



## AKTUÁLIS TARTALMAK



HÍREK



STATISZTIKAI ADATOK



RIASZTÁSOK

TÁJÉKOZTATÁSOK



CTI JELENTÉS, RIPORT



ESEMÉNY



BBA+ BESZÁMOLÓK

# HÍRLEVÉL

Nemzetközi  
IT biztonsági sajtószemle  
2026. 3. hét

## KONTAKT



edt@nki.gov.hu



FBC3 88A2 BF51 AD58  
A2D0 E9DD E078 ABD3  
E75D



nki.gov.hu



# HÍREK

## „Üriemberek” áldozatává vált a román erőmű – Gentlemen ransomware

(bleepingcomputer.com)

Zsarolóvírus-támadás érte a romániai Oltenia Energy Complexet (Complexul Energetic Oltenia), Románia legnagyobb szénelapú hőerőműjét karácsony második napján, amelynek következtében a vállalat IT-infrastruktúrája működésképtelenné vált. **Bővebben...**

## Egy kattintás és a rejtett Telegram proxy linkek felfedik az IP-címedet

(bleepingcomputer.com)

Gondolkodj mielőtt kattintasz! Egy látszólag ártalmatlan link is elindíthat a háttérben hálózati folyamatokat. Bizonyos esetekben ezek a folyamatok felfedhetik a felhasználó valódi IP-címét. A Telegramnál ez a jelenség a proxy linkek automatikus kezeléséhez kapcsolódik, és akár egyetlen kattintással is kihasználható. **Bővebben...**

## Kibertámadás bénította meg a belga AZ Monica kórház informatikai rendszerét

(bleepingcomputer.com)

2026. január 13-án reggel súlyos kibertámadás érte a belga AZ Monica kórház számítógépes rendszereit, ami miatt az intézmény mindkét telephelyén az összes szervert azonnal le kellett állítani, és az egészségügyi ellátás egyes részei részlegesen vagy teljesen leálltak. **Bővebben...**

## Hamis „Ingyenes Streaming Stick” ajánlatok

(idtheftcenter.org)

Az interneten és televíziós hirdetésekben egyre gyakrabban terjednek a hamis, „ingyenes streaming” szolgáltatást kínáló ajánlatok. Ezek az eszközök gyakran a megfélemlítésig hasonlítanak a valódi Amazon Fire TV Stickekhöz. Előfordul az is, hogy a csalók hitelesnek tűnő márkanevekkel próbálják meg legitimnek feltüntetni a termékeiket **Bővebben...**

## Új trükköt alkalmaznak a Facebook-csalók (bleepingcomputer.com)

A hackerek egyre inkább a browser-in-the-browser (BitB) módszerre támaszkodnak, hogy rávegyék a felhasználókat a Facebook-fiókjuk hitelesítő adatainak megadására. **Bővebben...**

# STATISZTIKAI ADATOK



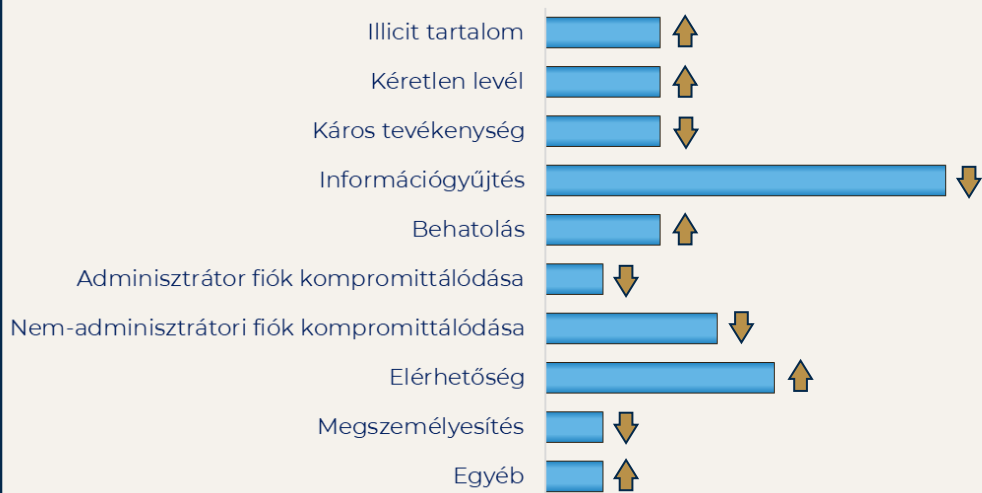
2026. 01. 09. — 2026. 01. 15.

Fenyegetettségi szint:

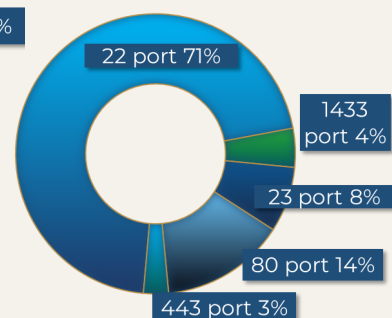
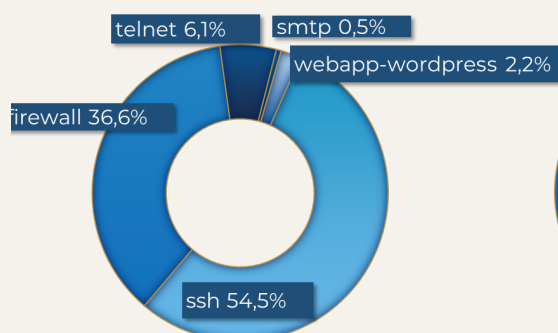


## Az NBSZ NKI által kezelt incidensekre vonatkozó statisztikai adatok

Az adatsorok melletti nyilak az előző héthez viszonyított változásokat mutatják.



## Az elosztott kormányzati IT biztonsági csapdarendszerből (GovProbe1) származó adatok



# RIASZTÁSOK TÁJÉKOZTATÁSOK



## Riasztás

### Microsoft termékeket érintő sérülékenységekről

A Nemzetbiztonsági Szakszolgálat Nemzeti Kiberbiztonsági Intézet (NBSZ NKI) **riasztást** ad ki a **Microsoft** szoftvereket érintő **kritikus kockázati besorolású** sérülékenységek kapcsán, azok súlyossága, kihasználhatósága és a szoftverek széleskörű elterjedtsége miatt.

A Microsoft tárgyhavi biztonsági csomagjában összesen **114** különböző **biztonsági hibát javított**, köztük **3 db nulladik napi (zero-day)** sebezhetőséget is amelyet a Microsoft tájékoztatása szerint támadók kihasználhattak, mivel már a javítás előtt publikálásra került.

[Elolvasom](#)

## Tájékoztatás

### Adobe szoftverek sérülékenységeiről

A Nemzetbiztonsági Szakszolgálat Nemzeti Kiberbiztonsági Intézet (NBSZ NKI) **tájékoztatót ad ki** az **Adobe** szoftverfejlesztő cég **termékeit érintő sérülékenységekkel** kapcsolatban, azok súlyossága, valamint az egyes biztonsági hibákat érintő aktív kihasználások miatt.

[Elolvasom](#)

# CTI JELENTÉS



## Az internet sötét oldala

A dokumentum **átfogó képet ad** a dark web működéséről, technológiai alapjairól és a hozzá kapcsolódó kiberbűnözési ökoszisztémáról. **Bemutatásra kerülnek** az anonimitást biztosító infrastruktúrák, a dark weben működő támadási piacok és szolgáltatások, valamint a legjellemzőbb kiberfenyegetési trendek és célzott támadási módszerek, különös tekintettel a dark web monitoring szerepére mint proaktív kiberhírszerzési eszközre. A cél az, hogy **a szervezetek és döntéshozók számára gyakorlati iránymutatást** nyújtson a tudatos, proaktív és hosszú távon fenntartható védelmi megközelítések kialakításához.

**Elovasom**

**Érdekesnek találta  
elemzésünket?**

**Szívesen olvasna  
hasonló témakörben?**

**Figyelmébe ajánljuk  
„A 2024-es CrowdStrike  
incidens elemzése” című  
jelentésünket!**



# ESEMÉNY



## EIVOK-72. Adatközpontok „kiskérdései” Szakmai Fórum

**2026. január 22. 16:30**  
**OTP Bank Nyrt.,**  
**1131 Budapest, Babér utca 9. és online (hibrid)**

A **HTE EIVOK 72.** információbiztonsági szakmai fórumán az **OTP Bank Nyrt.** és az **Omikron Kft.** kiemelkedő előadói tartanak előadást. **A fórum témája** rendkívül aktuális és érdekfeszítő, hiszen az **adatközpontok "kiskérdései"** lesznek fókuszban.

A rendezvényen való részvétellel **2 CPE pont** szerezhető.

**További érdekességekért  
és IT biztonsággal  
kapcsolatos tartalmakért  
látogasson el LinkedIn  
oldalunkra!**



**További információ**

**Regisztráció**

# HAVI CTI RIPORT



## 2025. december havi CTI riport

A Nemzetbiztonsági Szakszolgálat Nemzeti Kiberbiztonsági Intézet **havi rendszerességgel** ad ki **fenyegetéselemzést**, mely **összefoglalja a kibertér globális**, valamint **magyarországi helyzetét**. A riport megismerése megfelelő támpontot adhat az olvasó számára, hogy szervezete milyen IT biztonsági kihívásokkal nézhet szembe a közeli jövőben.

**Elo olvasom**

**Érdekli, hogyan formálhatják a jövőt a különböző kiberbiztonsági kihívások?**

**Fedezze fel velünk a legizgalmasabb témákat, a szakértői tippektől egészen a legújabb trendekig!**

**Kövesse podcastünket a legnépszerűbb felületeken!**



# BBA+ BESZÁMOLÓ

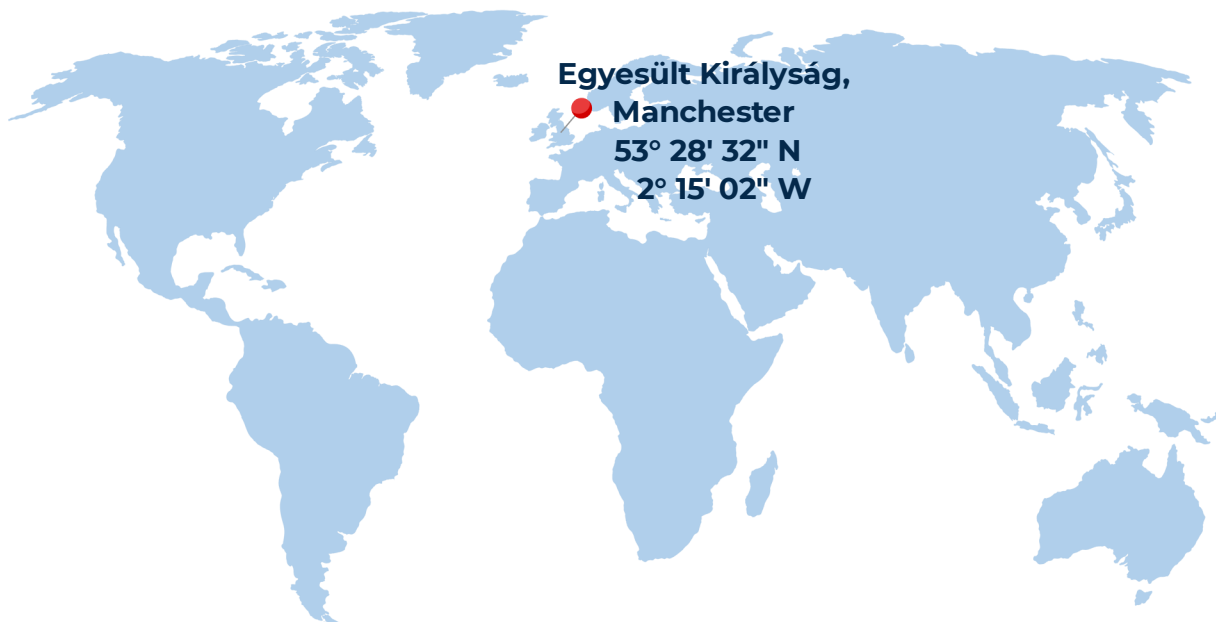
## NDC Manchester AI & Security



2025. december 3 - 4.

A **Nemzetbiztonsági Szakszolgálat Nemzeti Kiberbiztonsági Intézet** munkatársai **Széchenyi Terv Plusz pályázat** részeként szakmai ismeretbővítésen vesznek részt a súlyos és szervezett, határon átnyúló bűncselekmények elleni küzdelem, illetve ilyen jellegű bűncselekmények megelőzésének fejlesztése céljából.

**A projekt célja** a kiberfenyegetések elleni fellépéshez szükséges friss ismeretek gyűjtése és megosztása a hazai kiberbiztonsági szakemberekkel.



A Nemzeti Kiberbiztonsági Intézet munkatársai **2025. december 3-4.** között részt vettek a **Manchesterben** megrendezésre kerülő **NDC Manchester AI & Security** konferencián.

Az **NDC** (Now Developer Conference) egy **nemzetközi szoftverfejlesztői konferenciasorozat**, amely 2008 óta ad teret a .NET, a mesterséges intelligencia, a kiberbiztonság és a kapcsolódó technológiák szakértőinek.

#### **Kiknek ajánlott a beszámoló megismerése?**

-  Vállalati IT-biztonsági csapatoknak
-  Kiberfenyegetés-kutatóknak és elemzőknek
-  Biztonsági mérnököknek és fejlesztőknek
-  Kiberbiztonsági szakemberek számára

**A konferencia legfontosabb előadási az alábbiakban foglalhatók össze.**

# BBA+ BESZÁMOLÓ

Georges Bolssens

## The question is not when to start threat modeling. It's when to stop

Bolssens előadása a **threat modeling** „végtelenségének paradoxonát” vizsgálja: **a fenyegetések folyamatosan változnak**, így elméletben soha nem létezhet teljes, végleges fenyegetésmodell. A gyakorlati biztonsági munka azonban mégis megköveteli, hogy időről időre kijelentsük: „elegendő” állapotot értünk el ahhoz, hogy továbblépjünk a mitigációk implementálására. A prezentáció áttekintette a **klasszikus threat modeling lépéseket** (scope meghatározás, asset-azonosítás, útvonal- és trust boundary-analízis, STRIDE alapok), majd az előadó bevezetett olyan döntési kritériumokat, amelyek segítenek felismerni a **diminishing returns** pontját. Ide tartoznak a rendelkezésre álló erőforrások (idő, személyzet), az üzleti prioritások, a kockázati szint, és az, hogy a további modellezés hoz-e mérhető, kockázatcsökkentő értéket. Konkrét **esettanulmányok** mutatták be, hogyan lehet felismerni azt, amikor a további mélyítés már nem javítja érdemben a biztonsági helyzetet: például egy komplex microservice-architektúrában bizonyos komponensek fenyegetésmodellje annyira kevés új információt ad, hogy célszerűbb az implementációs oldalon haladni tovább. Bolssens kritikusan megjegyezte, hogy a csapatok gyakran „tökéletes” modellekre törekednek, miközben **a valós vállalati környezetben gyors iterációra és a legkritikusabb fenyegetések mihamarabbi mitigációjára van szükség**. A „good enough” fogalmát így nem minőségbeli kompromisszumként, hanem prioritás-vezérelt biztonsági döntésként kell értelmezni.



## Arohi Naik

### Beyond the Pink: Data Risks in Femtech You Can't Ignore

Arohi Naik előadása a **femtech alkalmazások adatbiztonsági és privacy-kockázataira** fókuszált, rávilágítva arra a kritikus problémára, hogy ezek a **termékek gyakran egészségügyi adatokat kezelnek**, mégsem tartoznak a szigorú egészségügyi adatvédelmi szabályozások alá. Részletesen bemutatta, hogy a femtech ökoszisztéma – menstruációs naptáraktól az ovulációs vagy hormonkövető rendszerekig – rendkívül érzékeny, reprodukív egészséghez kapcsolódó információkat gyűjt, mégis sokszor lifestyle-termékként pozicionálja magát. Naik elemezte a **femtech piac sajátosságait**: a gyors termékfejlesztést, a szabályozási kiskapukat, a minimális adatkezelési átláthatóságot és a felhasználói megtévesztést. Technikai oldalról rávilágított olyan **sebezhetőségekre**, mint a felhőben védtelenül tárolt logok, third-party adatszolgáltatások kontrollálatlan használata, gyenge API-védelem, hiányos titkosítás és pontatlan engedélykérések. Az előadás központi üzenete a **sociotechnical threat model megközelítés** fontossága volt: a femtech kockázatait nem lehet pusztán technikai szemszögből értékelni. A reprodukív egészségügyi adatok elvesztése vagy kiszivárgása számos országban komoly társadalmi, jogi vagy emberi jogi következményekkel járhat. A talk kifejtette, hogy a femtech szolgáltatók sok esetben agresszíven monetizálják a gyűjtött adatokat – például hirdetőhálózatok, profilozók és adatbrókerek felé –, miközben a felhasználók nincsenek tudatában adatuk értékének és kockázatának. Arohi Naik végül a design-szintű privacy-hez (Privacy-by-Design), a felhasználói kontroll visszaadásához és a szigorúbb iparági önszabályozáshoz szólított fel. A **femtech területének fejlődése** csak akkor lehet biztonságos, ha a termékfejlesztők és szabályozók felismerik a reprodukív egészségadatok különösen magas kockázati szintjét.

# BBA+ BESZÁMOLÓ

Aleksander Stensby

## 10 tips to level up your AI-assisted coding

Stensby a fejlesztői munka jövőjének egyik legdinamikusabban fejlődő területét, az **AI-asszisztált kódírást** vizsgálta. Előadásában **gyakorlati tanácsokkal** mutatta be, hogyan lehet az olyan eszközökből, mint a **Cursor**, **Claude Code** vagy a **GitHub Copilot**, nem csupán kódgenerátort, hanem valódi fejlesztőtársat faragni. A beszélgetés központi eleme a **prompting és context engineering** volt: a jó eredmények kulcsa a megfelelő kontextusbevitel, a célok pontos megfogalmazása és a feladat strukturált lebontása. Az előadó kiemelte a hosszú kontextusablakok jelentőségét, amely lehetővé teszi nagyméretű projektek koherens elemzését, teljes könyvtárstruktúrák vagy tesztkészletek értelmezését. Rámutatott, hogy a **modern AI-asszisztensek** a hibakeresést, kódelemzést és refaktorálást is képesek támogatni, de csak akkor, ha a fejlesztők **következetesen és átgondoltan** adják át a **szükséges információkat**. Stensby részletesen bemutatta a **Model Context Protocol** (MCP) jelentőségét is, amely lehetővé teszi, hogy az **AI közvetlenül kommunikáljon natív fejlesztői eszközökkel** – például GitHubbal, Slackkel, Figma-val vagy tesztkeretrendszerrel. Ez új workflow-kat hoz létre: az AI nem csak kódot generál, hanem build-folyamatot futtat, dokumentációt szinkronizál, vagy regressziós teszteredményeket elemez. Az előadás technikai fókusza mellett jelentős hangsúlyt kapott a **biztonság**: a nem megfelelően kezelt kontextus szivárgást okozhat, különösen ha privát repositoryk, API-kulcsok vagy tokenek is bekerülnek a promptokba. Stensby javasolta a minimális szükséges kontextus elvét, a sandbox környezetek használatát és a folyamatos log-ellenőrzést.



Jeff Watkins

## STOIC Security: Shielding Your Generative AI App from the Five Deadly Risks

Ágot Jeff Watkins előadása a **generatív AI-rendszerek biztonsági fenyegetéseinek** esszenciáját foglalta össze a STOIC-keretrendszerbe rendezve. A modell az **öt legkomolyabb kockázati kategóriát** emeli ki: **Stolen** (adatlopás, model leakage), **Tricked** (prompt injection), **Obstructed** (availability-elleni támadások), **Infected** (training vagy finomhangolás során bejutó káros tartalom) és **Compromised** (model corruption, supply chain manipuláció). Watkins célja az volt, hogy a fejlesztők, az архитеktek és a vezetők számára egy egyszerű, mégis robusztus gondolati keretet adjon a generatív AI-kockázatok kommunikálásához és kezeléséhez. Részletesen foglalkozott azzal, hogy a **hagyományos AppSec eszköztár** miért elégtelen a modern LLM-rendszerekhez. Az input nem determinisztikus, a modell állapota változó, a kontextus-kezelés pedig könnyen manipulálható — ezek miatt a **klasszikus sérülékenységi kategóriák** (pl. XSS, CSRF) analógiai szinten léteznek ugyan, de **teljesen új támadási felületek** nyílnak. Watkins hangsúlyozta, hogy a **generatív modellek** nem csak kódot futtatnak, hanem tartalmat „értelmeznek”, így támadási felületet jelentenek a természetes nyelvű utasítások és a rejtett parancsok is. **A STOIC keretrendszer minden kategóriájához konkrét védekezési stratégiákat is bemutatott.** A „**Tricked**” esetén például a rétegzett input-szűrés, a prompt-határok explicit definiálása, a kontextus szeparálása és az output sanitization jelentik az első védelmi vonalat. A „**Compromised**” esetében a modell supply chain integritása volt a fókusz: aláírt modellek, hash-ellenőrzés, kontrollált finomhangolási pipeline és hozzáférés-vezérlés. Watkins részletes példát mutatott arra, hogyan fertőződhet meg egy finomhangolt modell rosszindulatú vagy manipulatív tréningadatokkal, és hogyan válhat ez új, akár automatizált támadási eszköz alapjává. Kiemelte, hogy a generatív AI-t fejlesztő vállalatok gyakran csapnak át két végletbe: túlzott bizalom vagy irracionális félelem. A **STOIC célja** egy egyensúlyi modell:

# BBA+ BESZÁMOLÓ

**hatékony innováció** úgy, hogy közben strukturáltan mérhető és kezelhető marad a kockázat. Watkins szerint a következő évek biztonsági feladata az lesz, hogy a szervezetek AI-governance mechanizmusokat építsenek ki, ahol a modellfrissítések auditáltak, a prompt-környezetek kontrolláltak és a kimenetek ellenőrzöttek maradnak nem csak fejlesztés közben, hanem éles működés közben is.

**Jiachen Jiang**

## **Building Secure Infrastructure for Productive AI Agents**

Jiachen Jiang előadása arra fókuszált, hogyan lehet a **produktív AI-agenteket** biztonságos, mégis hatékony módon bevezetni vállalati környezetbe. Az AI-ügynökök egyre komolyabb feladatokba kapcsolódnak be — dokumentációírás, fejlesztés, tesztelés, pipeline-kezelés, rendszerkonfiguráció — de jogosultságaik gyors növekedése aránytalan kockázatot hoz létre. Jiang szerint a legnagyobb **hiba** az, hogy **sok szervezet „megbízható fejlesztőként” kezeli az AI-agentet**, miközben az valójában egy kontrollálatlan, statisztikai rendszer, amely ugyan kompetensnek tűnik, de nem érti a következményeket. Az előadás kiindulópontja egy egyszerű hasonlat volt: **az AI-agentet úgy kell kezelni, mint egy új junior fejlesztőt**, aki semmit nem tud a rendszer belső működéséről, és akit szigorú keretek közé kell terelni, mielőtt komoly feladatot kapna. Jiang bemutatta az úgynevezett **ephemeral sandbox environment** koncepcióját — ezek rövid életű, izolált környezetek, ahol az AI-agent kísérletezhet, kódot írhat, konfigurációt módosíthat, de a rendszer többi része tiltott számára. Az agent outputjai ezután manuális vagy automatikus review-n esnek át, mielőtt bármilyen éles művelet történne. Az előadás kitért a **jogosultságkezelésre** is. Jiang hangsúlyozta, hogy az AI-agentek esetében a **„least privilege”** még kritikusabb, mert a modellek nem értik a kockázatot, így nehezen különböztetik meg a biztonságos és veszélyes műveleteket. Példaként említette, amikor egy AI-agent hibás build-scriptet írt, amely törölte a teljes deploy-könyvtárat, vagy amikor a konfigurációs

állományok frissítése során véletlenül API-kulcsokat szivárogtatott naplófájlokba. Ezek nem rosszindulatból történnek — az **AI egyszerűen nem képes valós következmény-modellezésre**. Jiang részletesen bemutatta a **safe toolchain integration** architektúráját is. Ennek része a strict API-gateway, a titkosított credential store, az auditálható agent-tevékenység és a rollback-elhető munkamenet. A Coder/Anthropic laboratóriumi példák demonstrálták, hogy a kontrollált agent-környezetek 80–90%-kal csökkentik a hibás műveletek arányát, miközben az agent így is jelentős produktivitásnövekedést biztosít.

### **Pedram Hayati**

## **Skill Degradation: An Empirical Analysis of 400+ AI-Generated Security Fixes**

Pedram Hayati előadása egy nagyszabású, több mint **400 AI által generált biztonsági javítást vizsgáló empirikus kutatást** mutatott be. A vizsgálat célja kettős volt: egyrészt **megérteni**, mennyire megbízhatóak az LLM-ek által javasolt sérülékenységi javítások, másrészt **felmérni**, hogyan hat az AI-asszisztencia a fejlesztők hosszú távú szakmai tudására. A kutatás eredményei sokkolóak voltak: bár az AI-k által generált javítások első ránézésre korrektnek tűnnek, a valóságban jelentős részük hibás, hiányos vagy akár újabb sebezhetőséget vezet be. A vizsgálat során **a fejlesztők három csoportban dolgoztak**: AI nélkül, AI támogatással és AI + magyarázó mód használatával. A legérdekesebb megfigyelés az volt, hogy az AI-asszisztált csoport gyorsabban dolgozott, de súlyosabb hibákat hagyott a kódban. A fejlesztők döntő része nem tudta megmagyarázni, miért működik az AI által javasolt javítás, és sokan automatikusan elfogadták a patch-et anélkül, hogy megértették volna a hiba gyökerét. Ez a jelenség a kutatók szerint **skill degradation** – a szakmai kompetencia csökkenése a túlzott automatizációs reliance miatt. Hayati azt is bemutatta, hogy az AI által generált javítások gyakran elfedik az alapproblémát: például input-ellenőrzést adnak hozzá, de nem kezelik a logikai hibát; vagy

# BBA+ BESZÁMOLÓ

egy teljes függvényt átírnak, miközben egyetlen változó korrekt validálása is elegendő lenne. A kutatás több mint húsz esetben talált olyan AI-javítást, amely új sebezhetőséget hozott létre – például versenyhelyzetet, típusinkonzisztenciát vagy hibás jogosultságkezelést.

**Soroush Dalili & Pedram Hayati**

## **ToolShell, Patch Bypass, and the AI That Might Have Seen It Coming**

Az előadás a **2025-ös berlini Pwn2Own verseny** egyik legsúlyosabb incidensét, a **ToolShell** néven ismertté vált **SharePoint-sérülékenységi-láncot** (CVE-2025-49704 és CVE-2025-49706) dolgozta fel. A kutatók részletesen bemutatták, hogyan sikerült a támadónak **két különálló hiba összefűzésével** távoli kód futtatást (RCE) elérnie úgy, hogy megkerülte a SharePoint autentikációs mechanizmusait. A **technikai bemutató** során feltárták a root-okát annak, hogy a Microsoft által július 8-án kiadott első patch miért bizonyult hiányosnak, és miként vezetett ez gyors exploit-fejlesztéshez a fenyegető szereplők oldalán. Dalili és Hayati részletesen elemezték a sérülékenység eredetét, amely egészen a SharePoint 2010-es verziójáig visszanyúló **logikai és jogosultságkezelési konzisztencia problémákra** vezethető vissza. Kiemelték, hogy a modern vállalati rendszerekben a „legacy code path-ek” továbbra is jelentős kockázatot jelentenek, mert a fejlesztők és a biztonsági csapatok sokszor nincsenek tisztában azzal, hogy egy-egy régi API vagy modul milyen útvonalon hívódik meg új funkciók mellett. A **ToolShell** esetében ez lehetővé tette, hogy a támadó autentikációval nem rendelkező végpontokon keresztül olyan kódútvonalakra jusson be, amelyeknek már régen deprecated státuszban kellett volna lenniük. A prezentáció egyik **leginnovatívabb eleme** az volt, hogy bemutatták, **hogyan lehetett volna AI-alapú diff-analízissel már a patch kiadása napján felismerni a bypass lehetőségét**. Egy élő demó során megmutatták, hogy egy LLM képes volt összehasonlítani a patch előtti és utáni kódrészletet, és jelezni, hogy a logikai ellenőrzés még mindig sérülékeny, mivel a hitelesítési

ellenőrzés csak részben került be. Ez nem azt jelenti, hogy az AI „felfedezi az exploitot”, de képes olyan anomáliákra rámutatni, amelyeket egy manuális review könnyen kihagyhat — különösen nagy kódbázis esetén.

## Olena Kutsenko

### Deep dive into data streaming security

Olena Kutsenko előadása a **valós idejű adatfolyam-kezelő rendszerek** — különösen az Apache Kafka — **biztonsági kihívásairól** szólt. Kiemelte, hogy a **streaming platformok** a modern digitális rendszerek alapjai, de gyakran „implicit biztonságban” bíznak, miközben valójában alapból **nem biztonságosak**. Az incidensek — mint a Gamooga 17 GB-nyi kiszivárgott ügyféladata, vagy a GonnaOrder valós idejű logisztikai adatainak nyílt hozzáférése — jól szemléltetik, milyen súlyos következményei lehetnek akár néhány hibás konfigurációnak. Az előadás első része a Kafka **architekturális sérülékenységeire** fókuszált. Ezek közé tartozik a plain-text kommunikáció alapértelmezett használata, a vezetetlen ACL-ek, az elégtelen broker izoláció és az auditlog hiánya. Kutsenko hangsúlyozta, hogy egyetlen rosszul konfigurált broker **„time-bomb”** lehet, amelyhez ha valaki hozzáfér, az egész klaszter felett kontrollt szerezhet. A Kafka-clusterek jellemzően több tucat vagy száz mikroszolgáltatást szolgálnak ki, így a kompromittálás **dominóhatást** eredményez. A prezentáció második része a **védelmi mechanizmusok** mély technikai bemutatása volt. Kutsenko részletesen ismertette az end-to-end encryption, a disk encryption, a SASL/SCRAM autentikáció, a RBAC-alapú jogosultságkezelés és a mezőszintű titkosítás előnyeit-hátrányait. Kiemelte a **kulcskezelés** fontosságát, amely különösen szabályozott iparágakban kritikus: healthcare, pénzügy, insurance. Felhívta a figyelmet arra is, hogy sok szervezet ugyan titkosítja az adatfolyamot, de nem ellenőrzi a producer/consumer identitást, ami hamis biztonságérzethez vezet. A bemutató különösen értékes része volt a **patching és monitoring** realitásainak bemutatása. A CVE-2019-12399 például

# BBA+ BESZÁMOLÓ

egy memória-kezelési hibát érintett, és hónapokig észrevétlenül maradt sok szervezetnél, mert a Kafka cluster-ek monitoringja gyakran csak throughput és latency mérettel operál, nem pedig anomália-alapú viselkedésekkel. Kutsenko több **observability** mintát is ajánlott, köztük a broker-szintű auditlogolást, certificate rotation automatizálást, és a cross-cluster replication biztonságát.

**Dr. Neda Maria Kaizumi**

## **Your AI Is Still Biased (Even After You Checked)**

Dr. Neda Maria Kaizumi előadása arra világított rá, hogy a **mesterséges intelligenciák torzításai nem „egyszeri hibák”**, hanem folyamatosan változó, életciklus-szintű jelenségek, amelyek még akkor is újra megjelennek, amikor a szervezet már lefuttatta a fairness-teszteket, auditált, és kijavította a nyilvánvaló problémákat. A legtöbb vállalat ugyanis úgy kezeli a **bias-kérdést**, mint egy lezárható projektfeladatot, miközben valójában a torzítás inkább egy dinamikus folyamat, amely a tréning, a finomhangolás, a felhasználói visszajelzések és az adatgyűjtés során állandóan alakul. Kaizumi bemutatta, hogy a **torzítások visszatérése** gyakran a modellfrissítések nem szándékolt mellékhatása. Amikor a csapat új adatokat ad hozzá a tréningkészlethez, vagy finomhangolja a modellt egy új célra, könnyen előfordulhat, hogy a rendszer új, korábban nem létező vagy nem mért torzítást vesz fel. Ez különösen igaz a nagy nyelvi modellekre és ajánlórendszerekre, amelyek a felhasználói interakciók alapján folyamatosan tanulnak. A kutató **példát** is hozott: egy vállalati modell a fairness-javítás után kezdetben csökkentette a bizonyos demográfiai csoportokkal szembeni diszkriminációt, ám három hónap múlva – a felhasználói viselkedési adatok miatt – újratermelte a torzítást, ráadásul egy másik dimenzióban. Az előadás egyik hangsúlyos része az volt, hogy a **bias kezelése** nem lehet kizárólag a data science csapat feladata. A torzítás gyakran olyan döntésekbe van beágyazva, mint a UX flow, az adatgyűjtési stratégia, a termékprioritások vagy a jogi megfelelés. Ezért Kaizumi egy **multi-stakeholder bias governance modell**t javasolt, amelyben a termékmenedzser, a jogi csapat, a designer, a QA és a security szakértők együtt vesznek részt a rendszer fenntartásában. A szervezeteknek olyan

folyamatokat kell kialakítaniuk, mint a bias regression testing, a model lineage tracking, a finomhangolások dokumentálása és egy „bias incident response” mechanizmus, amely hasonlóan működik a security incident kezelésekhez.

## Avishay Balter OpenSSF Best Practices for Everyone!

Avishay Balter előadása átfogó és gyakorlatorientált bemutatót adott az **OpenSSF Best Practices Working Group** kezdeményezéseiről, valamint arról, hogyan lehet a **szoftverfejlesztési életciklus minden szakaszában növelni a biztonságot** anélkül, hogy a folyamat lassulna vagy a fejlesztők ellenállása nőne. Balter szerint a modern alkalmazásbiztonság legnagyobb kihívása nem maga a sérülékenység, hanem a kultúra, amelyben a biztonságot gyakran „plusz feladatnak” tekintik. Az **OpenSSF célja** ennek a szemléletnek az átformálása mérhető, fejlesztőbarát módszerekkel. Az előadó részletesen bemutatta a **Scorecard automatizált ellenőrző eszközt**, amely több mint egy tucat biztonsági dimenzióban értékeli a nyílt forrású projektek kockázatait, például függőségkezelés, build integritás, CI konfigurációk, branch protection, és titokszivárgás elleni védelem. A **Scorecard** egyik legfontosabb **előnye**, hogy **közvetlenül integrálható CI/CD pipeline-okba**, így a fejlesztők automatikusan visszajelzést kapnak a biztonsági gyengeségekről – még mielőtt azok éles környezetbe jutnának. Az előadás másik központi eleme a **Security Base Line és az OpenSSF oktatási programjai** voltak. Balter hangsúlyozta, hogy a fejlesztők többsége nem biztonsági szakember, ezért olyan „low friction” eszközökre van szükség, amelyek nem gátolják, hanem támogatják a mindennapi munkát. Az OpenSSF által ajánlott gyakorlatok között szerepel a kódalírás, a reprodukálható build-ek, a függőség-átvilágítás automatizálása, a biztonságos konfigurációk alapértelmezett beállítása, valamint a fenyegetésmodellezés könnyített formáinak bevezetése. Balter példákon keresztül mutatta be, **hogyan javult különböző nyílt**

# BBA+ BESZÁMOLÓ

**forrású projektek biztonsági érettsége**, miután integrálták az OpenSSF ajánlásait. Külön kiemelte, hogy a fejlesztők és a biztonsági csapat közötti feszültség csökkentéséhez elengedhetetlen a jól dokumentált, determinisztikus tooling és a „security-as-code” megközelítés, amelyben a biztonsági szabályok automatizált pipeline-elemek formájában valósulnak meg.

**Roman Zhukov**

## **The EU CRA Readiness for Contributors and Maintainers: A Survival Toolkit**

Roman Zhukov előadása az **Európai Unió Cyber Resilience Act (CRA) hatásait** vizsgálta a nyílt forráskódú szoftverek hozzájárulóira és karbantartóira. Zhukov hangsúlyozta, hogy a legtöbb figyelem a gyártókra és a szoftverfejlesztő alapokra összpontosít, míg az egyéni hozzájárulók és kisebb OSS-projektek fenntartói gyakran bizonytalanok abban, hogyan érinti őket a szabályozás. A **CRA célja** a fogyasztók védelme, de a globális OSS közösségre gyakorolt hatása miatt a fejlesztők számára világos útmutatásra van szükség, hogy biztonságosan folytathassák munkájukat. Az előadás részletezte a **CRA kulcsfontosságú követelményeit**, például a termékbiztonsági dokumentáció kötelezettségét, hibakezelési protokollokat, és a biztonsági kockázatok proaktív kezelését. Zhukov hangsúlyozta, hogy az egyéni fejlesztők nem lesznek automatikusan felelősségre vonva, amennyiben követik a bevált gyakorlatokat, dokumentálják hozzájárulásaikat, és aktívan használják a rendelkezésre álló eszközöket a biztonság érvényesítésére. A bemutatott **gyakorlati eszközök és módszerek** közé tartozott a Security Scorecard a projektbiztonság mérésére, az OSCAL szabványok használata a megfelelőségi dokumentációhoz, valamint a Security Base Line és az OpenSSF Global Cyber Policy Working Group irányelvei. Zhukov konkrét checklist-eket és sablonokat is bemutatott, amelyek segítségével a fejlesztők ellenőrizhetik, hogy hozzájárulásaik megfelelnek-e a CRA előírásainak, miközben nem lassítják a fejlesztési folyamatot. Külön figyelmet kapott az **OSS ökoszisztéma fenntarthatósága**, amely szoros kapcsolatban áll a jogi és biztonsági

megfelelőséggel. Zhukov kiemelte, hogy a CRA nem büntetni, hanem ösztönözni kívánja a biztonságos és ellenálló szoftverfejlesztést, és hogy a hozzájárulók számára az átlátható folyamatok, auditálható munkafolyamatok és eszközök (például automatizált tesztelés, dependency-scanning) kritikusak a sikeres alkalmazáshoz.

## Matteo Di Pirro

### Trust No Input: Taint Analysis at Compile Time

Matteo Di Pirro előadása a **nyelvspecifikus biztonsági mechanizmusok és a taint analysis** használatát tárgyalta a szoftverek biztonságosabbá tételére. A beszélgetés kiemelten foglalkozott azzal a problémával, hogy a modern, összekapcsolt alkalmazásokban a hagyományos hozzáférés-ellenőrzési módszerek gyakran nem elegendőek a kritikus adatok védelmére. A **taint analysis** lehetővé teszi, hogy a potenciálisan veszélyes, „szennyezett” adat útját kövessük a kódban, és a hibákat még fordítási időben, azaz a Productionba kerülés előtt detektáljuk. Az előadás első részében Di Pirro áttekintette a **language-based security** alapjait: hogyan lehet a programozási nyelvek típus- és annotációs rendszerét felhasználni biztonsági szabályok kényszerítésére. Bemutatta, hogyan modellezhető az adat érzékenysége, hogyan propagálható a taint státusz és milyen kompilációs ellenőrzésekkel akadályozható meg a nem kívánt adatfolyás. Példákat adott Java és Scala nyelveken, bemutatva, hogy a **statikus analízis** miként képes felfedezni SQL injection, XSS vagy bizalmas adatszivárgás kockázatokat anélkül, hogy futásidőben ellenőrzéseket kellene végezni. Di Pirro hangsúlyozta, hogy a generatív AI által írt kódok növelik a biztonsági hibák kockázatát, mivel a fejlesztők gyakran nem értik teljesen a generált kód működését. Itt a **taint analysis jelentősége** kiemelkedő, mert a fordítási időben történő ellenőrzés révén az esetleges veszélyforrásokat még a telepítés előtt azonosítani lehet. Az előadás során **gyakorlati példákkal szemléltette a statikus eszközök integrációját** a CI/CD pipeline-okba, lehetővé téve, hogy a fejlesztők automatikusan kapjanak visszajelzést a

# BBA+ BESZÁMOLÓ

biztonsági kockázatokról. Továbbá a foglalkozott a nyelvspecifikus biztonsági eszközök korlátaival és lehetőségeivel. Di Pirro kiemelte, hogy a **statikus elemzés** nem helyettesíti a dinamikus tesztelést, de jelentősen csökkenti az emberi hibákból és a túlzott AI-függőségből eredő kockázatokat.

A **nyelvbe épített biztonsági ellenőrzések kombinálása** a meglévő fejlesztési folyamatokkal lehetővé teszi a biztonság „by design” megközelítését, amely a későbbi incidensek számát minimalizálja. Az előadás kulcsgondolata, hogy a modern alkalmazások biztonságát a fordítási időben történő ellenőrzések és a nyelvi mechanizmusok integrálásával lehet jelentősen javítani, különösen, ha AI-generált vagy harmadik féltől származó kódot használunk. A résztvevők konkrét példákat és technikákat kaptak a taint analysis bevezetésére, amely közvetlenül hozzájárul a szoftverbiztonság erősítéséhez és a hibák korai elhárításához.

**Gary Lopez**

## **How to Break AI Systems (Before Someone Else Does)**

Gary Lopez és Amanda Minnich előadása az **AI rendszerek sebezhetőségeire és a proaktív tesztelés fontosságára** összpontosított. A talk kiindulópontja az volt, hogy a hagyományos szoftverbiztonsági tesztek nem képesek teljesen lefedni az AI-specifikus fenyegetéseket. Az előadók bemutatták, hogy az **AI rendszerek**, például chatbotok és autonóm AI-asszisztensek, alapvetően **különböznek a hagyományos programoktól**: nem képesek megkülönböztetni az utasításokat és az adatokat, így speciális támadásoknak vannak kitéve. A **prezentáció fő fókusza** az volt, hogy hogyan lehet kihasználni AI rendszerek gyengeségeit, mielőtt rosszindulatú szereplők tennék. **Konkrét támadási példákat** mutattak, köztük prompt injection, goal manipulation, adatlopás és rejtett instrukciók dokumentumokban, amelyek révén a rendszer hibásan viselkedik, vagy érzékeny információkat szivároztat. Az előadók hangsúlyozták, hogy ezek az események nem

elméletiek: a való világban gyakran fordulnak elő, és az AI rendszerek adaptív, önálló viselkedése miatt a hagyományos tesztelés nem elegendő. Részletezték a **red teaming** módszertanát, amelynek **célja az AI rendszerek célzott, támadás-orientált vizsgálata**. A résztvevők betekintést kaptak a gyakorlati eszközökbe, amelyek lehetővé teszik a hibák előzetes azonosítását és a sebezhetőségek automatizált feltérképezését. Lopez és Minnich bemutatta a vulnerable AI platform használatát, ahol a hallgatók kísérletezhettek a támadásokkal, így tapasztalati úton értékesítették a red teaming módszert. Az előadás kiemelte a **biztonsági tudatosság fontosságát AI rendszerek fejlesztésekor**. A hallgatóság megtanulta, hogy a sebezhetőségek felismerése és megelőzése nem egyszeri feladat, hanem folyamatos folyamat, amely magában foglalja a retraining, a felhasználói interakciók és az AI adaptív viselkedésének folyamatos monitorozását. A résztvevők konkrét, gyakorlatban alkalmazható stratégiákat kaptak arra, hogyan lehet AI rendszereket stressztesztelni, biztonságukat értékelni, és megelőző védelmi mechanizmusokat bevezetni.

**Lianne Potter**

**Are 'Friends' Electric?**

**What It Means to Be Human Now and Tomorrow in the Age of AI**

Lianne Potter előadása a **humán és mesterséges intelligencia közötti kapcsolat** kulturális, érzelmi és társadalmi aspektusait vizsgálta. A talk kiindulópontja az volt, hogy a technológiai fejlődés – az AI és a robotika – alapjaiban változtathatja meg az emberi interakciókat, társas kapcsolatokat és érzelmi kötődéseket. Potter a történelmi kontextust a sci-fi irodalom és a popkultúra példáin keresztül mutatta be, hangsúlyozva a Gary Numan 1979-es dalát és Philip K. Dick „Do Androids Dream of Electric Sheep?” című művét. **A prezentáció középpontjában az AI-alapú társas kapcsolatok és „elektromos barátok” jelensége állt.** Potter részletezte, hogy

# BBA+ BESZÁMOLÓ

a modern AI rendszerek képesek utánozó társas viselkedésre, érzelmi reakciók generálására és személyre szabott interakcióra. Felhívta a figyelmet azonban, hogy az AI nem rendelkezik valódi érzelmi intelligenciával, így a „barátság” vagy „társaság” élménye strukturált, adat-alapú és gyakran illuzórikus. A bemutató kiemelte a **kulturális és etikai kérdéseket**: milyen hatással lehet az AI-kompanionok elterjedése az emberi kapcsolatokra, a társadalmi izolációra, valamint az érzelmi egészségre? Potter hangsúlyozta, hogy a felhasználói függőség, a valós és virtuális kapcsolatok közötti elmosódás, valamint a társadalmi normák változása kritikus kockázatokat hordoz. Ugyanakkor az **AI társak potenciális előnyeit** is bemutatta: egyedülálló, idősek vagy szociálisan elszigetelt csoportok esetében csökkenthetik a magányt és támogatathatják a mentális jólétet. Az előadás hangsúlyozta a **multidiszciplináris megközelítést**, amely ötvözi a kultúratudományt, pszichológiát, AI fejlesztést és biztonságot. Potter részletesen beszélt arról, hogyan lehet a **design folyamatokba integrálni etikai és kiberbiztonsági szempontokat**, hogy az AI interakciók ne legyenek kizsákmányolóak vagy manipuláló jellegűek.

**Niall Merrigan**

## Exploiting the Supply Chain

AI Niall Merrigan előadása a **szoftverellátási lánc sebezhetőségeit és a támadások modern trendjeit** tárgyalta. A talk kiindulópontja az volt, hogy a hagyományos **supply chain** biztonságot gyakran azonosítják csak függőségekkel, azonban az AI és automatizált eszközök használata lehetővé teszi a támadóknak, hogy gyorsabb, nagyobb hatékonyságú és célzott támadásokat indítsanak. Merrigan bemutatta a **klasszikus módszerek és új trendek kombinációját**, például social engineering, adathalászat, insider threat és automatizált kódtámadás kombinációját. Kiemelte az AI alkalmazását a támadások skálázására, amely lehetővé teszi a több ezer potenciális célpont gyors elemzését és a támadás védelmi mechanizmusokkal való összehangolását. **Konkrét**



**példákon keresztül** bemutatta, hogyan használják a támadók az emberi tényezőt, valamint a percepciót, hogy hozzáférést szerezzenek szervezetekhez, gyakran az alkalmazottak tudta vagy közreműködése nélkül. Az előadás tartalmazott javaslatokat a **kiberbiztonsági védekezés fejlesztésére**, például a supply chain auditok, AI-alapú anomália detekció, és a folyamatos monitoring használatát.

Az **NDC Manchester AI & Security konferencia** előadásain való részvétel kiváló lehetőséget adott arra, hogy **naprakész ismereteket** szerezzünk az AI-alapú kiberbiztonsági megoldásokról. A szakmai előadások bemutatták a legújabb fenyegetésfelderítési technikákat, amelyekkel hatékonyabban lehet azonosítani és kezelni a kibertámadásokat. A konferencia rávilágított arra, hogyan alkalmazható a **mesterséges intelligencia** a bűncselekmények digitális nyomainak gyorsabb és pontosabb elemzésében. Kiemelten foglalkoztak az **automatizált incidenskezelési módszerekkel**, amelyek csökkentik a reagálási időt és növelik a szervezeti ellenállóképességet. A résztvevők **gyakorlati példákon keresztül** megismerhették a modern fenyegetésmodellezési és behatolás-megelőzési stratégiákat. Projektmenedzsment szempontból értékes, hogy a konferencia bemutatja, hogyan lehet AI-eszközökkel optimalizálni a szervezeti erőforrásokat és kockázatkezelési folyamatokat. Segítséget nyújtottak, a komplex biztonsági projektek menedzselésének jobb átláthatóságához, és megmutatták, hogyan támogathatják azokat az adatvezérelt elemzések. A szakmai közeg inspirálja az együttműködést, és lehetőséget teremt a biztonsági és projektmenedzsment gyakorlatok összehangolására. Összességében a konferencián való részvétel és az ott tapasztaltak megosztása **hozzájárul a szervezet biztonságos AI használatának a növeléséhez és a kapcsolódó bűncselekmények felderítési folyamatainak hatékonyabb támogatásához.**

# BBA+ BESZÁMOLÓ

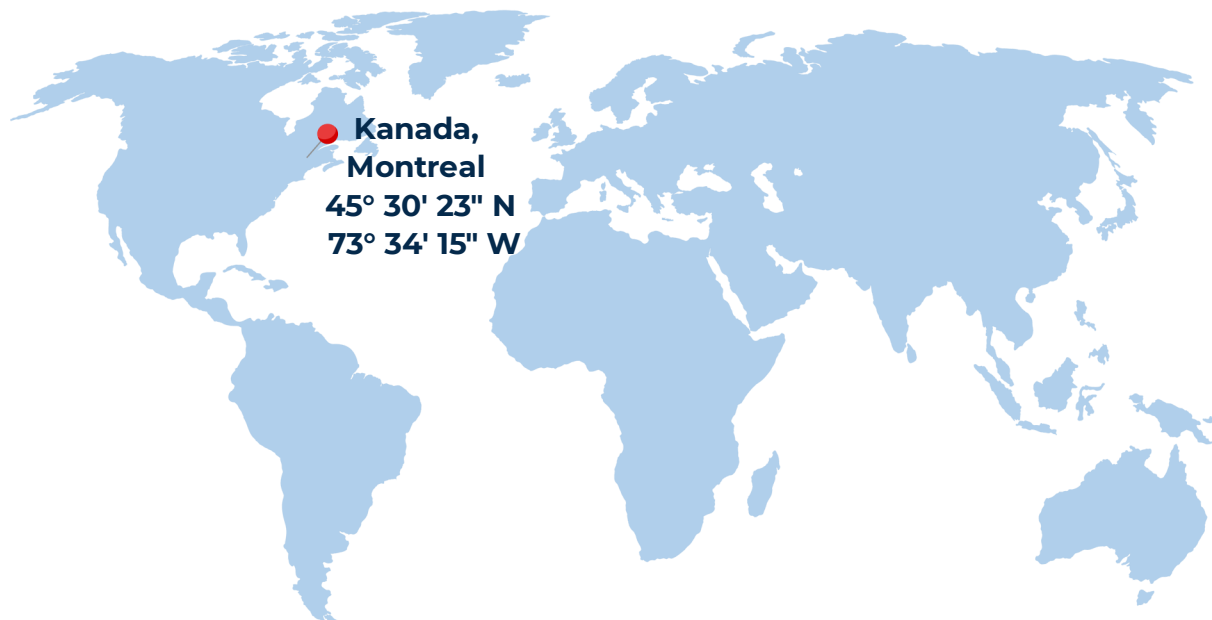
## SuriCon 2025 konferencia



2025. november 19 - 21.

A **Nemzetbiztonsági Szakszolgálat Nemzeti Kiberbiztonsági Intézet** munkatársai **Széchenyi Terv Plusz pályázat** részeként szakmai ismeretbővítésen vesznek részt a súlyos és szervezett, határon átnyúló bűncselekmények elleni küzdelem, illetve ilyen jellegű bűncselekmények megelőzésének fejlesztése céljából.

**A projekt célja** a kiberfenyegetések elleni fellépéshez szükséges friss ismeretek gyűjtése és megosztása a hazai kiberbiztonsági szakemberekkel.





A Nemzeti Kiberbiztonsági Intézet munkatársai **2025. november 19-21.** között részt vettek a **Montrealban** megrendezésre kerülő éves **SuriCon konferencián**.

A **SuriCon** 2015 óta évente megrendezett, **nemzetközi kiberbiztonsági konferencia**, amely a **nyílt forráskódú Suricata** köré épül. A Suricata egy **hálózatbiztonsági rendszer**, amely behatolás-észlelési (IDS), behatolás-megelőzési (IPS) és hálózati biztonsági monitorozási képességeket biztosít. A rendezvényt a **Open Information Security Foundation (OISF)** szervezi, és évente több szakembert, fejlesztőt és kutatót gyűjt össze, hogy megosszák a legújabb technikai eredményeket, gyakorlati tapasztalatokat és a Suricata jövőbeli fejlesztési irányait. A 2025-ös, Montrealban tartott esemény az 11. ilyen konferencia volt és három napon át tartott, kiegészítve gyakorlati képzésekkel és közösségi programokkal.

### Kiknek ajánlott a beszámoló megismerése?

-  Vállalati IT-biztonsági csapatoknak
-  Kiberfenyegetés-kutatóknak és elemzőknek
-  Biztonsági mérnököknek és fejlesztőknek
-  Kiberbiztonsági szakemberek számára

**A konferencia legfontosabb előadási az alábbiakban foglalhatók össze.**

# BBA+ BESZÁMOLÓ

## Ron Bowes & Glenn Thorpe (GreyNoise) Abusing HTTP Quirks to Evade Detection

Az előadás központi eleme, hogy a **HTTP-protokoll** rugalmas és meglepően toleráns felépítése számtalan olyan „szűrkezőnát” tartalmaz, amelyeket a támadók ügyesen kiaknázhathatnak. A **webkiszolgálók és a middleware-rétegek** — legyen szó Apache-ről, Nginxről, IIS-ről vagy épp PHP-ről és .NET-ről — gyakran megpróbálják „kijavítani” a hibás vagy szabálytalan kéréseket. Emiatt azonosnak tűnő, de valójában eltérően formázott HTTP-mezők (például duplikált host-header, atipikus kódolás vagy nem standard elválasztók) különböző komponensekben más-más jelentést kapnak. Ez a viselkedés a **megbízható működés fenntartása** szempontjából hasznos, de egyben lehetővé teszi, hogy a támadó egy olyan HTTP-konstrukciót küldjön, amelyet a célrendszer érvényesnek értelmez, miközben az IDS — például Suricata — teljesen másképp dolgozza fel ugyanazt a csomagot. Ezáltal a detektálás is ellehetetlenül. Példaként a path traversal lett felhozva, ahol nem csak a jól ismert „../” szekvencia variálható, hanem különböző kódolási formák, dupla százalékjeles dekódolás, Unicode-trükkök, illetve nem tipikus header-struktúrák is megkerülhetik az egyszerű mintakeresést. Hasonlóan, egy SQL- vagy shell-injekciós payload rejtve maradhat, ha az URL-kódolást, a chunked transfer encodingot vagy a sorvégeket eltérő módon használják — amelyek közül több olyasmi, amit a kommunikáló szerver automatikusan „helyrehoz” a támadó helyett. Ezek a különbségek elégségesek lehetnek ahhoz, hogy a Suricata egy jól megírt szabálya is hibásan értelmezze a forgalmat, és rávilágít arra, hogy a **protokollok implementációjának** esetei mennyi kikerülési lehetőséget rejtenek.

Az előadás második felében bemutatták azt is, hogyan lehet a Suricata-szabályok írását és validálását modern eszközökkel és automatizációval hatékonyabbá tenni. Felhívják a figyelmet arra,

hogya a szabályok tesztelésének nem szabad kizárólag statikus mintákra épülnie — szükség van fuzzolásra, különböző kódolások variálására, illetve ezek ellenőrzésére alkalmas tesztcsomagokra. Az **automatizált környezetek** segítségével gyorsan ellenőrizhető, hogyan reagál a Suricata különböző torzított vagy formabontó HTTP-kérésekre, és hol nem működik a detekció. A hangsúly a **szabályok folyamatos fejlesztésén** van: a támadói kreativitás és a protokoll engedékenysége miatt a védekezés csak akkor hatékony, ha folyamatosan teszteljük a rendszer viselkedését az atipikus valós támadási mintákkal szemben.

**Arezki Laga & Mohamed Amine Larabi**

## **A Signature-Based Approach to Detect Evolving Malware Communication Patterns and Behaviors**

Hagyományosan az Suricata-hoz és más IDS/IPS rendszerekhez használt **detekciós módszerek** úgynevezett „**signature-based**” **elven** működnek: a rendszer adatbázisában tárolt ismert támadási mintákkal veti össze a bejövő forgalmat, és ha egyezést talál, riasztást vált ki. Azonban az **egyik fő probléma**, hogy a rosszindulatú szoftverek (malware-ek) folyamatosan változtatják kommunikációs mintáikat, viselkedésüket, így a hagyományos, statikus alapon létrehozott signature-bázis egyre kevésbé képes lépést tartani a támadási technikákkal. Emellett a **signature-adatbázis** növekedésével együtt jár az IDS teljesítményének lassulása és a memóriahasználat növekedése — végső soron pedig a performancia csökkenése. Az ajánlott új megközelítés az „**approximate matching**” signature-alapon. Ez olyan signature-alapú eljárás, ami képes felismerni a korábban nem látott, de ismert malware-mintákhoz “hasznló” kommunikációt: azaz approximate string matching (ami általánosabb minták alapján történő

# BBA+ BESZÁMOLÓ

hasonlóságelemzés) segítségével azonosít új variánsokat. Ez azt jelenti, hogy nem kell minden egyes új malware-variánshoz külön szabályt írni — a módszer képes “távolabbi, de még releváns” egyezéseket is felismerni, ami nagyban csökkentheti a karbantartási terhet és javíthatja a detekció rugalmasságát. Ezzel a megközelítéssel hatékonyabban lehet az erőforrásokat felhasználni, mint a bonyolult, reguláris kifejezésekre épülő szabályok. Ez az **új, adaptív signature-alapú módszer** — amennyiben éles rendszerbe kerül — jelentős előrelépés lehet azoknak a biztonsági üzemeltetőknek, akik hagyományos IDS/IPS rendszereket (pl. Suricata) használnak, és szeretnék követni a malware-ek gyors evolúcióját anélkül, hogy folyamatosan bővítenék, finomítanák a szabályrendszert. Ez a technika a **„skálázhatóságot, hatékonyságot, karbantarthatóságot”** ígér a detekció lefedettségének növelése mellett. Ugyanakkor ez az irány is rávilágít egy fontos tendenciára: **a jövő IDS/IPS rendszereknek nem elég statikus alapon működniük** — szükség van adaptív, intelligens, „közel-egyezések” módszerekre, amely figyelembe veszi a malware-kommunikáció dinamikus, változó természetét. Ez pedig egyfajta evolúciót jelent a hagyományos signature-based megközelítés és az anomália- vagy viselkedésalapú módszerek között — középutat kínálva az ismertek biztos felismerésének és az ismeretlen variánsok észlelésének között.

## Markus Kont

### Pikksilm: A Tale of Unholy Alliance between Endpoint Agents and Suricata

A **hálózati alapú IDS/NSM-rendszerek** (mint a Suricata) nagyon jól követik a malware-szállítást (delivery), a Command-and-Control (C2) csatornákat, valamint az adat-exfiltrációt. Ugyanakkor a valódi problémák gyakran már korábban keletkeznek — az exploitáció és a malware telepítése a végponton (endpoint) zajlik, és ez a hálózati forgalomban akár teljesen „láthatatlan” lehet.

Emellett a támadók titkosítást és modern CDN-hálózatokat használhatnak, hogy elrejtsek a rosszindulatú forgalmat — így a hálózati monitorozás önmagában gyakran vakfolttal küzd. Ez azt jelenti, hogy pusztán a hálózati pakettek alapján nehéz vagy lehetetlen megbízhatóan megállapítani, mi is történik a végpontokon. A **Pikksilm egy nyílt forráskódú eszköz**, amelynek célja, hogy a **hálózat monitorozásából származó eseményeket** (pl. a Suricata által generált riasztásokat vagy session-logokat) **összekösse a végpontokon (endpoint) futó folyamatokról származó eseményekkel** (pl. folyamat-indítás, hálózati kapcsolat — Sysmon logból). Így látjuk, hogy melyik program milyen hálózati kapcsolatot kezdeményezett. Ezek alapján létrehozható egy közös **„entity\_id / community\_id”**, amely lehetővé teszi, hogy egy hálózati kapcsolatot visszakössünk a konkrét folyamatra, parancssorra, felhasználóra, hostra etc. Ezáltal a biztonsági elemzők azonnal átlátják a forgalom kommunikációinak összefüggéseit, így olyan támadások is láthatóvá válhatnak, amelyek a hálózati forgalom alapján elkendőzve, „normál” forgalomnak tűnnének. Az ilyen típusú **integrált megközelítés** — endpoint + network korreláció — segíthet minimalizálni azt a klasszikus „vakfoltot”, amit a pusztán hálózati monitorozás jelent: így a malware exploitálás, telepítés és a mögöttes folyamatok is láthatóvá válnak. A Pikksilm nyílt forráskódú és különösen **kis- vagy közepes méretű szervezeteknek** lehet kedvező. Azonban a megközelítés hatékonysága ipari méretű rendszereknél is használható lenne megfelelő fejlesztéssel.

# BBA+ BESZÁMOLÓ

Reid Wightman

## Suricata for ICS: Tips and Moar Research

Reid Wightman az **ipari irányítórendszerek** (ICS/SCADA) biztonságának ismert kutatója, korábban az ICS-CERT és a Dragos csapatának tagja. Nevéhez fűződik számos kritikus ICS-sebezhetőség feltárása. Előadásában a **Suricata hatékonyabb használatát** mutatta be ipari protokollok (ICS) elemzéséhez.

A fókusz egy olyan **„kerülőút” jellegű technikán** van, amelynél Suricata-szabályláncok (flow-változókat használó több egymásra épülő szabály) segítségével képes **kiolvasni protokollmezőket a válaszcsomagokból**, amiket a Suricata egyébként nem dekódolna. Valós példán keresztül mutatta be, hogyan detektálhatók egy Modbus vagy IEC-104 válaszban, a manipulált vagy értelmetlen értékek, pusztán IDS-szinten. **Az eljárás lényege**, hogy az első szabály kinyeri a releváns byte-pozíciót és elmenti egy flow-változóba, a következő szabályok pedig ezt a változót felhasználva vizsgálják az értéktartományt — gyakorlatilag egy mini **„preprocesszor”** jön létre csupán a létező Suricata szintaxis használatával. Ez a megközelítés anélkül teszi alkalmassá az IDS-t ipari rendszerek hatékonyabb védelmére, hogy magát a Suricata kódját nem kell módosítani. A gyakorlatban csupán **automatizálható szabálygenerátorokra** van szükség, amelyek a komplex szabályláncok írásával felkészítik az ipari protokollok vizsgálatára. Ennek segítségével képes lehet az ICS-rendszerek monitorozásakor olyan szintű részletességre, ami másként csak dedikált ICS-monitorozó eszközökkel lenne elérhető. Végeredményben a javaslat **egy hibrid út** — nem pusztán hagyományos IT-forgalom figyelése, hanem kifejezetten ipari forgalomra optimalizált Suricata-konfiguráció —, mely **kibővíti az IDS/NSM eszközök hatókörét a kritikus infrastruktúrák felé** is. Néhány konkrét ipari protokoll, amelyeknél az előadó szerint különösen hasznos lehet a Suricata-alapú, „soft-preprocesszoros” megközelítés — vagyis amikor szabályláncokkal és flow-változókkal próbálunk olyan mezőket is kiolvasni, amelyeket a natív Suricata nem képes vizsgálni:

### 1) Modbus/TCP

Az egyik legelterjedtebb **PLC-kommunikációs protokoll**. A standard Suricata ugyan felismeri a kereteket, de nem minden funkciókódot képes értelmezni. A **soft-processzoros módszerrel** az irreális vagy manipulált regiszterértékek riasztást váltanak ki.

### 2) DNP3

Elsősorban **energia infrastruktúrában** (alállomásokban, SCADA-rendszerekben) használatos. A **beágyazott struktúrái** miatt (header → object → variation → qualifier) nehéz preprocesszort írni hozzá. Ezzel a technikával detektálni lehet anomális I/O-értékeket vagy olyan visszatérő parancsokat, amelyek egy támadó felderítő tevékenységére utalnak.

### 3) BACnet

Jellemzően **épület automatizálási rendszerek** (HVAC, ajtóvezérlés, liftek, tűzvédelem) protokollja. A válaszcsomagok sokszor tartalmazznak változó hosszúságú mezőket és ezeket a Suricata alapértelmezetten nem mindig tudja dekódolni. Szabályláncsal például az „unexpected write property” kérésekre és azok válaszártékeire lehet figyelni.

### 4) IEC 104 / IEC-60870-5-104

Európában különösen gyakori **elektromos hálózatok SCADA-kommunikációjában**. Nagy hátrány, hogy a Suricatához nincs teljes értékű beépített preprocesszor, pedig a módszerrel különböző típusokat és azok paramétereit is monitorozni lehet.

### 5) Profinet / S7comm (Siemens)

Gyakran **PLC-k programozási és folyamatvezérlési** csatornái. Ezeknél tipikus probléma, hogy a parancsok és válaszok szerkezete több rétegben ágyazódik be. A különálló Suricata-szabályok összefűzésével még így is figyelhető például a memóriaterületekből kiolvasott értékek tartománya.

# BBA+ BESZÁMOLÓ

Konstantin Klinger

## Meerkat in the Sandbox: Turning Rule Hits into Verdicts

Az előadáson bemutatásra került, hogyan használta fel Konstantin Klinger a **Proofpoint sandbox-elemző** projektjében a **Suricata-t** és az **Emerging Threats (ET / ETPRO) szabálykészletet**: nem élő hálózati forgalmon, hanem sandbox-környezetben generált PCAP-ek elemzésére. A bemutatott környezet leginkább egy **on premise** Vírustotal működéséhez hasonlít, ami komplex pontozórendszer alapján értékeli a találatokat. Összességében az előadás azt bizonyítja, hogy a **szabály-metadata** nem pusztán „kiegészítő információ”, hanem stratégiai érték — különösen sandbox- és labor-környezetben, ahol a cél nem a valós forgalom blokkolása, hanem a behatolási láncok feltárása, a malware-kutatás és fenyegetés-intelligencia építése.

Lucas Aubard & Johan Mazel

## PYROLYSE: How to Burn Network Stack with Overlapping Data

A kutatók a **hálózati forgalom töredezettségéből adódó biztonsági kockázatokat** elemezték, ehhez fejlesztették a **PYROLYSE audit-eszközt**. Az ismert, hogy az **IPv4, IPv6 és TCP protokollok** lehetővé teszik, az eredeti adatcsomag több fragmentre (chunk-ra) szétszedve történő forgalmazását. Ezek a fragmentek részben vagy teljesen átfedhetik egymást (overlap), ami az adat újra-összeállítás (reassembly) során különböző viselkedést okozhat az eltérő hálózati szegmensekben. A **PYROLYSE eszközzel tesztelést végeztek**: több sorozatú overlapping fragment-kombinációt generáltak (több részletre bontva), és különböző OS-eket, NIDS-eket, beágyazott rendszereket vizsgáltak. A vizsgálat során 23 implementációt teszteltek, amelyek közül 14-20

különbéle reassembly-viselkedést azonosítottak protokoll és konfiguráció függvényében. Az **eredményeik** riasztóak: találtak nyolc olyan hibás viselkedést, amelyek lehetővé teszik, hogy a fragmentált és overlapping csomaggal támadó úgy manipulálja a forgalmat, hogy az IDS/NIDS másként rakja össze az adatot, mint a célrendszer: így a rule-matching (mintafelismerés) vagy teljes policy-ellenőrzés kikerülhet. Hangsúlyozzák **sok fragment összetett átfedése** esetén (3 vagy több) új, váratlan viselkedés léphet fel. Ezért a fragmentálásból adódó támadási vektorokat szisztematikusan, kell auditálni, nem csak néhány mintával.

## Ambre looss

### Shovel: Leveraging Suricata for Attack-Defense CTF

A **Shovel-t** a francia **ECSC Team** hozta létre, kifejezetten **CTF-versenyek** (attack-defense) közbeni **hálózati forgalom-elemzésre**. Ez egy ingyenes, nyílt forráskódú webes felület, amely a Suricata EVE (esemény-/flow-log) kimenetét vizuálisan, felhasználóbarát módon teszi hozzáférhetővé és kereshetővé. Tehát lényegében egy **frontend / vizualizációs réteg** a Suricata mögött. A Shovel **háromféle forgalom-feldolgozási módot** támogat: pcap-újrátjátszást (replay egyes pcap-ekkel), élő capture-t hálózati interfészről, vagy PCAP-over-IP-vel. Ez lehetőséget ad arra, hogy CTF- környezetben akár élő forgalmat figyeljünk, akár utólag pcap-ek alapján elemezzünk. A **frontend használata CTF-verseny közben** nagyban csökkenti az elemzéshez szükséges időt: gyorsan észlelhető, ha egy szolgáltatást exploitáltak és nem kell raw pcap-et böngészni vagy manuálisan adatokat kinyerni. Ez kritikus lehet idő szempontjából, kiváltképp nyomás alatt. A **Shovel segítségével** bemutatták, hogyan lehet egy profi NIDS / NSM-motor (Suricata) használatával nem csak üzemi környezetben, hanem esemény-orientált, gyors reagálást igénylő helyzetben **hatékonyan vizualizálni, szűrni és reagálni a hálózati**



NEMZETI  
KIBERBIZTONSÁGI  
INTÉZET

# BBA+ BESZÁMOLÓ

**eseményekre.** Ezáltal verseny helyzetben a résztvevők könnyebben írhatnak saját Suricata-szabályokat, izolálhatják a támadói forgalmat, vagy akár tűzfalakon keresztül blokkolhatják az exploitáló gépeket. A **projekt végső célja** azonban több, mint egy jól használható CTF-eszköz: **nyílt forráskódú bázist** kívánnak létrehozni a CTF-közösség számára, ahol bárki tanulhat, fejleszthet, és visszaadhat a közösségnek.