



Kognitív torzítások az AI-támogatott kriminalisztikában

Jogállami kockázatok és holisztikus mitigációs lehetőségek

Cognitive Biases in AI-Assisted Forensics Rule-of-Law Risks and Holistic Mitigation Approaches

Orbán József

Dr., PhD, szakértő
Nemzeti Média és Hírközlési Hatóság,
Védelmi és Rendészeti Frekvenciagazdálkodási Igazgatóság
josef.orban.mba@gmail.com



Absztrakt

Cél: A tanulmány célja, hogy bemutassa, a mesterséges intelligencia (AI) kriminalisztikai alkalmazása miként alakítja át a hibázás természetét. A szerző az AI-támogatott döntésekhez kapcsolódó automatizmus-torzítás, algoritmikus tekintélytorzítás, megerősítési torzítás, precizitási illúzió és felelősség-diffúzió kriminalisztikai jelentőségét vizsgálja.

Módszertan: A tanulmány a hazai és nemzetközi szakirodalomra, valamint a kriminalisztikai tendenciákra támaszkodva a bizonyíték teljes életciklusát holisztikus keretben elemzi és azonosítja a kognitív torzítások belépési pontjait.

Megállapítások: Az elemzés rámutat, hogy az AI-támogatott kriminalisztikában a kognitív torzítások nem kiküszöbölődnek, hanem új formában, láncolat-szerűen jelennek meg, és átalakítják a jogbiztonság érvényesülését, ha az algoritmikus eredmények technikai tekintélyként, magyarázat nélkül, intézményi elvárásként épülnek be a döntéshozatalba.

Érték: A tanulmány újdonsága, hogy az AI-támogatott kriminalisztikát a kognitív torzítások és a jogállami garanciák összefüggő, holisztikus kockázati ökoszisztémájaként értelmezi, és ehhez egy integrált mitigációs keretrendszert rendel. A gyakorlati ellenőrzőlistával kiegészített javaslat közvetlenül hasznosítható a nyomozati gyakorlatban, a szakértői munkában és a kriminalisztikai képzésben.

A szerző a kéziratot magyar nyelven nyújtotta be. Benyújtás ideje: 2026. 03. 10. Átdolgozás: 2026. 04. 07. Elfogadás: 2026. 05. 21.



Kulcsszavak: AI-támogatott kriminalisztika, kognitív torzítások, jogállami kockázatok, holisztikus jogi szemlélet

Abstract

Aim: The aim of this study is to demonstrate how the application of artificial intelligence (AI) in forensics transforms the nature of error. The author examines the forensic relevance of biases associated with AI-supported decision-making, including automation bias, algorithmic authority bias, confirmation bias, the illusion of precision, and the diffusion of responsibility.

Methodology: Building on both national and international literature, as well as contemporary forensic trends, the study analyses the entire lifecycle of evidence within a holistic framework and identifies the entry points of cognitive biases.

Findings: The analysis highlights that in AI-assisted forensics, cognitive biases are not eliminated but rather reappear in new forms and in chain-like patterns. These biases reshape the enforcement of legal certainty, particularly when algorithmic outputs are incorporated into decision-making as technical authority, without sufficient explanation, and under institutional expectations.

Value: The novelty of the study lies in interpreting AI-assisted forensics as an interconnected, holistic risk ecosystem of cognitive biases and rule-of-law safeguards, and in proposing an integrated mitigation framework. The recommendations, complemented by a practical checklist, are directly applicable in investigative practice, expert work, and forensic education.

Keywords: AI-assisted forensics, cognitive biases, rule-of-law risks, holistic legal approach

Bevezetés

A bírósági eljárás klasszikus római felfogás szerint nem öncélú formai rítus, hanem az igazság megtalálásának eszköze: Diocletianus rescriptuma szerint „*aliud nihil in iudiciis quam iustitiam locum habere debere*” (Fenyvesi, 2008), amely a folyamatot támogató kriminalisztikára is igaz kell, hogy legyen. Ennek elérése érdekében a szakemberek igyekeznek újabb és újabb eszközöket és módszereket bevonni az eljárásba.

A mesterséges intelligencia (AI) kriminalisztikai alkalmazása az elmúlt évtizedben látványos fejlődésen ment keresztül: az adatvezérelt nyomozás, a prediktív elemzés, a mintázatfelismerés és az automatizált döntéstámogatás ma már a bűnüldözés mindennapi eszközei közé tartoznak (Vári, 2024).

E technológiák megjelenése a bizonyítás megbízhatósága és a jogállami garanciák érvényesülése szempontjából alapvető jogelméleti és eljárásjogi dilemmát vet fel, különösen az emberi döntéshozatal kognitív korlátaival összevetve (Oerlemans & Royer, 2023).

A tanulmány három ponton járul hozzá az AI-támogatott kriminalisztika szakirodalmához.

- Egyrészt rendszerezi és AI-környezetben újraértelmezi a kriminalisztikai döntéshozatal öt kiemelt kognitív torzítását (automatizmus-torzítás, algoritmikus tekintélytorzítás, megerősítési torzítás, precizitási illúzió, felelősség-diffúzió), feltárva azok láncolatszerű működését a bizonyíték teljes életciklusában.
- Másrészt holisztikus jogi szemléletben elemzi, hogy e torzítások miként alakítják át a tisztességes eljárás, az indokolási kötelezettség, az elszámoltathatóság és a jogbiztonság érvényesülését AI-támogatott nyomozási környezetben.
- Harmadrészt egy technológiai, eljárási, szervezeti és kulturális elemeket integráló, gyakorlati ellenőrzőlistával kiegészített mitigációs keretrendszert javasol, amely a kognitív torzítások mérséklését a jogállami büntetőeljárás keretei között teszi lehetővé.

Elméleti háttér

A kriminalisztika tárgyát, alapelveit és a bizonyítás központi szerepét a 21. századi hazai szakirodalomban Fenyvesi dolgozta ki részletesen (Fenyvesi, 2013a; 2022). A kriminalisztika a büntetőeljárás azon területe, ahol az eljárási cselekmények döntő többsége nem mechanikusan végrehajtott normakövetés, hanem mérlegelést igénylő szakmai döntések sorozata.

A kriminalisztikai döntések közös jellemzője, hogy információhiányos környezetben, idő- és szervezeti nyomás alatt születnek, következményeik pedig utólag tipikusan nem visszafordíthatók, amelyben a hibakockázat nem véletlenszerűen előforduló rendkívüli jelenség, hanem hasonló strukturális feltételek mellett ismételten érzékelhető, a teljes egészt jellemző tulajdonság (Fenyvesi, 2013b). A holisztikus jogi szemlélet ebből kiindulva a kriminalisztikai hibát nem egyéni mulasztásra egyszerűsíti, hanem rendszerszintű kockázatként értelmezi. Ez a megközelítés illeszkedik ahhoz a Fenyvesi-tételhez, amely a kriminalisztika fejlődését a justizmordhoz vezető hibák csökkentésének történeteként írja le (Fenyvesi, 2013b).

A kriminalisztikában ezek a hibák és torzítások fokozott jelentőséggel bírnak, mivel a döntések normatív és ténykérdéseket egyszerre érintenek (Czebe & Kovacs, 2015; Fenyvesi et al., 2022).

Az AI mint döntéstámogató technológia a kriminalisztikában

Az úttörőnek tekinthető, a nagy kockázatú mesterséges intelligencia felhasználási szabályozását is célzó EU jogalkotás (2024/1689 EU rendelet) ellenére is számos kérdés maradt nyitott terület a kutatók számára, így a kriminalisztikában is. A kutatási rést a hazai végrehajtási rendelet sem szűkítette (2025. évi LXXV. tv.). Az AI kriminalisztikai alkalmazásai alapvetően döntéstámogató funkciót töltenek be. Ide sorolhatók az adatbázis-összefüggések feltárása, a mintázatfelismerés, kockázati és a valószínűségi becslések, a gyanúsított vagy az eseményrangsorok készítése. Hangsúlyosan említendő, hogy ezek a rendszerek olyan eszközök, amelyek az emberi döntéshozatalt befolyásolják. Az AI ebben az értelemben a döntés kognitív környezete, amely strukturálja az elérhető információkat, kiemel bizonyos mintázatokat, háttérbe szorít más adatokat, így nem csökkenti automatikusan a kognitív torzításokat, hanem átalakítja azok működését (Herke, 2021).

Amennyiben az AI-eredmények technikai (szak)tekintélyként jelennek meg, vagy magyarázat nélkül kerülnek alkalmazásra, ad abszurdum intézményileg elvárt iránymutatássá válnak, úgy az emberi döntéshozó szerepe formálissá válhat, miközben a jogi felelősség továbbra is őt terheli. A fentiek alapján megállapítható, hogy az AI-támogatott kriminalisztikában az emberi torzítások strukturális jellegűek (Fenyvesi, 2013a; 2014; Fenyvesi et al., 2022).

AI és kriminalisztikai tendenciák – technológiai fordulat

Az AI kriminalisztikai alkalmazása minőségileg új döntési környezetet hoz létre. Ebben a környezetben a kockázatok nem kizárólag technikai jellegűek, hanem az emberi mérlegelés, az intézményi működés és a jogi garanciák metszéspontjában keletkeznek (Lontai, Pamjav, & Petrétai, 2024a; 2024b). Az AI-támogatott nyomozás így illeszkedik abba a tendenciába, amelyet Fenyvesi a krimináltechnika tényeréseként és a kriminalisztika világtendenciáiként ír le (Fenyvesi, 2013a). A holisztikus jogi szemlélet ezért az AI-támogatott kriminalisztikát komplex kockázati ökoszisztémaként értelmezi. A működésének alapja az adatok nagy mennyiségű és gyors feldolgozása. Ez a gyakorlatban azzal jár, hogy a nyomozás fókusza részben adatközpontúvá válik: az releváns, ami mérhető, strukturálható és algoritmikusan feldolgozható.

Ez a szemléletváltás több kockázatot hordoz, mivel a nehezen formalizálható, kvalitatív információk háttérbe szorulhatnak, továbbá a bizonyítási relevancia technikai elérhetőséghez kötődhet. Az egyik legjelentősebb kockázata az átláthatatlanság. Számos algoritmus – különösen a komplex gépi tanulási modellek – esetében

a döntési út nem, vagy csak korlátozottan rekonstruálható. A forrásprobléma ott keletkezik, hogy már a nyomozó sem érti az eredmény keletkezését. Ebben a folyamatban a nyomozás irányát nem jogi, hanem adatlogikai szempontok befolyásolhatják. A relevancia ilyen átalakulása a bizonyítás szabadságának elvével kerülhet feszültségbe, különösen akkor, ha az AI által nem „észlelt” információk valójában kiesnek az értelmezési körből (Fenyvesi & Fábián, 2024).

Az AI gyakran intézményi szinten támogatott, sőt elvárt gyakorlat, ami azonban technológiai tekintélyt hoz létre, így az AI-eredmények implicit módon „jobb”, „objektívebb” döntésként jelennek meg. E környezetben a kritikus szakmai kétely kockázatos magatartássá válhat, ahol az egyéni mérlegelés szervezeten belül nem ösztönzött. A legérzékenyebb terület a felelősség megosztása. Bár az AI formálisan döntéstámogató eszköz, a gyakorlatban a döntési folyamat több szereplő között oszlik meg úgy, hogy a nyomozók és a jogalkalmazók mellett olyan laikus csapat is megjelenik, mint a fejlesztők, az adatgazdák és a rendszerüzemeltetők.

Ez a többszereplős modell felelősség-diffúzióhoz vagy részleges felelősség nélkülséghez is vezethet, ahol a hibák oka nehezen azonosítható, az egyéni jogi felelősség elmosódik, az intézményi elszámoltathatóság gyengül.

A kockázati környezet holisztikus értelmezése

Az AI-alapú kriminalisztikai rendszerek a korábbi adatokon tanulnak, ami növeli a hatékonyságot, ugyanakkor azt a kockázatot hordozza, hogy a múltbeli nyomozási gyakorlatok, torzítások és strukturális aránytalanságok beépülnek a rendszer működésébe (Herke, 2021). Ennek következménye lehet bizonyos mintázatok túltreprezentálása, továbbá más jelenségek láthatatlanná válása. Legnagyobb kockázati tényezőként a kriminalisztikai gondolkodás beszűkülését azonosíthatjuk. Ez a logika rokon a negatív indíciумok kriminalisztikai értékelésével, amelyre Fenyvesi és Fábián is felhívja a figyelmet (Fenyvesi & Fábián, 2024). E kockázatok nem izolált jelenségek, hanem a kognitív torzítások formájában sűrűsödnek össze.

Kognitív torzítások megjelenési formái az AI-támogatott kriminalisztikában

Az AI-támogatott kriminalisztikai környezetben a kognitív torzítások nem egyszerűen fennmaradnak, hanem új működési formát öltenek. Az algoritmusok által strukturált információs tér átalakítja az emberi észlelést, a döntési prioritásokat

és a felelősségérzetet. A következőkben bemutatott torzítások nem elszigetelt jelenségek, hanem egymással összefonódó torzításláncokat alkothatnak a bizonyítási folyamat során (Roth, 2017; Hefetz, 2025).

Automatizmus-torzítás, algoritmikus tekintélytorzítás és átláthatatlansági torzítások

Az automatizmus-torzítás az egyik leggyakrabban dokumentált jelenség (Jinad et al., 2024). Lényege, hogy az emberi döntéshozó – időnyomás vagy szervezeti elvárások esetén – hajlamos az algoritmus által javasolt eredményt kritikai vizsgálat nélkül elfogadni (Czebe & Kovacs, 2015; Martire et al., 2025). Kriminálisztikai környezetben ez megnyilvánulhat az AI által rangsorolt gyanúsítottak automatikus előtérbe helyezésében. Ide idézhetjük a hipotéziseket támogató, továbbá azok bekövetkezésének esélyeit csökkentő tények arányba állításának fontosságát (likelihood ratio), mint hibaszűrő eszközt.

Az algoritmikus tekintélytorzítás jelensége szorosan összefügg az automatizmus-torzítással, amely abból fakad, hogy az AI által előállított eredmények, az érvrendszer és a megfogalmazások objektívebbnek, tudományosabbnak vagy pontosabbnak látszódnak, mint az emberi megállapítások (Fenyvesi et al., 2022). Ez különösen erős akkor, ha az eredmény numerikus, grafikus vagy statisztikai formában jelenik meg. Ez az AI-szakértői vélemények kritikátlan elfogadásához, az AI-eredmények felülvizsgálatának elmaradásához, „a rendszer így számolta” típusú indokolásokhoz vezethet, amely érinti az indokolási kötelezettséget, mivel az indoklás technikai hivatkozássá redukálódhat, amely jogilag nem ellenőrizhető. Ez a jelenség gyengíti az emberi szakértői tekintélyt.

Az átláthatatlansági torzítás akkor lép fel, amikor az AI eredményét nem érti a döntéshozó és ezért nem is kérdőjelezi meg. A „feketedoboz” jellegű rendszerek ezt a torzítást strukturálisan erősítik. Ez a torzítás közvetlenül érinti az ellenőrizhetőség és elszámoltathatóság elvét. Ez a megfigyelés is indokolja a megmagyarázható AI (explainable AI) jogos igényét (Hefetz, 2025; Herke, 2021).

Azonban figyelembe kell venni, hogy az ember heurisztikus gondolkodásmódja képes kétségtelenül szokatlan, ugyanakkor jó megoldásokat találni.

Megerősítési torzítás, precizitási illúzió, felelősség-diffúzió és együttes hatás

A megerősítési torzítás klasszikus jelenség a kriminalisztikában, azonban AI-támogatott környezetben felerősödhet. Az algoritmusok gyakran előzetes paraméterezés, adatválasztás vagy hipotézis mentén működnek, így könnyen az eredeti

nyomozói feltételezések visszaigazolására szolgálnak. Ennek következménye, hogy a keresési tér beszűkül, de az érzékeny szem számára észlelhetővé válik a „releváns” találatok túlsúlya, továbbá az alternatív magyarázatok kiszorulása. A szakirodalom és kriminalisztikai tapasztalatok alapján az AI hajlamos a hibásan rögzített hipotézisek megerősítésére.

Az AI által szolgáltatott numerikus vagy valószínűségi kimenetek – például százalékos egyezések vagy kockázati pontszámok – hitelesebbnek tűnnek, mint amennyire ténylegesen azok. Itt említhető a bizonyítás túlzott mértékben egzaktként való értelmezése, a jogi mérlegelés mechanikussá válása, a bizonyítási küszöbök merev alkalmazása. A precizitási illúzió a jogi döntéshozatalt álbizonyosság felé tolhatja el. A felelősség-diffúzió lényege, hogy az AI-támogatott kriminalisztikai döntések több szereplő (fejlesztők, adatgazdák, üzemeltetők, nyomozók, jogalkalmazók) között oszlanak meg úgy, hogy a hibák okai elmosódnak, az egyéni jogi felelősség nehezen azonosítható, és ezzel az intézményi elszámoltathatóság is gyengül. Az automatizmus-torzítás összekapcsolódhat a tekintélytorzítással, az átláthatatlanság erősítheti a precizitási illúziót, míg a felelősség-diffúzió akadályozza a hibák feltárását.

A tartós AI-használat kedvezőtlen hatásait a későbbiekben részletesen elemezzük.

A kognitív torzítások hatása a bizonyíték holisztikus életciklusára

A holisztikus jogi szemlélet a bizonyítékot nem elkülönült eljárási mozzanatként, hanem dinamikus életcikluson át fejlődő jogi objektumként értelmezi. Ebben a megközelítésben a kognitív torzítások nem egyetlen döntési pillanatban jelennek meg, hanem láncolatszerűen hatnak végig a bizonyíték teljes életciklusán. AI-támogatott kriminalisztikai környezetben ez a hatás különösen felerősödik.

A bizonyíték életciklusának első szakasza a felismerés, amely során eldől, hogy egy információ egyáltalán bizonyítékként értelmezhető-e. Az AI-támogatott környezetben az információs tér jelentősen befolyásolja az emberi észlelést, az algoritmus által „kiemelt” adatok nagyobb figyelmet kapnak, így az AI által nem jelzett információk háttérbe szorulhatnak, az első AI-eredmény rögzítési (lehorgonyzási) hatást vált ki. A kognitív torzítás már itt meghatározza a bizonyítás irányát, mivel a későbbiekben a fel nem ismert bizonyítékok gyakran végleg kiesnek az eljárásból (Fenyvesi, 2021).

Az AI által relevánsnak jelölt elemek részletesebb dokumentálása mellett más, potenciálisan fontos információk alulrögzítésre kerülhetnek (Fenyvesi &

Fábián, 2024). Ennek következményeképp a bizonyítékállomány korán torzulhat és a védelem mozgásteret veszít.

A holisztikus kriminalisztikai szemlélet alapelvei szerint ez a torzulás nem utólag korrigálható technikai hiba, hanem eljárási minőségi probléma jelenik meg. Ez a szakasz jól mutatja, hogy a technikai integritás nem garantálja a jogi megbízhatóságot (Fenyvesi, 2013b).

A bizonyíték elemzési szakasza a holisztikus életciklus leginkább torzításérzékeny pontja. Az AI által szolgáltatott eredmények gyakran értelmezési keretnek adnak, amelyen belül az emberi döntéshozó gondolkodik. Alternatív magyarázatok így nem feltétlenül azért szorulnak háttérbe, mert cáfolhatók, hanem mert nem illeszkednek az AI által sugallt narratívába.

A bizonyíték bemutatása során a kognitív torzítások az AI-eredmények kiemelt vizuális kezelése, a bizonytalanságok háttérbe szorítása és az indokolás technikai nyelvre való leegyszerűsítése formájában jelenhetnek meg.

A bizonyíték életciklusa az ítélethozatallal nem zárul le. Jogorvoslatok, felülvizsgálatok és módszertani tanulságok révén visszahat a rendszer egészére. A kognitív torzítások azonban itt is akadályozhatják az AI kívánatos fejlődését. Gyakori problémaként jelenik meg az utólagos igazolási torzítás, intézményi önvédelem, az AI-rendszer kritikátlan további alkalmazása. Ennek következtében a rendszer nem tanul a hibákból, hanem újratermeli azokat, immár technológiailag megerősített formában (Kovács, 2022).

A kognitív torzítások hatása a bizonyíték holisztikus életciklusában nem lineáris, hanem egymást erősítő folyamatok sorozata. Kockázati tényezőként jelezzük, hogy az AI-támogatott kriminalisztikában a torzítások korábban lépnek be a folyamatba, mélyebben strukturálják azt, és nehezebben vagy egyáltalán nem korrigálhatók utólag.

Ez alátámasztja a tanulmány alapvető tézisé, miszerint a bizonyítás minősége nem kizárólag a technológia fejlettségétől, hanem attól függ, hogy a rendszer mennyire képes felismerni és kezelni saját kognitív korlátait.

Jogállami és eljárásjogi kockázatok

A kognitív torzítások technológiai közvetítése olyan eljárásjogi kockázatokat hoz létre, amelyek a tisztességes eljárás alapelveit érintik. E kockázatok sajátossága, hogy gyakran nem nyílt jogsértésként, hanem az eljárás minőségének fokozatos romlásaként jelennek meg.

A tisztességes eljárás különösen akkor problematikus, ha az AI által generált eredmények meghatározzák a nyomozás fókuszát, alternatív nyomozati irányok

de facto kizárásra kerülnek, az eljárás résztvevői eltérő mértékben férnek hozzá az AI működéséhez és eredményeihez. Ilyen esetekben a tisztességes eljárás formálisan fennáll, tartalmilag azonban sérülhet.

Az indokolási kötelezettség a jogállami büntetőeljárás egyik sarokköve. Az AI-támogatott kriminalisztikában fennáll annak veszélye, hogy az indokolás technikai hivatkozássá válik, amely kiüresedve nem teszi lehetővé az érdemi ellenőrzést. A jellemző problémák között említhetjük az AI-eredmények „adottként” való kezelését, az algoritmikus működés részleteinek hiányát az indokolásban, továbbá a döntési lánc megszakadását az emberi és technikai szintek között.

E folyamat a kontradiktórus eljárást tartalmilag gyengíti, még akkor is, ha annak formai keretei fennmaradnak. A büntetőeljárás egyik alapelve, hogy a döntésekért azonosítható jogalanyok viselnek felelősséget. AI-támogatott kriminalisztikában azonban a felelősség többszereplős technológiai lánc mentén oszlik meg. Ennek következménye a felelősség technológiára való áttolása, a döntéshozók pszichológiai leválása a következményekről, így az elszámoltathatóság gyakorlati nehézségekbe ütközhet. Jogállami szempontból különösen aggályos, ha egy technikai eszköz ténylegesen döntéshozóként jelenik meg, miközben jogilag nem kérhető számon. Ha a fegyverek egyenlőségének elve így torzul, az a büntetőeljárás legitimitációját hosszú távon alááshatja.

A jogbiztonság megköveteli, hogy a jogalkalmazás előrelátható és következetes legyen. AI-támogatott kriminalisztikában azonban a folyamatosan tanuló rendszerek működése dinamikusan változhat, ami a döntések kiszámíthatóságát csökkentheti. Ez problematikus a jellegükben hasonló ügyek ellenére eltérő algoritmikus viselkedés esetén, modellfrissítések dokumentálatlanságakor, valamint egységes módszertani keretek hiányában. A jogbiztonság ilyen módon nem jogszabályi, hanem technológiai okokból sérülhet (Gless et al., 2024).

Holisztikus enyhítési keretrendszer, módszertani javaslatok

A holisztikus jogi szemlélet abból indul ki, hogy a kognitív torzítások nem egyéni hibák, hanem az ember–technológia–intézmény kölcsönhatásából fakadó rendszerszintű jelenségek. Ennek megfelelően az enyhítési keretrendszernek egyszerre kell lefednie a technológiai, eljárási, szervezeti és kulturális dimenziókat (Fenyvesi, 2013a; 2013b).

A technológiai jobbítás célja nem az AI „objektivizálása”, hanem annak biztosítása, hogy az algoritmus ellenőrizhető, értelmezhető és korlátozott szerepű maradjon a bizonyítási folyamatban. A kulcselemek között említhető

a magyarázhatóság, a bizonytalanság explicit kezelése, az auditálhatóság és a verziókövetés.

Az AI által szolgáltatott eredményeknek olyan formában kell megjelenüniük, hogy azok értelmezése az eljárás résztvevői számára lehetséges legyen. A magyarázhatóság nem informatikai luxus, hanem eljárásjogi minimumkövetelmény.

Valószínűségi kimenetek esetén a bizonytalansági tartományok, feltételezések és korlátok kötelező megjelenítése szükséges a precizitási illúzió elkerülése érdekében. Az alkalmazott modellek változásait dokumentálni kell, hogy utólag rekonstruálható legyen, milyen technológiai állapotban keletkezett egy adott eredmény. Az eljárási jobbítás középpontjában az áll, hogy az AI-eredmények ne váltsák ki, hanem strukturált módon támogassák az emberi mérlegelést. Az AI által generált eredmények önmagukban nem alapozhatnak meg döntést; minden esetben kötelező az emberi felülvizsgálaton alapuló döntés, amit szintén dokumentálni kell (Cărcăle, 2025). Az eljárási rendnek ösztönöznie kell az alternatív magyarázatok vizsgálatát, különösen akkor, ha az AI egy domináns narratívát sugall (Casu et al., 2025). A bizonyíték életciklusának kritikus szakaszainál (felismerés, elemzés, prezentáció) előre meghatározott ellenőrzési és beavatkozási pontok szükségesek. A kognitív torzítások jelentős része nem egyéni, hanem szervezeti szinten keletkezik. Szükséges a felelősségi szerepek egyértelmű meghatározása és a döntésekért mindig azonosítható személynek kell felelősséget viselnie. Az AI-eredményekkel szembeni szakmai kétely nem deviancia, hanem elvárt magatartás kell legyen. Időszakos módszertani és jogi felülvizsgálatok szükségesek az intézményes torzítások feltárására.

A kriminalisztikai oktatásnak tudatosan kell foglalkoznia a döntési torzításokkal, különösen AI-környezetben. Az esetek utólagos elemzése, hibák feltárása és megosztása nem szankciós, hanem tanulási célt kell szolgáljon. Jogászok, kriminalisták, informatikusok és pszichológusok együttműködése elengedhetetlen.

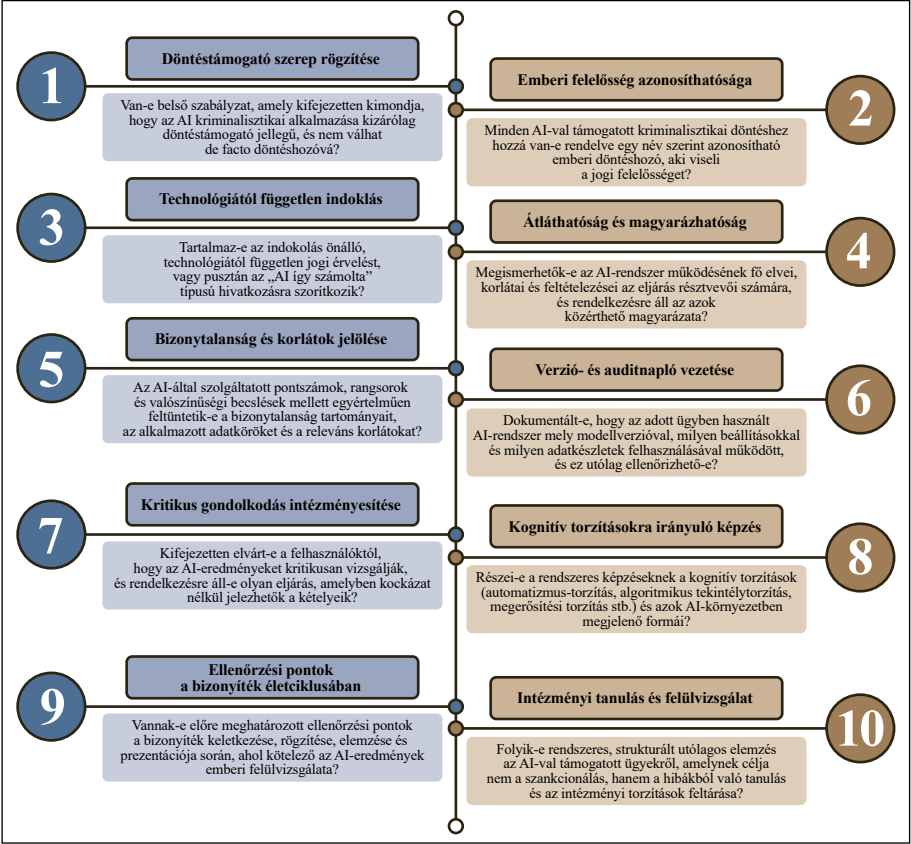
A bemutatott jobbítási szintek nem hierarchikusak, hanem egymást kiegészítő elemek, így a hatékony kockázatkezelés csak integrált keretrendszerként valósítható meg.

Az AI kognitív torzításának mérséklésére kiválóan felhasználható a vakteszt, a felismerésre bemutatás, a független ellenőrzés, a szabványos protokollok alkalmazása, a képzés és egy masszív minőségirányítási rendszer kiépítése.

A fentiek érdekében biztosítani kell, hogy az AI a konkrét azonosítási feladatról ne kapjon felesleges kontextusinformációt. A felismerésre bemutatásnál a mintákat olyan sorrendben mutatják be, hogy az AI nem ismeri, hogy melyik a „gyanúsított” minta. A független ellenőrzésnél más gyártótól származó AI szükséges, amely nem ismeri az első AI következtetését, és vizsgálatot végez. Nem elhanyagolható elemként kell kiemelni a kötelezően hibákra fókuszáló humán szakértő független véleményének figyelembevételét.

Az 1. számú ábra ellenőrzőlistája összegzi az AI-támogatott kriminalisztika javasolt ellenőrzési pontjait.

1. számú ábra
Gyakorlati minimumjavaslat az AI-támogatott kriminalisztikához



Forrás. A szerző saját szerkesztése.

Az ábrán látható, hogy a szabványos protokollok alkalmazása, szigorú metodológiai irányelvek követése, valamint a pontos dokumentálás a hibacsökkentési törekvések alapjának tekinthetők.

Az AI eszközök hiba-visszacsatolása, hibavizsgálat ROC görbéken keresztül (Orbán, 2019; Rotello & Chen, 2016), a szakértők rendszeres képzése a torzításokról, valamint az intézményi szinten beépített ellenőrzési mechanizmusok következetes alkalmazása jelentősen növelheti az eredmények megbízhatóságát.

Összegzés

Ezek után a kérdés nem az, hogy mit tud az AI, hanem az, hogy mit engedünk meg neki a büntetőeljárásban. Az AI alkalmazása a kriminalisztikában nem technológiai modernizációs, hanem eljárásminőségi kérdés. A tanulmány álláspontja szerint az AI kriminalisztikai alkalmazása csak akkor járulhat hozzá az igazságszolgáltatás minőségének javításához, ha az szigorú jogállami keretek között, holisztikus szemlélettel történik.

Felhasznált irodalom

- Cărcăle, V.-A. (2025). The Future of Artificial Intelligence (AI) Applications in Forensics. <https://doi.org/10.5281/ZENODO.15481632>
- Casu, M., Guarnera, L., Zangara, I., Caponnetto, P., & Battiato, S. (2025). A (Mid)journey Through Reality: Assessing Accuracy, Impostor Bias, and Automation Bias in Human Detection of AI-Generated Images. *Human Behavior and Emerging Technologies*, 25(1), 9977058. <https://doi.org/10.1155/hbe2/9977058>
- Czebe, A., & Kovacs, G. (2015). The impact of bias in latent fingerprint identification. 2015 6th IEEE International Conference on Cognitive Infocommunications (CogInfoCom), 569–574. <https://doi.org/10.1109/CogInfoCom.2015.7390656>
- Fenyvesi Cs. (2008). *Szembesítés Szemtől szembe a bűnügyekben*. Dialóg Campus Kiadó.
- Fenyvesi Cs. (2013a). A kriminalisztika alapkérdései. In Gaál I., & Hautzinger Z. (Szerk.), *Pécsi Határőr Tudományos Közlemények*, 14. (pp. 93–107). Magyar Hadtudományi Társaság.
- Fenyvesi Cs. (2013b). *A kriminalisztika tendenciái. A bűnügyi nyomozás múltja, jelene, jövője* (1. kiad.). Dialóg Campus Kiadó.
- Fenyvesi Cs. (2014). A justizmordhoz vezető kriminalisztikai hibák. *Belügyi Szemle*, 62(2), 45–60.
- Fenyvesi Cs. (2021). A felismerésre bemutatások hibái. *Iustum Aequum Salutare*, 25–39.
- Fenyvesi Cs., & Fábíán V. (2024). Negatív indíciumok üzenetei bírónak. *Bírák Lapja*, 1–16.
- Fenyvesi Cs., Herke Cs., & Tremmel F. (Szerk.) (2022). *Kriminalisztika*. Ludovika Egyetemi Kiadó.
- Gless, S., Lederer, F. I., & Weigend, T. (2024). *AI-Based Evidence in Criminal Trials?* William & Mary Law School.
- Hefetz, I. (2025). Evaluating bias in forensic evidence: From expert analysis to AI-based decision tools. *Forensic Science International: Synergy*, 11, 100645. <https://doi.org/10.1016/j.fsisyn.2025.100645>
- Herke Cs. (2021). A mesterséges intelligencia kriminalisztikai aspektusai. *Belügyi Szemle*, 69(10), 1709–1724. <https://doi.org/10.38146/BSZ.2021.10.2>
- Jinad, R., Gupta, K., Simsek, E., & Zhou, B. (2024). Bias and fairness in software and automation tools in digital forensics. *Journal of Surveillance, Security and Safety*, 5(1), 19–35. <https://doi.org/10.20517/jsss.2023.41>

- Kovács G. (2022). A kognitív torzítás veszélyei a szakértői véleményalkotásban. *Jog Állam Politika*, 14(3), 25–43.
- Lontai M., Pamjav H., & Petrétei D. (2024a). Artificial Intelligence in Forensic Sciences Revolution or Invasion? Part I. *Belügyi Szemle*, 72(4), 701–715. <https://doi.org/10.38146/bsz-ajia.2024.v72.i4.pp701-715>
- Lontai M., Pamjav H., & Petrétei D. (2024b). Artificial Intelligence in Forensic Sciences Revolution or Invasion? Part II. *Belügyi Szemle*, 72(8), 1513–1525. <https://doi.org/10.38146/bsz-ajia.2024.v72.i8.pp1513-1525>
- Martire, K. A., Neal, T. M. S., Gobet, F., Chin, J. M., Berengut, J. F., & Edmond, G. (2025). Psychological insights for judging expertise and implications for adversarial legal contexts. *Nature Reviews Psychology*, 4(4), 264–276. <https://doi.org/10.1038/s44159-025-00430-4>
- Oerlemans, J.-J., & Royer, S. (2023). The future of data-driven investigations in light of the Sky ECC operation. *New Journal of European Criminal Law*, 14(4), 434–458. <https://doi.org/10.1177/20322844231212661>
- Orbán, J. (2019). Evidence, Probability and ROC in Crime Cases. “Towards a Better Future: Democracy, EU Integration and Criminal Justice”, 55–64.
- Rotello, C. M., & Chen, T. (2016). ROC curve analyses of eyewitness identification decisions: An analysis of the recent debate. *Cognitive Research: Principles and Implications*, 1(1), 10. <https://doi.org/10.1186/s41235-016-0006-7>
- Roth, A. L. (2017). Machine Testimony. *Yale Law Journal*, (126), 1972–2053.
- Vári, V. (2024). The Intersection of Crime Geography and Police Work. *Criminal Geographical Journal*, (3–4), 39–66.

Alkalmazott jogszabályok

Az Európai Parlament és a Tanács (EU) 2024/1689 rendelete (2024. június 13.) a mesterséges intelligenciára vonatkozó harmonizált szabályok megállapításáról, valamint a 300/2008/EK, a 167/2013/EU, a 168/2013/EU, az (EU) 2018/858, az (EU) 2018/1139 és az (EU) 2019/2144 rendelet, továbbá a 2014/90/EU, az (EU) 2016/797 és az (EU) 2020/1828 irányelv módosításáról (a mesterséges intelligenciáról szóló rendelet).

2025. évi LXXV. törvény az Európai Unió mesterséges intelligenciáról szóló rendeletének magyarországi végrehajtásáról

A cikk APA szabály szerinti hivatkozása

Orbán J. (2026). Kognitív torzítások az AI-támogatott kriminalisztikában. Jogállami kockázatok és holisztikus mitigációs lehetőségek. *Belügyi Szemle*, 74(8), 2153–2166. <https://doi.org/10.38146/bsz-ajia.2026.v74.i8.pp2153-2166>

Nyilatkozatok

Összeférhetetlenség

A szerző nem jelentett összeférhetetlenséget.

Finanszírozás

A szerző nem kapott pénzügyi támogatást a kutatáshoz, a szerzőséghez és/vagy a cikk publikálásához.

Etikai nyilatkozat

Jelen cikkhez nem kapcsolódik adatkészlet.

Nyílt hozzáférésről szóló tájékoztatás

A közlemény nyílt hozzáféréssel, a Creative Commons Attribution–NonCommercial–NoDerivatives 4.0 International License (CC BY-NC-ND 4.0) feltételei szerint jelenik meg. A közlemény változatlan formában, nem kereskedelmi célból, bármely médiumban és formátumban szabadon megosztható és újraközölhető, feltéve, hogy az eredeti szerző(k), a közlés helye és a licenc hivatkozása megfelelően feltüntetésre kerül. Módosított vagy átdolgozott változat terjesztése, valamint kereskedelmi célú felhasználása nem megengedett.

Levelező szerző

A cikk levelező szerzője Orbán József, aki a jozsef.orban.mba@gmail.com e-mail címen érhető el.