

A mesterséges intelligencia toborzási integrációjának komplex vizsgálata a Z generációs attitűdök és a humán- algoritmikus döntéshozatal tükrében

A complex examination of AI integration in recruitment in light of Generation Z attitudes and human-algorithmic decision-making.

Ribánszki Nóra^{1*}, Dr. Karcsics Éva²

¹ Menedzsment és Üzleti Jog Tanszék, Gazdaságtudományi Kar, Neumann János Egyetem, Magyarország,
<https://orcid.org/0009-0009-1573-4131>

² Menedzsment és Üzleti Jog Tanszék, Gazdaságtudományi Kar, Neumann János Egyetem, Magyarország,
<https://orcid.org/0009-0001-4526-8794>
<https://doi.org/10.47833/2026.2.ECO.011>

Kulcsszavak:

Mesterséges intelligencia
Toborzás
Algoritmikus torzítás
TAM-modell
Ember-Robot Együttműködés

Keywords:

Artificial intelligence
Recruitment
Algorithmic bias
Technology Acceptance Model
Human-Robot Collaboration

Cikktörténet:

Beérkezett 2026. február 5
Átdolgozva 2026. augusztus 25.
Elfogadva 2026. augusztus 28.

Összefoglalás

Jelen tanulmány egy többmódszertanú kutatás keretében vizsgálja a mesterséges intelligencia (MI) toborzási alkalmazásának attitűdbeli és technikai dimenzióit, ahol az adminisztratív hatékonyság etikai aggályokkal párosul. A Technológia-elfogadási Modell (TAM) alapján végzett felmérés (N = 145) rámutat, hogy a Z generáció az MI-t „feltételesen elfogadja”: bár az operatív hasznosságot elismerik, a bizalmatlanság miatt az erős emberi felügyeletet elengedhetetlennek tartják. Egy kvantitatív kísérletben (N = 104 HR-szakember) az ATS-alapú algoritmus és a humán értékelők döntéseit is összevetettük. Bár a makroszintű egyezés erős ($r_s = 0,90$), az algoritmus demográfiailag semleges maradt, míg a humán döntéseket a szubjektív kockázatészlelés szignifikánsan módosította. Megoldásként az algoritmikus konzisztenciát és az emberi szociális érzékenységet ötvöző Human-Robot Collaboration (HRC) modellt javasoljuk.

Abstract

This study investigates the integration of artificial intelligence (AI) in recruitment through a mixed-methods approach. A survey (N = 145) reveals Generation Z's „conditional acceptance” of AI: while acknowledging its operational utility, they demand strong human oversight due to trust issues. Additionally, an experiment involving 104 HR professionals compared algorithmic and human decision-making in CV evaluations. Statistical analysis ($r_s = 0.90$) showed a strong macro-level correlation; however, the AI maintained demographic neutrality, whereas human decisions were significantly altered by subjective risk perception. We conclude that sustainable selection requires a Human-Robot Collaboration (HRC) model, balancing technical objectivity with social sensitivity.

* Kapcsolattartó szerző.
E-mail cím: ribanszki.norii@gmail.com

1. Bevezetés

A mesterséges intelligencia és az automatizált döntéshozatali rendszerek térnyerése paradigmaváltást indított el a humánerőforrás-menedzsmentben, alapjaiban írva felül a toborzási folyamatok hatékonysági mutatóit és időbeli kereteit. A technológiai transzformáció azonban komplex etikai és szociotechnikai kockázatokat hordoz: az algoritmikus torzítás szisztematikus jelensége és a döntéshozatali mechanizmusok átláthatatlansága, az úgynevezett „fekete doboz” (black box) effektus alapjaiban kérdőjelezi meg a mesterséges intelligencia alapú kiválasztási folyamatok méltányosságát és hitelességét.

A téma aktualitását az Európai Unió mesterséges intelligenciáról szóló rendelete (AI Act) is alátámasztja, amely a toborzási és kiválasztási célú szoftvereket a „magas kockázatú” kategóriába sorolja [3]. Ez a szabályozási környezet a kizárólagos technológiai implementáció helyett az algoritmikus elszámoltathatóság, a transzparens kockázatkezelési keretrendszerek és a humán felügyelet (human-in-the-loop) intézményesítését követeli meg a szervezetektől. A tudományos diskurzusban ugyanakkor továbbra is fennáll a feszültség a matematikai alapú algoritmikus objektivitás és a humán szubjektivitás között, amelynek feloldása kritikus a jövő munkaerőpiaci stabilitása szempontjából.

Jelen tanulmány célkitűzése, hogy két egymást kiegészítő kutatási pilléren keresztül elemezze a mesterséges intelligencia integrációját. Elsődlegesen a munkaerőpiacra belépő Z generáció attitűdjeit tártuk fel a Technológia-elfogadási Modell (TAM) kiterjesztett dimenziói mentén, különös tekintettel a bizalmi faktorokra. Másodsorban empirikus úton, egy kontrollált kísérleti elrendezésben vizsgáltuk az algoritmikus torzítás természetét, összevetve egy valós Applicant Tracking System (ATS) és 104 HR-szakember döntéshozatali mechanizmusait. Ebben a kontextusban az objektivitást operacionálisan a demográfiai semlegesség (nemi és életkori függetlenség) és a kritérium-konzisztencia (szakmai paramétereknek való matematikai megfelelés) mentén definiáljuk.

A vizsgálat az alábbi kutatási kérdésekre (RQ) és előfeltevésekre (E) épül:

- **RQ1:** Hogyan viszonyulnak a felsőoktatásban tanuló fiatal pályakezdekők a mesterséges intelligencia toborzásban történő alkalmazásához, milyen attitűdök, félelmek és elvárások jellemzik őket?
 - **E1:** A hallgatók jelentős része ambivalens attitűddel viszonyul az MI toborzási alkalmazásához, a bizalom és technológiai nyitottság egyensúlyban áll az adatvédelmi, méltányossági és emberi tényezővel kapcsolatos aggályokkal.
- **RQ2:** Kimutatható-e az algoritmikus torzítás a mesterséges intelligencia által végzett önéletrajz-értékelés során a HR-szakemberek értékelésével összehasonlítva?
- **RG3:** Mennyire eltérő rangsort állít fel egy mesterséges intelligencia és egy HR-szakember a jelöltek között az önéletrajzuk alapján?
 - **E2:** A mesterséges intelligenciát alkalmazó automatizált szűrőrendszerek a vizsgált minta alapján eltérő pontszámokat rendelnek azonos szakmai profilú, de eltérő demográfiai jellemzőkkel rendelkező pályázókhoz.
 - **E3:** A HR-szakemberek preferenciasorrendje és az MI-eszköz által képzett sorrend között szignifikáns eltérés mutatható ki, különösen nemi és etnikai dimenziók mentén.

Tanulmányunkkal arra keressük a választ, hogy a technikai precizitást és a szociális érzékenység miként integrálható a Human-Robot Collaboration, HRC modelljébe, megteremtve ezzel a fenntartható és méltányos toborzás alapjait.

2. Szakirodalmi áttekintés

A szakirodalom a mesterséges intelligencia integrációját a HRM folyamatokba jellemzően két, egymással dialektikus viszonyban álló dimenzió mentén tárgyalja [18]. Egyrészt a hatékonyság és az objektivitás növelésének eszközeként, másrészt az etikai kockázatok és az algoritmikus torzítás forrásaként [6].

2.1. Az MI transzformatív szerepe és technológiai eszköztára

A nemzetközi kutatások konszenzusa szerint az MI elsődleges hozzáadott értéke a rutinszerű, adminisztratív feladatok automatizálásában rejlik [16][17], és ez a transzformatív hatás a hazai HR-gyakorlatban is egyre markánsabban jelenik meg [12], amely lehetővé teszi a HR-szakemberek számára a stratégiai feladatokra való összpontosítást [2]. A toborzási folyamat digitalizációja három központi technológiai pilléren nyugszik.

Az intelligens jelöltkezelő rendszerek (Applicant Tracking System – ATS) jelentik a tömeges jelöltkezelés alapvető infrastruktúráját [12][9]. A természetes nyelvi feldolgozás (NLP) alkalmazásával az algoritmusok képesek a kompetenciák és kulcsszavak mentén automatizáltan szűrni és rangsorolni a pályázókat [19][8]. Ez a megoldás drasztikusan redukálja a közvetlen és közvetett költségeket és a kiválasztási időt, azonban a szakirodalom gyakran bírálja a rendszerek túlzott kulcsszó fókuszát, amely a nem konvencionális, de értékes tehetségek kizárásához vezethet [6].

A chatbotok és virtuális asszisztensek jelöltélmény fokozását célzó eszközök, amelyek 24 órás rendelkezésre állással támogatják az információnyújtást és az előszűrést [7][19][8]. Bár az azonnali és personalizált interakció erősíti a munkáltatói márkát, a technológia alkalmazása magában hordozza a folyamat elszemélytelenítésének kockázatát is [11].

Az aszinkron videóinterjú-elemzés (AVI) technológia legújabb generációja már nem kizárólag a verbális válaszokat, hanem a nonverbális kommunikációt (mimika, hanglejtés) is elemzi. Ez a terület képezi a legélesebb etikai viták tárgyát, különös tekintettel az érzelemfelismerő algoritmusok validítására és a diszkriminációs potenciálra [19].

1. Táblázat. A mesterséges intelligencia-alapú toborzás dialektikája: Racionális előnyök és humán-etikai korlátok összehasonlítása

Dimenzió	Előnyök	Aggályok
Hatékonyság és idő	Toborzási ciklusidő csökkenése	A mesterséges intelligencia képes lehet túlzottan a kulcsszavaktól függeni, ami miatt értékes jelöltek szűrődhetnek ki.
Objektivitás és méltányosság	A humán előítéletek (öntudatlan elfogultság) csökkenése a szűrés korai fázisában	A múltbeli diszkriminatív adatokból öröklődő előítéletek felerősítése.
Kompetenciafelmérés	Képes kizárólag a képzettségre és készségekre fókuszálni, így kizárva a demográfiai jellemzőket.	Képtelenség a puha készségek, a jelölt motivációjának és a vállalati kultúrába való illeszkedésének felmérésére.
Döntéshozatal	Döntéstámogatás prediktív analitikával a megtartási ráta és munkasiker előrejelzésére.	A döntési mechanizmusok átláthatatlansága, ami bizalomhiányt és az elutasítás okának meg nem értését okozza.
Adatkezelés	Nagy mennyiségű adat gyors és strukturált kezelése és tárolása.	Aggályok az érzékeny adatok gyűjtésével és nem transzparens felhasználásával kapcsolatban.

Saját szerkesztés, [1][7][10][15] alapján

2.2. Az algoritmikus torzítás és a “Fekete doboz” jelenség

Az MI-alapú döntéshozatal központi kritikája az objektivitás mítoszának megkérdőjelezése. Annak ellenére, hogy az algoritmusok mentesek az emberi előítéletektől, hajlamosak a betanításukhoz használt történeti adatokban rejlő társadalmi előítéletek reprodukálására és felerősítésére [13].

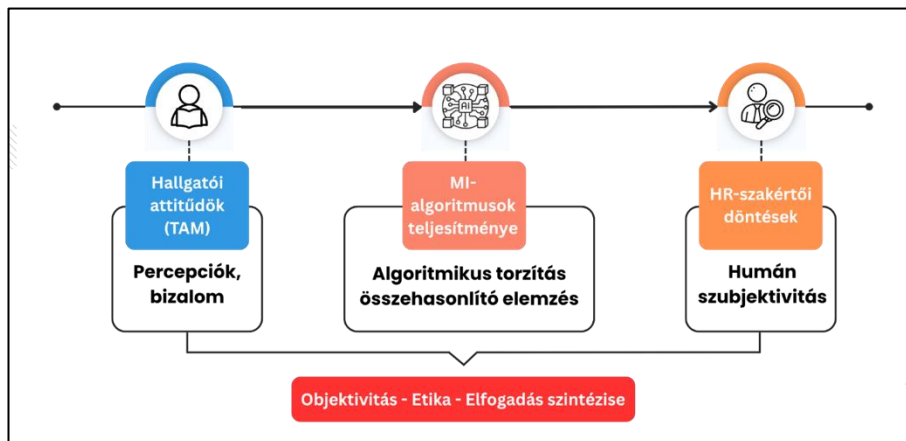
Az algoritmikus torzítás (bias) definíció szerint egy szisztematikus és ismétlődő tendencia, amely egy csoportot indokolatlanul előnyben részesít egy másikkal szemben [11]. Chen (2023) rámutat, hogy az algoritmusok nem semleges entitások, hanem a meglévő hatalmi és társadalmi viszonyokba ágyazott konstrukciók [8]. A problémát súlyosbítja a „fekete doboz” (black box) jelenség, vagyis a döntési mechanizmusok átláthatatlansága, amely lehetetlenné teszi annak megértését, hogy a rendszer miért utasított el egy adott jelöltet [21]. Ez a transzparencia-hiány alapvetően ássa alá a jelöltek bizalmát és jogi aggályokat vet fel [4], amelyre az Európai Unió AI Act szabályozása is reagál, magas kockázatúként besorolva a toborzási rendszereket [20].



1. ábra. Algoritmikus torzítás eredetének vizsgálata
Saját szerkesztés, [8] alapján

3. Módszertan

A kutatás célkitűzésének komplexitása indokolta a többmódszertanú (mixed-method) megközelítés alkalmazását. A vizsgálat egymásra épülő pilléreken nyugszik, egy kvalitatív fókuszú hallgatói attitűdvizsgálaton, egy kvantitatív kísérleti kutatáson.



2. ábra. Kutatási modell logikai felépítése
Saját szerkesztés, saját kutatás alapján

Az adatgyűjtési szakasz 2025 májusától 2025 szeptemberéig tartott, ezen időintervallumban a kutatás különböző komponenseihez tartozó adatfelvételek párhuzamosan zajlottak.

3.1. Hallgatói attitűdök és a TAM-modell

A kutatás első fázisának célja a munkaerőpiacra belépő Z generáció technológiaelfogadási hajlandóságának feltérképezése volt. A felsőoktatásban tanuló Z generáció tagjai számára a technológia, így a mesterséges intelligencia és a gamifikációs megoldások már az oktatási

környezetben is ismert jelenségek [5], ami alapjaiban határozza meg a toborzási szoftverekkel kapcsolatos elvárásait és elfogadási hajlandóságukat.

A vizsgálatban 145 felsőoktatásban tanuló vett részt, a minta összetétele a nemek arányát tekintve 37% férfi és 63% nő, akik dominánsan a Z generáció tagjai. A minta tanulmányi előrehaladását tekintve a válaszadók többsége, 65%-a a negyedik félévét (másodév) végezte az adatfelvétel idején. A képzési szint megoszlása szerint a résztvevők 78%-a alapképzésben (BSc), 22%-a pedig felsőoktatási szakképzésben vett részt. A munkaerőpiaci relevanciát erősíti, hogy a hallgatók döntő többsége rendelkezik munkatapasztalattal, a válaszadók 83%-a már szerzett gyakorlatot a munka világában (48%-uk egy évnél kevesebb, 35%-uk egy évnél több tapasztalattal bír), és mindössze 17%-uk nem rendelkezett még munkatapasztalattal a felmérés időpontjában.

Az adatgyűjtés módszere a strukturált esszéírás volt, a hallgatók 2000 karakter terjedelmű írásos véleményt fogalmaztak meg előre meghatározott kérdések mentén, amelyek a mesterséges intelligencia toborzásban való alkalmazására, az észlelt előnyökre és az etikai aggályokra vonatkoztak. A válaszok validitásának biztosítása érdekében az adatfelvételt megelőzően, 2025 májusában a résztvevők egy dedikált szakmai előadáson vettek részt, amely biztosította a szükséges fogalmi keretek ismeretét.

A beérkezett szöveges válaszokat deduktív tartalomelemzéssel vizsgáltuk. A kódolás során a válaszelemeket a Technológia-elfogadási Modell (TAM) kiterjesztett dimenzióihoz (észlelt hasznosság, könnyű használhatóság, bizalom, kockázatészlelés) rendeltük lehetővé téve a kvalitatív adatok gyakorisági vizsgálatát [9]. A kódolási folyamat tudományos átláthatósága és reprodukálhatósága érdekében egy előre rögzített kódkönyvet alkalmaztunk.

Az elemzés során a kódolási alapegységnek a tematikus kijelentéseket (mondatokat vagy logikai összefüggő tagmondatokat) tekintettük, amelyek önálló, értelmezhető állítást fogalmaznak meg az MI-alapú toborzással kapcsolatban

A deduktív kategóriarendszert a kiterjesztett TAM modell elméleti kerete határozta meg. A kódkönyv hat dimenziót és a hozzájuk rendelt horgonyszavakat tartalmazta:

1. **Észlelt hasznosság:** operatív előnyök és racionalizálás
2. **Észlelt könnyű használhatóság:** a technológia integrálhatósága és zökkenőmentessége
3. **Bizalom és etikai aggályok:** a döntéshozatal méltányosságával és transzparenciájával kapcsolatos megállapítások
4. **Kockázatészlelés:** az algoritmikus döntésekből fakadó potenciális károk
5. **Érzelmi attitűd:** a technológiával szembeni affektív viszonyulás
6. **Viselkedési szándék:** a jövőbeli használatra vonatkozó hajlandóság

A kódolási folyamat objektivitásának és megbízhatóságának biztosítása érdekében közösen vettünk részt a validálásban. Az esetleges kódolási eltéréseket konszenzusos szakmai megbeszélés útján oldottuk fel, ezzel növelve ezzel az elemzés következetességét.

3.2. Algoritmikus vagy humán döntéshozatal?

A második fázis célja az objektivitás tesztelése volt egy kontrollált kísérleti elrendezésben, összevetve egy valós ATS-rendszer és humán szakemberek rangsorolását.

A vizsgálathoz 5 darab, szabványosított fiktív önéletrajzot (CV) hoztunk létre, amelyek célzottan tértek el egymástól a demográfiai és szakmai jellemzők mentén, hogy mérhetővé tegyék a torzítás mértékét. A profilok az alábbiak voltak:

7. **Kovács Dániel (Alapteszt):** 26 éves férfi, 2 év releváns tapasztalattal, B2 angol nyelvtudással, ő szolgált referencia jelöltként (a tökéletes jelölt).
8. **Nagy Eszter (Nemi kontraszt):** 26 éves nő, Kovács Dániellel azonos kompetenciákkal és tapasztalattal, célja a nemi alapú diszkrimináció tesztelése volt.
9. **Gáspár Ramóna (Készség előny):** 30 éves nő, szintén 2 év tapasztalattal, de C1 felsőfokú angol és emellett német nyelvtudással, valamint projektmenedzseri háttérrel. Célja a magasabb kompetenciaszint hatásának vizsgálata.
10. **Tóth Gábor (Senioritás / Túlképzettség):** 57 éves férfi, 25+ év tapasztalattal, célja az életkori torzítás (ageism) és a túlképzettség kezelésének vizsgálata junior pozíciónál.

11. **Szabó Katalin (Kizáró ok és Senioritás):** 57 éves nő, 25+ év tapasztalattal, de csak alapfokú (B1) angol nyelvtudással. Célja a nyelvi szűrő és a túlképzettség együttes hatásának tesztelése.

A profilokat egy munkaerő-közvetítő vállalat valós ATS rendszerében futtattuk le, vizsgálva a szoftver által generált automatikus rangsort és illeszkedési pontszámokat. A kutatásban alkalmazott szoftver egy kétlépcsős, súlyozott kulcsszó- és kritériumillesztési logikát alkalmaz a jelöltek rangsorolásához. Azonnal diszkvalifikálja a kötelező alapkövetelményekkel nem rendelkezőket, valamint a továbbjutókat a hirdetés kulcsszavai és az önéletrajz közötti szemantikai hasonlóság alapján értékeli. A kísérlet során az algoritmus paraméterei kizárólag a szakmai kompetenciákra fókuszáltak, a demográfiai adatok (pl. nem, születési idő) kizárásával a rendszer „vak” előszűrést végzett, csökkentve a közvetlen demográfiai preferenciák hatását az előszűrés során. a k

A kutatás empirikus vizsgálatába bevont szakértői mintát 104 gyakorló humánerőforrás-szakember alkotta. A mintavétel hólabda-módszerrel, online kérdőíves formában történt 2025 szeptemberében. A részvétel alapfeltétele a toborzás-kiválasztás területén szerzett legalább egyéves szakmai gyakorlat megléte volt. A minta összetétele magas szakmai senioritást tükröz, a válaszadók átlagos HR-területen szerzett munkatapasztalata 12,77 év, míg a kifejezetten toborzási területen szerzett tapasztalat átlaga 7,46 év. A résztvevők domináns része HR generalista, toborzási specialista vagy HR vezetői munkakörben tevékenykedik. Szervezeti háttérüket tekintve a válaszadók többsége (52,9%) a közepes méretű vállalatokat képviselte, iparági besorolás szerint pedig elsősorban a kereskedelem és szolgáltatás (27,9%), valamint a gyártás és ipar (24,0%) szektorból érkeztek. A minta földrajzilag koncentrált, a kitöltők túlnyomó része Bács-Kiskun (65,4%) és Pest (25,0%) vármegyében végzi tevékenységét. A szakemberek feladata az volt, hogy ugyanazokat a fiktív önéletrajzokat értékeljék, mint a valós ATS rendszer.

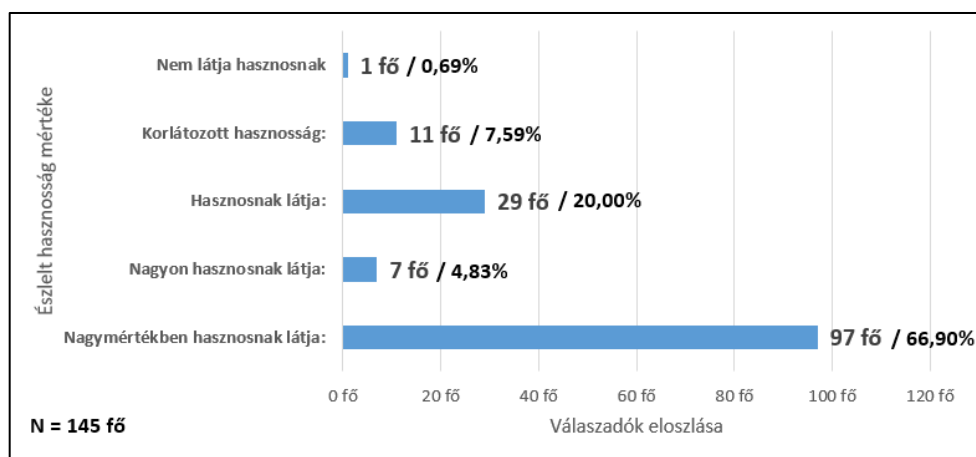
A két fázis összevetése lehetővé tette annak elemzését, hogy az algoritmikus és a humán döntéshozók hogyan kezelték a nemi különbségeket (Dániel vs. Eszter), a többlettudást (Ramóna) és a senioritást (Gábor és Katalin).

4. Eredmények

4.1. A Z generáció és a mesterséges intelligencia alkalmazása a toborzásban

A 145 fős hallgatói mintán végzett kvalitatív kutatás alapján megállapítható, hogy a fiatal munkavállalók viszonya a technológiához alapvetően ambivalens, a racionális hasznosság elismerése mellett megjelenik az érzelmi távolságtartás.

4.1.1. Észlelt hasznosság és hatékonyság



3. ábra. Észlelt hasznosság mértékének vizsgálata
Saját szerkesztés, saját kutatás alapján

A válaszadók többsége, 133 fő hasznosnak, nagyon hasznosnak vagy nagymértékben hasznosnak ítélte a mesterséges intelligencia integrációját a toborzási folyamatokba. A tartalomelemzés alapján a hallgatók a technológia legfőbb hozzáadott értékét az adminisztratív terhek csökkentésében, a folyamatok felgyorsításában és az objektivitás potenciális növelésében

látják. A "gyorsaság" és a "hatékonyság" kulcsszavak domináltak leginkább a pozitív visszajelzésekben.

4.1.2. A bizalom korlátai

A magas hasznosság-észlelés ellenére erős etikai aggályok fogalmazódtak meg a technológia alkalmazását illetően. A válaszadók a legnagyobb kockázatot (a kockázateszlelés dimenziója) abban látják, hogy az algoritmusok képtelenek a "puha készségek" (soft skills), a motiváció is a kulturális illeszkedés mérésére.

A hallgatók gyakran hivatkoznak a folyamat lélektelenségére, attól tartva, hogy a mesterséges intelligencia a komplex emberi profilt pusztá adathalmazzá redukálja és a döntésből hiányzik majd az empátia.

A transzparencia hiánya is visszatérő félelem volt a hallgatók körében, hiszen a jelöltek nem látják, hogy milyen szempontok alapján szűrt a rendszer, ami bizalmatlanságot szülhet.

2. Táblázat. A mesterséges intelligencia-alapú toborzás etikai aggályainak fő korlátai és jellemzői a hallgatók körében

Etikai aggály	Jelleg / aggály lényege	Példa az esszékből
Személytelenség	Empátia, motiváció, interakciók hiánya	„Mesterséges intelligencia nem értékeli az érzelmeket, motivációkat”
Algoritmikus elfogultság / diszkrimináció	Igazságtalanság, torz adatok	„Hátrányos helyzetű csoportok kizárása ²
Téves értékelés / megbízhatósági problémák	Rossz döntések hibás szűrés miatt	„Alkalmas jelölt kizárása”
Adatvédelmi és átláthatósági aggályok	Biztonság, transzparencia hiánya	„Érzékeny, személyes adatok kezelése, fekete doboz jelleg”

Saját szerkesztés, saját kutatás alapján

N = 145, a kvalitatív hallgatói esszék deduktív tartalomelemzése alapján azonosított főbb etikai kategóriák, a TAM dimenzióhoz rendelve, kvantitatív mérési skálát nem tartalmaz.

4.1.3. Demográfiai eltérések

Az eredmények szignifikáns különbséget mutattak a nemek között az aggályok jellegét tekintve. A férfiak inkább a rendszer technikai hibáitól, a téves elutasítástól és a kontrollvesztéstől tartanak (pl. értékes jelölt elvesztése technikai okokból).

A női válaszadók szignifikánsan érzékenyebbek az etikai dimenziókra, náluk dominánsabb az algoritmikus diszkriminációtól és a társadalmi előítéletek (bias) reprodukálásától való félelem.

4.1.4. Viselkedési szándék

A korábban említett tényezők eredőjeként a domináns attitűd a feltételes elfogadás (conditional acceptance). A hallgatók nem utasítják el a technológiát, de annak alkalmazását szigorú feltételekhez kötik, például, hogy a mesterséges intelligencia végezze az előszűrést, de a végső döntés minden esetben maradjon humán szakember kezében (human-in-the-loop).

4.2. Kísérleti eredmények

A kutatás második fázisában 5 fiktív önéletrajz értékelését hasonlítottuk össze egy valós piaci algoritmus és 104 gyakorló HR-szakember bevonásával. A vizsgálat célja az objektivitás és a torzítások (nem, kor, kompetencia) mérése volt.

4.2.1. Algoritmikus semlegesség és szigorú szűrés

A vizsgált rendszer a kísérleti környezetben demográfiai jellemzőktől független értékelést alkalmazott, amennyiben a kritériumok jól definiáltak.

A rendszer a referenciaként szolgáló férfi jelöltnek (*Kovács Dániel*) és a vele teljesen azonos kompetenciákkal rendelkező női jelöltnek (*Nagy Eszter*) egyaránt a maximális, 5/5-ös kompetenciapontszámot adta, ez az eredmény a vizsgált mintában az algoritmus nem mutatott kimutatható nemi alapú eltérést az értékelés során.

Rangsor alapján 3. helyre rangsorolta a technológia azt a jelöltet, aki a többi jelölthöz képest magasabb kompetenciákkal rendelkezett (*Gáspár Ramóna*), mert a kiemelkedő készségei mellett megjelent a vezetői tapasztalat, amit a rendszer túlképzettségnek definiált.

Az algoritmus a senior férfi jelöltet (*Tóth Gábor*) a junior pozícióhoz képest kockázatosnak ítélte. A rendszer "túlmenedzselő" profilként azonosította és alacsonyabb, 3/5-ös pontszámmal a 4. helyre rangsorolta, jelezve a túlképzettséget a pozícióhoz.

A nyelvi kritériumot (tárgylóképes angol) nem teljesítő senior jelöltet (*Szabó Katalin*) a rendszer automatikusan kizárta, függetlenül a jelentős szakmai tapasztalattól.

3. Táblázat. A mesterséges intelligencia-alapú önéletrajz-szűrő algoritmusok értékelésének összesítő elemzése

Jelölt	Végzettség	Tapasztalat	Tárgyalóképes angol nyelvtudás	Szűrés 1. fázis eredmény	Kompetencia-pontszám	Rangsor
Kovács Dániel	✓	✓	✓	Továbbjut	5/5	1. (Megfelel)
Nagy Eszter	✓	✓	✓	Továbbjut	5/5	2. (Megfelel)
Gáspár Ramóna	✓	✓	✓	Továbbjut	4/5	3. (Megfelel)
Tóth Gábor	✓	✓	✓	Továbbjut	3/5	4. (Túlképzett)
Szabó Katalin	✓	✓	✗	Kizárva	N / A	5. (Kizárva)

Saját szerkesztés, saját kutatás alapján

Az ATS-algoritmus kétfázisú értékelési logikájának eredményei. Az 1. fázis egy logikai, bináris szűrő a kötelező kritériumok meglétére (Továbbjut / Kizárva). A továbbjutó jelölteket a rendszer a 2. fázisban egy 1-től 5-ig terjedő kompetencia-skálán értékelte (5/5 = kiváló illeszkedés), amely meghatározta a végső rangsort.

4. Táblázat. Jelöltek preferenciasorrendje a humán szakemberek értékelése alapján

Rangsor	Jelölt	Összesített pontszám	Főbb megállapítás
1.	Kovács Dániel	182	A bírálók legmagasabb preferenciáját élvező jelölt.
2.	Gáspár Ramóna	224,64	Kiváló junior, de a direkt sorrendben megelőzte egy másik jelölt.
3.	Nagy Eszter	235,04	Erősen preferált junior jelölt, a top 3-ban helyezkedik el.
4.	Tóth Gábor	383,64	Túlképzett senior; elutasítás a pozíció profilja miatt.
5.	Szabó Katalin	486,18	Legkevésbé preferált; túlzott senioritás és kizáró ok (nyelvtudás) miatt diszkvalifikálva.

Saját szerkesztés, saját kutatás alapján

$N = 104$, a HR szakemberek szubjektív értékelése alapján felállított végső preferenciasorrend. A jelöltek helyezése 1-től 5-ig terjedő kényszerített rangsorolási skálán alapult, ahol az 1-es érték a legalkalmasabb, az 5-ös a legkevésbé alkalmas jelöltet jelölte. Az összesített pontszám ezen egyéni rangszámok aggregált eredményét tükrözní.

4.2.2. Humán szubjektivitás

A 104 HR-szakember értékelése megmutatta, hogy az emberi döntéshozatalban a szubjektív tényezők felülírhatják az objektív pontszámokat. Az emberi és az algoritmikus döntéshozatal közötti eltérések statisztikai validálása céljából rangkorrelációs vizsgálatot végeztünk. Az elemzés (Spearman-féle rangkorreláció, $r_s = 0,90$) igazolta, hogy bár a HR-szakemberek értékelése alapvetően erős együttmozgást mutatott a mesterséges intelligencia döntési irányával (amelyet a magas korrelációs együttható is jelez), a szubjektív emberi preferenciák szignifikánsan felülírták és módosították a jelöltek végső sorrendjét.

Az eredmények diszkrepanciát mutatnak az egyéni értékelések és a záró rangsor között, míg az egyéni pontozásos szakaszban Ramóna dominált 3,88-as átlaggal, a szakemberek végleges

preferencia-sorrendjében legtöbb esetben már a férfi jelölt, Kovács Dániel foglalta el az első helyet, megelőzve Ramónát és Esztert. A humán döntéshozók esetében a rangsorban enyhe eltérés jelent meg a férfi jelölt javára. Ez arra enged következtetni, hogy a humán döntéshozók körében a rangsor alapján feltételezhető preferencia mutatkozott a férfi jelölt iránt, vagy biztonságosabb választásnak ítélték őt.

A humán értékelők a mesterséges intelligenciához hasonlóan elutasították a senior jelölteket, azonban a döntés háttere eltért. A szakemberek nem kizárólag a szakmai illeszkedést hiányolták, hanem kockázatokat (alacsony motiváció, magas bérigény, fluktuáció) társítottak a túlképzettséghez.

A szakemberek 99,04%-a a tapasztalatot jelölte meg a legfontosabb döntési tényezőként. Érdekes, hogy bár a nemet csak a válaszadók 0,96%-a vallotta be befolyásoló tényezőként, a végső rangsorban mégis megmutatkozott a férfi jelölt enyhe preferenciája.

Az eredmények azt mutatják, hogy az algoritmus technikai értelemben a vizsgált mintában nagyobb demográfiai konzisztenciát mutatott a nemi kérdésekben (teljes egyenlőség), míg a humán döntéshozóknál fellelhető volt a szubjektív preferencia. Ugyanakkor a senioritás / túlképzettség mindkét döntéshozatali modellben erős negatív szűrőként funkcionált, megerősítve a junior pozíciókra jellemző életkori korlátokat.

5. Következtetések

A kutatás eredményeinek értelmezésre rávilágít arra, hogy a mesterséges intelligencia toborzási integrációja nem kizárólag technológiai implementációs kérdés, hanem egy mélyreható bizalmi és etikai dilemma is. Az elemzés arra enged következtetni, hogy az objektivitás a toborzásban nem abszolút, hanem relatív fogalom, természete alapvetően attól függ, hogy algoritmikus szigor vagy humán szubjektivitás vezérli-e a döntést.

5.1. A feltételes elfogadás, mint generációs válaszstratégia

A munkaerőpiacra belépő közgazdászok attitűdjeit nem a technofóbia, hanem egyfajta kockázatkerülő racionalitás jellemzi. A kutatás eredményei arra utalnak, hogy a hallgatók viselkedési szándékát a feltételes elfogadás dominálja, a technológia alkalmazását csak szigorú feltételekkel (átláthatóság és humán felügyelet) mellett támogatják.

Annak ellenére, hogy a válaszadók 91,7%-a racionálisan elismeri a mesterséges intelligencia gazdasági hasznosságát (gyorsaság, adminisztratív tehermentesítés), ez nem alakul át érzelmi bizalommal. A fekete doboz működés és a lélek hiánya miatt a jelöltek attól tartanak, hogy a rendszer a komplex emberi profilt pusztá adatpontokká redukálja, ami a soft skill-ek és a kulturális illeszkedés figyelmen kívül hagyásához vezet.

A hallgatói esszék rámutattak, hogy a hibás algoritmikus értékelés (pl. kulcsszóhiány miatt rosszabb értékelés) nemcsak egyéni sérelem, hanem gazdasági allokációs hiba is, hiszen a vállalatok potenciálisan értékes, de atipikus tehetségeket veszíthetnek el, ami csökkenti a versenyképességet.

5.2. Az objektivitás kettős arca

A kísérleti kutatás legfontosabb következtetése, hogy a „gépi objektivitás” és az „emberi döntéshozatal” eltérő dimenziókban hordoz kockázatokat.

A vizsgálat eredményei nem támasztották alá azt a gyakori előfeltevést, miszerint a mesterséges intelligencia automatikusan reprodukálja a nemi előítéleteket. A kísérletben alkalmazott algoritmus a megfelelően definiált kompetencia-kritériumok mentén a vizsgált esetekben nem mutatott kimutatható nemi alapú eltérést, az azonos szakmai profilú férfi és női jelöltet egyaránt maximális pontszámmal értékelte. Ez az eredmény arra enged következtetni, hogy a strukturált, kompetencia-alapú mesterséges intelligencia rendszerek valójában hozzájárulhatnak a strukturáltabb és konzisztensebb előszűréshez, mivel képesek csökkenteni a humán döntéshozókra esetenként jellemző szubjektív értékelési mintázatokat a szűrés korai szakaszában.

Ezzel szemben a HR-szakemberek döntéshozatala eltérést mutatott az objektív pontszámoktól. A preferencia-sorrend felállításkor a szakemberek a férfi jelöltet rangsorolták az első helyre, annak ellenére, hogy az objektív kritériumok (nyelvtudás, szakmai tapasztalat,

készségek) alapján egy női jelölt profilja erősebb értékelést kapott. Ez az eredmény arra enged következtetni, hogy a strukturált, kompetencia-alapú mesterséges intelligencia rendszerek valójában a méltányosság támogatói lehetnek, mivel képesek csökkenteni a humán döntéshozókra jellemző szubjektív értékelési mintázatokat a szűrés korai szakaszában, ezáltal erősítve az esélyegyenlőséget és a szervezeti diverzitás növelését [22].

5.3. Konszenzus a senioritás elutasításában

A kutatás alapján megállapítható egy rendszerszintű probléma, a senioritás és a túlképzettség kezelése. Mind az algoritmus, mind a humán döntéshozók elutasították a tapasztaltabb jelölteket, de eltérő okokból.

A mesterséges intelligencia a „fit” hiányát jelezte, mivel a jelölt profilja nem illeszkedett a junior specifikációhoz. A HR szakemberek kockázatok (motivációhiány, bérfeszültség, fluktuáció) társítottak a tapasztalathoz. Ez az eredmény abba az irányba mutat, hogy a túlképzettség mindkét döntési modellben erős negatív szűrőként funkcionál.

5.4. Az ideális megoldás: Human-Robot Collaboration (HRC)

A kutatás végső konklúziója, hogy bár az objektivitás a vizsgált mintán demográfiailag semlegesnek bizonyult a mesterséges intelligencia esetében, a toborzás fenntartható jövője nem a teljes automatizációban rejlik. A technikai pontosság és a humán szakértelem kiegészítik egymást. Ebből kifolyólag a legideálisabb megoldás a Human-Robot Collaboration (HRC) modelljének alkalmazása lehet.

Ez a hibrid megközelítés oldhatja fel a hallgatók bizalmi válságát és a HR-szakma hatékonysági kényszerét, megteremtve az egyensúlyt a technológiai innováció és az emberközpontú HR között.

6. Javaslatok

A kutatás eredményei alapján három fő irányvonal javasolt a mesterséges intelligencia etikus és hatályos jogszabályoknak megfelelő toborzási integrációjához.

6.1. Kétszakaszos, audit alapú toborzási modell bevezetése

A teljes automatizáció helyett a Human-Robot Collaboration elvére épülő, kétszakaszos modell implementálása javasolt, amely teljes mértékben illeszkedik az Európai Unió mesterséges intelligenciáról szóló rendeletének (AI Act) toborzási szoftverekre vonatkozó „magas kockázatú” besorolásához.

- **Első szakasz:** A folyamat kezdeti fázisában az MI-alapú rendszerek végzik a nagyméretű adathalmazok feldolgozását és a deduktív, objektív kritériumok mentén történő előszűrést, racionalizálva a tranzakciós költségeket.
- **Második szakasz:** A kritikus döntési pontokon, különösen a szűkített jelöltlista kialakításakor kötelező emberi felülvizsgálatot (auditot) kell beépíteni a nemzetközi standardok (pl. NIST AI Risk Management Framework [14]) iránymutatásai alapján. A humán toborzó feladata a gépi predikciók validálása, valamint az algoritmikusan nem kvantifikálható tényezők értékelése. Ez a fajta megközelítés tudná biztosítani az algoritmikus elszámoltathatóságot.

6.2. HR-szakemberek kompetenciafejlesztése és szerepvállalása

A technológia biztonságos adaptációjának feltétele a toborzók MI-írástudásának növelése. A HR-szakembereknek képessé kell válniuk a gépi predikciók kritikus értelmezésére és az algoritmikus torzítások azonosítására. Az adminisztráció redukálódásával a toborzói szerepkör fókuszának a készségalapú kiválasztás és a stratégiai tanácsadás felé kell mozdulnia.

7. Kutatás korlátai és további kutatási irányok

Bár jelen tanulmány átfogó empirikus adatokra támaszkodva lehetővé teszi a mesterséges intelligencia toborzási szerepével kapcsolatos következtetések levonását, a kutatás érvényességi körét árnyalják bizonyos módszertani korlátok, amelyek egyúttal kijelölik a jövőbeli vizsgálódás irányait is.

A technikai vizsgálat egyik legfőbb korlátja, hogy a kísérlet során egyetlen specifikus ATS (Applicant Tracking System) rendszert alkalmaztunk. Bár a választott szoftver valós piaci környezetben, élesben működő rendszer, az eredmények szükségszerűen ennek az egyetlen algoritmusnak a működési logikáját tükrözik. A jövőben több, eltérő fejlesztési háttérű algoritmussal történő párhuzamos vizsgálat árnyaltabb képet adhatna az objektivitás és a torzítás mértékéről, pontosabban megvilágítva az egyes technológiai megoldások közötti különbségeket.

A humán döntéshozatal vizsgálata kontrollált, kísérleti, ám hipotetikus szituációban zajlott, fiktív önéletrajzok és álláshirdetés felhasználásával. Ez a módszertan bár kiválóan alkalmas az izolált változók mérésére, nem képes teljes mértékben megragadni a valós szervezeti kontextus komplexitását. Olyan kritikus tényezők, mint a szervezeti kultúra nyomása, a valós időnyomás vagy a tényleges felvétellel járó felelősségvállalás súlya ebben a szimulált környezetben nem érvényesülhetnek. A további kutatások során érdemes lenne valós kiválasztási folyamatokba ágyazva vizsgálni, hogyan változnak a HR-szakemberek döntései éles helyzetben.

Szociológiai szempontból korlátot jelent, hogy a hallgatói attitűdök vizsgálata jelen tanulmányban a Z generációra koncentrált. Érdemes lenne a vizsgálatot más generációkra (X és Y generáció) is kiterjeszteni annak érdekében, hogy feltárható legyen a tanulmányban azonosított ambivalencia és bizalmi deficit generációs-specifikus jellemző-e vagy életkortól független, általános társadalmi jelenségként értelmezhető.

A kutatás eredményeire építve fontos és kihívást jelentő további kutatási irány a javasolt Human–Robot Collaboration (HRC) modell gyakorlati implementációjának vizsgálata. Szükséges feltárni, milyen konkrét módszerekkel és munkafolyamat-átalakításokkal valósítható meg operatív szinten az ember és a mesterséges intelligencia szinergikus együttműködése a toborzási folyamatokban.

Végezetül megkerülhetetlen jövőbeli kutatási irányt jelent az új szabályozási környezet, különösen az Európai Unió AI Act előírásainak hatásvizsgálata. Érdemes lenne feltárni, hogy az átláthatóságot és magyarázhatóságot elváró szigorú szabályok miként befolyásolják a HR-szakemberek bizalmát, döntési mozgásterét és a mesterséges intelligencia elfogadását. Ezen összefüggések megértése alapvető fontosságú ahhoz, hogy az MI alkalmazása a toborzásban felelősen, és hosszú távon is fenntartható módon valósuljon meg.

Köszönetnyilvánítás

Külön köszönet illeti a Neumann János Egyetemet a lehetőségért, hogy kutatásomat ilyen támogatott környezetben végezhettem, illetve a hallgatóknak, HR szakembereknek és a vállalatnak, akik részt vettek a kutatásomban és segítették a létrejöttét válaszaikkal.

Köszönöm az Egyetemi Kutatói Ösztöndíj Programnak a bizalmat, hogy jelen kutatás a 2025-2.1.1-EKÖP-2025-00021 projekt keretében jöhetett létre. A 2025-2.1.1-EKÖP-2025-00021 számú projekt a Kulturális és Innovációs Minisztérium Nemzeti Kutatási Fejlesztési és Innovációs Alapból nyújtott támogatásával, a 2025-2.1.1-EKÖP pályázati program finanszírozásában valósult meg.

Irodalomjegyzék

- [1] Aamer, A. K. A. és mtsai. (2023). The impact of artificial intelligence on the human resource industry and the process of recruitment and selection. DOI: https://doi.org/10.1007/978-3-031-26953-0_57
- [2] Anderson, J. W. és Visweswaran, S. (2025). Algorithmic individual fairness and healthcare: a scoping review. JAMIA Open, 8(1), ooae149. DOI: <https://doi.org/10.1093/jamiaopen/ooae149>
- [3] Artificial Intelligence Act (2024). What the Act means for staffing businesses. Elérhető: <https://artificialintelligenceact.eu/what-the-act-means-for-staffing-businesses/>
- [4] Asfari, A. & Gacek, J. (2024). The diversity-dissent paradox: Navigating police and university recruitment challenges amidst campus protest dynamics. DOI: [10.1177/26338076241297997](https://doi.org/10.1177/26338076241297997)
- [5] Babu, S. & Moorthy, A. D. (2023). Application of artificial intelligence in adaptation of gamification in education: A literature review. Computer Applications in Engineering Education, 32(1), 1-16. DOI: <https://doi.org/10.1002/cae.22683>
- [6] Berec, N. (2022). Szakdolgozat - Mesterséges intelligencia a HR-ben, valamint a gépi tanuló algoritmus alkalmazása
- [7] Blue Colibri (2024). A HR szoftver szerepe és előnyei a vállalati adminisztrációban. [online] Elérhető: <https://hu.bluecolibriapp.com/blog/hr-szoftver-ceges-adminisztracioban>
- [8] Chen, Z. (2023). Ethics and discrimination in artificial intelligence-enabled recruitment practices. Humanities and Social Sciences Communications, 10(1), 1-12. DOI: [10.1057/s41599-023-02079-x](https://doi.org/10.1057/s41599-023-02079-x)
- [9] Davis, F. D. (1989). Perceived usefulness, perceived ease of use, and user acceptance of information technology. MIS Quarterly, 319–340. DOI: <https://doi.org/10.2307/249008>
- [10] Dióssi, K. & Mikáczó, A. (2023). Mesterséges intelligencia a HR folyamatok, főként a toborzás támogatásában.
- [11] Kappel, G. (2024). Mesterséges intelligencia alapú döntéstámogató és döntéshozó rendszerek kockázatai a vállalatok vezetői szintű döntéshozatalában: Szakirodalmi áttekintés. Profuturo. DOI: <https://doi.org/10.26521/profuturo/2024/1/15166>
- [12] Kovács, L. (2023). A mesterséges intelligencia alkalmazása a HR-ben: Magyarországi helyzetkép. Vezetéstudomány, 54(1), 12-25.
- [13] McCombs School of Business. (2025). Algorithmic Bias | Ethics Defined. Elérhető: <https://youtu.be/cygpAm4ooGs?si=ndCfKqgiACIn09kH>
- [14] National Institute of Standards and Technology. (2023). Artificial Intelligence Risk Management Framework (AI RMF 1.0) (NIST AI 100-1) DOI: <https://doi.org/10.6028/NIST.AI.100-1>
- [15] Nawaz, N. és mtsai. (2024). The adoption of artificial intelligence in human resources management practices. International Journal of Information Management Data Insights, 4(1), 1-11. DOI: <https://doi.org/10.1016/j.ijime.2023.100208>
- [16] Norvig, P. & Russel, S. J. (2023). Mesterséges intelligencia – Modern megközelítésben I. kötet.
- [17] O'Donnellan, R. (2025). Machine learning by the numbers: Its impact on business. Elérhető: <https://www.intuition.com/en-sa/machine-learning-by-the-numbers-its-impact-on-business/>
- [18] Oracle (2019). AI in Human Resources: The Time is Now. [online] Elérhető: <https://www.oracle.com/a/ocom/docs/applications/hcm/oracle-ai-in-hr-wp.pdf>
- [19] Ribánszki, N. & Karcsics, É. (2025). Mesterséges intelligencia alkalmazása a toborzási folyamatokban. Gradus, 12(1). DOI: <https://doi.org/10.47833/2025.1.ECO.006>
- [20] Slattery, P. és mtsai. (2024). A Systematic Evidence Review and Common Frame of Reference for the Risks from Artificial Intelligence. DOI: [10.13140/RG.2.2.28850.00968](https://doi.org/10.13140/RG.2.2.28850.00968)
- [21] Trautwein, Y. és mtsai. (2025) 'Opening the "black box" of HRM algorithmic biases – How hiring practices induce discrimination on freelancing platforms', Journal of Business Research, 192, 115298. DOI: <https://doi.org/10.1016/j.jbusres.2025.115298>
- [22] Wilkens, U. és mtsai. (2025) 'Augmenting diversity in hiring decisions with artificial intelligence tools', The International Journal of Human Resource Management, 36(14), pp. 2585–2622. DOI: <https://doi.org/10.1080/09585192.2025.2492867>