

Protect Your MI-Sensitive

– Miért kell MI-szenzitivitást fejlesztenünk?



„Az AI-hoz kapcsolódó kockázatokra még nincs kiépült érzékenységünk” – hívja fel a figyelmet Keleti Arthur, ITBN alapító, jövőkutató. Amikor reggel rápillantunk a hőmérsékletre, a forgalmi helyzetre vagy éppen az árfolyamokra, valójában kockázatokat mérlegelünk. Mennyire van hideg? Milyen gyorsan jutunk el a munkába? Mekkora a veszteség, ha most váltunk pénzt? Tudatosan vagy tudat alatt, de folyamatosan reagálunk a környezetünk változásaira, és érzékenységet alakítunk ki ezekkel a tényezőkkel szemben. Van azonban egy új faktor, amelyre sokunknak még nincs megfelelő „érzékenysége”: a mesterséges intelligencia (AI).

Az inflációhoz hozzászoktunk, a devizapiac ingadozásai sem érnek teljesen váratlanul és a kiberbűnözők trükkjei sem ismeretlenek számunkra. Az AI viszont más. Nem pusztán adat vagy eszköz, de nem is ember. Gyors, okos,

és egyre több területen helyettesít minket. Ez az újdonság pedig egyszerre hoz lehetőségeket és fenyegetéseket – és közben felveti a kérdést: hogyan tudjuk megvédeni mindazt, ami számít számunkra, legyen szó munkánkról, tanu-

lásunkról, egészségünkéről vagy jövőnkről?

Mit jelent az AI-szenzitivitás?

Az AI-szenzitivitás az a képesség, hogy felismerjük, értelmezzük és megfelelően reagáljunk azok-

ra a kockázatokra és lehetőségekre, amelyeket a mesterséges intelligencia teremt.

Ellentétben más veszélyforrásokkal, itt nincsenek egzakt mérőszámok. Az UV-sugárzásnak van indexe, a szén-monoxidnak van érzékelője – az AI esetében azonban még nincs hasonlóan egyértelmű mutató, amihez igazíthatnánk félelmeinket vagy reményeinket. Ezért az AI-szenzitivitás fejlesztése nem luxus, hanem kényszerűség.

Az AI ráadásul más technológia, mint amivel eddig találkoztunk. Hasonlít ugyan az emberi intelligenciára, de számos ponton el is tér tőle. Éppen ezért fontos, hogy nyitottan és új megközelítéssel álljunk hozzá. Sokat kell tesztelnünk, próbálnunk, használnunk, hogy valóban megértsük a működését – és nem zárhatjuk ki, hogy időnként váratlanul, de értelmesen, összetetten válaszol, dönt vagy cselekszik. Ez a kiszámíthatatlanság és komplexitás szintén része kell legyen az érzékenységünk kialakításának, valamint a védelmi stratégiánknak.

Ez az érzékenység nem elvont fogalom: mindannyiunk személyes tapasztalatain, családjaink, gyerekeink, kollégáink biztonságán, valamint az általunk felügyelt rendszerek megbízhatóságán múlik. Minél jobban értjük, hogyan hat az AI a mindennapjainkra, annál felkészültebben reagálhatunk a fenyegetésekre.

Miért van rá szükség?

Az AI-alapú technológiák fejlődése nemcsak a hatékonyságot növeli, hanem a támadások számát és minőségét is megváltoztatja. Az Amazon biztonsági vezetője például arról számolt be, hogy a vállalatot érő kibertámadások száma az AI-automatizáció miatt hétszeresére nőtt.

Ez a világ, ahol mesterséges intelligencia segítségével megtévesztésig valóság-hű videókat, hangokat, képeket és e-maileket lehet előállítani, teljesen új szintre emeli a biztonsági kockázatokat. A szabályozás, az etikai irányelvek vagy a filozófiai viták önmagukban már nem elegendők – gyakorlati, közösségi és technológiai válaszokra van szükség.

Hogyan védekezhetünk?

- 1. Tudatosság fejlesztése** – az első lépés, hogy felismerjük: az AI által generált tartalom gyakran tökéletesen emberinek tűnik. Ennek tudatosítása nélkül könnyen áldozattá válhatunk.
- 2. Technológiai eszközök használata** – egyre több megoldás érhető el az AI által létrehozott tartalmak felismerésére, de ezek alkalmazása folyamatos figyelmet igényel.
- 3. Közösségi védelem** – A kiberbiztonságban az együttműködés mindig is kulcsfontosságú volt.



Az AI technológia hasonlít az emberi intelligenciára, de számos ponton el is tér tőle, éppen ezért fontos, hogy nyitottan és új megközelítéssel álljunk hozzá.

Az AI korszakában még inkább szükség van közös tapasztalatokra, esettanulmányokra és tudásmegosztásra.

- 4. Ellenálló képesség építése** – a cél nem az, hogy teljesen elzárkózzunk az AI-től, hanem hogy olyan stratégiákat alakítsunk ki, amelyekkel mérsékelhetjük a kitétettséget és növelhetjük a védekezés hatékonyságát.

ITBN – a közös tudás tere

Az AI-szenzitivitás védelme a kiberbiztonság egyik legösszetettebb feladata. Bár a szakértők már évtizedek óta szembenéznek új és ismeretlen kihívásokkal, az AI által jelentett veszélyek komplexitása páratlan.

Az ITBN (<https://itbn.hu/itbn-2025>) több mint húsz éve biztosít platformot a szakemberek és érdeklődők számára, hogy együtt dolgozzanak a legaktuálisabb kiberbiztonsági kérdések megválaszolásán. Az idei rendezvény mottója – Protect Your AI-Sensitive – pontosan azt fejezi ki, amire most a legnagyobb szükségünk van: közös megértésre, közös gondolkodásra és közös fellépésre.

Az ITBN nemcsak technológiai megoldásokat mutat be, hanem lehetőséget ad arra is, hogy közösen alakítsuk ki az AI-hoz fűződő érzékenységünket. Mindezt olyan szakmai közegben, ahol a világ vezető vállalatainak és hazai szakértőinek tapasztalatai egy helyen találkoznak.

A tavalyi hatalmas érdeklődés bizonyította, hogy ez a téma nemcsak a szakmát, hanem a szélesebb társadalmat is foglalkoztatja. Az idén az ITBN még szélesebbre tárja kapuit, és vár mindenkit, akinek fontos a biztonság és a jövő védelme – zárta gondolatait Keleti Arthur.