

Ki beszél, amikor az MI ír?

Média jog és -etika új perspektívában I.¹

Bayer Judit
Budapesti Gazdaságtudományi Egyetem

A generatív mesterséges intelligenciát mind gyakrabban használják a professzionális médiaiparban is. Ez a cikk a generatív mesterséges intelligencia által vagy annak segítségével létrehozott tartalmakkal kapcsolatos főbb etikai és jogi kihívásokról nyújt áttekintést. Ismerteti az Európai mesterségesintelligencia-törvény vonatkozó szabályozását és az elmúlt években felbukkanó újságírói szabványokat. Részletesen kitér a *deepfake* szabályozására, a transzparencia, a hitelesség és a pontosság újságírói követelményeire és arra, hogy miként érinti az újságírókat a szerkesztőségekben már kikerülhetetlenül felbukkanó mesterséges intelligencia.

Kulcsszavak: ChatGPT, generatív mesterséges intelligencia, média, újságírás, MI-törvény

Who Speaks When AI Writes?

Rethinking Media, Responsibility and Rights I

The rapid integration of generative artificial intelligence into the media ecosystem marks a structural transformation of public communication. What began as algorithmic assistance in data-driven reporting—sports, finance, weather—has evolved into systems capable of producing persuasive, human-like text, images and audiovisual content across virtually any domain.

This paper situates generative AI (GenAI) within the legal and ethical architecture of journalism, highlighting that the technology destabilises foundational assumptions about authorship, accountability and the nature of expression itself. The paper relies on a graduated model of AI involvement, distinguishing between assistive, transformative and fully generative uses. This spectrum is not merely descriptive but also normative, as legal responsibility and ethical expectations intensify with the degree of machine autonomy. Particularly in fully generative contexts, questions of liability, copyright and editorial accountability become acute. The apparent autonomy of AI, however, remains partly illusory: human input persists at every stage, from training via prompting to editorial oversight.

The analysis identifies three interrelated areas of ethical challenge. The first is connected to journalistic ethics and consist of three elements: transparency that aims to preserve trust and accountability both in terms of training data and output disclosure; the journalistic standard of accuracy is systematically challenged by the heightened risk of misinformation, hallucinations and unverifiable sources, raising the duty of verification to a whole new level. The second area of ethical concern is bias and discrimination that persist within and get extrapolated by generative systems, reflecting structural inequalities embedded in training data. Finally, labour implications raise concerns about the displacement of journalists and the erosion of professional autonomy.

Keywords: AI Act, ChatGPT, generative AI, journalism, media

¹ Ez a tanulmány a következő angol nyelvű cikk alapján készült, annak bővített és frissített megjelenése: Bayer, J. (2024): Legal implications of using generative AI in the media. *Information & Communications Technology Law*, pp. 1–20. <https://doi.org/10.1080/13600834.2024.2352694>.

1. Bevezetés

A mesterségesintelligencia-technológia (MI) elkerülhetetlenül a nyilvános kommunikációs tér átalakulásához vezet (Hepp et al. 2023). Az algoritmikus tartalomkezelés már évek óta meghatározza az internethasználók információhoz való hozzáférését. Algoritmikus úton generált híradásokat is közlésszerűnek már közel egy évtizede olyan témákban, ahol strukturált adatok állnak rendelkezésre, mint például a sport, az időjárás vagy a pénzügyek (Diakopoulos 2019). Ma már egyre több más típusú tartalmat is generál az MI (Gruber é. n.). Ezen túlmenően az automatizált beszélgetőrobotok – mint például az Alexa, a Siri vagy a Google Assistant – befolyásolják használóikat a hírek kiválasztásában is, ezáltal napirend-meghatározó funkciót töltve be (Resendez et al. 2025). Mivel a generatív modellek egyre inkább a médiatartalom-előállítás szokásos résztvevőivé válnak, számos etikai és jogi kérdést kell rendezni. Az Európai Unió első jelentős, átfogó mesterségesintelligencia-szabályozása, az úgynevezett MI-törvény (2024/1689 Rendelet, a továbbiakban: MI-törvény) tisztázta ezek közül a legfontosabbakat, azonban a médiaszektorban további kérdések merülhetnek fel. E cikk a fókusz az újságírói médiára korlátozza, és nem terjeszti ki a könyvekre vagy a filmekre.

A médiajog interdiszciplináris terület, amely több jogágot is érint. Különösen a szellemi tulajdon, a kártérítési polgári jogi felelősség, a büntetőjogi felelősség és az audiovizuális médiára vonatkozó ágazati szabályok szerinti felelősség olyan terület, amely elkerülhetetlen, ha generatív mesterséges intelligenciát (a továbbiakban: GenAI) alkalmaznak a médiában. Ugyanakkor az újságírást etikai szabályok – önszabályozás és társszabályozás – is irányítják. Ezek részben szerkesztői vagy iparág-specifikus szakmai szabványok formájában, részben pedig a digitális szolgáltatásokról szóló törvény (DSA) által meghatározott kódexek formájában jelennek meg. Pillanatnyilag a dezinformációra vonatkozó magatartási kódex a legrelevánsabb, azonban folyamatban van további uniós kódexek megalkotása is: ilyen az MI-generált tartalmak jelölésére vonatkozó Generatív MI-kódex, amely azonban fejlesztőkre vonatkozik, és kevésbé az alkalmazókra.

Ez a cikk elsőként azt vizsgálja, hogy a média szabadsága és a véleménynyilvánítás szabadsága vonatkozik-e a kizárólag MI-termékekből álló médiatermékekre vagy médiumokra, majd az etikai következményekre, végül pedig a konkrét jogi kihívásokra tér rá. Nem vállalkozik átfogó normatív következtetések vagy javaslatok megfogalmazására, hanem vállaltan leíró jellegű. Mivel azonban a szabályozás rendkívül dinamikusan alakul, és számos értelmezési kérdés még nyitott, néhány kérdéses pontban az értelmezés céljával következtetéseket is megfogalmaz. Ilyen elsősorban a véleménynyilvánítás szabadságára vonatkozó rész, valamint a háromszintű modell, amely a szerzőség és a felelősség mivoltát érinti.

A szabályozás célpontjaként szolgáló mesterséges intelligencia olyan gyorsan fejlődött, hogy szó szerint mozgó célpontot állított a szabályozás elé. 2022 novemberében, amikor az MI-törvény tervezetének jogalkotási folyamata már előrehaladott stádiumban volt, váratlanul megjelent a ChatGPT, és több száz új irányt jelölt ki a jövőbeli fejlesztések, felhasználási lehetőségek és szabályozási kihívások terén. Rögtön ezt követően több különböző vállalat is közzétette saját nyelvi modelljét (ilyen volt például a BERT² vagy a DALL-E³). Nem ez volt a történelem első olyan mesterséges intelligenciája, amely ténylegesen átmert a Turing-teszten,⁴ de az első olyan, amely a szélesebb közönség számára is elérhetővé vált. A ChatGPT a világ egyik leggyorsabban terjedő alkalmazásává vált, a Newman Report (2026) szerint 2026-ban világszerte heti 800 millió felhasználója volt. Egyebek között az aktuális hírekkel kapcsolatban is keresnek általa információkat. Ma már ez és a hasonló chatbotok (kérésre) leihivatkozzák a weboldalt is, ahonnan merítik az információt, de az átkapcsolás még elenyésző mennyiségű. A Google több kattintást hoz a sajtónak, azonban ez is csökkenő tendenciát mutat: a következő három évben a tartalomszolgáltatók a közvetített forgalom (átkattintások) 43 százalékos csökkenésére számítanak. Ennek az a fő oka, hogy a Google a Gemini nevű szöveggeneráló mesterségesintelligencia-chatbot

2 A BERT egy nyílt forráskódú nyelvi modell, a név jelentése: Bidirectional Encoder Representations from Transformers.

3 A DALL-E az OpenAI által kifejlesztett képgeneráló mesterségesintelligencia-modell.

4 A ChatGPT volt a második csevegőrobot, amely megfelelt a Turing-teszten. <https://www.mlyearning.org/chatgpt-passes-turing-test/>. További információk a Turing-tesztről: <https://www.techtarget.com/searchenterpriseai/definition/Turing-test>.

segítségével összefoglalókat nyújt, amelyek most a keresések mintegy 10 százalékában jelennek meg legfelül, a keresési találatok felett (valójában maguk a keresési találatok nem is mindig férnek a képernyőre ilyen esetben, csak lejjebb görgetve ismerhetők meg).

A GenAI felhasználási lehetőségeinek köre megszámlálhatatlan, akár egyének, akár vállalatok, akár más MI-al-kalmazások számára, mivel széleskörű, általánosnak tekinthető képességeket birtokol. Kérésre tartalmat hoz létre: szöveget, képeket vagy videót, az adott alkalmazástól függően. A legáltalánosabban elérhető képessége a szövegalkotás, amely olyan meggyőzően fejlődött az elmúlt években, hogy a felhasználó hajlamos elfelejteni: a válaszok valószínűségszámítás útján születnek, nélkülözve az arisztotelészi logika struktúráját, következtetéseit vagy a kimeneteket alátámasztó érdemi háttérkutatásokat.

2. A generatív MI-modellek etikai és jogi következményei a médiában

2.1. A véleménynyilvánítás és a média szabadsága

A véleménynyilvánítás szabadsága minden, egy személy által kifejezett tartalom védelmét szolgálja. Emberi jogként elsősorban az egyénekre vonatkozik, azonban jogi személyek is élhetnek ezzel a joggal (mint például a Sunday Times kontra Egyesült Királyság, 6538/74. sz. 1980. április 26. ügy, valamint az Observer és Guardian kontra Egyesült Királyság, 13585/88. sz. 1991. november 26. ügy esetében). E jog célja egyrészt annak biztosítása, hogy az egyének kifejezhessék gondolataikat; ebben a tekintetben az emberi méltóságban gyökerezik (Sajó 2004). Ha egy publikációt teljes egészében a GenAI hoz létre, ez a cél nem tekinthető megvalósultnak. A véleménynyilvánítás szabadságának másik fő célja a demokratikus közbeszéd biztosítása (Meiklejohn 1948, Barendt 2007, Post 2011). E cél megvalósulásának értékelése némi értelmezést igényel. Bár az MI által generált tartalom hozzájárulhat a nyilvános vitához, kérdéses, hogy ez a demokrácia nevében történne-e, mivel a demokrácia alkotóelemei az emberi egyének. Megállapítható, hogy a véleménynyilvánítás szabadságának sem az első, sem a második célja nem hozható fel olyan tartalom védelmére, amelyet teljes egészében MI generált, emberi beavatkozás nélkül.

Ugyanakkor a sajtószabadság a médiaszolgáltatókra is vonatkozik, függetlenül attól, hogy azok jogi személyiséggel bírnak-e vagy sem (A német alaptörvény 5. cikke, 1. szakasz, 2. mondat, BVerfGE 62, 230., 243. sor). Ezért még abban az esetben is, ha egy sajtótermék száz százalékban GenAI által készült, az továbbra is a sajtószabadság védelme alatt állna. Így egy mesterséges intelligencia által generált sajtótermék a sajtószabadság védelmét élvezné, de a véleménynyilvánítási szabadságét nem (Fechner 2025). Ez relatíve alacsonyabbra tenné a beavatkozás küszöbét, ha egy MI által generált publikációval valakinek a jogait sértenék, vagy törvényt sértenek, azonban a „sajtó” címkéje birtokában nem maradna híján a védelemnek. Hasonlóképpen, ha egy platform mint szerződő fél a szolgáltatási feltételei alapján eltávolítana ilyen tartalmat, elvileg nem kellene betartania a DSA részletes eljárási biztosítékait (DSA 2. és 3. fejezete), mivel az eltávolítás nem sérti az emberi jogokat (Bayer 2022). Ennek a koncepciónak a gyenge pontja abban rejlik, hogy emberi beavatkozás nélkül lényegében nem lehetséges teljes mértékben GenAI révén tartalmat előállítani (ma még). Az alkotás tágabb kontextusában Tim W. Dornis fogalmaz meg hasonló nézőpontot, aki kiemeli az MI autonómiájának illúzióját. Az emberi hozzájárulás a folyamat minden szakaszában jelentős, és a teljes irányítást (eddig még) sosem ruházták át MI-rendszerre. Igaz ugyan, hogy már működik néhány „teljesen” MI által vezérelt rádióállomás, de ezeket is felügyelik emberek (McLane 2023). Meghatározzák a hangvételt, a témákat, a márkaidentitást, minőségellenőrzést végeznek, gondoskodnak a szerzői jogok és a médiajogszabályok betartásáról, folyamatosan tanítják és finomhangolják az MI-rendszereket, és végső soron krízishelyzetben be tudnak avatkozni (BPR 2023). Például az elvileg teljesen MI-vezérelt, úttörő német Absolut Radiónál eleinte átnézték az MI-generált szöveget, mielőtt adásba került (BLM 2024), utóbb viszont már csak a blokkok promptjait nézték át, és ezt akár háromnapos blokkokban is megtehették. A hitelességet azzal próbálták elősegíteni, hogy a kért szövegekhez előre megadták a forrásokat, és azt a parancsot is, hogy más forrást nem használhat az MI. Így csak a szöveget kellett a célközönséghez igazítani, új információ gyűjtését nem bízták az MI-re.

Annak a kérdésnek az eldöntése, hogy pontosan milyen mértékű emberi hozzájárulás szükséges a védelem küszöbértékének eléréséhez, a bíróságok vagy a jogalkotók feladata, de az etikai határok kialakításán jelenleg is dolgoznak az újságírószervezetek, amint azt alább bemutatom. A tartalomalkotás teljes automatizálása nem felel meg az etikus újságírás szabványainak. Ezért, ha az újságíró felelősségteljesen használja az MI-t, akkor az újságíró által hozzáadott értéknek elegendőnek kell lennie a véleménynyilvánítás szabadságához való jog igazolásához (Conraths 2023).

2.2. Generatív MI-modellek az MI-törvényben

Az általános célú modellek váratlan megjelenése kész helyzet elé állította az uniós jogalkotót. Az akkor már szinte kész Európai mesterségesintelligencia-rendelet a felhasználás célja szerint határozta meg, hogy egy MI-rendszer magas kockázati faktorú-e. Ez azonban nyilvánvalóan nem volt alkalmazható az általános célú modellekre, hiszen azok bármilyen felhasználású rendszer részét képezhetik (MI-törvény 6. cikk). Natali Helberger és Nicholas Diakopoulos azt javasolta, hogy a generatív mesterséges intelligenciának és az általános célú mesterséges intelligenciának saját, önálló kategóriát kellene adni, és lényegében ezt a megoldást követte a jogalkotó (Helberger & Diakopoulos 2023). Általános célú mesterségesintelligencia-modellnek nevezik azokat, amelyeket nagy adatmennyiséggel, nagy léptékű önfelügyelet mellett képeztek ki, amelyek képesek sokféle különböző feladat elvégzésére, és amelyek integrálhatók számos alkalmazási rendszerbe (MI-törvény 3. cikk (66) bekezdés).

A törvényben az általános modelleket alapesetben korlátozott kockázatú modellként sorolták be, kivéve, ha rendszerszintű kockázatot jelentenek. Ezek esetében a magas kockázati szintet a képességük alapján határozzák meg, amely képességet a műszaki paraméterekből vezetik le. Feltételezhető a magas képesség, ha a modell betanításához felhasznált számítási teljesítmény összege meghaladja az 1025 FLOP-ot (lebegőpontos műveletet). Ez a betanítás összetettségére utal, és nem magának a modellnek a számítási kapacitására (MI-törvény (111) preambulumbekkezdés, valamint Hacker 2023). Kétségtelen, hogy a küszöbérték módosítása csak idő kérdése, de a Bizottság felhatalmazást kapott arra, hogy e célból felhatalmazáson alapuló jogi aktusokat fogadjon el.

Az általános kockázatú általános célú MI-modellek szolgáltatóinak is van néhány alapkötelességük. Kötelesek biztosítani a műszaki dokumentáció átláthatóságát, ideértve a modell képzését és tesztelését is, és kérésre az MI-hivatal és a nemzeti hatóságok rendelkezésére kell ezeket bocsátaniuk. Továbbá naprakészen kell tartaniuk és rendelkezésre bocsátaniuk a rendszer dokumentációját azon szolgáltatók részére, amelyek az általános célú MI-modellt be kívánják építeni MI-rendszereikbe; valamint tiszteletben kell tartaniuk a szerzői jogokat, és végül részletes összefoglalót kell készíteniük a képzési adatok tartalmáról az MI-hivatal által biztosított sablon alapján. A nyílt forráskódú modellek mentesülnek az első kötelezettség alól, de nem a második és a harmadik alól.

Ezen túlmenően azok az általános célú MI-modellek, amelyek rendszerszintű kockázatot jelentenek, további kötelezettségeknek vannak alávetve: értékelniük kell modelljüket szabványosított protokollok – például támadó szempontú tesztelés – szerint a rendszerszintű kockázat azonosítása és mérséklése érdekében; értékelniük és enyhíteniük kell az uniós szintű rendszerszintű kockázatokat – ideális esetben egy magatartási kódex segítségével –; nyomon kell követniük, dokumentálniuk és jelenteniük kell minden súlyos incidenst, és meg kell tenniük a korrekciós intézkedéseket; valamint biztosítaniuk kell a kiberbiztonság és a fizikai biztonság megfelelő szintjét (MI-törvény 55. cikke).

A szöveget, hangot és videót önállóan létrehozni képes generatív MI-modellek az általános célú MI-modellek legnépszerűbb alkalmazásai; ezek képezik e cikk témáját is (MI-törvény 50. cikk 4–5).

2.3. A médiában használt generatív MI-modellek (GenAI)

A nyilvánosan elérhető GenAI modelleket már korábban is használták tartalomfejlesztésre és -előállításra (Gruber é. n.). Ugyanakkor ezek bevonása, valamint az ember és a mesterséges intelligencia közötti interakció mértéke korábban túlnyomórészt rejtve maradt, mivel nem létezett széles körben elfogadott szabvány arra vonatkozóan, hogy ezt az információt egyáltalán meg kell-e jelölni, és ha igen, hogyan. Az első nyilvános leleplezéseket követően

az újságíró-szövetségek és az Európa Tanács elveket és irányelveket dolgoztak ki a mesterséges intelligencia felelősségteljes újságírásban való etikus használatára vonatkozóan (Morin 2023, CDMSI 2023), emellett ilyen célú a Partnership on AI „A PAI felelősségteljes gyakorlata a szintetikus médiában” című összetett ajánlása is (PAI 2023).

Az MI-törvény minimálszabályokat alkotott, elsősorban a fejlesztőre nézve: a mesterségesen generált hang-, videó-, kép- vagy szövegtartalmakat géppel olvasható formátumban „vízjelekkel” kell ellátni, kivéve, ha az MI-hozzájárulás csupán segítő jellegű, de nem változtatja meg lényegesen a bemeneti adatokat vagy azok szemantikáját. A vízjelek elhelyezése a rendszer szolgáltatójának kötelezettsége, és minden GenAI-rendszerre vonatkozik. A médiaiparban való használatára további konkrét kötelezettségek vonatkoznak. A közérdekű ügyekről való tájékoztatás céljából az alkalmazóknak – például az újságíróknak vagy a kiadóknak – közzé kell tenniük, ha egy szöveget mesterségesen generáltak vagy manipuláltak, azonban egy lényeges kivétellel: ha az MI által generált tartalmat emberi felülvizsgálatnak vagy szerkesztői ellenőrzésnek vetették alá, és a tartalomért szerkesztői felelősséget visel egy emberi személy. Ez utóbbi esetben el lehet tekinteni attól, hogy a nyilvánosságot tájékoztassák az MI bevonásáról. Amint később bemutatom, ez az előírás nincsen tekintettel a közönség jogaira: a médiatudatosságnak fontos feltétele, az emberi méltóság érvényesülésének pedig követelménye, hogy a közönséget informálják arról, hogy amit olvas, azt MI generálta vagy pedig ember írta. Ez a megengedő szabály azonban csak az írásos tartalomra (vagyis a szövegre) vonatkozik. A generált kép-, hang- vagy videótartalom nem mentesül a közzétételi kötelezettség alól, még közérdekű célokból sem. Emellett, ha a tartalmat egy nyilvánvalóan művészi, kreatív, szatirikus, fiktív vagy hasonló mű részeként használják fel, akkor a transzparencianyilatkozatot úgy lehet közzétenni, hogy az ne akadályozza a mű élvezetét (MI-törvény 50. cikk (4) bek.).

A jogi elbírálás kérdései – mint például a felelősség, a szerzőség és a szerzői jog – az MI részvételének mértékétől függenek. Álláspontom szerint a részvétel mélységét három szakaszra kellene osztani. Az elsőben az MI-t arra utasítják, hogy nyelvtanilag javítsa vagy stilizálja az írott szöveget, vagy javítsa az emberi felhasználó által alkotott fotót. A második esetben az MI feladata az lesz, hogy parafrázálja, bővítse vagy összefoglalja a szöveget, vagy hangot generáljon a szövegből, vagy átalakítsa és fejlessze a képet vagy a videót. Az eredmény egy másik termék lenne, de továbbra is szorosan kapcsolódna az eredeti bemeneti adatokhoz. A harmadik szinten az MI önállóan generálná a kívánt mennyiségű és minőségű szöveget, képet, videót vagy hangtartalmat a felhasználó utasításai alapján. Az első eset egyértelműen mentesül a vízjel-kötelezettség alól. A harmadik szakasz ugyan csak egyértelmű: maximális transzparenciát igényel, tájékoztatási és vízjel-kötelezettség alá esik, és izgalmas felelősségi és szerzői jogi kérdéseket vet fel, amelyeket a tanulmány jövőben megjelentetni tervezett második része részletesebben is elemez majd. A köztes, második eset az, amelyik a leginkább vita tárgyát képezheti. Komoly etikai kifogásokat eredményezhet, ezért az a biztonságos megoldás, ha a harmadikkal egy kategóriába soroljuk. Ugyanakkor képek, videók vagy nem képességmérésre használt szövegek esetében mind általánosabban elfogadott ez a használat, mind az újságírás, mind az egyéb munkavégzés során.

3. A GenAI újságírói felhasználásának etikai követelményei

A professzionális média működését nem csupán írott joggal szabályozzák, hanem jelentős szerepet töltenek be a szakmai etikai normák is. Ezek ráadásul határokon átívelőek, valamint jogi jelentőségük is lehet, amennyiben tartalommal töltik meg a jogi kereteket, mint például a szakma általánosan elfogadott gondossági követelményei.

Az etikai problémák tekintetében négy fő kategória azonosítható, amelyek az általánosabbaktól a konkrét ágazati problémákig terjednek.

3.1. Átláthatóság

Eleve súlyos etikai problémákat okoz, hogy az általános célú MI-modellek tanítását nem átlátható módon gyűjtött adatokkal végezték, és az adatok köre sem ismert (Johannisbauer 2023). Natali Helberger és Nicholas Diakopoulos (2023) az „*extraction fairness*” (az adatkinyerés méltányossága) kifejezést javasolja, ami azt jelenti, hogy

a tanításra használt adatoknak mind a jogi, mind az etikai előírásoknak meg kell felelniük. További aggály, hogy a modellek működése nem megmagyarázható (Opdahl et al. 2023), ezért a hibáik jellege, gyakorisága jelenleg nem tisztázott, és nem előre látható (Bommasani et al. 2022). Környezeti fenntarthatóságuk is komoly aggodalomra ad okot (Crawford 2021). Mindez azonban nem médiaszpecifikus, ezért itt nem tárgyaljuk részletesebben.

Az átláthatóság hiányának egy további aspektusa az, amikor az MI használatát eltitkolják, és úgy állítják be, mintha a tartalmat egy emberi személy generálta volna. Az MI által generált tartalomnak a felhasználó által alkotott tartalomként való felmutatása egyértelműen súlyosan etikátlan minden olyan kontextusban, ahol egyéni teljesítményt várnak el, például oktatási környezetben. A munkahelyi környezet esetében már nem ilyen egyértelmű a kérdés, hanem attól függ, hogy a munkaerőnek mi teszi ki a legfőbb értékét, mi az a hozzáadott érték, amelyért őt alkalmazzák. Egy vállalati titkárnő esetében feltehető, hogy a megírt és elküldött levelek, feljegyzések, jelentések esetében a végeredmény a lényeg, és minél magasabb színvonalú az, annál jobb. Ezzel szemben egyes újságírók esetében az is jellemző lehet, hogy az újságíró sajátos stílusa és alapossága az a fő érték, amely miatt a munkáltató őt alkalmazza, és nem fogadná el, hogy MI-vel írja vagy változtassa meg a közleményét.

Emellett az MI-szabályozás hajnalából származik az emberi méltóság védelméből fakadó „azonosítás elve”, amely tiltja a felhasználók félrevezetését a tekintetben, hogy MI-vel vagy emberrel van-e dolguk (Council of Europe 2020). Ez egyúttal gyakorlati biztosítékként is szolgál, amikor tájékoztatja a címzettet arról is, hogy az erkölcsi elvárásoknak alacsonyabbnak vagy legalábbis különbözőnek kell lenniük. A tartalomgyártás kontextusában is ez jelenti az egyik legfőbb érvet a transzparencia mellett: a médiatudatosság és a dezinformációval szembeni ellenállóképesség növelésére irányuló erőfeszítések részeként a közönségnek fontos megkülönböztetnie a szintetikus és az újságírói tartalmat, mivel az MI-tartalom megbízhatósága eltérhet az ember által generált tartalométól (Longoni et al. 2022). Ezen kívül azonban az információs környezet hitelessége érdekében is fontos a közönségnek tudnia, hogy egy tartalom MI-generált vagy eredeti, nemcsak szövegek, hanem – még inkább – a képek és a videók esetében is. Ebből következik, hogy az újságíró-társadalom magára nézve jóval szigorúbb szabályt alkotott, mint amit az MI-törvény előír.

3.2. A tartalomra vonatkozó újságírói szabványok

A média kontextusában különösen jelentős a szakszerűtlen, káros vagy illegális tartalom létrehozásának kockázata. Természetesen a tájékoztatás képessé teszi a közönséget arra, hogy felismerje ezeket, ugyanakkor az újságírói etika körében eleve meg kellene akadályozni ilyen tartalmak keletkezését. A médiaipar szereplői aktívan dolgoznak azon, hogy kidolgozzák az MI tartalomalkotásban való felhasználására vonatkozó irányelveiket. Mind több médiaszolgáltató teszi közzé az MI használatára vonatkozó nyilatkozatát, amelyekben kijelenti, mely feladatokhoz fogja használni az MI-t, és melyeknél tartózkodik ettől (például a Wired közzétett egy részletes MI-irányelvet, lásd Wired 2023). Egyes beszámolók szerint azonban néhány szerkesztőség fél közzétenni MI-irányelveit, mert attól tart, hogy bizalmatlanságot ébreszt ezzel az olvasóban, akik szkeptikusok az MI használatával szemben (AJP 2025).

Emellett újságíró-szövetségek és agytrösztök is kínálnak irányelveket a GenAI etikus használatához (Amditis 2023). A Tudósítók Határok Nélkül (RSF) szövetség nemzetközi bizottságot hívott össze, amely 2023 novemberében létrehozta a Párizsi Chartát a mesterséges intelligenciáról és az újságírásról, amely az első globális irányelv volt a mesterséges intelligencia etikus használatára vonatkozóan a médiaiparban (RSF 2024). A charta tíz alapelvet fogalmaz meg, amelyek között szerepel egyebek mellett az is, hogy az emberi cselekvésnek továbbra is központi szerepet kell játszania a szerkesztői döntésekben, a társadalomnak segítséget kell nyújtani a hiteles és a szintetikus tartalmak biztos megkülönböztetésében, valamint felelősséget kell vállalni a globális mesterségesintelligencia-irányításban való részvételért és a technológiai vállalatokkal az újságírás nevében folytatott tárgyalásokért.

A szakmai szabványokról kialakuló konszenzus hangsúlyozza, hogy a GenAI újságírói eszköz lehet, de nem szerző. Gyakorlati szerepe hasonlóvá válhat a Wikipédiához: soha nem hiteles, de hasznos a lehetőségek bővítésében (Marconi 2020), és további kutatások és alkotások kiindulópontjaként szolgál (David 2023). Hasonlították a Google-hez és a Google Translate-hez is, amelyek zökkenőmentesen épültek be a mindennapi újságírói munkamenetbe (Israely 2023). Mindezek az etikai követelmények szervesen illeszkednek az újságírás hagyományos szakmai-etikai

elvi közé. Mindig is ilyen volt például az a követelmény, hogy a sajtó köteles ellenőrizni minden hír valóságát, tartalmát és forrását a terjesztés előtt, lehetőleg két független forrásból. Tartalmazza továbbá az információk etikus felhasználását, valamint a szerkesztőség és a közönség felé tanúsított átláthatóságot. Az Ethical Journalism Network (EJN) álláspontja szerint az újságírás öt alapelve a következő: igazság és pontosság, függetlenség, tisztesség és pártatlanság, emberségesség és elszámoltathatóság. Az EJN, amelyet 2012-ben alapított Aidan White, és több mint hetven tagja között az EFJ (Európai Újságírók Szövetsége) is megtalálható, maga is aláírta a Párizsi Chartát.

Ahogy a GenAI használata egyre inkább a mindennapi élet, így az újságírói munkának is szokásos részévé válik, az etikus használatra vonatkozó szabványok döntő fontosságúvá válnak a felelősség megállapításában is, amikor egy újság vagy újságíró jogellenes publikáció által okozott károkért való felelősségéről kell döntenet. Lehet, hogy az újságírói szabványokat ki kell bővíteni az MI által generált szövegek tényellenőrzési eljárásainak beépítésével, valamint az ilyen szövegek közzététele előtti etikai ellenőrzésekre vonatkozó irányelvekkel. Ezek magukban foglalhatják az MI-rendszereknek adott utasítások alaposabb áttekintését és felügyeletét.

Az American Journalism Project 2025-ös tanulmánya szerint a hírszerkesztőségek keményen dolgoznak az MI-alkalmazás belső szabályainak kialakításán, ezért ajánlott néhány alapvető elvet, amelyet minden szerkesztőség hasznosítani tud: az MI tiltott és megengedett használatának tisztázását a szerkesztőség számára; az MI használatának nyilvánosságra hozatalát szerkesztői tartalom előállításakor, valamint annak leszögezését, hogy nem fogják a szerkesztőségi munkatársakat MI-asszisztensekre cserélni (AJP).

New York állam éppen e kézirat lezárásakor terjesztett be egy törvényt az MI-használat követelményeiről. Ennek értelmében a New York-i hírszolgáltató szervezetek kötelesek az újságírók és szerkesztői alkalmazottak tudomására hozni, hogyan és mikor alkalmaznak MI-t a munkahelyen, valamint minden generált MI tartalmat a közzététel előtt emberi alkalmazottnak kell szerkesztői kontrollt gyakorolva ellenőriznie. A törvény kiterjed az újságírói forrásvédelemre is (Fahy 2026). Egyébiránt New York államban már 2025 decemberében elfogadtak egy törvényt, amely kötelezi a reklámkészítőket, hogy jelezzék a mesterséges előadók alkalmazását a közleményekben és a reklámanyagokban.

A Center for News, Technology and Innovation (CNTI) 188 nemzeti és regionális MI-szabályt és -törvényt vizsgált meg abból a szempontból, hogyan érintik az újságírást (Goodson et al. 2025). Hét olyan szabályozási elemet azonosított, amelyek különösen nagy hatással lehetnek az újságírásra. Honlapján egy interaktív térkép segítségével meg lehet tekinteni az átvizsgált politikai dokumentumokat, illetve azt, hogy azok milyen, újságíráshoz kapcsolódó témákat tartalmaznak. A lista elemei mind részét képezik az EU MI-törvényének is, ám valójában nem mind kapcsolódnak szoros értelemben az újságírók munkájához. Az algoritmikus diszkrimináció és előítéletesség például kétségtelenül az egyik legkritikusabb problématerület, de nem specifikus az újságírást szempontjából – hacsak annyiban nem, hogy a GenAI által „termelt” elfogult tartalmakat tömegméretekben terjesztve a professzionális média súlyosabb ártalmakat okozhat, mint az egyéni chatbotok önmagukban. A leggyakrabban szabályozott terület az átláthatóság és az adatvédelem, a legkevésbé pedig a szólásszabadság kérdése merül fel. Összefoglalóan: a kutatás eredményeként azt állapították meg, hogy az MI-szabályozások rendkívül heterogének, ezenkívül egyes országokban – főleg az Egyesült Államokban – százával nyújtják be a törvényeket, ami lehetetlenné teszi az átfogó kép alkotását (Goodson et al. 2025).

a) Téves információk. A tartalmi problémák körében talán a legszembetűnőbb etikai kockázat a téves és a hamis információk generálásának kockázata. Mind a téves, mind a hamis információk szöveg, kép, hang, videó vagy ezek kombinációja formájában jelenhetnek meg. Egyre nagyobb nyugtalanságot kelt, hogy a generatív modellek ma már megtévesztő minőségben, ráadásul pillanatok alatt és szinte költségek nélkül generálnak képeket és akár videókat is (Wach et al. 2023). A deepfake-eket az MI-törvény ugyanolyan feltételekkel szabályozza, mint a szöveges tartalom generálását, párhuzamosan használva a „hoz létre” vagy a „manipulál” kifejezéseket (MI-törvény 52. cikk (3) bek.). A deepfake-technológiák bünygyi szempontból is egyre relevánsabbá válnak; ezeknek a felhasználásoknak azonban a média kontextusában kevésbé van jelentőségük. Az Európai Unió a Nők elleni erőszakról szóló irányelvében arra kötelezi a tagállamokat, hogy büntetőjogi eszközökkel tiltsák az „intim vagy manipulált tartalom” beleegyezés nélküli megosztását (Directive 2024/1385 2024).

Téves információ akkor keletkezik, amikor a felhasználó nincsen tudatában annak, hogy a generált tartalom valótlan vagy félrevezető. Gyakran a „hallucinál” kifejezést használják arra az esetre, amikor egy nyelvi modell hihetően hang-

zó, de teljesen valótlan vagy értelmetlen információt vagy tartalmat talál ki (Alkaissi & McFarlane 2023). Hogy ez termékminőségi problémaként értelmezhető-e, az attól függ, hogyan határozzák meg a generatív program funkcióit és képességeit. Jelenleg az ilyen modellek készítői kizárják a felelősséget a tartalmi tévedésekért. Ha minden felhasználó számára világos, hogy ezek a szoftvertermékek csupán a szórend statisztikai valószínűsége alapján képesek szöveget generálni, akkor azt is be kell látniuk, hogy az olykor értelmetlen vagy valótlan szöveg a szoftver velejárója, nem pedig hiba. A téves információk felismerése és kiszűrése az emberi felügyelő, jelen esetben az újságíró feladata. Korábbi esetek – mint például a CNET és a Men's Journal weboldalai, ahol ténybeli hibákkal teli, MI-generált cikkeket tettek közzé anélkül, hogy felfedték volna az MI részvételét – ugyancsak rávilágítanak erre a problémára (Christian 2023).

Alapelv tehát, hogy a GenAI terméke nem forrás, hanem csak egy eszköz. A BBC szabályzata szerint úgy kell kezelni, mintha anonim tipp vagy meg nem erősített információ lenne. Ezért MI-vel írt szöveg nem jelenhet meg közvetlenül, változtatás nélkül; egy szerkesztőségi munkatárnak fel kell dolgoznia, át kell írnia az alapanyagot. Minden állítást vissza kell vezetni egy eredeti cikkre vagy forrásra, és ha ez nem lehetséges, akkor törölni kell a cikkből, nem jelenhet meg.

b) *Dezinformáció.* A dezinformáció generálása egy kissé eltérő jelenség. Ilyen esetben az alkalmazó szándéka arra irányul, hogy célzott promptokkal vagy hamis információkat betáplálva hamis vagy félrevezető információt generáljon (Wardle & Derakshan 2017). A fő veszély abban rejlik, hogy rendkívül gyorsan és könnyedén lehet előállítani többféle célcsoportra optimalizált dezinformációs tartalmat, majd algoritmusok segítségével terjeszteni vagy felerősíteni. A promptok tudatos és komplex használatával olyan magas színvonalú és sokszínű információs csomag hozható létre, amely akár egy egyébként ellenálló közönséget is félrevezethet. (Ezzel szemben lásd Felix M. Simon és munkatársai [2023] érvelését arról, hogy a félrevezető információk fogyasztását elsősorban a kereslet – és nem a kínálat – korlátozza, ezért túlzottak azok az aggodalmak, hogy a generatív mesterséges intelligencia hatással lesz a félrevezető és hamis információk terjedésére).

c) *Verifikáció.* A tények ellenőrzése ma már szükséges tevékenysége minden felelős szerkesztőségnek. Az MI-korszakban azonban ez tovább differenciálódik. Nem csupán a tényeket kell ellenőrizni, hanem a forrást is verifikálni kell, hogy az valóban létező, visszakereshető, és nem MI-generált, tehát megbízható. Ezzel összefüggésben pedig verifikálni kell magát a folyamatot, és feltárni, hogy mennyi MI-használat történt az előállítás során, és hogy szükséges-e azt feltárni a nyilvánosság előtt.

3.3. Diszkrimináció

Egy harmadik etikai probléma az lehet, hogy a GenAI, más MI-khez hasonlóan, képes megerősíteni a meglévő társadalmi előítéleteket, és előfordulhat, hogy mérgező, ellenséges tartalmakat hoz létre – habár a betanítás erőteljes erőfeszítéseket tesz ennek elkerülésére. Opdahl és társai (2023) álláspontjával ellentétben én ezt a kettőt szorosan összefüggőnek tartom. A betanítást követően, az életciklus további szakaszaiban is javítható ez emberi felügyelettel: a fejlesztés, a finomhangolás és az alkalmazás során, végső soron a létrehozott tartalom emberi felülvizsgálatával. A mérgező tartalom (nem illegális gyűlöletbeszéd) problémája hasonló a félrevezető információk problémájához, egy jelentős különbséggel. A modell előítéletességre való hajlamát a fejlesztési és képzési fázisban lehet és kell kezelni, míg a dezinformációt ilyen módon nem lehet kizárni. Ez a követelmény, valamint a szigorú emberi felügyelet azonban csak a magas kockázatú alkalmazások esetében kötelező. Mindazonáltal a rendszerszintű kockázatokkal bíró általános célú modellek esetében kötelező a kockázatok értékelése és mérséklése a magatartási kódexre támaszkodva, ami az első lépés a probléma kezelésében (MI-törvény 55. cikk (1) bek. b) pont).

3.4. Munkaügyi kapcsolatok

A szakmai normáktól némileg eltérő jellegű az itt negyedikként tárgyalt etikai kérdés: a munkakörnyezet. Az újságírók védelme hagyományosan részét képezi a sajtószabadság kereteinek, túlnyomórészt nem a formális, hanem a puha jog területén alkotott elvek révén. Jelenleg az a legfőbb aggodalma az újságíró-társadalomnak,

hogy az újságírókat egyszerűen kiváltják a generatív MI-vel. Ezért az IPI 2025-ben arról foglalt állást, hogy az ember általi újságírást kell előnyben részesíteni (Humphries 2025). Újabb iparági sztenderdjei között az EFJ is az újságírók védelmére fókuszál elsősorban (EFJ 2025), az IFJ pedig még ennél is erőteljesebben állást foglal az érdekükben: hozzáteszi, hogy az újságíróknak részt kell venniük az MI szabályozásának kialakításában, és azt ajánlja, hogy a szerkesztőségek kollektív szerződést kössenek az MI használatának szabályozására (IFJ 2025).

Az Európa Tanács MI és újságírás irányelvei már 2023-ban összefoglalták a főbb elveket az újságírásban történő felelősségteljes mesterséges intelligenciarendszer-bevezetés körében (CDMSI 2023), és az említett szerkesztőségi irányelvek – más médiavállalkozásokéval együtt – valójában ezeket követik. Ezek strukturális szinten közelítették meg a kérdést, támogatást nyújtva a szervezeteknek a mesterséges intelligencia szakmai célú, tisztességes és felelősségteljes telepítésében és bevezetésében. A kérdés jogi problémaként is feltehető: jogosult-e, köteles-e vagy tilos-e az újságírónak a GenAI-t használnia a munkája során? A munkáltatónak minden bizonnyal joga van arra, hogy utasítsa az újságírót egy bizonyos eszköz vagy szoftver használatára (Kalbfus & Schöberle 2023). Mindazonáltal a kiadók politikája érvényesen megtagadhatja ezt, vagy szabályozhatja a használatát a fent leírt etikai normák szerint, lásd különösen az AP és a Wired irányelveit fentebb. A munkaügyi kérdések közé tartozik ezen kívül az újságírók rendszerhasználatra való felkészítése, képzése is. Ezek az elvek azonban jelenleg inkább tekinthetők az újságírói szakma normatív szándékú követeléseinek, mint érvényesíthető etikai elveknek.

4. Összefoglalás

Az iparági trendekből a következő irányelvek rajzolódnak ki:

1. Általánosan elfogadott a transzparencia követelménye, méghozzá szigorúbban annál, mint amit az uniós jogszabály megkövetel, mert minden MI-generált tartalomra kiterjed, akkor is, ha fennáll a szerkesztői kontroll. Ugyanakkor „szabadon” használhatnak MI-t az írás, a szerkesztés, a gyártás és a közzététel során, a közlési kötelezettség pedig csak a generált tartalomra terjed ki (Reuters and AI 2024), habár általában nincsen egyértelmű elhatárolás a parafrázálás, az összegzés és a generálás között. A szigorúbb iparági sztenderdek értelmében az újságírók nem állíthatnak elő közvetlenül publikálendő tartalmat a GenAI segítségével (Meir 2023). Kivételt képeznek a kevésbé kreatív tartalmak, mint a pénzügyi jelentések, az időjárásjelentés, a sporthírek, a bal-esetekről szóló híradások és hasonlóak (AP 2023).

2. Végző szerkesztői döntés szükséges ember részéről, tehát az MI nem tekinthető autonóm szerzőnek, és nem dönthet publikálásról.

3. Az újságíró és a médiavállalat megtartja az etikai és a jogi felelősséget. Ez magában foglalja a felelősséget az MI által szállított tények ellenőrzéséért, az elfogultság kiszűréséért, valamint a szerzői és a személyiségi jogok védelméért mind a promptként feltöltött információ, mind a generált információ esetében.

4. Végül az újságírók védelme mind munkajogi szempontból, mind pedig munkájuk védelme szerzői jogi szempontból. A munkajogi megfontolások keretében az irányelvek javaslatokat tartalmaznak arra nézve, hogy a szerkesztőségek deklarálják: nem bocsátanak el munkavállalót azért, hogy MI-vel helyettesítsék, pontosabban nem elfogadható az MI-re való hivatkozás az elbocsátások igazolására. Más szempontból pedig a dolgozók felé is legyenek transzparens az MI-használat módjai. Az újságírók műveiket pedig csak bizonyos feltételek teljesítése esetén engedik tréninghez használni, mint minden alkotói művet. (Ezekről a kérdésekről részletesebben fog szólni e cikk második része, amelyet a Médiakutató következő lapszámában tervezek közzéadni.)

A jogterület továbbra is a fejlődés stádiumában van, hiszen a technológia is folyamatosan fejlődik, és új meg új kérdéseket vet fel. Az is kirajzolódik, hogy az írás, a szerkesztés, a gyártás és a közzététel során igenis használnak MI-t, és csak akkor közlik ennek tényét a nyilvánossággal, ha elsősorban vagy kizárólag az MI-re támaszkodva állítanak elő hírértékű tartalmat.

(A következő lapszámában megjelenő második részben két, a mesterséges intelligencia által generált tartalommal kapcsolatos jogi kihívást tervezek tárgyalni: a tartalomért való felelősséget és a szerzői jogi kérdéseket.)

Felhasznált irodalom

- AJP, American Journalism Project (2025): Developing an AI Usage Policy in Your News Organization, <https://www.theajp.org/news-insights/insights/developing-an-ai-usage-policy-in-your-news-organization/>
- Alkaissi, Hussam & Samy I. McFarlane (2023): Artificial Hallucinations in ChatGPT: Implications in Scientific Writing. *Cureus*, vol. 15, no. 2, <https://doi.org/10.7759/cureus.35179>
- Amditis, Joe (2023): AI Throws a Lifeline to Local Publishers, <https://www.niemanlab.org/2022/12/ai-throws-a-lifeline-to-local-publishers/>
- AP (2023): AP, Open AI Agree to Share Select News Content and Technology in New Collaboration, <https://www.ap.org/media-center/press-releases/2023/ap-open-ai-agree-to-share-select-news-content-and-technology-in-new-collaboration/>
- Barendt, Eric (2007): *Freedom of Speech*. Oxford University Press, <https://doi.org/10.1093/acprof:oso/9780199225811.001.0001>
- Bayer, Judit (2022): Procedural Rights as Safeguard for Human Rights in Platform Regulation. *Policy & Internet*, vol. 14, no. 4, pp. 755–771, <https://doi.org/10.1002/poi3.298>
- BLM (2024): BLM Approves Germany's First AI Radio Station – Absolut Radio AI – Hosted by Artificial Intelligence, Managed by Humans, <https://www.worlddab.org/news/14545/blm-approves-germany%27s-first-ai-radio-station-absolut-radio-ai-hosted-by-artificial-intelligence-managed-by-humans>
- Bommasani, Rishi, Drew A. Hudson, Ehsan Adeli, Russ Altman & Simran Arora (2022): *On the Opportunities and Risks of Foundation Models*. Stanford University, Center for Research on Foundation Models (CRFM).
- BPR (2023): The World's First Totally A.I. Powered Radio Station Could Arrive Soon, <https://bprworld.com/news/the-worlds-first-totally-a-i-ota-radio-station-could-arrive-soon/>
- CDMSI (2023): *Guidelines on the Responsible Implementation of AI in Journalism*. Council of Europe, <https://rm.coe.int/cdmsi-2023-014-guidelines-on-the-responsible-implementation-of-artific/1680adb4c6>
- Christian, Jon (2023): CNET's Article-Writing AI Is Already Publishing Very Dumb Errors. *Futurism*, <https://futurism.com/cnet-ai-errors>
- Conraths, Timo (2023): Wege zum Roboterjournalismus? – Wie verändert KI die journalistische Arbeit? [Utak a robotújságíráshoz? – Hogyan változtatja meg a mesterséges intelligencia az újságírói munkát?] *Zeitschrift für Urheberrecht*, vol. 67, no. 8–9, p. 574.
- Council of Europe (2020): *CAHAI Feasibility Study*. CAHAI(2020)23. Council of Europe, <https://rm.coe.int/cahai-2020-23-final-eng-feasibility-study-/1680a0c6da>
- Crawford, Kate (2021): *The Atlas of AI: Power, Politics and the Planetary Costs of Artificial Intelligence*. Yale University Press, <https://doi.org/10.12987/9780300252392>
- David, Emilia (2023): The Associated Press Sets AI Guidelines for Journalists. *The Verge*, <https://www.theverge.com/2023/8/16/23834586/associated-press-ai-guidelines-journalists-openai>
- Diakopoulos, Nicholas (2019): *Automating the News*. Harvard University Press, <https://doi.org/10.4159/9780674239302>
- Directive 2024/1385 on Combating Violence against Women and Domestic Violence.
- Dornis, Tim W. (2021): Die “Schöpfung ohne Schöpfer” – Klarstellungen zur “KI-Autonomie” im Urheber- und Patentrecht [„Teremtő nélküli teremtés” – pontosítások az „AI-autonómia” kérdéséről a szerzői jog és a szabadalmi jog területén, GRUR (Ipari jogi védelem és szerzői jog)], *GRUR (Gewerblicher Rechtsschutz und Urheberrecht)*, vol. 123, no. 6, pp. 784–792.
- EFJ (2025): EFJ Publishes Its Position on Artificial Intelligence in Journalism, <https://europeanjournalists.org/blog/2025/09/30/efj-publishes-its-position-on-artificial-intelligence-in-journalism/>
- Fahy, Patricia (2026): Fahy, Rozic Introduce NY FAIR NEWS Act to Protect Journalists and the Public from Artificial Intelligence Jeopardizing News Reporting, <https://www.nysenate.gov/newsroom/press-releases/2026/patricia-fahy/fahy-rozic-introduce-ny-fair-news-act-protect>
- Fechner, Frank (2025): *Medienrecht [Médiatörvény]*. C. F. Müller.

- Goodson, Kayla, Jay Barchas-Lichtenstein, Samuel Jens, Emily Wright & Utsav Gandhi (2025): Journalism's New Frontier: An Analysis of Global AI Policy Proposals and Their Impacts on Journalism. CNTI, <https://cnti.org/reports/journalisms-new-frontier-an-analysis-of-global-ai-policy-proposals-and-their-impacts-on-journalism/>
- Gruber, Barbara (é. n.): Facts, Fakes and Figures: How AI Is Influencing Journalism, <https://www.goethe.de/prj/k40/en/lan/aij.html>
- Hacker, Philipp (2023): What's Missing from the EU AI Act, <https://doi.org/10.59704/3f4921d4a3fbeeec>
- Helberger, Natali & Nicholas Diakopoulos (2023): ChatGPT and the AI Act. *Internet Policy Review*, vol. 12, no. 1, <https://doi.org/10.14763/2023.1.1682>
- Hepp, Andreas, Wiebke Loosen, Stephan Dreyer, Juliane Jarke, Sigrid Kannengießer, Christian Katzenbach, Rainer Malaka, Michaela Pfadenhauer, Cornelius Puschmann & Wolfgang Schulz (2023): ChatGPT, LaMDA, and the Hype Around Communicative AI: The Automation of Communication as a Field of Research in Media and Communication Studies. *Human-Machine Communication*, vol. 6, no. 1, pp. 41–63, <https://doi.org/10.30658/hmc.6.4>
- Humphries, Rowan (2025): IPI General Assembly Resolution: In the Age of AI, Human-Made Journalism Must Be Prioritized and Protected. IPI, <https://ipi.media/ipi-general-assembly-resolution-in-the-age-of-ai-human-made-journalism-must-be-prioritized-and-protected/>
- IFJ (2025): IFJ Recommendations on the Use of Artificial Intelligence, <https://www.ifj.org/media-centre/news/detail/category/ai/article/ifj-recommendations-on-the-use-of-artificial-intelligence>
- Israely, Jeff (2023): How Will Journalists Use ChatGPT? Clues from a Newsroom That's Been Using AI for Years, <https://www.niemanlab.org/2023/03/how-will-journalists-use-chatgpt-clues-from-a-newsroom-thats-been-using-ai-for-years/>
- Johannisbauer, Ch. (2023): ChatGPT im Rechtsbereich. *MMR-Aktuell*, no. 455537.
- Kalbfus, Michael & Andreas Schöberle (2023): Arbeitsrechtliche Fragen beim Einsatz von KI am Beispiel von ChatGPT [Munkajogi kérdések mesterséges intelligencia használatával kapcsolatban, a ChatGPT példája]. *ArbRAktuell*, no. 251.
- Longoni, Chiara, Andrey Fradkin, Luca Cian & Gordon Pennycook (2022): News from Generative Artificial Intelligence Is Believed Less. In: *Proceedings of the 2022 ACM Conference on Fairness, Accountability and Transparency*, pp. 97–106. ACM, <https://doi.org/10.1145/3531146.3533077>
- Marconi, Francesco (2020): *Newsmakers: Artificial Intelligence and the Future of Journalism*. Columbia University Press, <https://doi.org/10.7312/marc19136>
- McLane, Paul (2023): Absolut Radio Gets a Head Start with AI. Radio World, <https://www.radioworld.com/news-and-business/news-makers/absolut-radio-gets-a-head-start-with-ai>
- Meiklejohn, Alexander (1948): Everything Worth Saying Should Be Said. *The New York Times*, 18 July 1948, <https://www.nytimes.com/1948/07/18/archives/everything-worth-saying-should-be-said-an-educator-says-we-talk-of.html>
- Meir, Nicole (2023): Standards around Generative AI. AP, <https://www.ap.org/the-definitive-source/behind-the-news/standards-around-generative-ai/>
- Morin, Olivier (2023): *AI Act: Journalists and Creative Workers Call for a Human-Centric Approach to Regulating AI*. European Federation of Journalists, <https://europeanjournalists.org/blog/2023/09/26/ai-act-journalists-and-creative-workers-call-for-a-human-centric-approach-to-regulating-ai/>
- Newman, Nic (2026): *Journalism and Technology Trends and Predictions 2026*. Reuters Institute for the Study of Journalism, <https://doi.org/10.60625/RISJ-PS1D-NP11>
- Opdahl, Andreas L., Bjørnar Tessem, Duc-Tien Dang-Nguyen, Enrico Motta, Vinay Setty, Eivind Throndsen, Are Tverberg & Christoph Trattner (2023): Trustworthy Journalism through AI. *Data & Knowledge Engineering*, vol. 146, no. 102182, <https://doi.org/10.1016/j.datak.2023.102182>
- PAI (Partnership on AI) (2023): Responsible Practices for Synthetic Media – A Framework for Collective Action, <https://syntheticmedia.partnershiponai.org/>
- Post, Robert C. (2011): Participatory Democracy and Free Speech. *Virginia Law Review*, vol. 97.

Resendez, Valeria, Theo Araujo, Natali Helberger & Claes de Vreese (2025): Hey Google, What Is in the News? The Influence of Conversational Agents on Issue Salience. *Digital Journalism*, vol. 13, no. 4, pp. 774–796, <https://doi.org/10.1080/21670811.2023.2234953>

Reuters (2024): Reuters and AI, <https://www.reuters.com/info-pages/reuters-and-ai/>

RSF (2024): Paris Charter on AI and Journalism, <https://rsf.org/en/paris-charter-ai-and-journalism>

Sajó, András (2004): *Freedom of Expression*. Institute of Public Affairs.

Simon, Felix M., Sacha Altay & Hugo Mercier (2023): Misinformation Reloaded? Fears about the Impact of Generative AI on Misinformation Are Overblown. *Harvard Kennedy School Misinformation Review*, <https://doi.org/10.37016/mr-2020-127>

Wach, Krzysztof, Cong Doanh Duong, Joanna Ejdyś, Rūta Kazlauskaitė, Paweł Korzyński, Grzegorz Mazurek, Joanna Paliszkievicz & Ewa Ziemia (2023): The Dark Side of Generative Artificial Intelligence: A Critical Analysis of Controversies and Risks of ChatGPT. *Entrepreneurial Business and Economics Review*, vol. 11, no. 2, pp. 7–30, <https://doi.org/10.15678/EBER.2023.110201>

Wardle, Claire & Hossein Derakhshan (2017): *Information Disorder: Toward an Interdisciplinary Framework for Research and Policy Making*. Council of Europe, <https://edoc.coe.int/en/media/7495-information-disorder-toward-an-interdisciplinary-framework-for-research-and-policy-making.html>

Wired (2023): How WIRED Will Use Generative AI Tools, <https://www.wired.com/about/generative-ai-policy/>

Jogforrások

Az Európai Parlament és a Tanács (EU) 2024/1689 rendelete (2024. június 13.) a mesterséges intelligenciára vonatkozó harmonizált szabályok megállapításáról (MI-törvény).

Az Európai Parlament és a Tanács 2022. október 19-i (EU) 2022/2065 rendelete a digitális szolgáltatások egységes piacáról és a 2000/31/EK irányelv módosításáról (DSA).

Német Szövetségi Köztársaság Alaptörvénye.

Bíróági döntések

Sunday Times kontra Egyesült Királyság, 6538/74. sz. ügy, 1980. április 26.

Observer és Guardian kontra Egyesült Királyság, 13585/88. sz. ügy, 1991. november 26.

BVerfGE 62, 230. – Német Alkotmánybíróság 1982. évi határozata.