

Válaszok a mesterséges intelligencia kockázataira

MOLNÁR ATTILA

RelNet Technológia Kft.

A mesterséges intelligencia terjedése az infokommunikációs infrastruktúrákban és a vállalati döntéstámogatásban átrajzolja a kiberbiztonság kockázati térképét. Az autonóm MI-ágensek megjelenése olyan új fenyegetési kategóriákat teremt, amelyekre a hagyományos, szabályalapú védelmi eszközök nem adnak kielégítő választ. Az MI-alapú kibertámadások és a szervezeteken belüli MI-kockázatok, továbbá a védekezés technológiai és szabályozási eszközei, végül a 2026. április végi AI Hungary Szimpózium, mind releváns szakmai tanulságokkal szolgálnak a hazai infokommunikációs szakemberek számára, akiknek egyre sürgetőbb feladatuk a mesterséges intelligencia kontrollált és biztonságos integrációja saját szervezeteikbe és ügyfeleik rendszereibe.

Bevezetés

Az elmúlt néhány évben a mesterséges intelligencia fejlődési üteme minden korábbi előrejelzést túlszárnyalt. Az infokommunikációs hálózatok anomáliadetektálásától [1] kezdve, az oktatáson és a munkapiacra zajló strukturális változásokon át [2], egészen a közlekedési infrastruktúra védelméig [3] az MI ma már jelenlévő valóság a Híradástechnika folyóirat tanúsága szerint is. A Híradástechnika 2024-es különszámának egyik áttekintése [4] joggal állapítja meg, hogy a technológia hatása az internet vagy az okostelefon elterjedéséhez hasonlítható átalakulást hoz, azzal a különbséggel, hogy ez az átalakulás minden technológiai (és társadalmi!) átalakulásnál gyorsabban zajlik.

E gyorsaság az informatikai biztonság szempontjából kettős kihívást jelent. Egyfelől az MI-eszközök a támadói oldal számára is elérhetőek, és már ma is valós fenyegetéseket generálnak. Másfelől a védekezés módszertana (például viselkedéselemzés, megfelelés, szervezeti folyamatok újratervezése stb.) szükségszerűen lemaradásban van a támadói innovációhoz képest. A RelNet Technológia Kft. ezt a kettős valóságot tapasztalja napi munkája során, és erre keres megoldásokat.

Apokalipszis vagy utópia?

A közbeszéd hajlamos a szélsőségekre: vagy a technológiai szingularitás borúlátó forgatókönyveit, vagy a minden problémát megoldó utópiát vizionálja. A valóság sokkal árnyaltabb. Dávid András, a RelNet Technológia Kft. alapító-ügyvezetője egy nyilvános interjúban [5] fogalmazta meg azt az iparági tapasztalatot, amely szerint az MI hatásának mértéke és iránya ma még nem látható át teljesen; aki azt állítja, hogy biztosan tudja, mire számíthat, az félrevezeti a szakmai közvéleményt.

A szingularitás hipotézise szerint az általános mesterséges intelligencia (angol rövidítésben: AGI) egy ponton túl önmaga fejlesztésére lesz képes, és ezzel a technológiai fejlődés emberi léptékkel nem modellezhető szakaszba lép. Ez a spekuláció hasznos keretező gondolat: jelzi, hogy a fejlődési görbe nem lineáris. Nem csupán az aktuális fenyegetésekre kell felkészülnünk, hanem egy folyamatosan változó kockázati térképet kell kezelni.

A következőkben bemutatott kockázatok és védelmi megközelítések a jelenlegi technológiai szinten már is fennállnak. Ugyanakkor kezelhetőek, feltéve, hogy a szervezetek és a szabályozói környezet kellő időben reagálnak.

Jó MI, rossz MI

Az MI-alapú kibertámadások tekintetében a legfontosabb fejlemény az automatizált, nagyszabású adathalászkampányok ugrásszerű növekedése. A nagy nyelvi modellek (ún. LLM-ek) képesek tömegesen előállítani olyan megtévesztő szövegeket, amelyek korábban idő- és munkaigényes előkészítést igényeltek. Ennél is komolyabb fenyegetést jelent az úgynevezett „prompt injection” jelenség, amelynek során a támadó rosszindulatú utasításokat csempész be az MI-rendszer bemenetébe, egy megosztott dokumentumba vagy egy e-mail szövegébe. Ezzel az MI-eszközt saját céljaira kényszeríti, például adatok kinyerésére, bizalmas tartalmak külső szerverre való továbbítására stb.

Az MI-eszközökkel történő védekezés ugyanakkor fejlődő területnek számít. A statikus, szabályalapú tűzfalak és szűrők az MI-alapú támadásokkal szemben kevésbé hatékonyak, ugyanis az MI-modellek nehezen választják le a végrehajtandó utasításokat a kapott adatokról, különösen, ha azok idegen nyelven, emoji-kkal vagy láthatatlan karakterekkel álcázottan érkeznek. A védekezés tehát szükségszerűen viselkedéselemzésen alapuló, dinamikus megközelítést igényel.

Külön kihívást jelent a *nem emberi belső szereplők* kategóriájának megjelenése. Amikor egy szervezet autonóm MI-ügynököket telepít belső folyamatainak automatizálására, ezek az ágensek vállalati hitelesítő adatokkal rendelkeznek, hálózati hozzáférésük van, de tevékenységük sok esetben nem kerül ugyanolyan biztonsági felügyelet alá, mint az emberi felhasználóké. Egy 2024-es iparági felmérés szerint a kiberbiztonsági szakemberek 74%-a nyilatkozott úgy, hogy szervezetén belül a rosszindulatú belső szereplők aggasztják legjobban. [7] Ebbe a kategóriába a nem megfelelően konfigurált MI-ügynökök is beletartoznak.

A fenyegetésre adott egyik válasz az MCP (Model Context Protocol), vagyis az MI-alkalmazások és külső rendszerek összekapcsolására szolgáló nyílt szabvány. Ez a protokoll önmagában hasznos infrastrukturális réteg, amely lehetővé teszi, hogy különböző MI-rendszerek és eszközök egységes felületen kommunikáljanak egymással, ugyanakkor új támadási felület nyílik, ha ez a kommunikációs csatorna nincs megfelelően védve. A piac már kínál olyan biztonsági megoldásokat, amelyek kifejezetten az MI-k közötti kapcsolatot ellenőrzik.

Jogi és szabályozási vonzatok

Az MI-kockázatok kezelésének szabályozói dimenziója az elmúlt két évben drámaian felgyorsult. Az Európai Unió MI-rendelete (EU AI Act, 2024/1689/EU) 2024. augusztus 1-jén lépett hatályba [8], és kockázatalapú megközelítéssel szabályozza az MI-rendszerek fejlesztését és alkalmazását. A rendelkezések fokozatosan válnak kötelezővé. 2026 augusztusától már a magas kockázatú rendszerekre vonatkozó megfelelőségi előírásokat is alkalmazni kell.

Az AI Act négy kockázati kategóriát különít el: elfogadhatatlan, magas, korlátozott és alacsony kockázatú MI-rendszereket. Az olyan, kibervédelmi szempontból leginkább érintett alkalmazások, mint a hozzáférés-menedzsment, az anomáliadetektálás vagy a viselkedéselemzés, jellemzően a magas kockázatú kategóriába esnek, ami részletes dokumentációs, auditnaplózási és megfelelőségi kötelezettséggel jár. A szabályok megsértése esetén kiszabható bírság elérheti a 35 millió eurót vagy a vállalat éves globális forgalmának 7%-át.

Különös jelentőséggel bír az elszámoltathatóság kérdése az MI-ágensek autonóm döntéshozatala kapcsán. Ha egy MI-ágens hibás döntése pénzügyi vagy adatvédelmi kárt okoz, a jogi felelősség megállapítása jelenleg még jogértelmezési *terra incognita*. A folyamatban lévő szabályozási munka ezért várhatóan az emberi felügyelet és a hamisításbiztos naplózás kötelező jellegét fogja erősíteni.

Az AI Hungary Szimpózium tanulságai

A HTE 2026. április 28-29-én rendezte meg az AI Hungary 2026 szimpóziumot. Deklarált célja volt az MI „átállítása távoli technológiai kockázatból a versenyképes gazdaság, a hatékony állam és a minőségi közszolgáltatások központi infrastrukturális elemévé”. [9] A RelNet Technológia Kft. nem csupán résztvevőként, hanem aktív szervezőként volt jelen: a vállalat az informatikai piacon elérhető, MI-kockázatokra specializálódott megoldásokat képviselő gyártók szakértőit hozta el, hogy a konferencia kiberbiztonsági diskurzusát konkrét technológiai tartalommal gazdagítsa.

Nem emberi ügynökök

A szimpózium egyik legfontosabb eredménye az autonóm MI-ügynökök szervezeti kockázatának szakmai felismerése volt. A plenáris előadások és workshopok egyaránt hangsúlyozták, hogy a szervezetek ma már olyan MI-ágenseket telepítenek belső munkafolyama-



Pillanatkép a HTE mesterséges intelligencia szimpózium 2026-ról

taikba, amelyek sokszor egyértelmű identitás vagy elszámoltathatóság nélkül mozognak a hálózaton. A biztonságtechnikai ipar egyik válasza erre az ágens alapú viselkedéselemzés. Ez a megközelítés a hagyományos felhasználói viselkedéselemzés (UEBA) kiterjesztését jelenti a nem emberi entitásokra, többek között valós idejű anomáliaérzékeléssel.

Az egyik workshopon bemutatásra került, hogy miként képes egy biztonsági naplómenedzsment (SIEM) rendszerbe integrált viselkedéselemző platform felismerni és automatikusan megakadályozni a nem kívánt MI-tevékenységeket, mielőtt azok tényleges adatvédelmi vagy biztonsági eseménnyé eszkalálódna. A legtöbb biztonsági csapat ugyanis még nem tudja megkülönböztetni az emberi tevékenységet az MI-ágensekétől.

Adatvédelem és hozzáférés-menedzsment

Az AI Hungary konferencián bemutatott megoldások alapelve az egységes adatbiztonsági életciklus volt, amibe beletartozik az adatok észlelése, osztályozása, figyelése is, egyetlen keretrendszerben. Az attribútum alapú hozzáférés-vezérlés (ABAC) és a kvantumrezisztens titkosítási szabványok alkalmazása a mesterséges intelligencia tanítási és futásidejű adatainak védelmére olyan megközelítések, amelyek az informatikai piacon ma már elérhető, termékkész megoldásokban öltenek testet.

A hagyományos szoftverek kiszámítható viselkedésével szemben az AI rendszerek a hozzáférés-kezelés területén is komoly kihívást jelentenek. A Gartner 2029-re vonatkozó előrejelzése szerint az MI-ágensek ellen indított sikeres kibertámadások több mint fele hozzáférés-vezérlési hiányosságokat fog kihasználni, jellemző

en közvetlen vagy közvetett „prompt injection” útján. [10] Ez különösen aggasztó a már említett nem emberi identitások tekintetében, mivel ezek az MI-ágensek által használt, sokszor túlzott jogosultságokkal rendelkező szolgáltatásfiókokat jelölik.

OT-biztonság

Az MI-alapú automatizáció az ipari rendszerek (OT) vonatkozásában sajátos kockázatot hordoz. Szemben az IT-hálózatokkal, az OT-környezetekben az üzemzavar fizikai biztonsági kockázatokkal is járhat. A konferencián bemutatott megközelítések a kiemelt jogosultságok kezelésének (angol rövidítéssel: PAM) kiterjesztésére összpontosítottak a nem emberi identitások irányába.

Az AI Hungary Manifesto

A szimpózium felsővezetői workshopjának résztvevői – köztük Dávid András, a RelNet Kft. ügyvezető igazgatója – rögzítették az AI Hungary Manifesto [9] legfontosabb stratégiai elveit. A kiáltvány négy pillére:

- a mérhető versenyképességi mutatókkal alátámasztott MI-transzformáció;
- az adatszuverenitás és a lokális (on-premise) modellek stratégiai értékének elismerése;
- a gyakorlatorientált szabályozás (az EU AI Act, a GDPR és a NIS2) egyetlen, innovációbarát keretrendszerként való értelmezése;
- valamint a szervezetekben szabályozatlanul terjedő MI-használat felszámolása biztonságos, auditálható vállalati alternatívák biztosításával.

Az együttműködés az eddig 2 független kisebb fogalmat elvivő, kevésbé kihasznált hálózat helyett egy közös, kezdetben magas kihasználtságú, összeségében kisebb sáv szélességet igénylő, ebből adódóan nagyobb spektrumhatékonyságú 2G hálózatot eredményez.

Óvatos transzformáció

A RelNet Technológia Kft. által képviselt megközelítés lényege az, hogy az MI-adaptáció elkerülhetetlen és kívánatos, de nem mehet a kontroll rovására. A szervezeteknek ügyelniük kell arra, hogy megőrizzék a kontrollt az MI-ügynökök és az adatvagyon felett. Ehhez azonban tudatos menedzselésre, amolyan kötélházi munkára van szükség, hiszen az innováció ritmusa sem lassulhat.

Az infokommunikációs szakemberek felelőssége ebben az összefüggésben kivételes. A technológia fejlődési üteme meghaladja a szabályozói és szervezeti alkalmazkodás természetes sebességét. Így, a szakértők feladata az, hogy elérhetővé tegyék a láthatóságot és a biztonságos MI-átalakulás eszköztárát.

Összefoglalás

A mesterséges intelligencia kockázatai az infokommunikáció területén három egymást erősítő síkon jelentkeznek: a technológiai fenyegetések szintjén (MI-alapú támadások, prompt injection, ágensek kompromittálása), a szervezeti irányítás szintjén (nem emberi identitások kezelése, Shadow AI, auditnaplózás hiánya) és a szabályozói megfelelés szintjén (EU AI Act, GDPR, NIS2 összehangolása). Ezek egyike sem kezelhető izoláltan.

A 2026. áprilisi AI Hungary szimpózium megerősítette, hogy a hazai szakmai közösség tudatában van e kihívások összetettségének, és aktívan keresi az integrált válaszokat. A konferencián bemutatott megoldások nem elszigetelt innovációk, hanem egy átfogó biztonsági architektúra építőkövei. A RelNet Kft. ezeket a megoldásokat aktívan közvetíti a hazai szervezetek felé, meggyőződve arról, hogy a versenyképes és biztonságos MI-adaptáció egymást kizáró célok helyett egymást erősítő feladatok.

Hivatkozások

- [1] Farkas Károly: „Mesterséges intelligencián alapuló eljárások alkalmazása infokommunikációs hálózatokban”, Híradástechnika LXXVIII/1 2023, HTE Infokom 2022, pp. 32–39.
- [2] Székely Levente: „Mitől félünk, ha a farkastól nem?”, Híradástechnika LXXIX/1 2024, HTE Infokom 2023, pp. 17–21.
- [3] Bódi Antal, Maros Dóra: „Az 5G-hálózat és a közlekedés információbiztonsági kihívásai”, Híradástechnika LXXVI/1 2021, HTE Infokom 2020, pp. 35–40.
- [4] Gyires-Tóth Bálint: „Mesterséges intelligencia – a harmincas években”, Híradástechnika LXXIX/2, 2024 Különszám – HTE 75 Infokommunikációs víziók, pp. 13–15.
- [5] Dávid András: „AI nagyobb változást hoz, mint az elmúlt 200 év összesen?” – videóinterjú. Elérhető: <https://www.youtube.com/watch?v=TNoFbBNsB4s>
- [6] AI Hungary Manifesto – HTE AI Hungary 2026 Szimpózium. Elérhető: <https://hte.hu/hu/web/guest/e/hirek1/10108/5506944>
- [7] Amid Rising Insider Threats, Most Companies are Vulnerable and Lack the Strategies to Protect Themselves, New Research Finds: https://www.securonix.com/press_release/2024-insider-threat-report/
- [8] AZ EURÓPAI PARLAMENT ÉS A TANÁCS (EU) 2024/1689 RENDELETE a mesterséges intelligenciára vonatkozó harmonizált szabályok megállapításáról: https://eur-lex.europa.eu/legal-content/HU/TXT/HTML/?uri=OJ:L_202401689
- [9] AI HUNGARY MANIFESTO infografika: https://www.hte.hu/documents/10180/5014263/AI_Hungary_Manifesto_v515.pdf
- [10] AI Cybersecurity Leadership: 5 Steps to Secure Enterprise Innovation: <https://www.gartner.com/en/articles/ai-cybersecurity-leadership>

A szerzőről



MOLNÁR ATTILA

több évtizedes tapasztalattal rendelkezik a digitális tartalomgyártás és a technológiai kommunikáció területén. Pályáját nyelvtudományi fókusszal kezdte, majd az informatikai szektor felé fordulva a két terület szinergiájára építette karrierjét. Szakmai útja során tapasztalatot szerzett technológiai ismeretterjesztő cikkek,

oktatási anyagok és B2B marketingtartalmak megalkotásában. Jelenleg a RelNet Technológia Kft. szenior tartalommarketing-specialistája, ahol fő feladata a hálózatbiztonsági megoldások hiteles szakmai közvetítése. (e-mail: amolnar@relnet.hu)