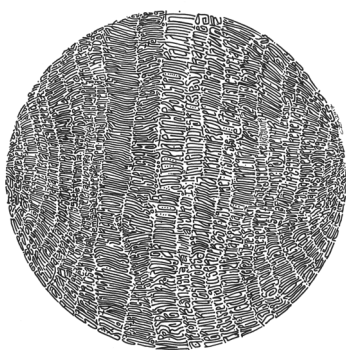


SZÉKELY JÁNOS – VERESS EMŐD

# A MESTERSÉGES INTELLIGENCIA ÉS A JOGALKALMAZÁS



**A mesterséges intelligencia megjelenése a társadalom konfliktusait nem fogja megszüntetni. Egyetlen technológia sem képes ilyen hatást kiváltani, a konfliktus léte az emberi társadalom szükséges velejárója.**

**A**sajtó és a hírfogyasztó közönség fantáziáját egyre inkább leköti a mesterséges intelligencia kérdése és e technológiai vívmány várható jövőbeni hatásai. Természetesen olykor napvilágot látnak visszafogottabb vélemények is, amelyek szerint a mesterséges intelligencia (rövidítve: „MI”; angol rövidítéssel: AI – azaz „artificial intelligence”) virágkora még a messzi jövőbe vész. Sőt felvetődött<sup>1</sup> a lehetőség, hogy e kérdésben a technológiai optimizmus mai, az 1950-es évek végét idéző fellángolása ez alkalommal is elhamarkodottnak bizonyul majd, a keserű csalódás mellett a további kutatások finanszírozásának szűkülését is előidézve. Elképzelhető az is, hogy a mesterséges intelligencia kiteljesedése – a magfúziós energiatermeléshez hasonlóan – a jövő ígéretes technológiája lesz, csak mindörökké az is marad... De ez a kevésbé valószínű forgatókönyv.

Az önvezető járművek és az általuk okozott súlyos balesetek (egy konkrét esetben mert a gépkocsi „úgy döntött”<sup>2</sup>, figyelmen kívül hagyja az általa észlelt gyalogost), ezekkel pedig a polgári jog szempontjából is jelentős felelősségi kérdések már időszzerűvé váltak, ahogyan az óriási adatbázisokból a személy legbensősebb gondolatai, akárcsak a saját maga számára is ismeretlen viselkedési mintázatai kifürkészésére, esetleg gazdasági vagy politikai kiaknázásra és sok minden másra alkalmas, gyakran a szakemberek számára sem átlátható működésű algoritmusok<sup>3</sup> is. A képfelismerés, a testtartásalapú azonosítás,<sup>4</sup> az arcfelismerés,<sup>5</sup> a hangfelis-

merés, a szem mozgásának (impliciten a személy figyelmének) követésére alkalmas rendszerek, a földrajzi helyzetmeghatározásra<sup>6</sup> és más, minden korábbinál szorosabb megfigyelésre használható technológiák már ma is komoly etikai és normaalkotási kérdéseket vetnek fel. A tömeges automatizált adatgyűjtés és -feldolgozás – legyen célja a reklám, a közbiztonságra veszélyes személyek kiszűrése vagy egyes önkényuralmi rendszerek fenntartása – igazán a mesterséges intelligencia fejlődésével válhat értelmessé, ugyanis ez utóbbi technológia az, amely a felhalmozott adattömegből kiolvasható összefüggések felfedezésére leginkább alkalmas.

Emiatt a nyilvános és titkos kutatásokhoz gyakorlatilag végtelen erőforrások<sup>7</sup> állnak rendelkezésre az MI fejlesztésére. Végső soron e technológia a termelékenység soha nem látott növekedéséhez, a közbiztonság szinte utópisztikus javulásához vagy a társadalmi kontroll minden korábbinál hatékonyabb megvalósításához, egy minden ediginél totálisabb totalitarizmushoz egyaránt vezethet, így nincs olyan állam (demokratikus vagy más berendezkedéssel rendelkező), sem nagyvállalat, ami ne tenne meg minden tőle telhetőt, hogy a mesterséges intelligencia gyümölcseihez hozzáférjen.

Nem meglepő tehát, hogy az MI egyre szélesebb körben foglalkoztatja államok, egyének és különböző gazdasági, rendészeti, katonai vagy éppen szakmai szervezetek, újságírók és jogalkotók képzelőerejét. A jogi szaksajtóra korlátozódva, csupán az elmúlt időszakban számos elmélkedés látott napvilágot a mesterséges intelligencia és az igazságszolgáltatás kölcsönhatásai kérdésében. A gépi döntéshozatal jogi kereteinek szabályozása és az európai térségben az e téren szükséges együttműködés, jogharmonizáció<sup>8</sup> és annak hiánya,<sup>9</sup> ahogyan a mesterségesen intelligens automatizált folyamatok által hozott döntések és az általuk működtetett gépek (üzemben tartóknak) esetleges büntetőjogi felelőssége<sup>10</sup> látszólag elméleti kérdései már most jogalkotói megoldás után kiáltanak. Ellentétben a genetika fejlődése által keltett, majdnem pánikszerű jogszabályalkotási hullámmal, a mesterséges intelligencia fejlődésével szemben a jogalkotók világszerte tanácstalanul állnak.

A jövő fürkészőinek egyik legjelentősebb felvetése, hogy a mesterséges intelligencia az igazságszolgáltatásban is felhasználható lesz, így esetleg egyes jogi hivatásrendek – elsősorban és mintegy „közkívánatra” – az ügyvédek, esetleg a bírók helyettesítését eredményezheti majd. Az ilyen nézetek egyik legfőbb hirdetője Richard Susskind, aki beszédes című írásaiban (*Az ügyvédség vége?*<sup>11</sup> vagy *Tomorrow's Lawyers: An Introduction to Your Future* [A holnap ügyvédei: Bevezető a jövődhöz]) firtatja a mesterséges intelligencia potenciális hatásait a jogi hivatásrendekre.

Egyik előadásában<sup>12</sup> összefoglalt álláspontja szerint már a 2020-as évek során a jogi szakmák tevékenységében az automatizáció és informatizáció a mainál is komolyabb jelentőséggel bír majd, más faktorok mellett amiatt is, mert küszöbön áll a számítástechnika újabb forradalma. Ennek során a számítógépek adatfeldolgozási sebességének és a rendelkezésre álló adattömegnek köszönhetően az alkalmazandó jogszabályok, sőt a jogviták kimenetelének megjósolása (ha nem is megoldása) tekintetében a gépi kapacitások az emberi képességeket elérhetik, azokat akár meg is haladhatják. Az okiratszerkesztés, az egyes jogvitákra alkalmazandó jogszabályok azonosítása, a megkötött jogügyletek átvilágítása mind szabványosítható, gépesített folyamatokká alakulnak, ezzel megfosztva állásától a kevésbé szakosodott jogászokat. Figyelemre méltó a szerző azon nézete is, miszerint ezzel párhuzamosan a magasan képzett jogászai szolgáltatások (pl. szakjogi tanácsadás) piaca is visszaszorul, és a nagyközönség számára nyíltan hozzáférhető jogi ismeretszolgáltatókkal (szako-

sodott online platformokkal) egészül ki. Susskind kutatásai, amelyek nem csupán az ügyvédi, hanem az orvosi, közjegyzői, mérlegképes könyvelői és más szellemi szabad szakmák tekintetében is mérvadó, mégsem mutatják ki egyértelműen azok művelőinek küszöbön álló gépi helyettesítését, csupán arra utalnak, hogy a jogi szakmák egyes jellemző tevékenységei terén az eddig emberi közreműködést is igénylő egyes munkafolyamatok várhatóan gépesítetté válnak. E fejlődési irány és potenciális hatása a jogi szolgáltatásokra már Romániából is<sup>13</sup> látható. A romániai Ügyvédi Kamarák Országos Egyesületének korábbi elnöke kifejezetten borúlátóan nyilatkozott<sup>14</sup> a mesterséges intelligencia által jelentett kockázatokról az igazságszolgáltatás számára.

Az e téren egyre komorabb jövőt előrevetítő elmélkedések a jogásztársadalmat annak kapcsán is foglalkoztatják, lehetséges-e a humán faktor (emberi bíró) géppel helyettesítése a jogviták megoldása során.

Mint a mesterséges intelligencia kapcsán oly sok más kérdés, ez is látszólag a valódságtól elrugaszkodott filozófiai elmélkedésnek tűnik, de valójában egy gyakran felmerülő súlyos probléma, a részrehajlás, az elfogultság leküzdésének technológiai lehetőségeit fürkészi. Ahogyan a labdarúgásban a videóbíráskodás bevezetésétől igazságosabb döntéseket remélt a nagyközönség (és ehelyett a bírói szigor nőtt<sup>15</sup> meg), úgy a bírók technológiai helyettesítése is (pl. Richard Susskind álláspontja szerint) egy olcsó, gyors és mindenekelőtt pártatlan, megvesztegethetetlen és elfogulatlan igazságszolgáltatás megvalósításához vezethet.

Az utópisztikus képzelgések e kérdésben két jelentős problémát vetnek fel: 1. az igazságszolgáltató mesterséges intelligencia jó eséllyel nem az emberi kogníció eszközeit fogja bevetni a jogvita megoldására, helyettük átláthatatlan informatikai algoritmusok vezetnek majd el a gépet döntéshez és 2. a patikamérleggel osztott, matematizált igazság, az ún. elfogultság teljes hiánya konkrét esetekben súlyos, az emberi igazságszolgáltatással és jogi tudattal egyaránt összeférhetetlen igazságtalanságokat szülhet.

1. A mesterséges intelligencia most ismert leghatékonyabb implementációja az ún. gépi tanulás (machine learning) módszere, ezen belül pedig a mélytanulási (deep learning)<sup>16</sup> algoritmusok. Az MI azonban nem azonos a gépi tanulás fogalmával,<sup>17</sup> hanem annál sokkal tágabb. Míg az „intelligencia” a rugalmas, kreatív gépi problémamegoldás jelentéstartalmára is utal, amit szokás *általános* mesterséges intelligenciának nevezni – ehhez a tudomány ma sem áll sokkal közelebb, mint 50 éve –, addig a gépi tanulás egy algoritmus működésére korlátozódik, ami nagy, de véges számú ismétlődő művelet gyors végrehajtásával bizonyos összefüggések felfedezésére alkalmas egy adott adattömegben. Ezáltal intelligens válaszokat produkál szűken definiált kérdésekre, például két fénykép tartalmának azonosságára nézve.

Minél nagyobb adattömegben alapul a gépi tanulás, minél több példa hozzáférhető számára a keresett adatokból, az algoritmus annál precízebben „tanulja meg” elhatárolni az adattömeg azon elemeit, amelyeket általa ki szeretnénk szűrni a többiek közül. A gépi tanulási algoritmusok működéséhez elsőnek tehát be kell tanítani a programot, meg kell ismertetni vele a lényegi kritériumokat, amiket az utóbb neki szolgáltatott adatokban fel kell ismernie. A gépi tanulás működése hasonló az elsőosztályos gyerek okulásához a szépírásgyakorlatok során: a tanító lerajzolja számára az elsajátítandó betűt, ezután a kisiskolás igyekszik azt legjobb tudása szerint utánozni, addig másolva egymásután az adott írásjelet, amíg annak írását kellő bizonyossággal el nem sajátítja.

Emiatt gépi tanulás útján a mesterségesen intelligensnek mondott gép csupán olyasmit képes megtanulni és utóbb felismerni, amire már számos példát szolgáltatunk neki. Olyasmit, aminek felismerését utóbb ő maga is gyakorolhatta úgy, hogy az általa elért helyes eredményt annak emberi felülvizsgálata (helyességének megállapítása) után pozitív megerősítés követi. Egy ilyen algoritmus például azonos képalkotási módszerrel, szövettani mintákról készült felvételek közül ki tudja válogatni a kimutatható daganatokat tartalmazókat,<sup>18</sup> akár egy radiológus képességeit jóval meghaladó pontossággal is, ha elég példát (fényképet) szolgáltatunk neki az azonosítandó jelenségről a betanítási szakaszban. Az, hogy az azonosítás pontosan hogyan történik, egyes algoritmusok esetén jobban, mások esetén kevésbé<sup>19</sup> nyilvánvaló a programozók számára. Innen ered a köznyelvi értelemben mesterséges intelligenciának nevezett algoritmusok igazságszolgáltatásban alkalmazásának két fő akadálya.

Egyrészt az igazságosság fogalmát az emberiségnek sem sikerült általánosan kielégítő módon meghatározni. John Rawls, a téma talán legjelentősebb filozófiai kutatója például az igazságos magatartás több eltérő elméletét<sup>20</sup> azonosította, amelyek a társadalmi együttműködést szolgálják. Az igazságszolgáltatás célja pontosan e társadalmi együttműködés (napjainkban már kevésbé divatos megnevezéssel az ún. társadalmi szerződés) megőrzése. Jelenleg nincs mód arra, hogy egy gép számára elegendő példát szolgáltatassunk egy jogvita társadalmilag igazságos megoldására ahhoz, hogy – amennyiben felismerné az illető jogvita meghatározó jellemzőit – azt ön maga megoldhassa. Adott ugyanis, hogy nem csupán az egyének értékelik és alkalmazzák az igazságosság fogalmát heterogén módon, hanem maguk a jogi normák is. Igazságos például éhező gyermekei számára ételmet lopó tolvajt azonos büntetéssel sújtani, mint egy jól szituált kleptomániást? Nyilván nem. Egy számítógép, amelynek betanítása a véletlenszerűen kiválogatott lopások ügyében született ítéletekkel történik, ezt a különbséget nem biztos, hogy észlelni lenne képes.

A gépi tanulási algoritmusok mindenkor a betanításukra szolgáló adathalmazból önmaguk számára érthető (bár mások számára nem mindig nyilvánvaló) kritériumok alapján igyekeznek egyes elemeket közös jellemzőik alapján elhatárolni. Mivel az emberi élet számos mozzanata napjainkban már digitalizálható, elképzelhető, hogy egy személy minden megismerhető adata (genetikai, egészségügyi, pénzügyi, tanulmányi információk, szülei, testvérei személyazonossága, pártállása, banki adatai, írásos és szóbeli rögzített megnyilvánulásai, büntetett előélete stb.) egyetlen, az adott személyt teljes mértékben leíró adathalmazba begyűjthetővé válik. Ebből a számítógép megtanítható arra, hogy mondjuk az egyes bűncselekményekben elmarasztalt vagy bizonyos betegségeken szenvedő személyek közös magatartási jellemzőit azonosítsa. Utóbb a gép bizonyos személyről meg tudja majd állapítani, hogy bizonyos bűncselekmény elkövetésére vagy bizonyos betegségekre hajlamos vagy sem, csupán a személy „digitális lábnyomából” kiindulva. Az igazság szolgáltatásához – a ma ismert gépi tanulás módszerével – gép ennél közelebb nem juthat: valamely (jellemzően statisztikai) módszerrel történő profilalkotással valószínűsítheti, hogy egy személy bizonyos magatartásokra hajlamos, vagy elképzelhető, hogy ezeket tanúsította is. Ez a megállapítás azonban nem ítélethez, csupán előítélethez vezet a gépet, olyan nézetalkotáshoz, amely a humán faktort is jellemzi, de amelyet a számítógépek alkalmazásával szeretnénk kiküszöbölni az igazságszolgáltatás során. A gépi tanulási algoritmusok ilyen működése jó eséllyel napjaink totalitárius társadalmában talál majd először gyakorlati felhasználásra.

Másrészt a gép kognitív folyamatai, belső érvelése nem mindig ellenőrizhető az igazság megismerésének szemszögéből. Az informatikusok nem mindig tudják megállapítani, hogy bizonyos gépi tanulási algoritmusok miért éppen egy adott eredményt produkálnak, még akkor sem, ha az eredmény az algoritmus rendeltetésszerű működéséből származik. Ezt nevezik az algoritmusok átláthatósága vagy elszámoltathatósága<sup>21</sup> problémájának a gépi döntéshozatal során. Ha nem tudjuk pontosan meghatározni, hogy egy algoritmus miért bizonyos döntésre jutott (vagy esetleg csak azzal tudjuk meghatározni, hogy amiatt, mert az a *legvalószínűbb* a gép számára), egyúttal nem tudjuk ellenőrizni a határozat igazságos (törvényes) jellegét sem. Ez a probléma<sup>22</sup> már most létezik,<sup>23</sup> bár még kevesek számára tűnik<sup>24</sup> jelentősnek. A jogban azonban az elszámoltathatóság, a határozathozatal jogi és ténybeli indoklása elengedhetetlen. Csupán az indoklás ismeretében lehetséges megállapítani, hogy egy határozat jogszerű és egyben igazságos-e. Csupán az indoklás alapján lehetséges a jogorvoslat vagy perorvoslat lehetőségét biztosítani a jogkeresők számára egy jogszerűtlen, igazságtalan határozattal szemben.

2. A tág értelemben vett jog semmiképpen sem egyszerűsíthető le matematikai, logikai folyamatokká, mert a jogalkalmazás sokkal emberhez kötöttebb és szubjektívebb, semhogy teljes egészében algoritmusok vezéreljék. A jog olyan állami alkotás, amely domináns értékeket és politikai célokat közvetít. Képzeljük el Batthyáneum vagy a Székely Mikó kollégium ügyét, ahol az MI a tények és bizonyítékok objektív és komoly mérlegelése alapján visszaszolgáltatási határozatában megállapítaná: a Római Katolikus vagy a Református Egyház és nem a Román Állam a vitatott sorú ingatlanok tulajdonosa.

El merjük képzelni, hogy a jogalkalmazást bármely állam mint érdekközösség átengedi a teljes objektivitásnak? Az alternatíva az, hogy a mesterséges intelligenciának is lesz nemzetiisége, és megjelenik a nacionalistává programozott mesterséges intelligencia? A társadalom egyre bonyolultabbá válásával a jogi szabályozás iránti szükséglet csak nő.

A mesterséges intelligencia és általában az informatika a jogalkalmazás szolgáltatásban egyre nagyobb szerepet játszik, és e szerepkör egyre csak bővülni fog. Már ma a mesterséges intelligencia az egyik legnehezebb új szabályozási terület, hiszen ennek működése és elterjedése precíz jogi szabályozás nélkül elképzelhetetlen.

A mesterséges intelligencia veszélyeit jogi preventív eszközökkel kell korlátozni. Megfelelő jogi keretek nélkül – maguk azon műszaki szakemberek szerint, akik fejlesztésének élvonalában járnak<sup>25</sup> – egyre veszélyesebb hazardjátékká látszik válni.

A valóság az, hogy a jogalkalmazással szemben az a valós állami elvárás, hogy adott mértékben igenis szubjektív legyen, és ami szubjektív, az emberi előítéletek által vezérelt. Ennek egyidejűleg van negatív és pozitív oldala. Ha mesterséges intelligenciával végeztetünk komplexebb jogi feladatokat, akkor ezeket az érzelmeket és előítéleteket a mesterséges intelligenciának is el kell sajátítania, és működésének alapjává kell tennie. Ez a folyamat egyébként lehetséges: a mesterséges intelligencia óriási adatmennyiség feldolgozásával fejlődik, tökéletesedik. Ha a mesterségesen intelligens folyamat kialakításához előítéletekkel terhelt adatok kerülnek betáplálásra, a folyamat mintegy örökl<sup>26</sup> az adatokban tükröződő előítéleteket. De akkor nem várhatunk a mesterséges intelligenciától sem igazságot.

Nyilván nem szabad kizárni, sőt észszerű is, hogy egyszerű jogi munkák automatizálhatók lesznek, és ez a jogászok számának bizonyos mértékű csökkenését fogja

eredményezni. A jogi szakmák eltűnését jósló álláspontok mégis alaptalanok. A jogi szakmáért nem kell aggódni, akkor sem, ha a mesterséges intelligencia szerepe megnő, és egyre intenzívebben használunk informatikai eszközöket a jogalkalmazásban. Az „emberi tényezőt” csupán az emberi tényező leképezésével lehet a jogalkalmazásba bevezetni, azaz a mesterséges intelligencia nem helyettesíti, hanem csak segítheti a jogászt. Érzelmek nélkül: rá tudjuk bízni például a büntetőjog alkalmazását algoritmusokra? Vagy képzeljük el az iszlám vallási jogot, a sariát alkalmazó mesterséges intelligenciát, amint objektíven jóváhagyja a prostitúció tilalmának kijátszására alkotott meghatározott idejű (önmegsemmisítő) házasság intézményét (az akár kiskorúval megkötött házassági szerződést).

A mesterséges intelligencia megjelenése a társadalom konfliktusait nem fogja megszüntetni. Egyetlen technológia sem képes ilyen hatást kiváltani, a konfliktus léte az emberi társadalom szükséges velejárója. Ha a jövőt fürkésszük, inkább a konfliktusfelületek növekedésétől kell tartani.

Képzeljünk el egy jogvitát, amikor a felek ügyvédjei mindkét oldalon mesterséges intelligencia segítségét veszik igénybe. A mesterséges intelligenciák konfliktusa és versenye is kialakul. Valamelyik álláspontnak mégis győznie kell, máskülönben a jogvita lezáratlan marad. A jog alkalmazásakor a bíró mindkét oldalon fennálló részigazságokat mérlegel, és a jogvita megoldásának, megoldhatóságának érdekében valamelyik részigazságot abszolutizálja. Ezért van az, hogy a jogi és erkölcsi igazság sok esetben eltérhet.

Dönthet két mesterséges intelligencia konfliktusában egy harmadik mesterséges intelligencia? Az emberi tényezőt leképező mesterséges intelligencia és a mesterséges intelligenciát segítségül hívó ember között igenis van, és nem is kicsi különbség. A mesterséges intelligencia számos jogi folyamatban fontos szerepet fog játszani, például a szerződéskötések esetén jelentősen csökkentheti a tranzakciós költségeket. Azonban a jogviták esetén a szerepe inkább előkészítő, tanácsadó feladatok ellátására alkalmazható.

A mesterséges intelligencia általános veszélye hosszú távon az emberi elbutulás. Túlzó használata a jogi kultúrát és tudást, az ember konfliktusmegoldó képességének megfelelő szintjét veszélyezteti. A mesterséges intelligenciát csak a kontrollra képes, önmagában magas felkészültségű jogász tudja teljes mértékben felhasználni.

Az innovációt, a jogot mint az emberi társadalom értékközpontú irányítási és fejlesztési rendszerét, az úttörő erkölcsi folyamatokat az ember kezéből „kiengedni” valóban hazárdjáték.

## ■ JEGYZETEK

1. Oscar Schwartz: *‘The discourse is unhinged’: how the media gets AI alarmingly wrong*. The Guardian 2018. 07. 25. sz. <https://www.theguardian.com/technology/2018/jul/25/ai-artificial-intelligence-social-media-bots-wrong> (utolsó megnyitás: 2020. január 10.).
2. Sean O’Kane: *Uber reportedly thinks its self-driving car killed someone because it ‘decided’ not to swerve*. The Verge 2018. 05. 07. sz. <https://www.theverge.com/2018/5/7/17327682/uber-self-driving-car-decision-kill-swerve> (utolsó megnyitás: 2020. január 10.).
3. Andrew Smith: *Franken-algorithms: the deadly consequences of unpredictable code*. The Guardian 2018. 08. 29. sz. <https://www.theguardian.com/technology/2018/aug/29/coding-algorithms-frankenalgorithm-program-danger> (utolsó megnyitás: 2020. január 10.).
4. Jim Giles: *Cameras know you by your walk*. New Scientist 2012. 09. 19. sz. <https://www.newscientist.com/article/mg21528835-600-cameras-know-you-by-your-walk/> (utolsó megnyitás: 2020. január 10.).
5. Arielle Pades: *Racial Recognition Tech Is Ready for Its Post-Phone Future*. Wired 2018. 10. 09. sz. <https://www.wired.com/story/future-of-facial-recognition-technology/> (utolsó megnyitás: 2020. január 10.).

6. Stuart A. Thompson – Charlie Warzel: *How to Track President Trump*. The New York Times 2019. 2019. 12. 20. sz. <https://www.nytimes.com/interactive/2019/12/20/opinion/location-data-national-security.html> (utolsó megnyitás: 2020. január 10.).
7. Matt Stroud: *The Pentagon is getting serious about AI weapons*. The Verge 2018. 2018. 04. 12. sz. <https://www.theverge.com/2018/4/12/17229150/pentagon-project-maven-ai-google-war-military> (utolsó megnyitás: 2020. január 10.).
8. Irina Rîmaru: *Declarația UE de cooperare cu privire la inteligența artificială*. Juridice 2018. 2018. 05. 24. sz. <https://www.juridice.ro/582875/declaratia-ue-de-cooperare-cu-privire-la-inteligenta-artificiala.html> (utolsó megnyitás: 2020. január 10.).
9. Cristian Bălan: *Roboții nu vor înlocui judecătoria*. Juridice 2018. 2018. 08. 27. sz. <https://www.juridice.ro/598527/robotii-nu-vor-inlocui-judecatorii.html> (utolsó megnyitás: 2020. január 10.).
10. Adrian Șandru: *Despre necesitatea sau nu a introducerii răspunderii penale a roboților*. Juridice 2018. 2018. 08. 31. sz. <https://www.juridice.ro/598476/despre-necesitatea-sau-nu-a-introducerii-raspunderii-penale-a-robotilor.html>
11. Madocsai Kinga Blanka: *Az ügyvédség vége – a civilizáció vége?* Jogászvilág 2014. 2014. 10. 20. sz. <https://jogaszvilag.hu/uzlet/az-ugyvedseg-vege-a-civilizacio-vege/> (utolsó megnyitás: 2020. január 10.).
12. Richard Susskind: *Artificial Intelligence and the Law Conference at Vanderbilt Law School*. <https://www.youtube.com/watch?v=xs0iQSYBoDE> (utolsó megnyitás: 2020. január 10.).
13. Maxim Macovei: *Înlocuiește inteligența artificială juristul?* Juridice 2018. 2018. 08. 26. sz. <https://www.juridice.ro/598482/inlocuieste-inteligenta-artificiala-juristul.html> (utolsó megnyitás: 2020. január 10.).
14. Gheorghe Florea: *Inteligența artificială s-ar putea întoarce împotriva umanității!* Juridice 2017. 2017. 11. 16. sz. <https://juridice.ro/essentials/1828/gheorghe-florea-inteligenta-artificiala-s-ar-putea-intoarce-impotriva-umanitatii> (utolsó megnyitás: 2020. január 10.).
15. Jochim Spitz – Pieter Moors – Johan Wagemans – Werner F. Helsen: *The impact of video speed on the decision-making process of sports officials*. Cognitive Research: Principles and Implications 2018. 3/16. sz. <https://doi.org/10.1186/s41235-018-0105-8> (utolsó megnyitás: 2020. január 10.).
16. Robert D. Hof: *Deep Learning*. MIT Technology Review 2013. 2013. 04. 23. sz. <https://www.technologyreview.com/s/513696/deep-learning/> (utolsó megnyitás: 2020. január 10.).
17. Bernard Marr: *What Is The Difference Between Artificial Intelligence And Machine Learning?* Forbes 2016. 2016. 12. 06. sz. <https://www.forbes.com/sites/bernardmarr/2016/12/06/what-is-the-difference-between-artificial-intelligence-and-machine-learning/>
18. Hu Zilong – Tang Jinshan – Wang Ziming – Zhang Kai – Zhang Ling – Sun Qingling: *Deep learning for image-based cancer detection and diagnosis - A survey*. Pattern Recognition 2018. novemberi sz. 134–149. <https://www.sciencedirect.com/science/article/pii/S0031320318301845> (utolsó megnyitás: 2020. január 10.).
19. Andrew Smith: *Franken-algorithms: the deadly consequences of unpredictable code*. The Guardian 2018. 2018. 08. 30. sz. <https://www.theguardian.com/technology/2018/aug/29/coding-algorithms-frankenalgos-program-danger> (utolsó megnyitás: 2020. január 10.).
20. John Rawls: *The Main Idea of the Theory of Justice*, <https://www.csus.edu/indiv/c/chalmersk/econ184sp09/johnrawls.pdf> (utolsó megnyitás: 2020. január 10.).
21. *Algorithmic Accountability. Applying the concept to different country contexts*. World Wide Web Foundation, 2017. [https://webfoundation.org/docs/2017/07/WF\\_Algorithms.pdf](https://webfoundation.org/docs/2017/07/WF_Algorithms.pdf) (utolsó megnyitás: 2020. január 10.).
22. Nicholas Diakopoulos – Sorelle Friedler: *How to Hold Algorithms Accountable*. MIT Technology Review 2016. 2016. 11. 17. sz. <https://www.technologyreview.com/s/602933/how-to-hold-algorithms-accountable/> (utolsó megnyitás: 2020. január 10.).
23. Nicholas Diakopoulos: *Algorithmic Accountability. Journalistic investigation of computational power structures*. Digital Journalism 3/2015. <https://www.tandfonline.com/doi/abs/10.1080/21670811.2014.976411?journalCode=rdij20> (utolsó megnyitás: 2020. január 10.).
24. Jakko Kemper – Daan Kolkman: *Transparent to whom? No algorithmic accountability without a critical audience*. Information, Communication & Society 2019. <https://www.tandfonline.com/doi/full/10.1080/1369118X.2018.1477967> (utolsó megnyitás: 2020. január 10.).
25. James Vincent: *Elon Musk, DeepMind founders, and others sign pledge to not develop lethal AI weapon systems*. The Verge 2018. <https://www.theverge.com/2018/7/18/17582570/ai-weapons-pledge-elon-musk-deepmind-founders-future-of-life-institute> (utolsó megnyitás: 2020. január 10.).
26. James Zou – Londa Schiebinger: *AI can be sexist and racist — it's time to make it fair*. Nature 2018. <https://www.nature.com/articles/d41586-018-05707-8> (utolsó megnyitás: 2020. január 10.).