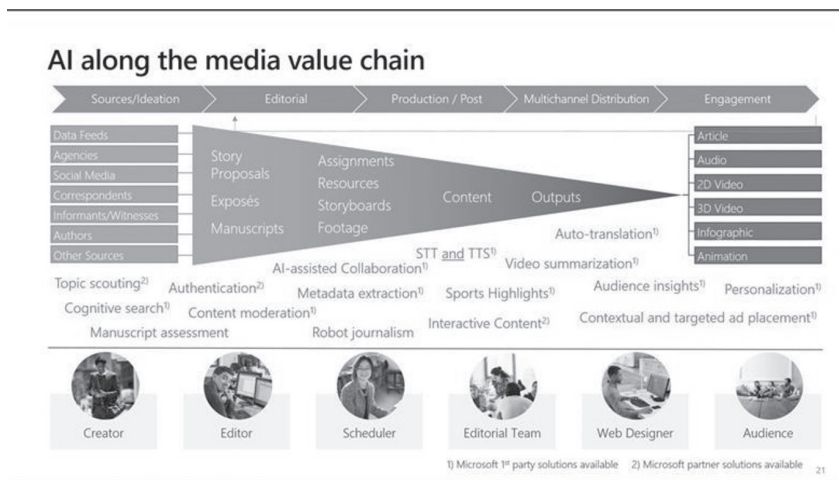


Abbildung 1: KI entlang der Medien-Wertschöpfungskette (Quelle: www.ibc.org)

Die SWOT-Analyse dient zur Identifikation von Stärken, Schwächen, Chancen und Risiken von Unternehmensstrategien und -ressourcen und wurde in diesem Fall auf die Implementierung von KI-Technologien im Mediensektor entlang der Wertschöpfungskette bezogen. Anbei ein Beispiel für den Bereich „Creation“, wie dieser sich aus Sicht der befragten Interviewpartner:innen darstellt:



Abbildung 2: SWOT-Bereich „Creation“

Demnach können Stärken des vermehrten Einsatzes von KI im Medien-Workflow etwa eine Effizienzsteigerung durch die Automatisierung repetitiver Aufgaben, eine Freisetzung kreativer Ressourcen bei der Contenterstellung durch KI-gestützte Tools sowie eine Erleichterung datenintensiver Prozesse wie Recherche, Marktanalyse und Musterdetektion sein. Schwächen hingegen sehen die Befragten beim zeit- und kostenintensiven Training der KI-Systeme, der schwierigen Überprüfung der erstellten Inhalte, dem Zugriff auf große Nutzerdatenmengen (Stichwort Datenschutz) und der Neigung zu Halluzinationen. Die befragten Interviewpartner:innen sehen durchaus das Risiko der Verbreitung falscher Informationen durch fehlerhafte oder verzerrte KI-Ausgaben und einen möglichen *bias* der Trainingsdaten, etwa bei der Restrukturierung großer Datenmengen.

Chancen ergeben sich durch neue kreative Möglichkeiten durch die Integration von KI in den Produktionsprozess (zum Beispiel KI-generierte Musiktexte, automatisierte Nachrichtenaggregation), durch die Unterstützung des Innovationspotenzials von Medienunternehmen im internationalen Konkurrenzzumfeld und durch Geschäftsmodellinnovation und Erschließung neuer Märkte. Auch die Schaffung neuer, KI-zentrierter Jobs wird als Möglichkeit wahrgenommen. Risiken ergeben sich durch die mögliche Erosion eigener kreativer Kompetenzen und die Gefahr zunehmender Abhängigkeit von KI-Lösungen großer Plattformen und deren proprietären Lösungen, die Marginalisierung assistierender Tätigkeiten am Arbeitsmarkt sowie das Potenzial für gesellschaftliche Desinformation und Manipulation, mangelnde Qualitätskontrollen und Vertrauenswürdigkeit. Gerade für das

Feld des Journalismus wurden auch die Gefahren des „automation creep“ als fortschreitender Ersatz menschlicher Arbeit durch Maschinen und KI, enthumanisierter Journalismus und algorithmusgetriebener Sensationalismus genannt.

Ein zweites Beispiel für den Bereich „Editing“ und wie dieser sich aus Sicht der befragten Interviewpartner:innen im Rahmen einer SWOT-Analyse darstellt:



Abbildung 3: SWOT-Bereich „Editing“

In der Medienproduktion wird die Rolle von KI im Bereich des Editing vielschichtig betrachtet. KI-Technologien bieten das Potenzial, traditionelle Arbeitsschritte zu ersetzen, man sieht aber das Bedürfnis, dass journalistische Kernkompetenzen gewahrt werden müssen. Während die automatisierte Weiterverarbeitung von Inhalten als wichtig angesehen wird, bleibt die Kennzeichnung von KI-erstellten Elementen unerlässlich. Nur repetitive Aufgaben oder etwa das Lektorat werden erleichtert, nicht jedoch das Verständnis komplexer Zusammenhänge überflüssig; besondere Verantwortung kommt hier den Führungskräften in Medienunternehmen zu, die sicherstellen müssen, dass Qualitätskriterien der journalistischen Arbeit gewährleistet bleiben. Die sorgfältige redaktionelle Beurteilung und die Beachtung des sozialen Kontexts bleiben wahrgenommene Schwächen der KI-Tools.

KI-Tools erleichtern die Themenfindung, indem sie wichtige Themen aus umfangreichen Dokumenten identifizieren und Vorschläge für Inter-

viewfragen generieren können, womit sie in der journalistischen Prozesskette unterstützen können. Hinsichtlich neuer KI-basierter Formate gibt es unterschiedliche Chancen, von der Entwicklung einer datenschutzkonformen KI-Lösung über die Verbesserung der Überprüfung nicht-textlicher Inhalte bis hin zur *deep fake detection*.

Die Darstellung der Einschätzung der Expert:innen betreffend möglicher Effizienzgewinne wird freilich nicht von allen Forscher:innen geteilt. So schätzt zum Beispiel Acemoglu (2024), dass zwar niedrig-qualifizierte Arbeiten produktiver vonstatten gehen würden, diese aber nur 5 Prozent aller Aufgaben in den nächsten 10 Jahren umfassen. Trabelsi (2024) stützt das Argument der Effizienzsteigerung durch KI durch nun besser mögliche *big data*-Analysen, betont aber das Risiko der Polarisierung und steigender Ungleichheiten auf Arbeitsmärkten. Auch Brynjolfsson (2025) sieht potenzielle Gewinne eher für den Niedriglohnsektor und bei seltener auftretenden Problemen, wenn Menschen wenig Vorerfahrungen besitzen. Keine Evidenz hingegen sieht er für Produktivitätsgewinne auf Industriebene (unabhängig vom Sektor).

Es zeigt zusammenfassend (siehe zu den ausführlichen Ergebnissen Belinskaya et al. 2024) in dieser Selbsteinschätzung und der SWOT-Analyse, dass Medienunternehmen (in unserem Beispiel journalistisch arbeitende, unterhaltende und Technologieplattformen nutzende Medien und KI-Entwickler:innen) die Vorteile und Nachteile KI-gestützter Wertschöpfungsformen auf Medienmärkten recht klar benennen können.

Die Vorteile der Integration von KI- beziehungsweise Großen Sprachmodellen (*LLM*) in die Wertschöpfungskette werden vor allem in der Automatisierung repetitiver Aufgaben im Medienprozess gesehen. Beispielsweise werden Recherchetätigkeiten, Datenanalysen und kreative Content Produktion unterstützt. Viele Systeme sind in der Lage, Muster, Trends und thematisch relationale Tendenzen sichtbar zu machen und bieten somit den Anwender:innen effizientere *workflows* und neue Zugänge zur Dateninterpretation.

Allerdings sind den teilnehmenden Expert:innen die möglichen Nachteile der Technologienutzung durch zum Beispiel größere Abhängigkeit, weniger Verständlichkeit (Stichwort „Explainability“ der KI: die zunehmende Abstraktion und Intransparenz der Prozesse führt zu Verständniseinbußen bei weniger technikaffinen Akteur:innen), möglicherweise vermehrte Falschinformationen und durch *bias* verzerrte Entscheidungen der KI bewusst. Dazu kommen datenschutzrechtliche Bedenken sowie mögliche unbewusste Einflussnahmen auf die eigenständige (menschliche) Kreativar-

beit. Im Bereich journalistisch operierender Medienunternehmen ist die Hauptherausforderung das Erfüllen journalistischer Qualitätskriterien der Überprüfung, Datenqualität und Einhaltung des *human-in-the-loop*-Prinzips; im Bereich digitaler Plattformen geht es vor allem um algorithmische Transparenz. Nicht zuletzt sehen die befragten noch abzuwartende Umwälzungen auf dem Arbeitsmarkt und bei den erforderlichen Qualifikationen ihrer Mitarbeiter:innen.

Die Wahrnehmung der KI-Integration ist somit ambivalent: Während operative Vorteile durchaus anerkannt werden, dominiert eine kritische Grundhaltung hinsichtlich längerfristiger gesellschaftlicher und branchenspezifischer Risiken. Da immer wieder der Wunsch nach klaren Regulierungsrichtlinien aufkam, sollen nun für diese einige medien- und KI-ethische Ansätze auf ihre Fruchtbarkeit für künftige Regulierungsmaßnahmen analysiert werden.

3. Aktuelle medien- und KI-ethische Ansätze

Die ethische Reflexion des KI-Einsatzes im Medienbereich orientiert sich in der Literatur an mehreren etablierten theoretischen Strängen, die je nach Schwerpunkt der Autor:innen (siehe Abbild 3) von diesen detaillierter ausgearbeitet werden.

Theoriestrang/Konzept	Autor:innen	Jahr	Keywords
Tugendethisch-teleologisch Tugendethische Nutzung	Ess Spiekermann Cohen Vallor	2020 2019 2012 2024	Welche Art Mensch muss ich werden, in der ständigen Ausübung meiner technologischen Interaktionen, um zufrieden/glücklich zu sein? → <i>Eudaimonia</i> ; <i>Wertebewusstsein</i> im Umgang mit (digitaler) Technologie; Frage nach dem „Warum“ neuer technologischer Entwicklungen; vernünftige <i>Wertepriorisierung</i> und <i>Wertebalancierung</i> , so dass jede(r) ihr/sein Telos im Leben erreichen kann; Wertebewusste Programmierung
Kantianisch-deontologisch	Floridi et al.	2018	Digitale Technik muss einerseits die Kantische <i>Menschenwürde</i> , andererseits die möglichst komplette Selbstrealisierung („ <i>self-realization</i> “) unterstützen
Narrative Ethik und Werte	Grimm et al.	2019	Privatheit, Autonomie, Sicherheit als grundlegende Werte und Rechte; <i>narrative Ethik</i> der Digitalisierung
Autonomiediskussion	Thimm & Bächle Coeckelbergh	2019 2020	<i>Autonomie</i> als komplexe (und interdisziplinär zu klärende) Zuschreibung, also als <i>Bedingung für die Möglichkeit</i> der Übernahme ethischer Verantwortung (nicht als reine Beschreibung der technischen Entscheidungsmöglichkeiten von algorithmischen Plattformen)
Designprinzipien	Dignum	2019	<i>Accountability</i> , <i>Responsibility</i> und <i>Transparency</i> für die Produktion und Anwendung neuer Technologien
Gerechtigkeitszugänge, öff. Vernunftgebrauch	Sen, Rawls Habermas	2010 2001 1991	<i>Gerechtigkeit</i> der Datengebarung/des Datenzugangs; <i>Gerechtigkeit als Fairness</i> ; <i>Komparative Gerechtigkeit</i> ; Kommunikative Vernunft

Abbildung 4: Einige medien- und KI-ethische Zugänge

Eine ausführliche Diskussion der Ansätze würde den Rahmen sprengen; zudem ist diese Übersicht natürlich keine vollständige Liste. Sie zeigt aber trotz der unterschiedlichen Herkunftsdisziplinen (vor allem Medienethik, KI-Ethik, Ökonomie und Philosophie) wichtige Kernaussagen, die sich auch in der Regulierungsdebatte zur KI wiederfinden, beziehungsweise auch, welche Theoriestränge sich womöglich noch nicht so prominent bemerkbar gemacht haben, aber für die Regulierung fruchtbar gemacht werden könnten.

Tugendethische und teilweise teleologisch ausgerichtete Ansätze (zum Beispiel Spiekermann 2019; Ess 2020 und Vallor 2024) legen einen Fokus auf Wertebewusstsein, das Ziel der „*Eudaimonia*“ (in etwa: „Glückseligkeit durch ethische Vernunft und Ausüben von Tugenden“) und verlangen nach einer „Wertepriorisierung“ bei der Entwicklung digitaler Technologien. Kantianisch-deontologisch geprägte Zugänge (zum Beispiel Floridi et al. 2018) betonen die Menschenwürde und das Ermöglichen der Selbstrealisierung durch digitale Technik. Einen Schwerpunkt auf die Werte Privatheit, Autonomie und Sicherheit und die Nutzung einer „narrativen“ Ethik setzen zum Beispiel Grimm et al. (2019). Coeckelbergh (2020) und Thimm/Bächle (2019) beschäftigen sich unter anderem mit dem Thema Autonomie und betrachten diese als gesellschaftlich und interdisziplinär auszuhandelndes Konstrukt und als Bedingung der Möglichkeit der Übernahme ethischer Verantwortung (was Maschinenethik in einem spezifischen Sinn, nämlich

dass Maschinen „autonom“ entscheiden könnten, ausschließt). Einen eher pragmatischen Zugang findet man bei bestimmten Anforderungen an das Design von KI-Tools, etwa an *accountability*, *responsibility* und *transparency* (vgl. Dignum 2019). Viele Richtlinien und Ethik-Kodizes für KI-Entwickler:innen versuchen, solche und weitere Werteanforderungen zu integrieren.

In der KI-Ethik Debatte geht es nicht zuletzt auch um Gerechtigkeitszugänge und den öffentlichen Vernunftgebrauch für eine mögliche gesellschaftliche Akzeptanz neuer KI-Entwicklungen. Dabei wird zwar auf einige diesbezüglichen klassischen Ansätze (vgl. zum Beispiel Habermas 1991; Rawls 2001 und Sen 2010) verwiesen, aber aus Sicht des Autors werden diese Ansätze bisher zu wenig ausgearbeitet (vgl. hierzu Litschka 2025). Mögliche Themenfelder wären hier die Gerechtigkeit des Datenzugangs und der Datennutzung (vgl. Litschka/Saurwein/Pellegrini 2024) sowie Fairnessfragen bei Entscheidungen durch KI-(Sprach-)Modelle, aber auch generell die Rawlsianische Frage nach der gesellschaftlichen Grundstruktur und der Institutionen, die durch die vermehrte KI-Verbreitung und Nutzung verändert werden.

Wie könnte man also künftig diese gerechtigkeitsorientierten Ansätze verwenden, um zum Beispiel gerechte Zugangsmöglichkeiten, Wertediversität, Pluralität und prozedurale Transparenz im KI-gestützten Mediensektor zu sichern / zu erlangen und künftige Regulierungsmaßnahmen noch mehr auf Gerechtigkeitskonzepte zu stützen?

4. Welche KI-Ethik für die Regulierung? Rawls und Sen revisited

Eine naheliegende Frage an dieser Stelle lautet, warum Ethiker:innen bei Fragen der KI-Regulierung und Governance nach einer bestimmten (prozessorientierten) Gerechtigkeitsethik suchen sollten, und nicht mit anderen oben beschriebenen (oder weiteren verfügbaren) Ansätzen auskommen können. Die Antwort darauf ist, dass insbesondere zwei Probleme die Integration normativer Werte in Regulierungsdokumente und -akte kompliziert machen:

- *Das Problem der Werteaggregation*: Aus der klassischen ökonomischen Forschung ist bekannt, dass das Treffen gesamtgesellschaftlicher Entscheidungen aufgrund individueller Präferenzen bestimmten logischen Problemen unterliegt. Schon Arrows Unmöglichkeitstheorem (vgl. Arrow 1951), auf dem auch Ergebnisse anderer *Social Choice*-Theoreti-

ker wie Sen (2017) aufbauen, zeigt, dass es kein komplett konsistentes (Wahl-)System gibt, das wünschenswerte Kriterien wie Nicht-Diktatur, Pareto-Effizienz und Unabhängigkeit von irrelevanten Alternativen gleichzeitig erfüllen kann, wenn die Präferenzen aller Wähler berücksichtigt werden.

- *Das Problem des Wertpluralismus*: Soziale Entscheidungen unterliegen immer auch politischen Problemen und Differenzen, die sich durch unterschiedliche Kulturen, Bildungszugänge und Traditionen ergeben. Ein bekanntes Beispiel hierfür war die kaum aufzulösende Debatte über die widerstreitenden Werte Freiheit und Sicherheit bei der Bekämpfung der *Covid*-Pandemie. Für beide demokratisch akzeptierten und der Gesellschaft wichtigen Werte gab und gibt es gute Argumente; nur beide Werte gleichzeitig zu maximieren, funktioniert selten.

Umgelegt auf unser Problem der Regulierung bedeuten diese Probleme, dass es ohne eine allgemeine Zustimmung zu einem Prozess für Regulierungsvorgaben schwierig werden wird, die unterschiedlichen Präferenzen der Unternehmen, Nutzer:innen und Behörden abzustimmen. Freilich bleiben uns am Ende eines solchen Weges Werteentscheidungen kaum erspart, werden aber womöglich durch einen als gerecht empfundenen Prozess dorthin erleichtert. Gerechtigkeitsorientierte Zugänge versuchen nämlich oft, individuelle Präferenzen und politische Differenzen zu umgehen, indem sie unter anderem fragen, welche *Prozesse* zu einer gesellschaftlich akzeptablen Entwicklung und Nutzung von KI führen und welche *Institutionen* Kooperation innerhalb einer Gesellschaft und ein gemeinsames Gerechtigkeitsverständnis sichern. Zwei bekannte Ansätze hierzu stammen von John Rawls und Amartya Sen.

Rawls' (2001) Theorie der Gerechtigkeit als Fairness und seine bekannten Gerechtigkeitsprinzipien für Hintergrundinstitutionen der Gesellschaft könnte man folgendermaßen für das oben angesprochene Problem anwenden: Da Gleichheit nicht nur formaler Natur sein, sondern auch zu realen, vergleichbaren Erfolgsaussichten führen soll, muss in der KI-Wirtschaft sogenannte *substanzielle* Chancengleichheit hergestellt werden. Dies meint die aktive Beseitigung von Ungleichheiten statt rein prozeduraler Gleichheit (wie sie zum Beispiel vor dem Gesetz herrscht). Man kann dies dadurch begründen, dass algorithmische Diskriminierung einerseits im Gegensatz zum Recht aller auf gleiche Bürgerschaft steht, andererseits eine Überwachung durch KI-Systeme ein „gutes Leben ohne unerwünschte Einflussnahme“ (wie es Rawls ausdrückt) behindern kann. Faire Chancengleichheit

ist eben nicht nur formale Chancengleichheit, sondern Herstellung einer fairen Chancenverteilung (ähnliche Fähigkeiten führen zu ähnlichen Erfolgsaussichten).

Für Medienunternehmen, die mit KI operieren, bedeutet das, dass die Notwendigkeit öffentlicher Rechtfertigung, insbesondere durch öffentliche Deliberation über KI-Modelle, ernster als bislang genommen werden muss. Dies gilt nicht nur für öffentlich organisierte Unternehmen wie öffentlich-rechtliche Anbieter, sondern auch für private (Plattform-) Unternehmen; öffentliche Rechtfertigung sticht private Begründungen, zum Beispiel bezüglich optimaler und effizienter, gewinnorientierter Geschäftsmodelle, aus. Die oft geäußerten Bedenken vieler privat organisierter Medienbetriebe über mögliche Einschränkungen der unternehmerischen Freiheit, die solche Kommunikationsstrategien angeblich mit sich bringen, können so entkräftet werden: Autonomie und Freiheit werden nicht durch Rechtfertigungsprozesse und Regulierungen unterminiert, sondern durch mögliche ungerechte Einflussnahmen einer Technologie wie KI (vgl. Gabriel 2022: 12). Die bei öffentlich-rechtlichen Sendern beispielsweise immer schon stärkere Tradition der Rechtfertigung ihrer Methoden (der Themenauswahl, der Recherche, der verwendeten Technologie, etc.) sollte somit vermehrt in den privaten Sektor Eingang finden, etwas, das wohl nur durch stärkere Regulierung stattfinden kann.

Sowohl die selbstregulierenden *Governance*- und *Accountability*-Maßnahmen der Medienunternehmen selbst, als auch die rechtlichen Vorgaben der Medien- und Technologiepolitik, und nicht zuletzt auch die europäischen Regulierungsrichtlinien bezüglich des Designs neuer KI-Anwendungen sollten also entsprechend angepasst werden. Dies betrifft neben der Notwendigkeit der öffentlichen Deliberation (unabhängig der richtigerweise stattgefundenen Deliberation solcher wegweisender Verordnungen wie dem *AI-Act*) auch die aktive Beseitigung aufgetretener Ungleichheiten; hier könnte man an unterschiedlich vorhandene KI-Kompetenzen des Publikums denken und Maßnahmen, die diese Differenzen ausgleichen.

Sen entwickelte seinen „Capability Approach“ weiter zu einem eigenen Gerechtigkeitsansatz („Comparative Justice“; vgl. Sen 2010; Litschka 2015 und Litschka 2019). Im Unterschied zu Rawls verzichtet Sen auf absolute, universelle Gerechtigkeitsprinzipien und betont stattdessen die Bewertung konkreter institutioneller Wirkungen und realer menschlicher Handlungsmöglichkeiten im internationalen Vergleich. Diese realen Handlungsmöglichkeiten (eben „Capabilities“) müssen unter anderem von der Politik (mittels verschiedener Konversionsfaktoren) hergestellt werden. Nahelie-

gend in unserem Zusammenhang und direkt anschließend an die empirischen Ergebnisse, die eine Überforderung der Nutzer:innen von KI befürchteten, wäre ein Fokus auf KI-Capabilities der Bürger:innen in der Bildungspolitik.

Betreffend der Medienlandschaft sieht Sen sowieso eine starke Rolle unabhängig operierender Medien(unternehmen). Wie er an einer Stelle (vgl. Sen 2010: 201) formuliert: Global agierende Medien und transnationale Organisationen können die so genannten „Grenzen der Gerechtigkeit“ erweitern. Denn der Fokus gerade in technologischen Umbruchszeiten muss immer auch auf Bewohner:innen anderer Länder statt nur eine Nation gelegt werden, womit vermieden werden soll, dass provinzielle (stark von einer Nation oder Kultur geprägten) Werte zu stark in den Vordergrund treten („open“ statt „closed impartiality“). Der „Impartial Spectator“ dient hier als metaphysisches Modell für eine objektive Beurteilung und die Integration der „distant voices“, um eine offene, pluralistische und global legitimierte Wertebasis zu schaffen. So können laut Sen verschiedene und divergierende Gerechtigkeitsprinzipien durch den öffentlichen und unparteilichen Vernunftgebrauch vereint werden.

Konkret bedeutet dies für Medien- und KI-Unternehmen und die zuständigen Regulierungsbehörden, dass vermehrt öffentliche Stakeholder-Dialoge mit Entwickler:innen, Behörden, Unternehmen und User:innen anzustreben sind, um diese Werteintegration zu erreichen; für Regulierungsprinzipien, wie sie grundlegend für Regelwerke wie den DMA (*Digital Markets Act*) / DSA (*Digital Services Act*) / AI-Act und deren Nachfolger sein sollten, sind internationale (-kulturelle) Gesichtspunkte einzubeziehen.

Mit Sen sind auch viele Medienökonom:innen und Medienethiker:innen der Meinung, dass utilitaristische Denkweisen die normativen Probleme der KI-Weiterentwicklungen nicht abdecken können (vgl. beispielsweise Christians 2007; Sen/Williams 1982 und Karmasin/Litschka 2013) und verlangen nach Analysen mittels kommunikativer Rationalität (ähnlich wie Habermas 1991) und deontologischen Konzepten. Ökonomisch rationale Argumente in utilitaristischen Theorien wie dem *Uses-and-Gratifications*-Ansatz (also was genau „nutzt“ eine KI-Anwendung spezifischen User:innen?) sind nicht in der Lage, wichtige normative Konzepte wie die Pflichten von Individuen, die Einbettung von Individuen in eine reaktionsfähige Mediengesellschaft oder Massenmedien als soziale Institutionen, die nicht durch isolierte individuelle Entscheidungen verändert werden können, zu erfassen (vgl. Christians 2007). Nur deontologische (und diskursive) Theorien und, wenn man den Argumenten dieses Beitrags folgt, *capability*-

orientierte Denkmodelle können eine solche vollständige normative Sicht vermitteln.

Der *Capability*-Ansatz selbst betont vielerorts die Bedeutung deontologischer Kategorien und von „Prozessen“ auf dem Weg zur Erreichung bestimmter Ziele, zwei für den Utilitarismus unwichtige Konzepte. Sen (1985: 4) hat beispielsweise ein prozedurales Verständnis der Bedeutung von Märkten in einer Gesellschaft und Rechten als unveräußerlichen Aspekten einer Person. Was den Markt für Künstliche Intelligenz betrifft, so findet sich in einem Großteil der wirtschaftswissenschaftlichen Literatur, aber auch in Äußerungen von *big-tech-Managern* der Branche (anschaulich dazu Vallor 2024) utilitaristisches Denken und ein starker Glaube an das Funktionieren von Märkten oder Marktplätzen für Ideen (Karmasin/Litschka 2013). Während diese Theorien aus der Perspektive der Betonung des innovativen Potenzials funktionierender Technologiemarkte und der Abschreckung von zu viel staatlichem Einfluss vernünftig sind, scheinen sie nicht in der Lage zu sein, die vielschichtigen Dilemmata zu bewältigen, mit denen sich Organisationen und Bürger:innen bei der aktuellen Nutzung von KI konfrontiert sehen. An anderer Stelle haben der Autor und Kollegen beispielsweise beschrieben, welche Probleme die Marktkonzentration auf Plattformmärkten und in der Social-Media-Industrie verursachen kann (siehe zum Beispiel Litschka et al. 2024). Eine utilitaristische Ethik scheint nicht dazu beizutragen, diese Probleme zu lindern.

Deontologische Ansätze betonen neben der Bedeutung unveräußerlicher Rechte und der Rolle autonomer und universalisierbarer Entscheidungen, bei denen neben den erwarteten Ergebnissen von Interaktionen auch Verfahren und Prozesse berücksichtigt werden müssen, vor allem eine bestimmte Art von Rationalität. Öffentliche Argumentation – unter anderem ermöglicht durch ein funktionierendes System globaler Medien – ist nicht nur der Grundpfeiler der Demokratie, sondern auch eines möglichen universellen Gesellschaftsvertrags. Für die künftige Entwicklung der KI-Ethik ist es wichtig, diese Art von Vernunft einzubeziehen: Sie verlangt, dass wir unsere Argumente vor allen anderen rechtfertigen können. Wenn wir unterschiedliche Kulturen und Traditionen in die KI-Entwicklung und -Nutzung einbeziehen wollen (siehe auch Ess 2020), um zumindest teilweise universalisierbare Vereinbarungen über KI-Regelungen zu erreichen, ist dieses diskursethische Prinzip weiterhin gültig.

Während also Rawls insbesondere auf substanzielle Gleichheit im Zugang zu Chancen pocht und eine öffentlich nachvollziehbare Legitimation von KI-basierten Entscheidungen fordern würde, betont Sen den Pluralis-

mus und die Notwendigkeit internationaler Vergleichsperspektiven. Beide Ansätze bieten somit komplementäre Orientierungspunkte für eine ethisch begründete und gesellschaftlich akzeptierte Regulierung von KI-Technologien: Rawls als Grundlage für interne Chancengleichheit und institutionelle Rechenschaftspflicht, Sen als Impuls für transnationale Offenheit, Partizipation und Diversität der Wertvorstellungen. In der Medienpraxis empfiehlt sich eine verbindende Perspektive, welche substanzielle Gleichheit und globale Pluralität abbildet.

5. Zusammenfassung und Ausblick

Die empirischen Erkenntnisse der durchgeführten Studie bestätigen eine kritische Selbstsicht der Akteur:innen im Hinblick auf die Integration von KI-Technologien im Medienbereich. Im Medienproduktionsalltag haben sich schon an vielen Stellen effizienzsteigernde und innovationsfördernde Prozesse durch eine Anwendung von KI gezeigt. Auf der anderen Seite werden die gesellschaftlichen und ethischen Herausforderungen, die durch Automatisierung, Datenmonopolisierung und algorithmische Steuerung entstehen, kritisch gesehen. Während die Medien- und KI-Ethik schon vielerorts Probleme aufgezeigt und Lösungswege vorgeschlagen hat, sind manche gerechtigkeitsbasierten Theorieimpulse, wie zum Beispiel jene von Rawls (substanzielle Fairness und öffentliche Rechtfertigung) und Sen (vergleichende Gerechtigkeit und Pluralität), bislang noch weniger in politische Maßnahmen und Regulierungen eingeflossen. Der Beitrag zeigt einige erste Möglichkeiten auf, diese Impulse für die Theorieentwicklung einerseits und die praktische Implementierung in Regulierungsvorhaben andererseits aufzunehmen. Unter anderem wären demnach in dieser Sichtweise für eine menschengerechte und gesellschaftlich akzeptable Entwicklung von KI im Medienbereich Prozesse und Institutionen notwendig, die pluralistische Prinzipien offen einbeziehen und substanzielle Chancengleichheit gewährleisten.

Literatur

Acemoglu, Daron (2024): The Simple Macroeconomics of AI, in: Economic Policy, 5. April 2024, (online unter: <https://economics.mit.edu/sites/default/files/2024-04/The%20Simple%20Macroeconomics%20of%20AI.pdf> – letzter Zugriff: 17.11.2025).

- Arrow, Kenneth Joseph (1951): Social Choice and Individual Values (= Cowles Commission for Research in Economics, Monograph No. 12), New York.
- Belinskaya, Yulia et al. (Hg.) (2024): KI in der Medienwirtschaft, RTR Studienreihe zu Künstlicher Intelligenz, Wien.
- Brynjolfsson, E. et al. (2025): Generative AI at Work, in: The Quarterly Journal of Economics 140 (2/2025), S. 889–947.
- Christians, Clifford G. (2007): Utilitarianism in media ethics and its discontents, in: Journal of Mass Media Ethics 22 (2–3/2007), S. 113–131.
- Coeckelbergh, Marc (2020): AI Ethics, Cambridge.
- Dignum, Virginia (2019): Responsible Artificial Intelligence. How to Develop and Use AI in a Responsible Way, Cham.
- Ess, Charles (2020): Digital Media Ethics, 3. Aufl., Cambridge.
- Floridi, Luciano et al. (2018): An Ethical Framework for a Good AI Society: Opportunities, Risks, Principles, and Recommendations, in: Minds and Machines 28 (4/2018), S. 689–707.
- Gabriel, Iason (2022): Towards a Theory of Justice for Artificial Intelligence, in: Daedalus 151 (2/2022), S. 1–12.
- Grimm, Petra et al. (Hg.) (2019): Digitale Ethik. Leben in vernetzten Welten, Paderborn.
- Habermas, Jürgen (1991): Erläuterungen zur Diskursethik, Frankfurt am Main.
- Karmasin, Matthias / Litschka, Michael (2013): Normativität in der Medienökonomie, in: Matthias Karmasin / Matthias Rath / Barbara Thomasz (Hg.), Normativität in der Kommunikationswissenschaft, Berlin, S. 191–207.
- Litschka, Michael (2015): Medien-Capabilities als polit-ökonomisches Konzept in der Medienethik: Theoretische Grundlagen und mögliche Anwendungen, in: Communicatio Socialis 48 (2/2015), S. 190–201.
- Litschka, Michael (2019): The Political Economy of Media Capabilities: The Capability Approach in Media Policy, in: Journal of Information Policy 9 (63/2019), S. 79–110.
- Litschka, Michael (2025): AI Ethics and the Capability Approach, in: Genealogy+Critique 11 (1/2025), S. 1–16.
- Litschka, Michael / Saurwein, Florian / Pellegrini, Tassilo (2024): Open Data Governance und digitale Plattformen, Ethische, ökonomische und regulatorische Herausforderungen und Perspektiven, Wiesbaden.
- Pinzolits, Robert et al. (2025): KI im Mediensektor: Eine SWOT-Analyse entlang der Wertschöpfungskette und ihre regulatorischen Implikationen, in: Rimscha, M. Björn von / Ehrlich, Gianna / Riemann, Robin (Hg.), Alles rational? Der menschliche Faktor in Medienorganisationen: Proceedings zur Jahrestagung der Fachgruppe Medienökonomie der DGpuK 2024, Mainz, S. 106–120.
- Pinzolits, Robert et al. (2025): KI und Medienvertrauen. Mediennutzung und Medienvertrauen in Österreich im Spannungsfeld von KI und Sozialen Medien. RTR-Studienreihe: Künstliche Intelligenz in der Medienwirtschaft, 14. Oktober 2025 (online unter: https://www.rtr.at/medien/aktuelles/publikationen/Publikationen/Publikationen_2025/Studie_KI_und_Medienvertrauen.de.html – letzter Zugriff 17.11.2025).

- Rawls, John* (2001): *Justice as Fairness: A Restatement*, Harvard.
- Sen, Amartya* (2010): *The Idea of Justice*, London.
- Sen, Amartya* (2017): *Collective Choice and Social Welfare*, New York.
- Sen, Amartya / Williams Bernard* (Hg.) (1982): *Utilitarianism and Beyond*, Cambridge.
- Spiekermann, Sarah* (2019): *Digitale Ethik: Ein Wertesystem für das 21. Jahrhundert*, München.
- Thimm, Caja / Bächle, Thomas C.* (2019): *Autonomie der Technologie und autonome Systeme als ethische Herausforderung*, in: Matthias Rath / Friedrich Krotz / Matthias Karmasin (Hg.), *Maschinenethik. Normative Grenzen autonomer Systeme*, Wiesbaden, S. 73–87.
- Trabelsi, Mohamed Ali* (2024): *The impact of artificial intelligence on economic development*, in: *Journal of Electronic Business & Digital Economics* 3 (2/2024), S. 142–155.
- Vallor, Shannon* (2024): *The AI Mirror*, Oxford.

