



HÍRLEVÉL

Nemzetközi
IT-biztonsági sajtószemle
2025. 23. hét



HÍREK

- Új backdoor kampány ASUS routerek ellen
- Ha a hacker böngészőjét használja, sokba fog kerülni!
- Letöltés helyett fertőzés: MI-csalik a célkeresztben
- Több mint 1700 GitHub-projektet érintő sérülékenység
- Az OpenSSH mint támadási vektor



STATISZTIKAI ADATOK

- Incidensek eloszlása típus és kockázati besorolás szerint
- Események eloszlása csapdatípusok alapján
- Támadott port szerinti eloszlás



IT BIZTONSÁGI TIPP

- Az online fájlkonvertáló alkalmazások veszélyei



BESZÁMOLÓ

- FIRST Technical Colloquium 2025



KONTAKT

edt@nki.gov.hu

PGP kulcs

FBC3 88A2 E465 BF51
AD58 A2D0 E9DD E078
ABD3 E75D



További hírekért, látogasson el **weboldalunkra!**

NEWS

IT biztonsági HÍREK

Ha a hacker böngészőjét használja, sokba fog kerülni!
(thehackernews.com)

A hagyományos Man-in-the-Middle (MitM) típusú támadásokkal szemben egyre nagyobb teret nyer egy új, kifinomultabb technika, a Browser-in-the-Middle (BiTM) támadás. A BiTM célja továbbra is a kliens és a szolgáltatás közötti adatfolyam kompromittálása. **Bővebben...**

Letöltés helyett fertőzés: MI-csalik a célkeresztben
(thehackernews.com)

Hamis telepítőprogramokat alkalmaznak népszerű mesterséges intelligencia (MI) alapú eszközökhöz, mint például az OpenAI ChatGPT vagy az InVideo AI, amelyekkel különféle rosszindulatú szoftverek — így a CyberLock és Lucky_Gh0\$t zsarolóprogramok, valamint az újonnan azonosított Numero malware — terjesztése történik. **Bővebben...**

Több mint 1700 GitHub-projektet érintő
sérülékenység
(gbhackers.com)

Egy friss, nagyszabású kutatás súlyos biztonsági problémát tárt fel a Node.js-alapú statikus fájlkiszolgálók körében: egy széles körben elterjedt útvonalbejárás (path traversal) sebezhetőség, amely több mint 1756 nyílt forráskódú projektet érint a GitHubon. **Bővebben...**

Az OpenSSH mint támadási vektor
(gbhackers.com)

Egy friss kutatás súlyos biztonsági hiányosságokat tárt fel a modern webinfrastruktúrában. A vizsgálat rámutat, hogy a HTTP/2 szerver push és a Signed HTTP Exchange (SXG) technológiák kihasználhatók a Same-Origin Policy (SOP) megkerülésére. **Bővebben...**



Új backdoor kampány
ASUS routerek ellen
(greynoise.io)

2025 márciusában a GreyNoise kutatói egy kifinomult támadási kampányt azonosítottak, amely során ismeretlen támadók több ezer, internetre csatlakoztatott ASUS routerhez szereztek jogosulatlan, tartós hozzáférést. **Bővebben...**



További hírekért, látogasson el **weboldalunkra!**

Statisztikai Adatok

2025.05.30.-2025.06.05.

Az NBSZ NKI által kezelt incidensekre vonatkozó statisztikai adatok:



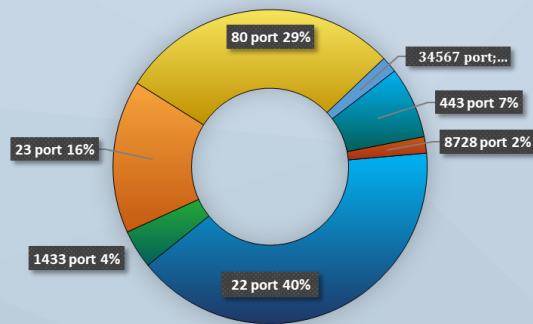
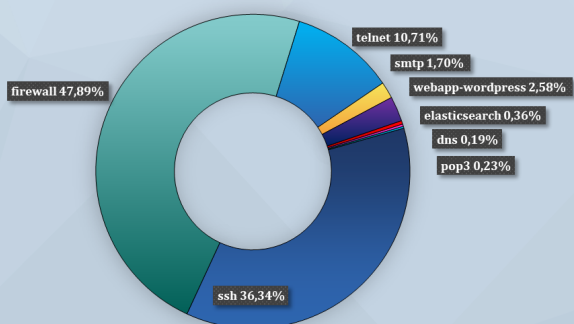
Fenyvegetettségi szint: magas



■ Alacsony ■ Közepes ■ Magas ■ Kritikus

Incidensek eloszlása típus és kockázati besorolás szerint

Az elosztott kormányzati IT-biztonsági csapdarendszerből (Gov1probe) származó adatok:



További érdekességekért, látogasson el [weboldalunkra!](#)



Aktuális tartalmak



Az online fájlkonvertáló alkalmazások veszélyei

Valószínűleg sokak számára ismerős az a helyzet, amikor olyan fájlformátumban kapunk egy dokumentumot vagy képet, amelyről korábban még csak nem is hallottunk – nemhogy megtudnánk azt nyitni. Ilyen **például**, amikor a megszokott **JPEG vagy PNG** formátumú képek **helyett HEIC formátumban** érkezik egy fotó.

Ilyenkor a legtöbben **online fájlkonvertáló szolgáltatásokhoz fordulnak** segítségért. Ezek az eszközök valóban **gyors és kényelmes megoldást** kínálnak – ám nem árt tisztában lenni azzal, hogy **használatuk** bizony komoly **kockázatokat** is **rejt**enek magukban.

[Elovasom](#)

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



További hírekért, látogasson el **weboldalunkra!**

Aktuális tartalmak



Malware elleni védekezés: a láthatatlan ellenség SANS OUCH!

Megjelent a **SANS** és a Nemzetbiztonsági Szakszolgálat **Nemzeti Kibervédelmi Intézet** közös kiadványának **2025. júniusi száma**, melyben bemutatjuk azt a rosszindulatú szoftvert (**malware**), amely lehetővé teszi a **kiberbűnözők** számára, hogy **hozzáférjenek számítógépes rendszereinkhez és mobileszközeinkhez**.

Emellett olyan egyszerűen alkalmazható **módszereket** is ismertetünk, amelyekkel **megelőzhetők** az **ilyen támadások**.

Elovasom

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



További érdekességekért, látogasson el **weboldalunkra!**



FIRST Technical Colloquium 2025



2025. március 25-27.

A Nemzetbiztonsági Szakszolgálat Nemzeti Kibervédelmi Intézet munkatársai **Széchenyi Terv Plus 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 Kibervédelmi Intézet** munkatársai részt vettek a 2025. március 25-27. között Amsterdamban megrendezésre kerülő éves **FIRST Technical Colloquium** konferencián.

A **Forum of Incident Response and Security Teams** (FIRST) műszaki kollokviumai és szimpóziумai ingyenes, szakmai párbeszédre alkalmas fórumot biztosítanak a FIRST tagsággal rendelkező csapatok, valamint meghívott vendégeik számára. A rendezvényeken a résztvevők megoszthatják egymással tapasztalataikat sebezhetőségek, biztonsági incidensek, eszközök és minden olyan kérdés kapcsán, amelyek hatással vannak az incidensreagáló és kiberbiztonsági csapatok napi működésére.

Az amszterdami TC (Technical Colloquium) mindig is technikai fókuszú volt, különös tekintettel az incidensreagálásra és kiberfenyegetés felderítésre. Az esemény egy egynapos Dark Web témájú képzéssel indult, majd ezt követte a két napos előadássorozat. Többek között szó esett a **mesterséges intelligencia** szerepéről **az incidensészlelésben és -kezelésben**, a fenyegetési szereplők (pl. APT csoportok) technikáiról, kreatív válaszcádási módszerekről, valamint a felhő- és konténeralapú környezetek biztonsági kihívásairól. A rendezvény célja az volt, hogy technikai mélységgel és valós tapasztalatokon alapulva **segítse a közösség tudásmegosztását és fejlődését** a folyamatosan változó fenyegetési környezetben.

Kiknek ajánlott a beszámoló megismerése?

- ▶ Vállalati IT-biztonsági csapatoknak
- ▶ Kiberfenyegetés-kutatóknak és elemzőknek
- ▶ Biztonsági mérnököknek és fejlesztőknek
- ▶ Minden kiberbiztonsági szakember számára



Gabriel Cirlig (HUMAN Security, GB), Lindsay Kaye (HUMAN Security, US): Hogyan szörfözz a Dark Weben, mint egy profi (vagy mint egy kelet-európai)

„Az infrastruktúrát támadások érik, és a helyzet napról napra romlik. Úgy tűnik, az akciók jól szervezettek – de hol is kezd el a támadók felkutatását?” Ebben a 4 órás **OSINT- és nagyléptékű adatgyűjtési workshopon** az előadók bemutatták, hogyan tudunk **„igazi kiber-nyomozóként” információhoz jutni** a szervezett támadások, árusított technikák és „know-how”-okat tartalmazó platformokon.

A workshop során az előadók gyakorlati példákon keresztül mutatták be a **Dark Web** világát, továbbá bemutattak olyan erőforrásokat is, amelyek segítenek megtalálni az általunk keresett információkat, mindezt úgy, hogy közben megőrizzük anonimitásunkat.

Betekintést nyertünk a **TOR hálózat** használatába, felkerestük a legnépszerűbb Dark Webes információmegosztó oldalakat, feltártuk ezen fórumok megtekintésének potenciális veszélyeit, illetve olyan technikákat tanultunk az információk biztonságos (akár automatizált) nagyméretű kinyerésére a különböző platformokról. A workshop részeként továbbá említésre került az ilyen jellegű felderítő és információgyűjtő tevékenységekkel járó potenciális jogi következmények, illetve ezen jogi vonatkozásoknak a „kezelése”.

Marina Bochenkova (Corbion Group BV, NL): A hívószótól a csatamezőig: az okosvárosok kiberbiztonsági kihívásai

Az előadó bevezetésében a **„smart city”** (okosváros) kifejezés eredetét mutatta be, amely immár több mint egy évtizede divatos hívószó a politikusok, várostervezők és technológiai vállalatok körében – a csillogó ígéretek azonban gyakran elfedik a **mögöttes kockázatokat**.

Az előadás során ezután **önkormányzatokat érő kibertámadások**, leállások és károk rendszeres médiafigyelméről volt szó, azonban az IT biztonsági szakembereknek elsősorban az IT rendszerek biztonságára és helyreállítására kell összpontosítania.





Az okosvárosok természetüknél fogva **az IT és OT** (informatikai és hagyományos mechanikai vezérlő) **rendszerek keverékére épülnek**, viszont ezek együttes, átfogó kockázatkezelésére jelenleg nincs kialakult gyakorlat. A **biztonság szempontjából kritikus következmények**, amelyek az okosvárosok tervezésében és megvalósításában részt vevő különféle szereplőkből adódnak, gyakran észrevétlenek maradnak. Az emberi tényezők, a sebezhető eszközök és az adatkezelési problémák csak tovább fokozzák a támadási felületet, új és kreatív lehetőségeket kínálva a fenyegetési szereplők számára.

Példaként az előadó megemlítette, hogy Amsterdam 2010-ben tartott egy fejlesztői versenyt, amely során a cél az amszterdami polgárok életét megkönnyítő alkalmazások fejlesztése volt. Amsterdam az „**Open Data Projekt**” keretében mindenki számára elérhetővé teszi a várossal kapcsolatos információk hatalmas tömegét. Ilyen a valós idejű parkolók foglaltsági információja, házépítési tervek, földalatti kábelek helyét, méretét és mennyiségét, illetve a metróvonalak pontos elhelyezkedése. A verseny keretében egy fejlesztő a korábban felsorolt információkat felhasználva fejlesztett egy alkalmazást, amely a közvilágítás, a legdrágább házak, és a rendőrségtől való távolság információit felhasználva megmutatja melyik környéken a legkifizetődőbb a lakásokba betörni. Habár triviálisnak tűnik, ez egy nagyszerű példa arra, hogy a nyilvánosan elérhető információkat kihasználva egy támadó milyen **extra támadási felülethez juthat**.

Az okosvárosok egyedülálló és mindenre kiterjedő kockázati kombinációt képviselnek, amelyet csak átfogó elemzéssel lehet kezelni annak érdekében, hogy növeljük az eljárásrendek és a működés biztonságát, megbízhatóságát és ellenálló képességét. Ha újraértelmezzük, mit is jelent valójában egy okosváros, már meglévő, gyakorlatban is alkalmazható stratégiákat tudunk bevonni a védelem megerősítésére – a világiárványoktól kezdve az állam által támogatott kibertámadásokig. Ahogy a politikailag motivált támadások hatóköre és mellékhatásai növekednek, fel kell készülnünk arra, hogy nem intézményeink, hanem **akár egész városaink válhatnak a következő célponttá**.

Az előadás összességében bővítette az okosvárosokról alkotott fogalmainkat, rávilágított az ezekhez kapcsolódó adat-, emberi és technológiai kockázatokra, és olyan módszertanokat mutatott be, amelyek segíthetnek ezek kezelésében.

Julie Agnes Sparks (Datadog, US), Juvenal Araujo (Datadog, PT): CI/CD biztonság a gyakorlatban: stratégiák a fenyegetések észlelésére és kezelésére

Az előadó azt a kérdést járta körbe, hogy mire lehet képes egy támadó, ha bejut egy olyan környezetbe, amely tele van érzékeny adatokkal és szellemi tulajdonnal. Az előadásban korábbi, **CI/CD** (Continuous Integration/Continuous Delivery, azaz folyamatos integráció és kiadás) rendszerek elleni támadások elemzésén keresztül mutatta be az előadó, hogyan segíthet egy **jól kialakított CI/CD fenyegetési modell a veszélyek felismerésében és a hatékony reagálásban**. Az előadás során végig vettük, hogy **milyen logolási eszközeink vannak** a fenti specifikus környezetben, illetve **milyen best practiceket tudunk alkalmazni** a logolás megtervezése és felállítása során. Említést tett az incidenskezelési metodológiáról, illetve végig vettük a **leggyakrabban kihasznált támadási módszereket és vektorokat** (pl.: Acces Token viselkedés és ezek exploitációja).

Mackenzie Jackson (Aikido Security, NL): Shadow Patching: mesterséges intelligencia segítségével azonosított, nem közölt biztonsági javítások nyílt forráskódú szoftverekben

Az előadó azzal a felvetéssel indított, hogy a **silent patching** – azaz a biztonsági sebezhetőségek nyilvános bejelentés nélküli javítása – **komoly vakfoltot jelent a szoftverellátási lánc biztonságában**. Mivel minden hatodik sebezhetőséget, ilyen módon, „csendben” (bejelentés nélkül) javítanak, a hagyományos, nyilvános sérülékenységi adatbázisokra (pl. CVE, NVD) támaszkodó biztonsági eszközök gyakran nem tudják észlelni ezeket a hibákat, így **a szervezetek rejtett kockázatoknak vannak kitéve**. Ezenfelül a kitettségéről nem tudva a szervezet arról sem kap adott esetben információt, hogy mennyire **fontos** az adott **szoftver/eszköz frissítése**.

Az előadás egy teljesen új megközelítést mutatott be, amely nyelvi modellek (LLM-ek) erejét használja fel rejtett javítások azonosítására nyílt forráskódú szoftverekben.



Az előadó bemutatta, hogyan működik a saját fejlesztésű, erre a célra létrehozott kettős LLM-architektúrájuk, amely **nyilvánosan elérhető változásnaplók** (changelogok) **elemzése alapján képes felismerni és kategorizálni a csendben javított sebezhetőségeket**. Egy demón keresztül került bemutatásra, hogyan segített ez a mesterséges intelligencia-alapú megközelítés több száz korábban ismeretlen sérülékenységi javításának feltárásában nagyobb, nyílt forráskódú projektekben – amelyek 20%-a kritikus vagy magas kockázatú besorolást kapott.

Főbb témák:

- A silent patching veszélyei és hatása a szoftverellátási lánc biztonságára
- A dual-LLM modell részletes felépítése és működése
- Valós esetekből származó eredmények és azok jelentősége az IT biztonsági szakma/közösség számára
- A Human-in-the-Loop (HITL) szerepe a mesterséges intelligenciával támogatott folyamatok ellenőrzésében
- Összehasonlító elemzés a hagyományos biztonsági kutatási módszerekkel
- A jelenlegi megközelítés korlátai és a jövőbeni fejlesztési lehetőségek



Nikolas Dobiasch (SAP SE, AT): Engineering the Loop: az SAP útja a folyamatos viselkedésalapú tesztelés felé

Az előadó egy folyamatos biztonsági validációs program összetett vállalati környezetben való elindításához szükséges alapos tervezésre, világos használati esetek feltárására és szisztematikus érvényesítési módszerekre épülő folyamatot járt körbe. Az előadás bemutatta, **hogyan alakította ki az SAP a saját Breach and Attack Simulation (BAS) alapjait**, különös tekintettel az **észlelési képességek** validálására, a **biztonsági kontrollok** értékelésére és a fejlett **támadási forgatókönyvek** tesztelésére.

Bemutatta az SAP módszertanát a központosított **Detection Lab** kiépítésére, a többlépcsős validációs keretrendszerüket, valamint azt, hogyan készülnek fel a különböző üzleti területekre történő kiterjesztésre. A kezdeti implementációból származó gyakorlati példákon keresztül betekintést nyerhettünk abba, **hogyan lehet skálázható biztonsági validációs programot építeni**, amely együtt tud fejlődni a vállalati igényekkel.

Megmutatták, hogyan közelítették meg az infrastruktúra kialakítását, a használati esetek prioritizálását és az érintett szereplők összehangolását annak érdekében, hogy stabil alapot teremtsenek **egy átfogó, vállalati szintű biztonsági validációs rendszerhez**.



A Phish 'n' Ships módszer feltárása: A módszer, amellyel több tízmillió dollárt loptak el online vásárlóktól

Az előadó betekintést nyújtott egy rendkívül komplex, összehangolt és üzletszerűen elkövetett, majd végső soron értékesített (**PaaS – Phishing as a Service**) csalási módszerbe. Felfedezte azt a kutatási metódust, amely során sikerült beazonosítani a csalást és betekintést nyújtott a további felderítés egy szeletébe. Az elkövetőcsoport metodológiája összehangoltan alkalmazott sérülékenységek kihasználásából eredő technikai megoldásokat, phishing és search engine „hacking/optimization” technikákat a csalás véghezviteléhez.

Az eredeti felfedezés egy ritka, niche termék keresése során látott napvilágot, a csoport kizárólag ritka, réteg termékekre automatizált módon hozott létre termék oldalakat több ezres nagyságrendben.

A létrehozott termék oldalak nagyon hasonlítottak legitim webshopok felületeire. Azonban a csalás kivitelezéséhez különböző sérülékenységeket kihasználva kompromittáltak tárhely szolgáltatók szervereit és az azokon futó webshopokat. Miután sikeresen bejutottak egy scripttel átirányították a termék listázásokat az általuk létrehozott hamis webshopok oldalaira, amelyen a pénzügyi adatainkat megadva már áldozatul is estünk ennek a csalásnak.

Az előadó által végzett kutatások során **több mint 1000 weboldalt fertőztek meg** és **több mint 121 hamis webshopot hoztak létre**, amely több tízmillió dolláros károkat okozott abban a megközelítőleg 5 évre kiterjedő időszakban, amíg a z IT biztonsági kutatók vissza tudták követni a csoportot.

Diego Matos Martins (IBM, BR): Az identitás őrzői: tanulságok személyazonosság-lopási incidensekből

Ez az előadás egy valós **Adversary-in-the-Middle (AiTM)** típusú támadás elemzését mutatja be, amely során **egy fenyegetési szereplő sikeresen megkerülte a többtényezős hitelesítést (MFA) egy globális vállalat Microsoft 365-fiókjánál**. A jogosulatlan hozzáférés után a támadó **Business Email Compromise (BEC) módszerrel** folytatta a támadást, majd másodlagos AiTM és BEC akciókat indított az eredeti áldozat partnerei és kontaktlistáján szereplő további célpontok ellen. Ezek a további „**terjeszkedő**” támadások különösen veszélyesek voltak, mivel a már kompromittált, legitim vállalati partnertől eredtek, így könnyebben **képesek voltak áthatolni a védelmi rendszereken**.

Az előadás során az alábbi témákat érintettük:

- A kiberbiztonsági világban a támadási fókusz közelmúltbeli eltolódásáról
- Az Adversary-in-the-Middle módszerről és annak alkalmazásáról
- A támadók által használt elkerülési és adatkinyerési technikákról
- Egy konkrét esettanulmányról
- Védekezési, észlelési és vizsgálati lehetőségekről



Stephan Berger (InfoGuard AG, CH): A Linux rootkitekről mélyrehatóan: fejlődésük, észlelés és védekezés

Az előadás átfogó képet nyújtott a **Linux rootkitek** fejlődéséről a kezdeti megjelenésüktől egészen a napjainkban ismert, kifinomult változatokig.

Az előadáson belementünk **rootkit technikák** működésébe, hatékony észlelési stratégiákba, valamint a védekezés jövőbeli lehetőségeibe. A történeti háttér, a jelenlegi módszerek és a feltörekvő fenyegetések bemutatásán keresztül az előadás tudást és eszközöket adott a **Linux rendszerek rootkit-támadások elleni védelméhez**.

Az előadás bevezetésként **ismertette a Linux rootkitek alapvető működését**, majd végigkövette a **rosszindulatú eszközök fejlődését** az egyszerű programoktól a ma már rendkívül kifinomult technikákig. A rootkiteket kategorizáltuk: kernel-szintű, felhasználói módú és hibrid típusokra, ismertetve ezek működését néztük meg. Ilyen működési módszer volt például a **kernelfunkciókra való „ráakaszkodás” (hooking)**, a **felhasználói folyamatok „elfogása”** és a rendszer ilyen módú kompromittálása, vagy a két módszer kombinációja. Kitértünk a rootkit által alkalmazott **rejtve maradási biztosító technikákra** és a **tartós jelenlét (persistencia) megvalósítására** is.

Ezt követően az előadás a **detekciós stratégiákra fókuszált**. Kezdeként bemutatta az aláírásalapú (signature) észlelést, kiemelve annak előnyeit és korlátait, majd rátért a rendszer viselkedéselemzésére, amely rendszeranomáliák figyelésével segít a rejtett rootkitek felfedésében.

Konkrét esettanulmányokkal demonstrálták a módszer gyakorlati hatékonyságát. Szó esett az **integritásellenőrzés** fontosságáról is, különös tekintettel a rendszerfájlok és binárisok megbízható összehasonlítására.

Az előadás végül bemutatta a fejlettebb észlelő eszközöket és keretrendszereket is, kiemelve a legelterjedtebb rootkit-detektáló eszközöket, gyakorlati példákon keresztül szemléltetve azok működését. Az előadás során nyújtott elemzés rámutatott a rootkit-fejlesztők és a kiberbiztonsági szakemberek közötti folyamatos technológiai versenyre, és hangsúlyozta az észlelési és elhárítási módszerek folyamatos fejlesztésének szükségességét. Összeségében egy hatalmas témát a rendelkezésre álló idő korlátai között igyekezett minél részletesebben bemutatni.

Az előadás prezentációja az alábbi [linken](#) érhető el.

Vanja Svajcer (Cisco Talos, HR): BYOVD és zsarolóvírusok

Az előadó rámutatott, hogy az utóbbi időben a **zsarolóvírus-csoportok** egyre gyakrabban **alkalmazzák** a „Hozd a saját sebezhető illesztőprogramodat” (**Bring Your Own Vulnerable Driver – BYOVD**) **technikát** – „de mit is jelent ez pontosan?” Felvetette azt a kérdést, hogy

- „Milyen típusú illesztőprogramok alkalmasak erre?”
- „Mely drivereket használják a támadók, és miért épp azokat?”
- „Milyen sebezhetőségeket használnak ki, és mi a céljuk ezek kihasználásával?”
- „Hogyan történik mindez, és kik az ismert fenyegetési szereplők, akik alkalmazzák?”

Az előadó elsődlegesen ezen kérdés mentén mutatta be a zsarolóvírusok ezen aspektusát.



Az előadás bevezetett a BYOVD technikába, és részletesen bemutatta, **hogyan használják ki a zsarolóvírus-csoportok az illesztőprogramok sebezhetőségeit**. Elsőként három fő sebezhetőségi kategóriát vizsgáltunk meg a régebbi Windows illesztőprogramok esetében, amelyeket a támadók gyakran kihasználnak: tetszőleges **MSR-írások** (Model Specific Register), tetszőleges **kernelmemória-írások**, valamint **elégtelen hozzáférés-ellenőrzés**. Ezek a sérülékenységek lehetővé teszik a támadók számára a jogosultságok kiterjesztését, nem aláírt kódok betöltését, EDR-szoftverek megkerülését, valamint egyéb, a végső kártékony tevékenységhez vezető lépések végrehajtását.

Az előadás második részében azon zsarolóvírus-csoportokra fókuszáltunk, amelyek BYOVD-t használnak támadásaikhoz – például a Kasseika, Akira, Qilin, BlackByte és RansomHub csoportokra. Megvizsgáltuk ezen csoportok 2024-es BYOVD-hez kapcsolódó taktikáit, technikáit és eljárásait (TTP-k).

Röviden áttekintettük azokat a Windows által bevezetett **exploitvédelmi mechanizmusokat** is, amelyek **célja a BYOVD típusú támadások megelőzése**. Ide tartozik például a virtualizáción alapuló biztonság (VBS), a hypervisorral védett kódintegritás ellenőrzés (HVCI), a kernel-szintű vezérlésáramlás-védelem (kCFG) és a kernel „shadow stack” funkció, amelyek mind jelentős szerepet játszanak a rendszer biztonságának erősítésében.

Végül kitérünk arra is, **hogyan észlelhetők és előzhetők meg a BYOVD technikára épülő támadások** a blue team (védekező oldal) nézőpontjából, adatforrásokkal és észlelési stratégiákkal alátámasztva.

Habár az előadás technikai és operatív oldalról egyaránt bemutatta a BYOVD működését, ezzel hasznos ismereteket nyújtva a digitális kriminalisztikai szakértőknek, incidenskezelőknek és kiberfenyegetés-kutatóknak, ugyanazzal a problémával küzdött, mint a korábbi rendkívül technikai előadások, hogy egy hatalmas és nagyon mély területet próbált bemutatni egy rendkívül rövid időkeretben.



Mahdi Alizadeh (Databricks, NL): Incidenskezelés Kubernetes környezetben

Az előadó a **Kubernetes** bemutatásával kezdte az előadást, amely a modern, éles IT üzemeltetési környezetek egyik alapvető építőkövévé vált, köszönhetően skálázhatóságának, rugalmasságának és a konténerorchesztráció egyszerűsítésében betöltött szerepének. (Lényegében egy nagy ipari felhasználásra fejlesztett rugalmas és skálázható konténerizációs keretrendszer).

Ugyanakkor összetettsége és dinamikus jellege **egyedi kihívásokat jelent a biztonsági incidensek kezelése során**. Egy feltört Kubernetes környezet komoly számítási kapacitást és hozzáférést biztosíthat a támadók számára, lehetőséget adva például adatlopásra, szellemi tulajdon eltulajdonítására vagy kriptovaluta bányászatra.

A Kubernetesben történő incidenskezelés **speciális tudást igényel**, mivel a hagyományos biztonsági gyakorlatok gyakran nem képesek megfelelően kezelni a konténerizált rendszerek sajátosságait. A konténerek mulandó jellege (bármikor a semmiből elindíthatunk egy konténert és ugyan ilyen könnyen le is állíthatjuk, vagy teljesen törölhetjük), a korlátozott naplózás és monitorozás, valamint az észlelőeszközök hiányosságai miatt különösen nehéz az incidensek felismerése, izolálása és kezelése. Sok biztonsági csapat nincs tisztában a Kubernetes-specifikus támadási vektorokkal, és hiányzik náluk az ilyen környezetekben szükséges szakértelem.

Az előadás először példákat mutatott be **Kubernetes támadási láncokra**, és olyan technikákat ismertetett, mint például a **jogosultságkiterjesztés „rosszindulatú podok”** segítségével, amelyek kifejezetten erre a környezetre jellemzőek. Ezután áttekintette azokat a kulcsfontosságú naplókat (logokat), amelyeket érdemes gyűjteni, és bemutatta, **hogyan segíthet a lemez- és memória forensic vizsgálata az incidenskezelésben**.





Kitért az egyes módszerek gyakorlati alkalmazására (mikor melyikhez érdemes nyúlni), mikor érdemes egyszerre többet alkalmazni és azok előnyeire és hátrányaira a védekező csapat szempontjából. Végül szó volt azokról a kihívásokról is, amelyekkel az elemző csapatok szembesülhetnek a vizsgálat során.

Az előadás prezentációja az alábbi [linken](#) érhető el.

Kirill Boychenko (Socket, US), Philipp Burckhardt (Socket, US): Biztonsági szakemberek útmutatója a rosszindulatú csomagokhoz

Az előadás bemutatta, **hogyan használják ki a támadók az elírásokat (typosquatting), a szerzők megszemélyesítését és az innovatív malware-kampányokat a szoftverellátási láncok megfertőzésére.** Az előadó kitért arra, hogy manapság már „szinte senki” nem készít 100%-ban eredeti kódot és az összes szoftver, szolgáltatás nagymértékben támaszkodik rengeteg, már korábban megírt, bevált és 3. fél által karbantartott (jellemzően open source) repositorykra. Az előadás során **gyakorlati fenyegetésfelderítési módszereket** ismerhettünk meg, valamint lépésről lépésre végigvezetett útmutatókat kaptunk arra vonatkozóan, hogyan lehet felismerni, elemezni és kivédeni ezeket a rosszindulatú fenyegetéseket. Az előadás érdekessége, hogy az egyik legtöbbet használt és **elsődleges támadási vektor ismét a felhasználó,** (esetünkben szoftver és szolgáltatásfejlesztők) „figyelmetlenségét” használja ki.

Arda Büyükkaya (EclecticIQ, NL): A Scattered Spider APT felhős taktikái: A zsarolóvírus-terjesztés életciklusának megértése

A mai **felhőalapú üzleti környezetben** a kibertámadók egyre gyakrabban célozzák meg a felhő infrastruktúrákat, hogy nagy hatású **zsarolóvírus-támadásokat** hajtsanak végre. Az előadás a **Scattered Spider** néven ismert fenyegetési szereplő **taktikáit, technikáit és eljárásait (TTP-k) vizsgálta**, különös tekintettel arra, hogyan zajlik náluk a zsarolóvírusok terjesztésének teljes életciklusa a felhőalapú környezetekben.

Részletes kutatásokra és valós esettanulmányokra támaszkodva – különösen a **biztosítási és pénzügyi szektor ellen irányuló támadásokból** – megtudhattuk, hogyan alkalmaz a Scattered Spider kifinomult **social engineering módszereket**, mint például a hangalapú adathalászat (vishing) és az SMS-alapú adathalászat (smishing), hogy magas jogosultságú fiókokat – például IT helpdesk vagy identitáskezelő adminokat – kompromittáljanak. Az előadás részletesen tárgyalta, hogyan használják a **SIM-cserét a többfaktoros hitelesítés (MFA) megkerülésére**, és hogyan szereznek **jogosulatlan hozzáférést kritikus felhőszolgáltatásokhoz és SaaS (Software as a Service) platformokhoz**.

Az előadás azt is bemutatta, **hogyan élnek vissza a támadók a legitim felhőfunkciókkal** – például a Microsoft Entra ID „Cross-Tenant Synchronization” funkciójával vagy az identitásszolgáltatókkal – a tartós hozzáférés és jogosultságkiterjesztés érdekében. Az előadásban szó volt arról is, **hogyan alkalmaznak nyílt forráskódú eszközöket a felhő-felderítéshez**, hogyan bénítják meg a védelmi megoldásokat, illetve milyen észlelés-elhárítási módszereik vannak – például távoli menedzsment eszközök (RMM), protokoll-alagút technikák vagy nem menedzselt virtuális gépek létrehozása.

Elemezésre kerültek a Scattered Spider zsarolóvírus-kampányait, különös tekintettel a felhőalapú IaaS (Infrastructure as a Service) környezetek – például a VMware ESXi – ellen irányuló stratégiákra.



Az előadás részletesen kitért az **automatizált terjesztési taktikákra** is és a **felhő-natív eszközök használatára a zsarolóvírus payload-ok hatékony futtatásának céljából**. Ezen technikák alkalmazása azért is különösen alattomos, mert jelentősen megnehezítik az áldozatok helyreállítási kísérleteit.

Az előadás végigvezetett egy **Scattered Spider támadás** életciklusán, a kezdeti fiókkompromittálástól a zsarolóvírus végrehajtásáig. A prezentáció végül megelőzési lehetőségekkel zárult: **best practiceket** ismerhettünk meg a hitelesítés, fiókvédelem, és a kompromittálódás észlelésére szolgáló lekérdezések terén.

Az előadás prezentációja az alábbi [linken](#) érhető el.

Emanuele Mezzi (Vrije Universiteit Amsterdam / Ethikon Institute, NL):

A nagy nyelvi modellek megbízhatatlanok a kiberhírszerzés területén: gyenge teljesítmény, következetlenség és rossz kalibráció CTI-feladatokban

Az előadó felvezetésében kitért arra, hogy több friss kutatás is azt állítja, hogy a **nagy nyelvi modellek** (Large Language Models – LLM-ek) segíthetnek kezelni a kiberbiztonságban használt hatalmas adatáradatot azáltal, hogy **javítják a kiberfenyegetettségi hírszerzés** (Cyber Threat Intelligence – CTI) **feladatainak automatizálását**.

Az előadásban olyan **értékelési módszertant** ismerhettünk meg, amely lehetővé teszi **LLM-ek tesztelését CTI-feladatokon „zero-shot”, „few-shot” és „finomhangolt” tanulás esetén**, valamint alkalmas a modellek következetességének és önbizalom-szintjének kvantitatív mérésére is.



A kutatás során három élvonalbeli LLM-et vizsgáltak meg egy 350 fenyegetési jelentésből álló adatkészleten, és új következtetéseket mutattak be arra vonatkozóan, hogy **milyen biztonsági kockázatot jelenthet az LLM-ekre támaszkodni CTI-munkafolyamatokban**.

A kutatás konklúziójaként megállapítható, hogy az **LLM-ek nem garantálnak elegendő teljesítményt** valós jelentések esetén, miközben **válaszaik következtelenek** és **indokolatlanul magabiztosak**. A few-shot tanulás és a finomhangolás ugyan részben javítja az eredményeket, de nem elégségesek ahhoz, hogy megalapozzák az LLM-ek használatát olyan CTI-szenáriókban, ahol nincs megfelelő, strukturált, címkézett adathalmaz.

Az előadás prezentációja az alábbi [linken](#) érhető el.

Régi idők csalása 2024-ben: Zsarolóvírus előkészítés zsarolóvírus nélkül

Az előadás két kiberbiztonsági kutató/elemző közös munkáját mutatta be, amely során egy ismeretlen csoportot követve derítették fel annak egy akcióját és végső soron a módszertanukat. A csoport a **Latrodectus nevű kártevőt** – annak minden modulját kihasználva – egy **kémkedésre épülő támadáshoz használta fel**, amely végül **pénzlopással zárult**. Megmutatták, hogyan segítheti a **kiberhírszerzés (Threat Intelligence)** az incidenskezelési folyamatok felgyorsítását, valamint más áldozatok és még aktív rosszindulatú infrastruktúrák azonosítását.

