

Mesterséges intelligencia a kibervédelemben

A mesterséges intelligencia a nemzeti biztonságban és a kibertámadásokban

Az AI használata a nemzeti biztonságban kézenfekvő. Hozzáférhető tanulmányok sora vizsgálta az elmúlt pár év során a terület tendenciáit és eredményeit. Találhatunk közöttük általános összefoglaló,¹ a katonai területre összpontosító tanulmányt,² a használhatóság osztályozására koncentráló írást,³ de a kínai⁴ vagy akár az orosz vonatkozásokat is feldolgozó cikkeket⁵ is.

Ehhez szeretnék hozzájárulni egy részterület, a mesterséges intelligencia alapú kibervédelem helyzetének vázlatos áttekintésével. A téma apropója a 2020. szeptember 8-án bemutatott Magyarország Mesterséges Intelligencia Stratégiája 2020–2030, amely szerint „az MI új támadási felületeket és módszereket jelent a kibervédelem területén”,⁶ valamint Magyarország új Nemzeti Biztonsági Stratégiája⁷ is, amely nemrég lépett hatályba, és korszerűen kezeli korunk biztonsági kihívásait.

A veszély világos: ha ügyesen megírt hagyományos malware-ek is olyan sebességgel tudnak fertőzni, amivel ember nem veheti fel a versenyt, mennyivel inkább veszélyes lesz egy AI által támogatott kibertámadás a maga változékonyságával és nagyságrendileg nagyobb sebességével? Szakmai körökben ma már evidencia a kijelentés,⁸ hogy az AI-t szükséges alkalmazni a kibervédelemben is. Az alábbiakban elsősorban a nem szakmabeliek számára szeretnék kis körképet adni a kibervédelmi AI fontosságának megértéséhez.

Két probléma merülhet fel: biztonságos-e az AI használata, és megtérül-e a jelentős többletköltség. A válasz mindkét kérdésre: igen – ennek indoklásául azonban picit alaposabban végig kell néznünk a témát. Kezdjük ezt technikai áttekintéssel, hogy lássuk,

¹ Daniel Hoadley – Nathan Lucas: *Artificial Intelligence and national security*. Washington, Center for New American Security, 2018.

² Porkoláb Imre – Négyesi Imre: A mesterséges intelligencia alkalmazási lehetőségeinek kutatása a haderőben. *Honvédségi Szemle*, 147. (2019), 5. 3–19.

³ Gregory Allen – Taniel Chan: *Artificial Intelligence and national security*. Cambridge, Belfer Center Study, 2017.

⁴ Gregory C. Allen: *Understanding China's AI Strategy: clues to Chinese strategic thinking on Artificial Intelligence and national security*. Washington, Center for New American Security, 2019.

⁵ Koós Gábor – Szternák György: A mesterséges intelligencia katonai felhasználásának lehetőségei. *Felderítő Szemle*, 18. (2019), 4. 117–135.

⁶ A tanulmány megírásakor sajnos ez még nem volt hozzáférhető, ezért itt csak utalok rá.

⁷ 1163/2020. (IV. 21.) Kormányhatározat a Magyarország Nemzeti Biztonsági Stratégiájáról 1. sz melléklet.

⁸ Ez bármelyik hivatkozott tanulmányból kiderül.

mivel is állunk szemben, ezután összefoglaljuk hogyan állnak a témához a különböző szektorok 2020 nyarán.

Technikai áttekintés

A védelemi technikák általában kétféle módszert követnek, nincs ez másképp a kibertérben sem.

1. Új paradigmák felállítása: ekkor már nem folyhat a védekezés az eddig megszokott módszerek és elvek mentén, jelentős újításokra van szükség, nem elegendő a régit csiszolni.
2. A támadás módszerének felhasználása. Átalakítva sok módszer védelemre is használható.

Alább két külön fejezetben nézünk példákat mind a két típusú védelmi elvre.

Új paradigmák

Múltalapú védelem – már (mindenhol) a múlté

A múltalapú védelem azon az elven nyugvó régi paradigma, hogy az a biztos, ami már megtörtént – a jövőt senki sem látja előre. Tervezéskor abból indultak ki, hogy a dolgok hasonlóan fognak történni, mint a múltban, ha pedig közbejött valami, akkor B vagy C terv lépett életbe. A múltalapú védelem elégtelensége nemcsak az informatikában, hanem sok más területen is korunk sajátosságává vált, érdemes tehát általánosítva megvizsgálnunk.

Napjainkban a statisztikai módszerek, a pszichológia és szociológia fejlődésével jól képzett elemzők egyre inkább képesek lettek megjósolni tömeges jelenségeket. Az egyre pontosabb gazdasági szimulációk átalakították a tőzsdéket, a jövőmodellek kihatottak a politikára és az élet számos területére, a védelmi tudományok is jó ideje használnak különféle szimulációs modelleket. Ezekhez már a kezdeteknél kézenfekvő volt az akkor megjelenő számítógépekre támaszkodni. Ott még ugyan nem tartunk, ahol Isaac Asimov híres *Alapítvány*ában, de a tömeges jelenségek lehetséges kimeneteleinek az esélyét egyre pontosabban meg tudjuk határozni.

Ám az ember nem így működik. A személy továbbra is saját tapasztalatának érzelmi emlékeire támaszkodik. Mindennapi életünket nem statisztikai valószínűségek kiértékelésében éljük. Nem vagyunk képesek a környezetünket folyton elemezve figyelembe venni minden kockázatot; és nem is akarjuk. Ki több, ki kevesebb veszélyről tud – szakértettsége szerint –, és senki nem az összesről.

Általános jelenség, hogy hiába ismert egy fenyegetés, sőt annak elhárítási módja, másnak talán még mondjuk is, hogy vigyázzon – de saját életünkben soha nem teszünk meg minden biztonsági lépést. Esetleg nem érünk rá, vagy nincs kedvünk megtenni.

A jelenség lélektani háttere, hogy a mindennapokban az ember a „velem nem történhet meg” alapattitúddal él, és csak akkor számol igazán komolyan egy-egy fenyegetéssel, miután ő maga vagy közvetlen környezete megtapasztalta – akkor is csak egy ideig, mivel nem képes folyamatos stresszhelyzetben élni. Természetes emberi reakcióként a (szerinte) kevésbé valószínű veszélyekről a tudatunk nem vesz tudomást, hiszen egy egészséges lélek nem akar folyamatos rettegésben élni.

Ez az oka, hogy a mai napig működhetnek a klasszikus betörések és az ezeréves átvetések, amelyekről minden gyerek tud, és az ellenszer is rég rendelkezésre áll. Az óvintézkedések oktatásával a csalók dolga nehezebbé válik, de a nagy számok törvénye alapján mindig akadnak, akik épp akkor figyelmen kívül, nincs kedvük vagy energiájuk az óvintézkedésekre, vagy épp nem jut eszükbe az.

A fenti jelenség hatványozottan van jelen a kibertérben. Egyrészt azért, mert ott senkinek nem lehet fizikailag baja, és ez tudat alatt mindenkit hamis biztonságérzettel tölt el. Másrészt, mert a támadóknak sokkal könnyebb dolguk van, mintha fizikailag támadnának, mivel százmilliónyi vagy milliárdnyi próbálkozást lehet tenni gyorsan és titokban.

A védelmi programozók sokáig csak annyit tudtak elérni, hogy egy-egy sikeres felderítés után feketelistára tették a támadót. Ám egy ilyen vírus képes volt térben és időben még sokáig fertőzni, mivel az ellenszert (a frissítést) sokan nem telepítették a „csak más gépe lesz vírusos” elv alapján. Később az internet terjedésével automatizálódott a frissítés, ám a közelmúltig folyamatos problémát jelentett, hogy a beazonosított kártevő módosított változatához újabb frissítés kellett. Gyorsabb volt sok-sok vírusváltozatot legyártani, mint azokhoz a frissítéseket. A védelmi fejlesztők ezért már jóval az AI megjelenése előtt próbáltak szélesebb védelmet programozni, amelyben nemcsak konkrét vírusokat képes a pajzs leleplezni, hanem a gyanús tevékenységre utaló kódokat is érzékeli. Így jöttek létre az úgynevezett heurisztikus módszerek, amelyeket ahhoz a kábítószer-ellenes szabályozáshoz lehet hasonlítani, amikor nem a drog konkrét képletet tiltották be, hanem a hatóanyagot magát. A kifejezés az ismert ógörög „Heuréka!” felkiáltásból ered, amely állítólag akkor hangzott el, amikor Arkhimédész rájött híres törvényére a kádban. Bizonyos törvényszerűségek mentén ezek az algoritmusok is rájönnek egy adott tevékenység vagy kód gyanús voltára. A kortárs bűnügyi pszichológia is alkalmazza ezt a módszert, hiszen az emberi mozgás- és viselkedésminták elemzésével a ráutaló magatartások sok típusa ismerhető fel. Például egy kritikus rendezvényen megfigyeléssel keresik azokat az embereket, akik szerepelnek a veszélyes személyeket tartalmazó adatbázisban, miközben az elhárítás figyeli a gyanús viselkedésmintákat az adatbázisban nem szereplő embereken is. Ezek alapján a tömegből sok potenciális szabotőr kiemelhető. Hasonlóan működik együtt a múltalapú kibervédelem vírusadatbázisa a heurisztikus támadáselemzővel.

Ám ha ezek a szabotőrök egy különleges tréning során megtanulják, hogy hogyan kerüljék el ezeket a felfedezhető viselkedési sémákat, akkor nehezen lepleződnek le. Az AI-alapú támadási kódok is megtanulnak alkalmazkodni, és az álcázást is hatékonyan segíti az AI. Például a rajvírusok egymásnak küldhetnek információkat lelepleződéskor, így tanítva egymást, fejlesztve a rajtudást, amelynek új, mutálódott egyedei legközelebb

megpróbálják kikerülni az előző hibát.⁹ Vagyis már nem elegendőek a múlton alapuló védelmek, még ügyes tippelő (heurisztikus) rendszerrel megerősítve sem.

A gépi tanulás azonban szárnyakat adhat ezeknek a heurisztikus védelmi módszereknek, amikor a felhasználási módok és a támadási szokások elemzésével képes megkülönböztetni az ügyfelek tevékenységét és a rosszindulatú fenyegetések által indított hasonló tevékenységet. A felhasználó számára ez úgy érzékelhető, hogy kevesebbszer téved, mint elavult társai. Az AI is a múltra alapoz, azonban nem csupán eltárolja az információkat: az ilyen rendszer nem vár arra, hogy programozó szakemberek majd egyszer átírják (frissítsék) a védelmet szolgáló algoritmusokat. Az alapvető különbség a korábbi adatbázis + heurisztika kettőshöz képest, hogy itt az adatbázist is az AI-rendszer építi folyamatosan, és a heurisztikát is ő „találja ki”, nem pedig programozók. Más szóval nem statikusan a múltra épül, hanem dinamikusan a múltból „tanul”. Így képessé válik egyszerű dolgokat megjósolni, sőt önmaga módosítására javaslatot tenni. Egy jól betanított AI a hagyományos támadó algoritmusok közel 100%-a ellen képes pajzsot képezni úgy, hogy közben a rendszer a jogosult humán felhasználókat nem nézi támadónak. Tehát megvan a szinte tökéletes ellenszer! Csakhogy, amennyiben nem egy hagyományos vírus, hanem egy másik AI támad, azt igen nehéz lesz megkülönböztetni az emberi mintáktól, hiszen a támadó AI is a humán aspektusokat fogja másolni. Tehát az AI-val végzett támadás esetén pont úgy kell eljárunk, mint régen, az eseményt utólag kell elemeznünk, és az alapján javítani vagy módosítani saját védelmünket Vagyis az AI-alapú védelem egy sokrétegű modern páncélhoz hasonlatos, amely jól véd minden régi fegyver (nyílvesszők, dárdák) ellen, sőt egy sor, „olcsó” lőfegyver ellen is – de nem csodapajzs minden, a modern fizika és kémia eredményeit felhasználó löveggel szemben. Lesz, ami áttöri!

Valós idejű védelem

Az informatikában a valós idő olyasmit jelent, hogy ember által (szinte) érzékelhetetlen késéssel történik valami. Mindenki tapasztalt már zavaró, pár másodperces hangcsúszást egy lassú hálózaton zajló IP-beszélgetés esetén. Ám a kibervédelemben többnapos, akár többhetes eltérés volt a támadás, annak észlelése, majd a megszüntetése között. A valós-idejűség azért nagy lépés, mivel ez a reakcióidő csökkent milliszekundumnyira. A közel-múltig biztonsági szakemberek vizsgálták át folyamatosan a rendszereket, figyelték a védelmi szoftverek naplózását, a gyanús eseményekről szóló riasztásokat. Azonban még ez a drága munkaerő sem elég gyors: számtalan sikeres támadás mögött az áll, hogy az emberi elemzőképeségek legjava sem képes megközelíteni a hatékony vírusok fertőzési sebességét.¹⁰ Mire megtalálják és kiirtják a támadást, már sok kár keletkezett.

⁹ Ivan Zelinka et alii: Swarm virus – Next-generation virus and antivirus para-digm? *Swarm and Evolutionary Computation*, 43. (2018). 207–224.

¹⁰ Max Heinemeyer: *The next paradigm shift – AI-driven cyber-attacks*. DarkTrace Research, White Paper, 2. 2018.

A hagyományos, múltalapú rendszerelemző algoritmusok csupán segédeszközök voltak a biztonsági team kezében.

Ezeknek a rendszerelemzéseknek a nagy része átadható a mesterséges intelligenciának. Így a védelem valós időben képes a tevékenységsorozatokat elemezni, és azokban a fenyegetésre utaló mintákat felfedezni. A jó szakember ráadásul fárad és hibázik – az AI-rendszer nem. Képtelenség volt minden hálózati forgalmat és rendszertevékenységet jól megvizsgálni, és csak akkor engedni, ha rendben van. Eddig egyfajta „ártatlanság véelme” elv alapján engedélyezni kényszerültünk mindent. Míg nem derült ki valamiről, hogy ártalmas, addig nem volt letiltva a rendszerben. Ez is hozzájárult, hogy a hagyományos védelem mellett a jól álcázott támadások sokáig kifejtették káros tevékenységüket.

Eddig két lehetőség volt: ha a védelem automatikusan és alaposan megvizsgált mindent – alaposan lelassította a rendszert. Ha pedig folyton kérdezgett, például egy munkaasztali védelmi modul, azt a felhasználó általában megunt, és úgy állította át, hogy engedje, amit csak lehet – utat nyitva a támadásoknak (központi védelemnél pedig szellemi munkaerőt foglalt le). Az AI használatával lehetőség nyílt arra, hogy kevesebb emberi beavatkozással különböztesse meg a védelem, hogy mi az ártalmas, és mi nem. A védelem AI-alapú fejlődése hasonlítható a beléptető vagy átvilágító kapuk adta biztonsági szintugráshoz, amely alig csökkentve az átengedési sebességet képes mindenkiről megállapítani, van-e nála veszélyesnek minősülő holmi, vagy lázas-e. Vagy hasonlítható a kamerákat figyelő jó biztonsági őrhöz, aki észreveszi a gyanús viselkedésű személyt az épületben.

A klasszikus védekezésnél az is komoly gond, hogy abban nem optimális a humán-erőforrás-kihasználás hatékonysága, hasonlóan ahhoz, amikor a madáchi falanszterben Michelangelók faragnak széklábat. Nagy tehetségű, sokáig képzett és komoly szakmai gyakorlattal rendelkező szakemberek több műszakba osztva, idegőrlő bejegyzés-bogarászást végeznek, ami jól fizető, ám unalmas, hosszabb távon élvezhetetlen munka. Ráadásul korunkban ezer területen hiányzik ezeknek a tehetséges és magasan képzett szakembereknek a kreativitása.

Ez változik most meg. A jövőben a biztonsági szakemberek munkájának része marad ugyan, hogy felülvizsgálják azokat az automatika által gyanúsnak vélt eseményeket, amelyekben az AI-rendszer nem tud dönteni – ám ezek nem rutinfeladatok lesznek, hanem sok esetben szakmai kihívásnak megfelelő, sajátos teszteket igénylő eljárások. Emellett egy másik, izgalmas munkát is biztosít számukra ez a technológia, mivel az AI folyamatos finomhangolását és fejlesztését, új rendszerek betanítását végezhetik. Hiszen az AI-rendszer csak az állandó tanítás által marad hatékony, sőt egyre gyorsabban és nagyobb arányban képes a bejutott kártevőket ártalmatlanítani.¹¹

¹¹ Mirko Zorz – Chris Morales – Mike Banic: Automating the hunt for cyber attackers. A Black Hat konferencia interjúja, 2017.

AI-detektorok

A hagyományos detektorok alapvetően különböztek az AI által működtetett érzékelőktől. Régebben értékhatárok figyelésére programozott szűrőket lehetett használni, vagy karaktérsor-egyezéseket vizsgálni. Az AI viszont képes összefüggéseket felismerni a detektorai által szolgáltatott adatok összességéből.

Nagy szükség is van a fejlődésre, hiszen az AI-alapú támadások fokozódó „kiterő képességei” különösen nagy problémát okoznak. Elsőpró hatásfokkal cselezik ki a hagyományos, szabályalapú biztonsági eszközöket. Az AI megtanulhatja a szabályokat, és kijátszhatja azokat, lehetővé tesz új támadásmódokat, ráadásul óriási sebességgel támad. Példaként hozhatjuk a Deeplockert, ami egy többrétegű, AI-alapú feketedobozba zár mindent, amit egy víruskereső technológia veszélyként érzékelhetne, így kerüli ki a védelmet.¹² A fentebb említett rajvirus pedig a rajintelligenciában rejlő öntökéletesítési módszert kiaknázva képes észrevétlen maradni. Mindkét megoldás az AI-ban lévő rejtőzködési lehetőségek terén újszerű, csak eltérő módon (egy tanulmányban majd részletesen bemutatom ezeket). A hasonló vírusok könnyedén kijátsszák a hagyományos detektálást, az AI-detektorok használata adhat csupán esélyt megtalálásukra.

Ahogy az emberi elmének szüksége van az érzékszerveire a tanuláshoz, így az AI-rendszerek is attól lesznek kellően gyorsak, hogy saját maguk által szerzett információk elemzésével fejlesztik magukat. Ehhez szükségük van az alkalmazások, felhasználói eszközök viselkedésének figyelésére – ezeket elemezve képesek megkülönböztetni a váratlan eseményeket, észlelni a gyanús műveleteket. Ez a felderítési taktika segíthet a jövőben a kibertámadások gyors azonosításában, beleértve az AI-bázisú próbálkozásokat is.

Egyéb ötletek

A fentiek alapján láthatóak a főbb irányok, ám hely hiányában itt csupán megemlíteni tudunk pár további hasznos példát a számos egyéb terület közül, ahol a kibervédelem része lett az AI. Már a 2000-es évek közepétől használják az AI-t a hálózati behatolás-felismerés és megelőzés hatékonyabbá tételére, sőt újabb és újabb módszerek vetődnek fel,¹³ olyan sajátos ötletekkel, mint a biológiai immunitást másoló rendszerelképzelés.¹⁴ Végül megemlíthetjük a Botnetek elleni küzdelmet, ahol az AI jól használható a kiberbiztonsági besorolások automatizálásához, vagy a támadási események előrejelzéséhez.¹⁵

¹² Marc Ph. Stoecklin – Jiyong Jang – Dhilung Kirat: DeepLocker: How AI can power a stealthy new breed of malware. *Security Intelligence*, 2018. augusztus 8.

¹³ Selma Dilek – Hüseyin Çakır – Mustafa Aydın: Applications of Artificial Intelligence techniques to combating cyber crimes. *International Journal of Artificial Intelligence & Applications*, 6. (2015), 1. 21–39.

¹⁴ Nabil Ali Alrajeh – Lloret, J.: Intrusion detection systems based on Artificial Intelligence techniques in wireless sensor networks. *International Journal of Distributed Sensor Networks*, 2013. október 20.

¹⁵ Machine Learning in Cyber Security Domain –2: Rating and Incident Forecasting. NormShield blog, 2015.

A támadáshoz használható módszerek védekezésre

Amikor a biztonsági megoldások készítői a támadások új szintjével néznek szembe, a pajzs megalkotása érdekében alaposan meg kell ismerniük a támadó mechanizmust. Egy orvosi területről kölcsönvett analógiával azt mondhatjuk, hogy meg kell vizsgálnunk a vírust, hogy létrehozzuk a „oltást” – tehát ha tanulmányozzuk az AI-alapú támadásokat, számos új tulajdonságot azonosíthatunk a hagyományos támadásokhoz képest, amelyekkel aztán saját védelmi rendszerünket betaníthatjuk, mint ahogyan a legyengített vírussal teszik ellenállóvá az emberi test immunrendszerét.

Ám ez az ötletlopás fordítva is működik: a védelmi célra létrehozott ötleteket a támadók is használhatják a védelmi mechanizmusok kikerülésére. A mesterséges intelligencia bevetése akár támadó, akár védelmi oldalon a klasszikus módszerekhez képest körülbelül olyan mérvű előnyhöz juttatja használóját, mint az ókori fegyverek ellen alkalmazott mai lőfegyverek. A fegyverpárhuzam annyiban nem jó, hogy teljesen eltérő programokat alkalmaz a támadás és a védelem, viszont bizonyos modulokat és módszereket lehet támadási vagy védekezési célra egyaránt felhasználni, teljesen eltérő módon alkalmazva őket. Nézzünk meg néhány példát arra, hogy a támadási ötletekből hogyan lehet védekezési megoldásokat gyártani!

AI a pszichológiai befolyásolás (social engineering) ellen

Sok ismert social engineering módszer ellen az AI-t hatásosan lehet alkalmazni. Ugyanis egy ilyen rendszer képes megkülönböztetni a visszaélések és a normális tevékenységek közötti árnyalt és kifinomult különbségeket, a rendellenességek elemzése és az adatok változásának kiértékelése alapján. Az AI ezen tulajdonságát éppúgy lehet támadásra, mint védekezésre használni. Más helyen részleteztük, hogy az AI-alapú technológia hogyan képes a szélhámosság új dimenzióit megnyitni – most csak a védekezésről ejtünk néhány szót.

Ez idáig csupán szabályok szigorításával és a folyamatok bonyolításával lehetett a csalások ellen biztosítani a különféle online szolgáltatásokat. Ezzel a platformok távolodtak attól, hogy felhasználóbarátnak lehessen őket nevezni. A felhasználóbarát megjelenés pedig nem csupán üzleti szempontból fontos tényező. Az az igazi jelentősége, hogy csökkenti az emberek (a humán erőforrás) túlterheltségét, így növeli a munkavégzés hatékonyságát és gyorsaságát. Ez az állami vagy katonai szolgáltatásoknál is igen fontos tényező. Ráadásul sokszor épp a túlzott szabályozottság vezet a felhasználó „lázadásához”, vagyis a szabályok megszegéséhez, és biztonsági rések megnyitásához. Amikor a felhasználók nem érzik életszerűnek a szabályokat, akkor valószínűleg lesz, aki nem tartja be. Így épp az erős szabályozottság lesz a támadók segítségére.

Tehát nem szabadna egymással szembeállítani a biztonságos (de nehézkes) kezelést és a felhasználóbarát felületet. Az AI nélkül azonban nem volt más lehetőség, valamelyik kárára kellett döntenünk. Az AI-n alapuló védelemmel azonban bizakodhatunk abban, hogy a kettő összebékíthető, nem bonyolódnak tovább a felhasználókra vonatkozó eljárásrendek

és szabályok, mégis biztonságosabbá válik a személy azonosítása. Ez a tény önmagában is hozzájárul a biztonsági szint növekedéséhez, hiszen nem akarják majd kijátszani a felhasználók. Emellett ugyanennek a védelemnek az elemzőképessége jobban felismeri a rengeteg felhasználói hibát és csalási próbálkozást, „értve” azok mintázatát. Ráadásul folyamatosan megtanítható a csalás vagy hibázás újabb variációinak figyelembevételére is. Ha pedig komolyabb azonosítás az elvárás, akkor használható olyan többlépcsős autentikációval működő modell, ami AI használatával válik igazán hatékonyá.¹⁶

Sérülékenységi adatbázisok

A szoftvergyártók a hibafoltozásokhoz jó ideje használnak különféle sérülékenységi adatbázisokat. Ám ezek rég túlnőtték az ember által átlátható mértéket. Ezért lehetett és kellett az AI-t segítségül hívni a sérülékenységi lehetőségek tesztelésére. A megtalált réseket kétféleképp lehet felhasználni: egy támadó behatol általuk – míg egy védelmi szakember befoltozza azokat.¹⁷

A védelem a termék szoftver kiadása előtt jelentősen csökkentheti a biztonsági rések számát, mivel AI-kereséssel azokat is meg lehet találni és befoltozni, amelyeket a hagyományos tesztelés nem vett észre. Később, a termék kiadása után is jól használható védelmi módszer, hiszen jobban lépést tud tartani a kiberbűnözőkkel, akik folyamatosan elemzik az egyre bővülő sérülékenységi adatbázist. A foltozások egy része tehát automatizálhatóvá vált – bár folyamatosan szükség van a réskereső algoritmusok további tanítására, finomhangolására. Ezenfelül az AI javaslatokat is tehet arra, milyen új információk kerüljenek a sérülékenységi adatbázisba. Annyit meg kell még említenünk, hogy a szakirodalom szerint radikálisan újra kell gondolni eddigi használatukat, csak így kerülhetjük el a támadásra való alkalmazásukat.

Spamek izolációja

A támadók AI nélkül is képesek olyan programot írni, ami sok kis részlet és minta segítségével állítja össze a leveleket, hihetőbbé változtatják az érzékelhető feladót. Aki látott már versíró programot, vagy véletlenszerű borkóstolászöveg-készítőt, nem nehéz elképzelnie, hogy sematikus szövegfordulatok keverésével milyen könnyű eltérő e-maileket generálni. Ha ehhez hozzátesszük az AI szövegalkotó képességeit, akkor sokkal kevésbé sematikus e-maileket kapunk.

Viszont, ha ezt a szövegértő képességet egy AI-alapú spamvizsgálatba integráljuk, az megtaníthatóvá válik sokféle árulkodó jel azonosításra. Így képes egy fejlett automatika által összeállított levelet egy igazi levéltől megkülönböztetni. Emellett az e-mail láthatatlan részében található információk elemzésével megtalálja a finomabb jeleket is,

¹⁶ Phiri Jackson et alii: Using Artificial Intelligence techniques to implement a multifactor authentication system. *International Journal of Computational Intelligence Systems*, 4. (2011), 4. 420–430.

¹⁷ Chartis Research Staff: *The state of AI in risk management* Chartis, 2019. október 22.

amelyek csalo küldőre utalnak, így a levél szövegének elemzése nélkül is pontosabban képes a kéretlen küldőket kizárni.

A mesterségesintelligencia-alapú kibervédelem helyzete

Ami az üzleti szektorban megtérül, tehát elterjed, az nagy valószínűséggel az állami szektorban is követendő technológia. Ezért érdemes ezt a szektort is megvizsgálnunk. (Ideális esetben ez fordítva van, és a katonai vagy egyéb állambiztonsági felhasználás messze az üzleti lehetőségek előtt jár – ám a hazai és európai realitások esetében ez nem mondható el.)

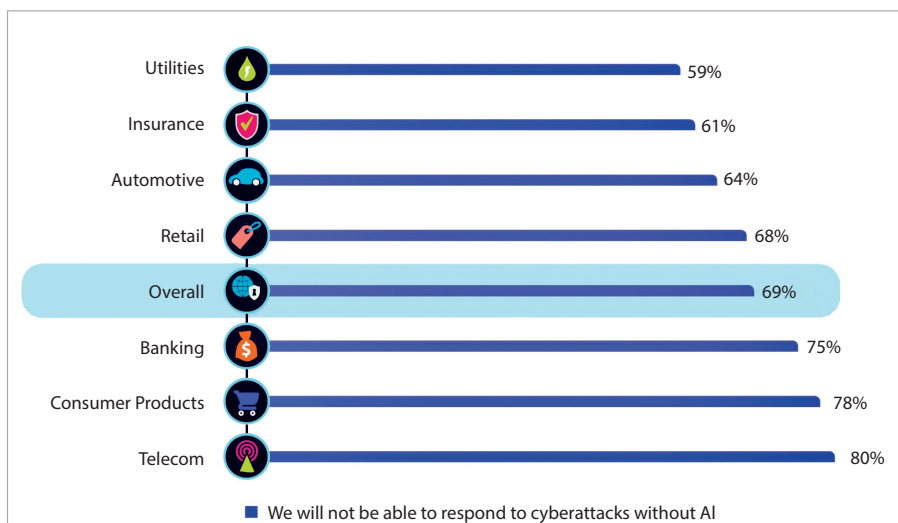
Ha a bevételt nem termelő szervezeteket nézzük, ott is fontos szempontok például a biztonság, a hatékonyság és a törvényességtámogatás. Ezen három szempont együttese arra mutat, hogy „megtérül” az AI használata a nem üzleti szférában is, persze nem nagyobb bevétel formájában. A szabályozás azonban ebben a szektorban nehezebb, hiszen eltérő érdekeltségű politikai erők konszenzusán kell hogy alapuljon, nem pedig a pusztá hatékonyságon, mint egy cég esetében.

Az üzleti világ tapasztalatai

Az AI terjed mindenhol, ahol eddig is jelen volt a klasszikus informatika. Például az üzleti elemzések készítése, a különféle bankkártya-tranzakciók (eddig emberek által végzett) előjegyzése, a vásárlási biztonság növekedése, és sok más egyéb piaciilag hasznos területen meghatározóvá vált ez a technológia. Ebből látható, hogy az AI nem csupán a játékipar, az egészségügy vagy okosdolgok terén hódít, hanem egészen biztosan be fogja hálózni a technikai élet minden területét. Így viszont azok védhetősége sem lesz megoldható a hagyományos módszerekkel. Ezért a védelmi rendszerek fejlesztőinél már régóta megjelentek az egyre finomabban tanuló és tanítható megoldások. Érdemes ezt a fejlesztési irányt mérések alapján is megvizsgálni, így jól kivehető a tendencia, amely az AI irányába mutat.

A Capgemini Kutatóintézet számos kimutatást készített az AI piaci használatáról. Egy alapos tanulmányban kimutatták, hogy ez a technológia számos ágazatban meghatározó kibervédelmi módszer. 850 különböző méretű vállalat kibervédelmi szakembereit kérdezték a fejlett országokban 2019 első felében arról, hogy cégüknél AI alkalmazása hol és merre tart.¹⁸ A felmérés sok nézőpontot megvizsgálva kimutatja, hogy a cégek több mint fele ilyen képességgel felszerelt megoldást vesz igénybe a hálózatbiztonság különféle területein (lásd a mellékelt ábrán), sőt bizonyos aspektusokban a 80%-ot is eléri a kibervédelmi AI használata. A megkérdezett vállalatok jó része komoly összegeket vesztett kibertámadások során, ezért nagy összegeket investál a fejlesztésbe, és 70%-ban az AI-ban látják a megoldás esélyét.

¹⁸ Capgemini Research Institute: *Reinventing cybersecurity with Artificial Intelligence*. Párizs, Capgemini Research Institute, 2019.



1. ábra: Nem tudunk megfelelni a kibertámadások kihívásainak AI nélkül – mondja a válaszadók majd 70%-a!

Forrás: Capgemini-felmérés, 2019

A Capgemini-felmérés egészéből „megjósolható” (valószínűsíthető) volt – bár a kutató-intézet nem tesz ilyen kijelentést –, hogy a védelmi megoldások piacán lassan nem lehet majd olyan terméket kapni (vállalatok számára), amely valamilyen formában ne használná az AI-t. Jelenleg mindenki kínál ilyen megoldást, és a fejlesztések ebbe az irányba mutatnak, a marketing erről szól, tehát várhatóan a klasszikus működésű védelem visszaszorul az olcsó személyieszköz-védelem területére (ezek is fel fogják használni az AI-módszereket, illetve az AI-szerverek eredményeit).

El vagyunk-e késve a szabályzásokkal?

Azért mondható el, hogy sem a technológia bevezetésével, sem a szabályozásával még nem vagyunk elkésve, mivel úgy tűnik – legalábbis a kibertámadások tekintetében –, kissé lelassult a fejlődés. 2017-ben még rengeteg szakportál átvette azt a véleményt, hogy a kibertámadások 62%-a mögött 2018-ban az AI fog állni. Ez azonban csupán BlackHat 2017 konferencián részt vevő szakemberek közötti felmérésen alapult, a portál hivatalos jóslataiban sem szerepelt ilyen adat.¹⁹ Sőt ennek beteljesüléséről még a 2019-es évben sem igen jöttek hírek és mérések.²⁰

¹⁹ A hivatalos USA- és EU-összesítésekben nem szerepel ilyen adat. Lásd: www.blackhat.com/docs/us-17/2017-Black-Hat-Attendee-Survey.pdf, valamint www.blackhat.com/docs/eu-17/Black-Hat-Attendee-Survey.pdf. A 2018-as jelentésben sincs szó az AI-fenyegetésről. Lásd: www.blackhat.com/docs/us-18/black-hat-intel-where-cybersecurity-stands.pdf.

²⁰ 2020-as felmérés a cikk írásakor nem lelhető fel.

A dokumentálatlanság talán nem zárja ki, hogy létezik a veszély. Elképzelhető, hogy a jóslat bevált, csak olyan jól álcázott, hogy senki nem vette észre, hogy évek óta zajlik, vagy aki észrevette, nem akarja kiszivároztatni, hogy így járt.

A mesterségesintelligencia-alapú kibervédelem nemzetközi és hazai szabályozása

Az európai AI-szabályzás egészéről elmondható, hogy az AI 2010-es fellendüléséhez képest kissé későn, a világ többi részétől azonban nem lemaradva kezdődött meg 2018-ban. Azóta gőzerővel folyik a terület szabályozásának kidolgozása, számos megállapodás, szakmai és politikai megnyilvánulás született. A kiadott dokumentumokat egy külön tanulmányban tekintettem át, azt összefoglalva sajnos úgy tűnik, hogy földrészünk lemaradóban van az AI-technológia kibervédelmi aspektusait illetően. A szabályozókat ugyanis jobban átítatják napjaink politikai és elvi divatjai, mint az indokolható lenne, a jogász megközelítés elnyomja a szakmai szempontokat. Például hosszasan taglalják a környezetvédelem, a rasszizmus témáját, vagy hogy kevés nő dolgozik a területen – de az AI és a kibervilág összefüggéseit a legtöbb megnyilatkozásban nem, vagy alig említik.

Az ilyen jövőbe mutató, kiforratlan területeken, mint az AI, a helyzetet nehezíti az igen ambivalens EU–USA viszony, amely közös érdekek és ellenérdekeltségek szövevénye. Bár katonai szövetségesek vagyunk, az USA technológiai fölénye elsöprő, és előnyt feltéve őrzi. A processzorgyártásban, az operációs rendszerekben, de a nagy szoftverek vagy az online szolgáltatások (például felhők, a keresőoldalak, e-mail stb.) terén is egyeduralkodók a tengerentúli cégek, lényegében egymással versenyeznek.²¹ Ez várható az AI vagy a kvantumszámítógépek fejlesztése területén is. Az EU-dokumentumok statisztikai adatai szerint az AI-ra fordított összeg csupán töredéke az USA-ban ráfordított összegnek, sőt még a kínaitól is jelentősen elmarad: „Európa erre összesen 2,4–3,2 milliárd EUR-t fordított 2016-ban, szemben a 6,5–9,7 milliárd EUR-nak megfelelő ázsiai és a 12,1–18,6 milliárd EUR-nak megfelelő észak-amerikai összeggel.”²² Az arányok azóta sem javultak jelentősen.

Ezekből az következik, hogy az Európai Unió optimista hangvételű dokumentumai, amelyek például saját AI- vagy kvantumszámítógép-fejlesztéseket vízionálnak, nem reálisak. A saját európai operációs rendszer, saját processzorgyártás sikertelenségéből le kellene vonni a tanulságokat.

Ezzel szemben a nemrég napvilágot látott *NATO Science and Technology Trends 2020–2040* kiadványt²³ teljesen átszövik a kibertér és a mesterséges intelligencia kapcsolataira történő utalások. Minden témával kapcsolatban legalább említik, sok helyen

²¹ Ezt friss hírek is alátámasztják. Szoftverek terén és processzorok terén lásd például: <https://fxssi.com/most-profitable-it-companies>, valamint www.marketreportsworld.com/TOC/16092890. Szolgáltatások esetében például <https://alertify.eu/top-10-information-technology-it-companies-in-world-2019/>

²² Európai Bizottság COM (2018)795: *Mesterséges Intelligenciáról szóló összehangolt terv*, 11. lábjegyzet.

²³ NATO Science and Technology Organization: *NATO Science and Technology Trends 2020–2040. Exploring the S and T Edge*, Brüsszel, NATO Science and Technology Organization, 2020.

pedig oda vonatkozó technikai megoldásokat is vázolnak. A mesterséges intelligencia NATO-beli terjedéséről több tanulmány is megjelent.²⁴ Ezekben is számos konkrét programról esik szó, amelyekben részben vagy teljesen a mesterséges intelligencia védelme, vagy kibervédelemben való használhatósága áll a középpontban. Ezek közül a legszellemesebb, amikor kétféle AI teszteli – és közben tanítja – egymást.

Az EU- és a NATO-megközelítések összevetéséből levonható az a tanulság, hogy a jelenlegi európai politikai erők még nem látták be a dolog igazi jelentőségét és veszélyét, viszont szerencsére a katonai felső vezetés és a gazdasági élet teljesen tudatában van ennek.

Ebben nyilván közrejátszik az is, hogy a katonai és gazdasági életben világos dominanciája van az USA-nak, aki az AI (egyben az AI-alapú kibertevékenység) zászlóshajója. Ráadásul az európai politikai erők polgárjogi megközelítéséből adódó hátrány az USA előnyére is válik,²⁵ ezért nyilván nem kívánja a katonai téren elért szemléletét rájuk erőltetni.

Tehát csupán remélhetjük, hogy az AI és kiberszakterület a közeljövőben jobban egymásra talál hangsúlyok szempontjából minden európai szakdokumentumban, és a biztonsági szempontokat nem előzik majd meg a jóléti társadalomból fakadó jogászokodó aspektusok.

A magyar szabályozás csak a vázolt nemzetközi erőterben értelmezhető, bővebb elemzésére itt nincs mód, de annyit megállapíthatunk, hogy jó irányban halad. Az új Mesterséges Intelligencia Stratégia²⁶ nagyon sok helyen és összefüggésben tárgyalja a két témát, de még a Nemzeti Biztonsági Stratégia is az EU-s dokumentumok előtt jár, hiszen több helyen egymás mellett említi őket.²⁷

Válaszok és konklúzió

A technikai áttekintésben érzékelhettük, hogy ezek a rendszerek messze nem olyan autonómok, hogy tanulási folyamataik során az ember ellen fordulhassanak. Viszont rajtuk kívül nincs más esélyünk az AI-alapú kibertámadásokkal szemben, amelyek küszöbön vannak. Még nem léptünk be abba a korba, amikor ez elsődleges probléma lenne – részben épp azért, mivel nagy ütemben fejlődik a védelmi oldal is. A biztonságra a fő garancia az az alapelv, hogy ilyen rendszereknél mindig vannak „biztonsági gombok” az emberi kezelőknél. Emellett a technológia fejlődésével egyre jobban megakadályozható, hogy téves információkat tanuljon meg a rendszer. A fontos döntéseket pedig mindig csak emberi jóváhagyással hozhatja meg egy mesterséges döntéshozó.

A nagyobb probléma nem azzal van, ellenünk fordul-e a védelmünk. A gond, hogy a különféle AI-megfigyelések (illetve már a betanítás is) mindig nagy mennyiségű adattal

²⁴ Négyesi Imre: A mesterséges intelligencia és a hadseregek. *Hadtudomány*, 29. (2019), 3. 71–79.

²⁵ Tegyük hozzá, hogy sajnos a paletta többi komoly résztvevője (Kína, Oroszország) előnyére is válik.

²⁶ Mesterséges Intelligencia Koalíció: Magyarország Mesterséges Intelligencia Stratégiája 2020–2030. 2020.

²⁷ Mesterséges Intelligencia Stratégia (2020) i. m. 106, 160–162. pont

dolgoznak. Ezen adatok tárolásbiztonsági problémái mellett jogi (például a magánszférát érintő) kétségek is felmerülnek. Így érkezhetünk el oda, hogy a védelemből személyiségi-jogi veszélyforrás lesz. Ennek feloldására egyelőre nem alkottak sehol elégséges szabályozást vagy megoldást. De az biztos, hogy új, nemzetközileg elfogadott alapelvekre, majd szabályozókra, valamint azok betartatására lesz hozzá szükség. Ezek alapja például a térfigyelő rendszerek és hatósági megfigyelések szabályozása lehet.

Miért éri meg a drágább AI-alapú kibervédelem?

Az üzleti oldalt vizsgálva megkaptuk az igenlő válasz indoklását: a mesterségesintelligencia-alapú kibervédelem nem drága abban az értelemben, hogy nagyobb biztonság mellett nagyobb bevételt tesz lehetővé már közepes cégek esetén is. Az ilyen cégeknél sok esetben összegszerűen is kimutathatók egyes veszteségek, amelyeket a nem megfelelő kibervédelem okozott. De a cég presztízsének vesztesége, vagy az alkalmazottak stresszességének jelentős növekedése még itt sem számszerűsíthető. A fentebb említett fontos szempontok (biztonság, hatékonyság, törvényesség támogatás) mellett ilyen tényezőket is figyelembe véve, a nemüzleti szférában is egyértelműen indokoltak az AI használatából adódó többletkiadások,

Gyakorlati megfontolások

A fentebb részletezett technikai indokok alapján, az EU, az USA és a NATO hozzáállását figyelembe véve jelen pillanatban hazánk – sőt igazából minden európai állam számára – biztonsági szempontból optimális megoldás olyan kész, amerikai kibervédelmi rendszerek beszerzése, amelyek fejletten használják az AI-t. Akár az állami gerinchálózatok, akár az állami fenntartású intézmények védelmében ilyen megoldások különböző erejű verziói fogadhatók el. A minősített adatok védelménél nyilván nem elegendő kizárólag a piacon kapható termékekre támaszkodni. A szabályozókban hangoztatott együttműködés keretében lesz szükséges egyik piaci rendszer magyar módosítása. Ilyen saját fejlesztéssel kiegészítve az erős piaci termék egy ideig megfelelő védelmet jelenthet.

Felhasznált irodalom

- 1163/2020. (IV. 21.) kormányhatározat a Magyarország Nemzeti Biztonsági Stratégiájáról 1. sz. melléklet.
 Allen, Gregory – Taniel Chan: *Artificial Intelligence and national security*. Cambridge, Belfer Center Study, 2017.

- Allen, Gregory C.: Understanding China's AI Strategy: clues to Chinese strategic thinking on Artificial Intelligence and national security. Washington, Center for New American Security, 2019.
- Alrajeh, Nabil Ali – Lloret, J.: Intrusion detection systems based on Artificial Intelligence techniques in wireless sensor networks. *International Journal of Distributed Sensor Networks*, 2013. október 20. Online: <https://journals.sagepub.com/doi/10.1155/2013/351047>, 2013.
- Capgemini Reaserch Institute: *Reinventing cybersecurity with Artificial Intelligence*. Párzis, Capgemini Research Institute, 2019. Online: www.capgemini.com/wp-content/uploads/2019/07/AI-in-Cybersecurity_Report_20190711_V06.pdf
- Chartis Research Staff: *The state of AI in risk management* Chartis, 2019. október 22. Online: www.chartis-research.com/technology/artificial-intelligence-ai/state-ai-risk-management-10976
- Dilek, Selma – Hüseyin Çakır – Mustafa Aydın: Applications of Artificial Intelligence techniques to combating cyber crimes. *International Journal of Artificial Intelligence & Applications*, 6. (2015), 1. 21–39. Online: <https://doi.org/10.5121/ijaia.2015.6102>
- Európai Bizottság COM(2018)795, *Mesterséges Intelligenciáról szóló összehangolt terv*, 11. látjegyzet. Online: <https://eur-lex.europa.eu/legal-content/HU/TXT/HTML/?uri=CELEX:52018DC0795&from=ES>
- Heinemeyer, Max: *The next paradigm shift – AI-driven cyber-attacks*. DarkTrace, Research White Paper, 2. 2018. Online: www.oixio.ee/sites/default/files/the_next_paradigm_shift_-_ai_driven_cyber_attacks.pdf
- Hoadley, Daniel – Nathan Lucas: *Artificial Intelligence and national security*. Washington, Center for New American Security, 2018.
- Koós Gábor – Szternák György: A mesterséges intelligencia katonai felhasználásának lehetőségei. *Felderítő Szemle*, 18. (2019), 4. 117–135. Online: www.knbsz.gov.hu/hu/letoltes/fsz/2019-4.pdf
- Mesterséges Intelligencia Koalíció: *Magyarország Mesterséges Ingelligencia Stratégiája 2020–2030*. 2020. Online: <https://digitalisjoletprogram.hu/files/6f/3b/6f3b96c7604fd36e436a-96a3a01e0b05.pdf>,
- NATO Science and Technology Organization: *NATO Science and Technology Trends 2020–2040. Exploring the S&T Edge*. Brüsszel, NATO Science and Technology Organization, 2020.
- Négyesi Imre: A mesterséges intelligencia és a hadseregek. *Hadtudomány*, 29. (2019), 3. 71–79. Online: www.mhtt.eu/hadtudomany/2019/2019_3/2019eA%20mesters%c3%a-9ges%20intelligencia%20%c3%a9s%20a%20hadseregek_N%c3%a9gyesi%20Imre.pdf
- Machine Learning in Cyber Security Domain –2: Rating and Incident Forecasting*. NormShield blog, 2015. Online: www.normshield.com/machine-learning-in-cyber-security-domain-2-cyber-security-rating-and-incident-forecasting/
- Jackson, Phiri – Tie-Jun Zhao – Cong Hui Zhu – Jameson Mbale: Using Artificial Intelligence techniques to implement a multifactor authentication system. *International Journal of Computational Intelligence Systems*, 4. (2011), 4. 420–430. Online: <https://doi.org/10.1080/18756891.2011.9727801>
- Porkoláb Imre – Négyesi Imre: A mesterséges intelligencia alkalmazási lehetőségeinek kutatása a haderőben. *Honvédségi Szemle*, 147. (2019), 5. 3–19.
- Marc Ph. Stoecklin – Jiyong Jang – Dhilung Kirat: DeepLocker: How AI can power a stealthy new breed of malware. *Security Intelligence*, 2018. augusztus 8. Online: <https://securityintelligence.com/deeplocker-how-ai-can-power-a-stealthy-new-breed-of-malware/>

- Zelinka, Ivan – Roman Senkerik – Lubomir Sikora – Swagatam Das: Swarm virus – Next-generation virus and antivirus para-digm? *Swarm and Evolutionary Computation*, 43. (2018). 207–224. Online: <https://doi.org/10.1016/j.swevo.2018.05.003>
- Zorz, Mirko – Chris Morales – Mike Banic: *Automating the hunt for cyber attackers*. A Black Hat konferenciainterjúja, 2017. Online: www.helpnetsecurity.com/2017/08/08/automating-hunt-cyber-attackers/; magyarul: <https://blackcell.hu/automatizalt-mesterseges-intelligencia-kibertamadasok-ellen-interju/>