



Dampak Algoritma AI terhadap Komunikasi Publik: Memahami Manipulasi Informasi dan Realitas

Ilham Nurfajri*, Erlangga Teguh Hadi Pratama, Gilang Septian Tupamahu, Ridwan Saputra, Yusi Erwina

Program Studi Ilmu Komunikasi, Universitas Muhammadiyah Tangerang

Abstrak: Algoritma kecerdasan buatan (AI) memainkan peran penting dalam membentuk informasi yang diterima oleh publik melalui platform digital, seperti media sosial dan mesin pencari. Dengan menyaring konten berdasarkan preferensi pengguna, algoritma ini dapat menciptakan "gelembung informasi", yang membatasi paparan terhadap sudut pandang berbeda dan memperburuk polarisasi sosial. Penggunaan AI dalam komunikasi publik juga menimbulkan tantangan etis terkait bias, manipulasi informasi, dan ketidakadilan dalam penyebaran informasi. Artikel ini mengkaji dampak algoritma AI terhadap komunikasi publik, termasuk dampak sosialnya, serta perlunya regulasi untuk memastikan penggunaan AI yang adil dan transparan dalam ruang publik.

Kata kunci: Algoritma AI, Komunikasi Publik, Gelembung Informasi, Polarisasi Sosial, Manipulasi Informasi, Regulasi

DOI:

<https://doi.org/10.47134/converse.v1i3.3543>

*Correspondence: Ilham Nurfajri

Email: ilhamnurfajri110302@gmail.com

Received: 02-01-2025

Accepted: 04-01-2025

Published: 31-01-2025



Copyright: © 2024 by the authors. Submitted for open access publication under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).

Abstract: Artificial intelligence (AI) algorithms play an important role in shaping the information received by the public through digital platforms, such as social media and search engines. By filtering content based on user preferences, these algorithms can create "information bubbles," limiting exposure to differing viewpoints and exacerbating social polarization. The use of AI in public communications also raises ethical challenges related to bias, manipulation of information, and unfairness in the dissemination of information. This article examines the impact of AI algorithms on public communications, including their social impacts, as well as the need for regulation to ensure fair and transparent use of AI in the public sphere.

Keywords: AI Algorithm, Public Communication, Information Bubble, Social Polarization, Information Manipulation, Regulation

Pendahuluan

Algoritma AI mengasumsikan fungsi penting dalam penyaringan informasi yang dapat diakses oleh pengguna, yang secara mendalam membentuk perkembangan opini publik. Algoritma ini menyesuaikan konten sesuai dengan kecenderungan pengguna, sering mengakibatkan munculnya "gelembung filter" yang memperkuat keyakinan yang sudah ada sebelumnya dan membatasi paparan ke berbagai perspektif. Mekanisme ini berpotensi mendistorsi wacana publik dan mempengaruhi partisipasi demokrasi.

Mekanisme penyaringan informasi Algoritma personalisasi: Algoritma ini menekankan konten yang sesuai dengan interaksi pengguna historis, sehingga menghasilkan umpan berita yang disesuaikan yang dapat mengaburkan sudut pandang

subjek (Powers, 2017). Gelembung Filter: pengguna sering menemukan diri mereka terjatuh dalam ruang gema di mana mereka sebagian besar menemukan ide-ide yang kongruen, yang dapat mengintensifkan polarisasi, terutama dalam skenario politik (Liu et al., 2021) (Park & Park, 2024).

Dampak Pada Opini Publik Efek Kognitif: Prevalensi berita palsu dan informasi bias dapat mengurangi kapasitas kognitif publik, yang menyebabkan warga negara yang salah informasi dan merusak proses demokrasi (Valle et al., 2024). Homogeni Perspektif: Kecenderungan politik yang intens menyebabkan berkurangnya paparan terhadap berbagai konten, sehingga melanggengkan bias yang ada dan membatasi wacana sipil (Liu et al., 2021).

Meskipun algoritma kecerdasan buatan memiliki potensi untuk meningkatkan pengalaman pengguna melalui pengiriman konten terkait, mereka secara bersamaan menghadirkan bahaya dengan mendorong polarisasi dan membatasi keragaman informasi, yang pada akhirnya dapat membahayakan partisipasi demokratis dan kewarganegaraan yang terinformasi dengan baik.

AI secara mendalam mengubah lanskap komunikasi, terutama dalam kaitannya dengan konsep transparansi dan inklusivitas. Meskipun AI memiliki potensi untuk meningkatkan efisiensi komunikasi, ia secara bersamaan memperkenalkan seluk-beluk yang dapat membatasi akses informasi, sehingga mempengaruhi kualitas wacana. Bagian selanjutnya akan memeriksa interaksi yang rumit ini.

Peningkatan transparansi mengenai keselamatan dan tujuan AI dapat menumbuhkan kepercayaan dan mendorong upaya kolaboratif. Meskipun demikian, transparansi mengenai kemampuan dapat mengakibatkan kerugian kompetitif, yang dapat menghalangi inovasi dan membatasi akses ke teknologi yang menguntungkan (Bostrom, 2017). AI menambah metodologi inovasi terbuka konvensional, memfasilitasi kolaborasi yang lebih efektif. Namun, kemajuan ini dapat menimbulkan perbedaan dalam akses, karena entitas yang memiliki kemampuan AI canggih dapat memonopoli bidang inovasi (Holgersson et al., 2024).

Integrasi teknologi AI untuk tujuan komunikatif ditandai dengan kurangnya keseragaman; variabel demografis seperti usia dan status sosial ekonomi membentuk hambatan yang signifikan untuk mengakses. Perbedaan ini dapat mengakibatkan kesenjangan digital yang nyata, di mana hanya populasi terpilih yang menuai manfaat dari kemajuan dalam kecerdasan buatan (Goldenthal et al., 2021). Batasan yang dikenakan pada akses agen ke informasi yang diperlukan dapat menghambat kemajuan komunikasi, karena protokol yang ditetapkan dapat membatasi penyebaran informasi penting (Kim & Song, 2023).

Sebaliknya, meskipun kecerdasan buatan memiliki potensi untuk meningkatkan komunikasi, secara bersamaan dapat mengintensifkan ketidaksetaraan yang sudah ada sebelumnya, yang berpuncak pada situasi di mana sejumlah individu istimewa dapat sepenuhnya memanfaatkan teknologi. Skenario ini mendorong pertanyaan kritis mengenai masa depan akses yang adil ke informasi dalam paradigma AI-sentris.

Penggabungan *Artificial Intelligence* (AI) dalam bidang komunikasi publik menimbulkan dilema etika yang cukup besar, terutama dalam kaitannya dengan bias, transparansi, dan akuntabilitas. Karena teknologi AI semakin meresapi jurnalisme dan lanskap media, menjadi penting untuk menghadapi tantangan ini untuk mendorong penggunaan yang bertanggung jawab dan etis. Bagian selanjutnya menggambarkan masalah etika dan metodologi penting yang bertujuan mengurangi bias dalam algoritma AI.

Masalah Etis dalam AI dalam Komunikasi Publik Bias dan Diskriminasi: Sistem AI memiliki potensi untuk menyebarkan bias yang sudah ada sebelumnya yang tertanam dalam kumpulan data pelatihan, menghasilkan dampak diskriminatif dalam penyebaran berita dan wacana publik (Sahoo, 2024). Transparansi dan Akuntabilitas: Karakteristik algoritma AI yang tidak dapat dipahami menghambat pemahaman kerangka pengambilan keputusan mereka, sehingga memperumit upaya untuk menganggap akuntabilitas untuk output mereka (Sahoo, 2024). Pengawasan Manusia: Penentuan otomatis yang dijalankan oleh sistem AI mungkin rentan terhadap ketidakakuratan, yang memerlukan pengawasan manusia untuk menjamin hasil etis (Sahoo, 2024).

Meminimalkan Bias dalam Algoritma AI Data Pelatihan Beragam: Jaminan bahwa kumpulan data pelatihan beragam dan representatif dapat secara signifikan berkontribusi pada pengurangan bias dalam output yang dihasilkan AI (Xu, 2024). Transparansi Algoritmik: Perumusan kerangka kerja yang meningkatkan interpretabilitas algoritma AI dapat meningkatkan akuntabilitas dan menimbulkan kepercayaan (Xu, 2024). Pedoman Etis: Pembentukan dan penegakan pedoman etika dalam organisasi dapat secara efektif mengintegrasikan prinsip-prinsip jurnalistik ke dalam sistem AI (Porlezza & Schapals, 2024).

Sementara pendekatan ini memiliki potensi untuk secara substantif mengurangi masalah etika, tantangan yang melekat tetap ada dalam penerapan pragmatis pedoman ini, karena kemajuan cepat teknologi AI sering melampaui kerangka peraturan yang ada.

Manipulasi informasi oleh kecerdasan buatan (AI) menimbulkan konsekuensi sosial-politik yang signifikan, terutama memfasilitasi polarisasi masyarakat dan degradasi proses demokrasi. Fenomena ini dikatalisis oleh algoritma yang secara aktif membentuk sentimen publik dan mempengaruhi interaksi sosial, yang berpuncak pada dampak penting untuk

kohesi sosial dan tata kelola. Bagian selanjutnya menggambarkan contoh utama dan implikasi terkaitnya.

Polarisasi Masyarakat Divergensi Sosial-Politik: Konsekuensi AI di pasar tenaga kerja menunjukkan korelasi dengan polarisasi sosial-politik, di mana individu yang terkena dampak negatif oleh otomatisasi cenderung condong ke arah ideologi yang lebih konservatif, berbeda dengan mereka yang memperoleh manfaat dari AI, yang cenderung ke perspektif liberal (Jacobs, 2024). **Echo Chambers:** Algoritma yang digunakan oleh platform media sosial sering menimbulkan ruang gema, sehingga memperkuat keyakinan yang sudah ada sebelumnya dan mengurangi paparan terhadap berbagai sudut pandang, yang memperburuk keretakan sosial (Spina et al., 2023).

Dampak pada Opini Publik Penyebaran Disinformasi: Penyebaran misinformasi yang merajalela melalui platform yang digerakkan oleh AI merusak kepercayaan publik terhadap kerangka kelembagaan, mengakibatkan berkurangnya kapasitas untuk kewarganegaraan yang terinformasi dan melemahnya proses demokrasi secara bersamaan (Valle et al., 2024). **Manipulasi Kehendak:** Teknologi AI memiliki kemampuan untuk memanipulasi sentimen publik dengan menargetkan individu dengan misinformasi yang disesuaikan, sehingga menimbulkan persepsi yang menyimpang tentang realitas dan melemahkan keterlibatan warga (Ienca, 2023).

Konsekuensi dari dampak sosial Erosi Legitimasi Demokratis: Proliferasi informasi yang salah dan polarisasi masyarakat dapat memicu krisis legitimasi dalam administrasi publik, karena warga menjadi semakin kecewa dengan pemerintahan (Valle et al., 2024). **Kerusakan Hubungan Sosus:** Komunikasi yang dimediasi AI memiliki potensi untuk mengubah dinamika interpersonal, menumbuhkan persepsi positif dan negatif, yang memperumit interaksi sosial (Hohenstein et al., 2023).

Sebaliknya, sementara peran AI dalam manipulasi informasi menghadirkan risiko yang cukup besar, ia secara bersamaan menawarkan solusi prospektif, seperti meningkatkan proses pemeriksaan fakta dan menumbuhkan literasi media, yang dapat mengurangi beberapa dampak buruk.

Keterlibatan badan pemerintah dan badan pengatur dalam pengawasan kecerdasan buatan (AI) sangat penting untuk mempromosikan kesetaraan dan pelestarian keragaman informasi dalam domain publik. Langkah-langkah peraturan yang memadai dapat mengurangi potensi bahaya yang terkait dengan teknologi AI sekaligus menumbuhkan kepercayaan publik dan merangsang inovasi. Bagian selanjutnya menggambarkan elemen-elemen utama dari kerangka peraturan ini.

Kerangka dan Standar Peraturan Pemerintah sedang dalam proses merumuskan kerangka peraturan holistik, dicontohkan oleh Undang-Undang AI UE, yang menggarisbawahi perlunya standar yang koheren untuk menegakkan hak-hak

fundamental dan prinsip-prinsip demokrasi (Gamito, 2024). Organisasi Pengembangan Standar (SDO) sangat penting dalam menetapkan kriteria penting untuk sistem AI berisiko tinggi, sehingga memfasilitasi pemeliharaan kelayakan teknis dan akuntabilitas (Gamito, 2024). Mengatasi Masalah Etis Langkah-langkah peraturan harus mencakup dilema etika, termasuk masalah bias, transparansi, dan akuntabilitas, untuk mendorong integrasi teknologi AI yang adil (Pantanowitz et al., 2024) (Almeida et al., 2023). Dianjurkan bagi pembuat kebijakan untuk merancang kerangka peraturan yang dapat disesuaikan yang dapat berkembang seiring dengan kemajuan teknologi sambil memprioritaskan keselamatan publik dan perlindungan hak asasi manusia (Mendes et al., 2023). Tantangan dalam Regulasi AI Ciri-ciri khas AI, yang ditandai dengan ketidakpastian dan kerumitan, menghadirkan hambatan signifikan bagi paradigma regulasi konvensional, yang memerlukan strategi inovatif yang menggabungkan konsolidasi otoritas dan kemampuan intervensi yang cepat (Judge et al., 2024). Ada potensi bahaya bahwa ketergantungan pada standar yang diselaraskan dapat membahayakan akuntabilitas demokrasi jika proses penetapan standar tidak memiliki inklusivitas dan transparansi (Gamito, 2024).

Sementara regulasi sangat penting untuk mendorong keadilan dalam AI, masih ada wacana berkelanjutan mengenai keseimbangan antara inovasi dan pengawasan peraturan. Beberapa berpendapat bahwa peraturan yang terlalu ketat dapat menghambat kemajuan teknologi, sehingga menekankan perlunya kerangka kerja peraturan yang fleksibel yang responsif terhadap evolusi cepat yang melekat dalam pengembangan AI.

Metode

Metode penelitian ini mengadopsi pendekatan penelitian perpustakaan, yang berarti pengumpulan data dilakukan dengan mencari informasi dari berbagai sumber, seperti jurnal ilmiah, buku referensi, dan bahan bibliografi. Dalam pendekatan ini, peneliti tidak hanya mengandalkan buku, tetapi juga memanfaatkan internet untuk mengakses jurnal, teori, penelitian sebelumnya, dan pandangan yang relevan dengan topik penelitian. Tujuan utamanya adalah untuk menemukan teori-teori yang dapat mendalami pemahaman tentang masalah yang sedang diteliti.

Penggunaan internet sangat penting karena menyediakan akses ke sejumlah besar informasi yang berkaitan dengan topik penelitian, yang dapat memperkaya studi dengan literatur dari berbagai penelitian sebelumnya di seluruh dunia. Keuntungan utama dari internet adalah kemudahan akses dan antarmuka yang ramah pengguna, yang memungkinkan pencarian data lebih efisien dan efektif dalam proses pengumpulan data untuk penelitian ini.

Hasil dan Pembahasan

Peran Algoritma dan AI dalam Pembentukan Diskursus Publik

Gelembung informasi dan ruang gema mewakili fenomena yang muncul dari proses kurasi konten yang digerakkan secara algoritmik, terutama lazim dalam platform media sosial dan mesin pencari. Gelembung informasi bermanifestasi ketika individu menemukan diri mereka terpapar pada spektrum sudut pandang yang terbatas, seringkali secara tidak sengaja, sedangkan ruang gema sengaja mengecualikan dan mendelegitimasi perspektif yang berbeda. Misalnya, algoritma rekomendasi yang digunakan oleh YouTube memiliki kapasitas untuk menumbuhkan ruang gema dengan mendukung konten yang kongruen secara ideologis, sehingga memperburuk polarisasi politik (Park & Park, 2024).

Definisi Gelembung Informasi: Konsekuensi dari paparan selektif terhadap informasi, di mana perspektif terkait secara tidak sengaja diabaikan.

Definisi Echo Chambers: Terlibat dalam pengecualian sistematis dan delegitimisasi sudut pandang yang berlawanan, sehingga menumbuhkan ketidakpercayaan terhadap sumber eksternal (Nguyen, 2020).

Dampak Jangka Panjang Perspektif yang menyempit: Individu mungkin menghadapi tantangan dalam memahami masalah rumit karena tidak mencukupi berbagai sudut pandang yang beragam. Peningkatan polarisasi: penguatan terus-menerus dari keyakinan analog dapat menimbulkan fraktur sosial dan menghambat wacana konstruktif (Park & Park, 2024) (Nguyen, 2020). Disosiasi Kognitif: menghadapi informasi yang kontradiktif dapat membangkitkan ketidaknyamanan, mendorong individu untuk lebih memperkuat keyakinan mereka yang sudah ada sebelumnya (Nguyen, 2020).

Sebaliknya, beberapa sarjana berpendapat bahwa paparan konten yang disesuaikan dapat meningkatkan keterlibatan dan kepuasan pengguna, menunjukkan hubungan beragam antara kurasi algoritmik dan konsumsi informasi. Namun demikian, sudut pandang ini sering mengabaikan konsekuensi buruk bagi pemikiran kritis dan kohesi komunitas.

Dimensi Ontologis dalam Komunikasi Berbasis AI

Munculnya kecerdasan buatan dalam praktik komunikatif menimbulkan pertanyaan ontologis yang mendalam mengenai prinsip-prinsip transparansi dan kebebasan aksesibilitas informasi. Meskipun AI memiliki kemampuan untuk meningkatkan efisiensi komunikasi, ia secara bersamaan mengkurasi konten yang didasarkan pada data khusus pengguna, yang secara tidak sengaja dapat membatasi variasi informasi yang dapat diakses. Dualitas yang melekat ini memerlukan evaluasi yang ketat tentang apakah komunikasi yang dimediasi oleh AI benar-benar dapat dianggap terbuka, atau jika itu memaksakan kendala pada otonomi pengguna.

AI sebagai Mediator Komunikasi Sistem kecerdasan buatan secara sistematis mengkurasi konten melalui algoritma yang meneliti perilaku pengguna, menghasilkan lingkungan informasi yang dipersonalisasi namun berpotensi terbatas (Wang, 2023). Ketergantungan pada AI untuk moderasi konten dapat menumbuhkan ruang gema, di mana pengguna sebagian besar menghadapi perspektif yang sesuai dengan perspektif mereka sendiri, sehingga membatasi ruang lingkup wacana (Wang, 2023). Implikasi Etis Ketidakmampuan AI untuk terlibat dalam rasionalitas menghakimi memicu kekhawatiran mengenai pengaruhnya terhadap pengambilan keputusan etis dalam bidang komunikasi (O'Regan & Ferri, 2024). Pengasuh mengartikulasikan kekhawatiran bahwa agen percakapan dapat membahayakan otonomi manusia dan proses perkembangan, menunjukkan preferensi untuk interaksi manusia yang otentik daripada keterlibatan AI yang dimotivasi oleh keuntungan (Kucirkova & Hiniker, 2024). Pendidikan Dialogis dan Keterbukaan Pendidikan dialogis menganjurkan keterbukaan dan keterlibatan di berbagai konteks budaya, menyatakan bahwa AI dapat memungkinkan dialog yang lebih inklusif jika dirancang dengan cermat (Wegerif et al., 2019). Meskipun demikian, tantangan terus-menerus terletak pada jaminan bahwa AI tidak melanggar identitas monologis, tetapi sebaliknya menumbuhkan dialog dan pemahaman otentik (Wegerif et al., 2017).

Sebaliknya, beberapa orang berpendapat bahwa AI memiliki potensi untuk meningkatkan komunikasi dengan memberikan informasi yang disesuaikan yang memenuhi persyaratan individu, sehingga memperkaya pengalaman pengguna. Sudut pandang ini menggarisbawahi perlunya menyelaraskan kemampuan AI dengan pertimbangan etis untuk menjaga komunikasi tetap terbuka dan inklusif.

Tantangan Etis dalam Algoritma AI

Kecerdasan Buatan (AI) memiliki potensi yang cukup besar untuk meningkatkan banyak domain; Namun, secara bersamaan menghadirkan kebingungan etis, terutama dalam kaitannya dengan bias dan diskriminasi. Tanpa pengawasan yang tepat, mekanisme AI dapat memperkuat bias yang sudah ada sebelumnya, sehingga memfasilitasi penyebaran informasi yang salah dan fragmentasi kohesi masyarakat. Wacana ini meneliti konsekuensi buruk prospektif dari AI yang tidak diatur dalam domain kurasi informasi, metodologi untuk memastikan hasil yang adil, dan strategi untuk menambah transparansi algoritma AI.

Konsekuensi Buruk dari AI yang Terabaikan Penguatan Bias: AI memiliki kapasitas untuk memperbesar bias sosial, yang berpuncak pada hasil yang merugikan di sektor-sektor seperti perekrutan dan peradilan pidana (Saleh Hamed Albarashdi, 2024). Penyebaran Misinformasi: Sistem AI mandiri dapat mengkurasi informasi dengan cara yang salah menggambarkan persepsi publik, mengintensifkan polarisasi masyarakat

(Floridi, 2024). Pelanggaran Privasi: Agregasi kumpulan data yang luas dapat melanggar hak privasi individu, menimbulkan kesulitan etis (Bellaby, 2024). Mempromosikan Kesetaraan dalam Algoritma AI Data Pelatihan Bervariasi: Integrasi kumpulan data yang beragam dapat membantu mengurangi bias dalam output AI (*"Exploring the Ethical Implications of AI-Powered Personalization in Digital Marketing,"* 2024). Audit Sistematis: Penerapan penilaian rutin sistem AI dapat memfasilitasi identifikasi dan perbaikan bias (Saleh Hamed Albarashdi, 2024). Keterlibatan Pemangku Kepentingan Kolaboratif: Dimasukkannya kelompok heterogen dalam fase pengembangan dapat meningkatkan pertimbangan etis (Floridi, 2024). Meningkatkan Transparansi Standar Pengungkapan Eksplisit: Sistem AI diharuskan untuk menyampaikan metodologi dan sumber datanya secara transparan kepada pengguna akhir (Cuartielles et al., 2024). Protokol yang Dapat Diakses: Pembentukan pedoman terbuka untuk pemanfaatan AI dapat menumbuhkan budaya akuntabilitas (Cuartielles et al., 2024). Keterlibatan Masyarakat: Melibatkan publik dalam dialog mengenai aplikasi AI dapat memastikan bahwa banyak perspektif diperhitungkan (Floridi, 2024).

Meskipun pendekatan ini memiliki potensi untuk mengurangi risiko yang terkait dengan AI, kekhawatiran tetap ada bahwa transparansi saja mungkin terbukti tidak memadai. Kerumitan sistem AI dapat mengaburkan pemahaman, sehingga menghambat kemampuan pengguna untuk sepenuhnya memahami implikasi keputusan yang didorong oleh AI.

Dampak sosial dari Manipulasi Informasi oleh AI

Konsekuensi sosial dari manipulasi informasi yang difasilitasi oleh kecerdasan buatan sangat signifikan, terutama dalam hal memperburuk polarisasi dan menghalangi pemahaman bersama di antara spektrum sudut pandang. Algoritma yang mengkurasi informasi dapat menghasilkan ruang gema, di mana individu sebagian besar menemukan konten yang beresonansi dengan keyakinan mereka yang sudah ada sebelumnya. Keterlibatan selektif ini dapat mengintensifkan perpecahan sosial, membuatnya semakin menantang bagi individu untuk berinteraksi dengan perspektif alternatif. Bagian selanjutnya menyelidiki dimensi kritis dari fenomena ini.

Informasi dan Kualitas Banjir AI generatif memiliki kapasitas untuk membanjiri ekosistem informasi, terutama selama peristiwa penting, menghasilkan masuknya konten yang mungkin menipu atau salah (Lucas et al., 2024). Kesulitan membedakan antara konten asli dan yang dihasilkan AI semakin mempersulit kepercayaan publik dan keterlibatan kritis dengan informasi (Sanchez-Acedo et al., 2024). Personalisasi dan Penargetan Mikro Algoritma sering menggunakan metodologi penargetan yang canggih, menyesuaikan konten dengan preferensi individu, yang dapat memperkuat bias dan membatasi paparan

ke beragam sudut pandang (Schmitt & Flechais, 2024). Penargetan mikro semacam itu dapat menimbulkan persepsi realitas yang miring, karena pengguna mungkin sebagian besar menemukan informasi yang memvalidasi keyakinan mereka, sehingga memperdalam kesenjangan (Ienca, 2023). Implikasi untuk Harmoni Sosus Manipulasi informasi melalui AI dapat membahayakan harmoni sosial dan pluralisme dengan menumbuhkan suasana di mana dialog dan pemahaman terhambat (Tomassi et al., 2024). Ketika individu menjadi semakin tidak mampu terlibat dengan sudut pandang yang berbeda, kohesi komunitas dapat memburuk, yang berpuncak pada konflik dan perpecahan yang meningkat.

Sebaliknya, beberapa orang berpendapat bahwa AI memiliki potensi untuk meningkatkan akses ke perspektif yang beragam jika direayasa dengan prinsip-prinsip transparansi dan inklusivitas di garis depan. Namun demikian, lintasan saat ini menunjukkan bahwa, tanpa tindakan proaktif, bahaya polarisasi dan pembagian cenderung melampaui keuntungan prospektif ini.

Peran kebijakan dalam Mengatur AI untuk komunikasi publik

Tata kelola kecerdasan buatan dalam komunikasi publik sangat penting untuk menegakkan prinsip-prinsip kesetaraan dan integritas dalam penyebaran informasi. Kerangka kebijakan yang kuat dapat mengurangi bahaya yang terkait dengan kecerdasan buatan, termasuk manipulasi sentimen publik dan penyebaran informasi yang bias. Tata kelola ini dapat direalisasikan melalui metodologi berorientasi risiko, seperti yang diilustrasikan oleh Undang-Undang AI Uni Eropa, yang berupaya merekonsiliasi inovasi dengan langkah-langkah perlindungan yang diperlukan. Bagian selanjutnya menggambarkan aspek utama dari skema peraturan ini.

Regulasi Berbasis Risiko Undang-Undang AI Uni Eropa menggarisbawahi kerangka kerja yang berorientasi risiko untuk menghindari regulasi yang berlebihan sekaligus menjaga kesejahteraan masyarakat (Ebers, 2024). Metodologi ini memfasilitasi pembentukan peraturan yang ditargetkan yang sesuai dengan risiko berbeda yang terkait dengan aplikasi kecerdasan buatan dalam komunikasi publik.

Transparansi dan Akuntabilitas Klausul transparansi dalam Undang-Undang AI sangat penting dalam melengkapi pengguna untuk membedakan konten yang dihasilkan AI dan dalam mengurangi kemungkinan manipulasi (Piasecki et al., 2024). Pedoman eksplisit yang mengatur pembuatan konten AI dapat berkontribusi pada pelestarian integritas informasi yang disebarluaskan di domain publik.

Pertimbangan Kepentingan Umum Sistem kecerdasan buatan harus dibangun dengan penekanan pada memajukan kebaikan publik, yang mencakup keharusan etis dan aplikasi pragmatis (Züger & Asghari, 2024). Legislatur harus memastikan definisi kepentingan publik dalam penerapan kecerdasan buatan untuk menjamin akses informasi yang adil.

Sebaliknya, peraturan yang terlalu ketat dapat menghambat inovasi dan membatasi keuntungan prospektif kecerdasan buatan dalam meningkatkan komunikasi publik. Keseimbangan antara langkah-langkah regulasi dan kebutuhan untuk kemajuan teknologi tetap menjadi tantangan penting.

Simpulan

Algoritma AI telah menjadi elemen penting dalam membentuk komunikasi publik, terutama dalam mengelola arus informasi dan memengaruhi persepsi realitas. Teknologi ini memiliki kemampuan untuk mempersonalisasi konten berdasarkan preferensi individu, namun juga dapat digunakan untuk memanipulasi informasi demi tujuan tertentu. Manipulasi ini mencakup penyebaran berita palsu (*fake news*), amplifikasi bias, dan penciptaan realitas alternatif melalui *deepfake* dan konten sintetik lainnya.

Tantangan utama yang dihadapi adalah bagaimana menjaga transparansi dan akurasi dalam komunikasi publik di tengah dominasi algoritma AI. Regulasi yang ketat, literasi digital, dan pengembangan algoritma yang lebih etis menjadi solusi penting untuk mengurangi dampak manipulasi informasi. Dengan pendekatan yang bertanggung jawab, algoritma AI dapat digunakan secara positif untuk mendukung komunikasi yang inklusif dan berbasis fakta, sekaligus membangun kepercayaan masyarakat terhadap informasi yang mereka terima.

Daftar Pustaka

- Almeida, V., Mendes, L. S., & Doneda, D. (2023). On the Development of AI Governance Frameworks. *IEEE Internet Computing*, 27(1), 70–74. <https://doi.org/10.1109/MIC.2022.3186030>
- Bellaby, R. (2024). The ethical problems of ‘intelligence–AI.’ *International Affairs*, 100(6), 2525–2542. <https://doi.org/10.1093/ia/iiae227>
- Bostrom, N. (2017). Strategic Implications of Openness in <sc>AI</sc> Development. *Global Policy*, 8(2), 135–148. <https://doi.org/10.1111/1758-5899.12403>
- Cuartiellas, R., Mauri-Ríos, M., & Rodríguez-Martínez, R. (2024). Transparency in AI usage within fact-checking platforms in Spain and its ethical challenges. *Communication & Society*, 257–271. <https://doi.org/10.15581/003.37.4.257-271>
- Ebers, M. (2024). Truly Risk-based Regulation of Artificial Intelligence How to Implement the EU’s AI Act. *European Journal of Risk Regulation*, 1–20. <https://doi.org/10.1017/err.2024.78>
- Exploring the Ethical Implications of AI-Powered Personalization in Digital Marketing. (2024). *Data Intelligence*. <https://doi.org/10.3724/2096-7004.di.2024.0055>

- Floridi, L. (2024). Introduction to the Special Issues. *American Philosophical Quarterly*, 61(4), 301–307. <https://doi.org/10.5406/21521123.61.4.01>
- Gamito, M. C. (2024). The role of ETSI in the EU's regulation and governance of artificial intelligence. *Innovation: The European Journal of Social Science Research*, 1–16. <https://doi.org/10.1080/13511610.2024.2349627>
- Goldenthal, E., Park, J., Liu, S. X., Mieczkowski, H., & Hancock, J. T. (2021). Not All AI are Equal: Exploring the Accessibility of AI-Mediated Communication Technology. *Computers in Human Behavior*, 125, 106975. <https://doi.org/10.1016/j.chb.2021.106975>
- Hohenstein, J., Kizilcec, R. F., DiFranzo, D., Aghajari, Z., Mieczkowski, H., Levy, K., Naaman, M., Hancock, J., & Jung, M. F. (2023). Artificial intelligence in communication impacts language and social relationships. *Scientific Reports*, 13(1), 5487. <https://doi.org/10.1038/s41598-023-30938-9>
- Holgersson, M., Dahlander, L., Chesbrough, H., & Bogers, M. L. A. M. (2024). Open Innovation in the Age of AI. *California Management Review*, 67(1), 5–20. <https://doi.org/10.1177/00081256241279326>
- Ienca, M. (2023). On Artificial Intelligence and Manipulation. *Topoi*, 42(3), 833–842. <https://doi.org/10.1007/s11245-023-09940-3>
- Jacobs, J. (2024). The artificial intelligence shock and socio-political polarization. *Technological Forecasting and Social Change*, 199, 123006. <https://doi.org/10.1016/j.techfore.2023.123006>
- Judge, B., Nitzberg, M., & Russell, S. (2024). When code isn't law: rethinking regulation for artificial intelligence. *Policy and Society*. <https://doi.org/10.1093/polsoc/puae020>
- Kim, T., & Song, H. (2023). Communicating the Limitations of AI: The Effect of Message Framing and Ownership on Trust in Artificial Intelligence. *International Journal of Human-Computer Interaction*, 39(4), 790–800. <https://doi.org/10.1080/10447318.2022.2049134>
- Kucirkova, N., & Hiniker, A. (2024). Parents' ontological beliefs regarding the use of conversational agents at home: resisting the neoliberal discourse. *Learning, Media and Technology*, 49(2), 290–305. <https://doi.org/10.1080/17439884.2023.2166529>
- Liu, P., Shivaram, K., Culotta, A., Shapiro, M. A., & Bilgic, M. (2021). The Interaction between Political Typology and Filter Bubbles in News Recommendation Algorithms. *Proceedings of the Web Conference 2021*, 3791–3801. <https://doi.org/10.1145/3442381.3450113>
- Lucas, J. S., Maung, B. M., Tabar, M., McBride, K., & Lee, D. (2024). The Longtail Impact of Generative AI on Disinformation: Harmonizing Dichotomous Perspectives. *IEEE Intelligent Systems*, 39(5), 12–19. <https://doi.org/10.1109/MIS.2024.3439109>

- Nguyen, C. T. (2020). ECHO CHAMBERS AND EPISTEMIC BUBBLES. *Episteme*, 17(2), 141–161. <https://doi.org/10.1017/epi.2018.32>
- O'Regan, J. P., & Ferri, G. (2024). Artificial intelligence and depth ontology: implications for intercultural ethics. *Applied Linguistics Review*. <https://doi.org/10.1515/applirev-2024-0189>
- Pantanowitz, L., Hanna, M., Pantanowitz, J., Lennerz, J., Henricks, W. H., Shen, P., Quinn, B., Bennet, S., & Rashidi, H. H. (2024). Regulatory Aspects of Artificial Intelligence and Machine Learning. *Modern Pathology*, 37(12), 100609. <https://doi.org/10.1016/j.modpat.2024.100609>
- Park, H. W., & Park, S. (2024). The filter bubble generated by artificial intelligence algorithms and the network dynamics of collective polarization on YouTube: the case of South Korea. *Asian Journal of Communication*, 34(2), 195–212. <https://doi.org/10.1080/01292986.2024.2315584>
- Piasecki, S., Morosoli, S., Helberger, N., & Naudts, L. (2024). AI-generated journalism: Do the transparency provisions in the AI Act give news readers what they hope for? *Internet Policy Review*, 13(4). <https://doi.org/10.14763/2024.4.1810>
- Porlezza, C., & Schapals, A. K. (2024). AI Ethics in Journalism (Studies): An Evolving Field Between Research and Practice. *Emerging Media*, 2(3), 356–370. <https://doi.org/10.1177/27523543241288818>
- Powers, E. (2017). My News Feed is Filtered? *Digital Journalism*, 5(10), 1315–1335. <https://doi.org/10.1080/21670811.2017.1286943>
- Sahoo, M. (2024, November 4). Ethics in AI – Critical Skills for the New World. ADIPEC. <https://doi.org/10.2118/222249-MS>
- Saleh Hamed Albarashdi. (2024). Discrimination Associated with Artificial Intelligence Technologies. *EVOLUTIONARY STUDIES IN IMAGINATIVE CULTURE*, 637–645. <https://doi.org/10.70082/esiculture.vi.2099>
- Sanchez-Acedo, A., Carbonell-Alcocer, A., Gertrudix, M., & Rubio-Tamayo, J.-L. (2024). The challenges of media and information literacy in the artificial intelligence ecology: deepfakes and misinformation. *Communication & Society*, 223–239. <https://doi.org/10.15581/003.37.4.223-239>
- Schmitt, M., & Flechais, I. (2024). Digital deception: generative artificial intelligence in social engineering and phishing. *Artificial Intelligence Review*, 57(12), 324. <https://doi.org/10.1007/s10462-024-10973-2>
- Spina, D., Sanderson, M., Angus, D., Demartini, G., McKay, D., Saling, L. L., & White, R. W. (2023). Human-AI Cooperation to Tackle Misinformation and Polarization. *Communications of the ACM*, 66(7), 40–45. <https://doi.org/10.1145/3588431>

- Tomassi, A., Falegnami, A., & Romano, E. (2024). Mapping automatic social media information disorder. The role of bots and AI in spreading misleading information in society. *PLOS ONE*, 19(5), e0303183. <https://doi.org/10.1371/journal.pone.0303183>
- Valle, V. C. L. L., Fernández Ruiz, M. G., & Buttner, M. (2024). Fake news, influência na formação da opinião pública e impactos sobre a legitimidade da decisão pública. *A&C - Revista de Direito Administrativo & Constitucional*, 24(95), 73–97. <https://doi.org/10.21056/aec.v24i95.1898>
- Wang, S. (2023). Factors related to user perceptions of artificial intelligence (AI)-based content moderation on social media. *Computers in Human Behavior*, 149, 107971. <https://doi.org/10.1016/j.chb.2023.107971>
- Wegerif, R., Doney, J., Richards, A., Mansour, N., Larkin, S., & Jamison, I. (2019). Exploring the ontological dimension of dialogic education through an evaluation of the impact of Internet mediated dialogue across cultural difference. *Learning, Culture and Social Interaction*, 20, 80–89. <https://doi.org/10.1016/j.lcsi.2017.10.003>
- Xu, X. (2024). Research on Algorithmic Ethics in Artificial Intelligence. 2024 6th International Conference on Internet of Things, Automation and Artificial Intelligence (IoTAAI), 499–503. <https://doi.org/10.1109/IoTAAI62601.2024.10692746>
- Züger, T., & Asghari, H. (2024). Introduction to the special issue on AI systems for the public interest. *Internet Policy Review*, 13(3). <https://doi.org/10.14763/2024.3.1802>