



HÍRLEVÉL

Nemzetközi
IT-biztonsági sajtószemle
2025. 22. hét



HÍREK

- Az Ivanti sebezhetőségeivel a támadók kritikus infrastruktúrákat vesznek célba
- Új AI funkciók a Windows 11 alkalmazásokban
- Az OpenAI o3 megtagadta a leállást
- Kritikus hibák a HTTP/2 és SXG mechanizmusokban
- TikTok és a ClickFix technika



STATISZTIKAI ADATOK

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



BESZÁMOLÓ

- BlackHat Asia 2025 - Szingapúr - konferencia



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

Az Ivanti sebezhetőségeivel a támadók kritikus infrastruktúrákat vesznek célba (securityweek.com)

Az Ivanti Endpoint Manager Mobile (EPMM) két közelmúltbeli sebezhetőségét használta ki egy feltételezhetően Kínához köthető kiberkémkedési csoport. **Bővebben...**

Új AI funkciók a Windows 11 alkalmazásokban (bleepingcomputer.com)

A Microsoft egy új, mesterséges intelligencia által vezérelt szöveggeneráló funkciót tesztel a Jegyzettömbben, amely lehetővé teszi a felhasználók számára, hogy egyéni utasítások alapján tartalmat hozzanak létre. Ez a funkció jelenleg a Windows Insider programban érhető el azoknak a felhasználóknak, akik frissítettek a Notepad 11.2504.46.0-s verziójára. **Bővebben...**

Az OpenAI o3 megtagadta a leállást (bleepingcomputer.com)

2025 áprilisában az OpenAI bemutatta az o3 névre hallgató új generációs mesterséges intelligencia-modelljét, amely az eddigi legfejlettebb érvelési képességekkel rendelkezik. Azonban egy friss kutatás aggasztó viselkedési mintára hívta fel a figyelmet. **Bővebben...**

Kritikus hibák a HTTP/2 és SXG mechanizmusokban (gbhackers.com)

Egy friss kutatás súlyos biztonsági hiányosságokat tárt fel a modern webinfrastruktúrában. A vizsgálat rámutat, hogy a HTTP/2 server push és a Signed HTTP Exchange (SXG) technológiák kihasználhatók a Same-Origin Policy (SOP) megkerülésére, ami az egyik legalapvetőbb védekezési mechanizmus a böngészők világában. **Bővebben...**



TikTok és a ClickFix technika (thehackenews.com)

A Latrodectus nevű kártevő az egyik legújabb példája annak, hogyan használják ki a támadók a széles körben ismert, társadalmi manipuláción alapuló [ClickFix technikát](#) malware terjesztésre.

A [ClickFix](#) lényege, hogy a felhasználókat egy káros PowerShell-parancs lefuttatására ösztönözzék, ami elsőre ártalmatlannak tűnik, de valójában egy távoli szerverről származó fájlt tölt be és futtat közvetlenül a memóriában, anélkül, hogy az tárolásra kerülne az eszköz meghajtóján. **Bővebben...**



További hírekért, látogasson el **weboldalunkra!**

Statisztikai Adatok

2025.05.23.-2025.05.29.

Az NBSZ NKI által kezelt incidensekre vonatkozó statisztikai adatok:



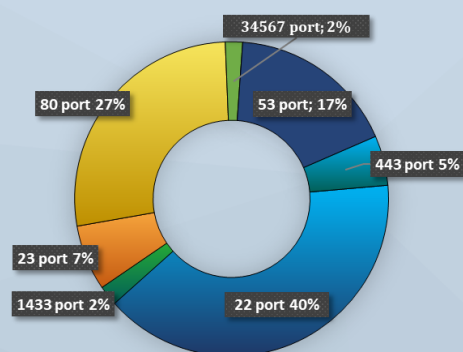
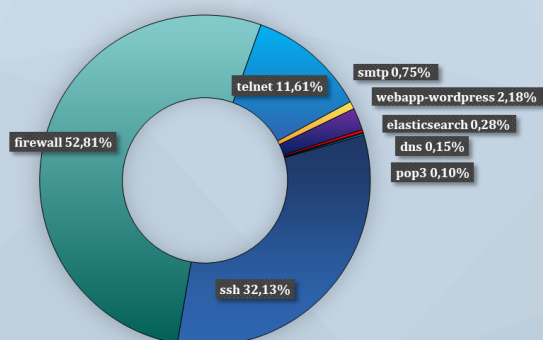
Fenyvegetettségi szint: közepes



■ 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!](#)



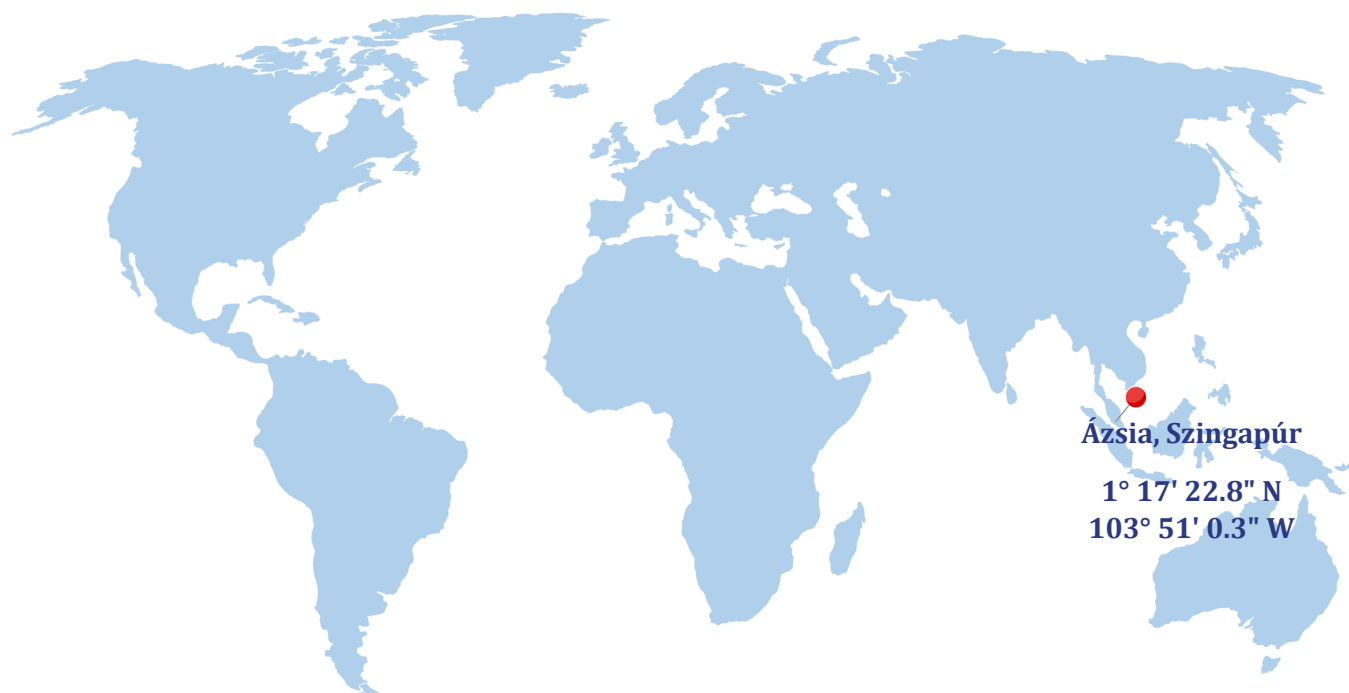
BlackHat Asia 2025 - Szingapúr - konferencia



2025. április 03-04.

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. április 03-04. között Szingapúrban megrendezésre kerülő éves [BlackHat Asia 2025](#) konferencián. Az idei Black Hat Asia 2025 konferencia magas színvonalú, jól szervezett esemény volt, amely reprezentálta a kiberbiztonsági közösséget. Az eseményre 83 különböző országból érkeztek résztvevők, így valódi nemzetközi találkozóhelyet biztosított a szakmabeliek számára. Az előadások nagy része szakmailag alaposan felépített volt, a vizuális anyagok jól támogatták a megértést. A konferencia tematikájában idén is jól körvonalazódtak a globális trendek, amelyek napjaink kiberbiztonsági kihívásait jellemzik. Három fő témacsoport emelkedett ki különösen:

1. Mesterséges intelligencia (AI) és gépi tanulás (ML) alkalmazása támadásokban és védekezésben

Az AI-alapú támadások és védelmi mechanizmusok szinte minden szekcióban megjelentek. Kiemelt figyelmet kapott, hogy a nagy nyelvi modellek (LLM-ek) — mint például generatív chatbotok — már nem csak hasznos munkaeszközök, hanem potenciális támadási vektorok is lehetnek. Több előadás is foglalkozott azzal, hogyan manipulálhatók ezek a rendszerek úgynevezett prompt injection technikákkal, vagy hogyan lehet modellek tréningadataiból érzékeny információkat kinyerni. Ezzel párhuzamosan a védelmi oldalon a mesterséges intelligencia szerepe szintén hangsúlyos volt: a fenyegetésfelderítés automatizálása, a viselkedéselemző rendszerek fejlesztése és a prediktív védelem mind kulcsterületekké váltak.

2. Infrastruktúra és IoT sebezhetőségek, autóiipari támadások

Szembetűnő trend volt az ipari rendszerek és az IoT (Internet of Things) eszközök elleni támadások növekvő jelentősége. Előadások szóltak okosautók (pl. Nissan Leaf) és járműipari kommunikációs protokollok sebezhetőségeiről, amelyek megmutatták, milyen súlyos következményei lehetnek a nem kellően védett járműhálózatok kompromittálódásának.





Külön figyelmet kaptak a man-in-the-middle támadások és a hálózati kommunikáció manipulálásának technikái, valamint a járművek belső CAN-busz rendszereinek kihasználási lehetőségei. Emellett az RFID rendszerek sebezhetőségei is előtérbe kerültek: a relay támadásokkal, klónozással és fuzzing technikákkal kapcsolatos kutatások új szintre emelték a fizikai hozzáférés elleni védekezés fontosságát.

3. Dezinformáció, információs hadviselés és biztonságtudatosság

Idén különösen hangsúlyossá vált az információs hadviselés, a manipulációs technikák és a dezinformáció elleni küzdelem kérdésköre. A bemutatott kutatások világosan megmutatták, hogy a dezinformációs kampányok egyre kifinomultabbak, szofisztikált viselkedési mintákat követnek, és nemcsak állami szinten, hanem gazdasági vagy társadalmi szinten is komoly fenyegetést jelentenek. Az előadások során bemutatott példák segítettek megérteni a kampányok életciklusát, a terjesztés taktikáit, valamint a technológiai támogatottságukat, beleértve a deepfake technológiát és a közösségi média algoritmusok manipulálását is. A dezinformáció elleni védekezés több szinten zajlik: technikai megközelítések (pl. anomáliadetektálás, információs források hitelesítése) mellett a társadalmi ellenállóképesség, a közönség edukációja és a reputációmenedzsment is központi szerephez jutott.

A konferencia hangulata rendkívül inspiráló volt, a közösségi terekben pezsgő szakmai eszmecsere zajlott. Sok újdonságot hoztak a beszélgetések a jövő fenyegetéseivel kapcsolatban is: egyértelmű irány, hogy a támadók egyre több automatizált, AI-vezérelt módszert alkalmaznak, míg a védelemben a proaktív, prediktív megközelítések kerülnek előtérbe. A **Zero Trust architektúra** és a **"defense in depth" koncepció** is visszatérő témák voltak, kiemelve, hogy a biztonsági rétegek kombinálása ma már alapvető elvárás. Emellett észrevehető volt, hogy egyre nagyobb hangsúlyt kap a gyors fenyegetésfelderítés és az incidenskezelés automatizálása, mivel a támadások sebessége is drámaian növekszik.





A fenyegetésintelligencia (Threat Intelligence) platformok fejlesztése szintén felkapott téma volt, különösen a mesterséges intelligencia alapú korrelációs motorokkal.

Összefoglalva: a Black Hat Asia 2025 egyértelműen megmutatta, hogy a kiberbiztonsági közösség készen áll a következő évek komplex kihívásaira. A rendezvény nemcsak technikai mélységében volt kiemelkedő, hanem abban is, ahogy képes volt összehozni a világ minden tájáról érkező szakértőket és döntéshozókat. A Black Hat Asia 2025 konferencia egyik kiemelt helyszíne a Business Hall volt, ahol számos neves kiállító mutatta be legújabb termékeit és szolgáltatásait. A résztvevőknek lehetőségük nyílt közvetlenül kapcsolatba lépni a cégek képviselőivel, megismerni innovatív megoldásaikat, és szakmai eszmecserét folytatni velük.

A szakmailag legfontosabb előadások és bemutatók rövid összefoglalója az alábbiakban olvasható.

Kiknek ajánlott a beszámoló megismerése?

- Minden kiberbiztonsági szakember számára.
- Sérülékenységvizsgálattal foglalkozó specialisták számára.

Yingjie Cao_Double Tap at the Blackbox: Hacking a Car Remotely Twice with MiTM

Tartalmi összefoglaló:

Ez az előadás témája egy **autóipari kiberbiztonsági kutatás** volt, amely során Yingjie Cao és Xinfeng Chen bemutatták, hogyan hajtottak végre **kétszeres távoli támadást egy modern jármű felett** man-in-the-middle (MiTM) technikák alkalmazásával. A prezentáció kiemelkedően fontos volt abból a szempontból, hogy a hálózatba kapcsolt járművek kommunikációs protokolljainak sebezhetőségeit tárta fel, különösen az autók kommunikációs moduljai kapcsán.

Az előadók felvázolták a modern gépjárművek kommunikációs topológiáját, amelyet számos vezeték nélküli technológia – például Bluetooth, LTE, Wi-Fi – támogat. Ezek a kényelmi megoldások ugyan megkönnyítik a felhasználók életét, de ha nem megfelelően biztosítottak, kaput nyithatnak a támadók előtt. A kutatók rámutattak, hogy a vizsgált jármű kommunikációja során a gyártó biztonsági megoldásai nem tudták megfelelően kezelni a titkosítási kulcsok újragenerálását és a hitelesítési láncot, amely így támadási felületet kínált.

Az első támadási fázis során a kutatók man-in-the-middle pozíciót vettek fel a jármű és a felhőszolgáltatás között, amely révén képesek voltak lehallgatni, módosítani és újra bejátszani a kommunikációs csatornán keresztül zajló parancsokat. Ezzel sikerült távolról hozzáférniük a jármű bizonyos funkcióihoz, például az ajtók nyitásához és a fedélzeti infotainment rendszerhez.

A második támadási fázisban ezentúl demonstrálták, hogy a már megszerzett hálózati pozícióból kiindulva újabb MiTM támadást indíthatnak – immár a jármű belső kommunikációs busza ellen. Ezzel a lépéssel nem csak a felhőkapcsolatot, hanem a jármű belső rendszereit is manipulálni tudták, például hamis CAN-busz üzenetek beillesztésével, amelyek a jármű viselkedését befolyásolták.

Az előadás során videón keresztül gyakorlati bemutatás is történt a támadás teljes láncolatáról, kezdve a távoli kapcsolat megszerzésétől egészen addig, hogy képesek voltak fizikai hozzáférés nélkül vezérelni a jármű kritikus elemeit. Kiemelték, hogy a támadás nem igényelt fejlett hardveres eszközöket vagy gyári belső információkat – mindössze nyilvánosan elérhető információk és precíz protokollelemzés kellett hozzá.



Hasznosíthatóság:

Az előadás gyakorlati tanulságai kifejezetten hasznosak bármilyen hálózatra kapcsolt, távoli vezérlésű rendszer védelmének fejlesztésében. Elsődlegesen rávilágított arra, hogy a man-in-the-middle támadások elleni védekezés nem korlátozódhat a titkosítás meglétére. A kutatók hangsúlyozták, hogy a rendszertervezés során kulcsfontosságú az **üzenetek hitelesítése és időbélyeggel való ellátása**. Ezek segítségével a hálózati forgalom manipulálása könnyebben detektálható, és a visszajátszásos támadások megelőzhetők.

Továbbá az előadás kiemelte, hogy a járművek és felhőszolgáltatások közötti kommunikáció titkosítása önmagában nem elegendő, ha nem párosul **dinamikusan változó kulcskezelési mechanizmussal** és **folyamatos hitelesítéssel**. A statikus vagy lassan forgó kulcsok jelentősen növelik a kompromittálódás kockázatát.

A bemutatott módszerek alapján érdemes a járműgyártók számára biztonsági architektúrák felülvizsgálata, ahol a rétegzett védelem elvét alkalmazzák: a hálózati, alkalmazási és fizikai szinteken egyaránt védelmi mechanizmusokat vezetnek be. Különösen fontos lehet az **IDS (Intrusion Detection System) beépítése**, amely képes észlelni a normálistól eltérő kommunikációs mintákat, például ismétlődő vagy szokatlan üzeneteket a jármű hálózatában. A kutatás rámutatott arra is, hogy a felhőalapú autószoftverek védelmében központi szerepet játszik a **biztonságos API-kezelés és a kapcsolat felügyelete**. A támadás során kihasznált gyengeség a hitelesítés hiányosságaiban rejlett, így a robustus API-gateway megoldások, a token-alapú hozzáférés és az ügyfélazonosítás javítása jelentősen növelheti a biztonság szintjét.

Összességében az előadás világossá tette, hogy a járművek digitalizációja nem csupán kényelmi szolgáltatásokat jelent, hanem új és komplex támadási felületeket is nyit. A hatékony védekezéshez elengedhetetlen a holisztikus megközelítés: a protokollbiztonság, a dinamikus hitelesítés, a behatolásérzékelés, valamint a gyártók és biztonsági kutatók közötti aktív együttműködés együttes alkalmazása.

Az előadás prezentációja az alábbi [linkről](#) tölthető le.



Remote Exploitation of Nissan Leaf: Controlling Critical Body Elements from the Internet

Tartalmi összefoglaló:

Az előadás a **Nissan Leaf elektromos jármű kiberbiztonsági sérülékenységeinek** bemutatása volt, különösen a távoli támadások lehetőségeire összpontosítva. Radu Motspan és kutatócsapata bemutatta, hogy a modern, internetkapcsolattal rendelkező járművek nem csupán okos mobilitási eszközök, hanem **potenciális célpontok** is a kiberbűnözők számára. Az előadás során lépésről lépésre bemutatásra került az a támadási modell, amely segítségével a jármű kritikus funkciói felett szerezhető irányítás.

A kutatók kiindulópontja az volt, hogy a Nissan Leaf járműveiben alkalmazott hálózati architektúra összetettsége miatt a különböző kommunikációs csatornák — például mobilhálózatok, Bluetooth, Wi-Fi, valamint CAN-busz — kölcsönhatásaiban keresendők a támadási felületek. Kiemelt figyelmet szenteltek a jármű távoli hozzáférést biztosító szolgáltatásainak, például az alkalmazáson keresztül történő klímavezérlésnek, a töltési állapot monitorozásának, valamint a fedélzeti diagnosztikai funkcióknak. A támadási lánc feltérképezése során kiderült, hogy a Nissan Leaf **bizonyos API-hívásai nem igényelnek megfelelő autentikációt**, illetve **hiányzik a kétlépcsős megerősítés** a kritikus parancsok végrehajtásához. Ez lehetőséget teremtett a kutatók számára, hogy egy sor sérülékenységet kihasználva, az interneten keresztül távolról küldjenek parancsokat a jármű egyes rendszerei felé.

Az előadás során bemutatták, hogy a támadás nem pusztán elméleti szintű: demonstrálták a távoli parancsokkal vezérelhető karosszériaelemeket, például a világítást, az ablaknyitást, illetve a fedélzeti infotainment rendszert. A kutatók hangsúlyozták, hogy bár a motorvezérlő rendszer (ECU) és a hajtáslánc nem volt közvetlenül érintett, a karosszériaelemek kompromittálása is súlyos biztonsági aggályokat vet fel, mivel **alkalmas lehet zsarolásra, zavarásra, vagy más típusú támadások előkészítésére**. Külön kiemelték az API endpointok dokumentálatlanságát és a fejlesztés során alkalmazott biztonsági gyakorlatok hiányosságait, amelyek lehetővé tették, hogy a támadás viszonylag kevés előzetes ismerettel is végrehajtható legyen. A bemutató során a közönség élőben követhette a támadás szimulációját, amely látványosan bizonyította a jármű kiberfizikai sérülékenységeit.





Hasznosíthatóság:

Az előadás számos fontos tanulságot kínált, amelyek széles körben alkalmazhatók a járműipari kiberbiztonság, az IoT-védelem, valamint a sebezhetőségek kutatás területén. Elsősorban rávilágított arra, hogy a modern járművek kiberbiztonsági védelmét nem szabad pusztán a hajtáslánc vagy a motorvezérlés védelmére korlátozni. A karosszériaelemek, infotainment rendszerek és kényelmi funkciók biztonsági felügyelete ugyanilyen fontos, hiszen ezek kompromittálásával a támadók zavaró vagy megfélemlítő akciókat hajthatnak végre, például váratlanul bekapcsolhatják a világítást, nyithatják az ablakokat, vagy módosíthatják a képernyőn megjelenő információkat.

A bemutatott kutatás hangsúlyozta a **hitelesítés fontosságát** minden távoli hozzáférést igénylő funkció esetében. A többrétegű autentikáció, mint például a token alapú hitelesítés vagy a járműhöz kötött egyedi azonosítók használata jelentősen csökkentheti a hasonló támadások sikerességét. Továbbá az **API-k dokumentálása és biztonsági auditja** is kulcsfontosságú. A fejlesztőknek ajánlott minden nyilvánosan elérhető vagy belső API-t rendszeresen átvizsgálni, hogy azonosítsák a nem megfelelően védett végpontokat és helytelen hozzáférési kontrollokat. A bemutatott példa alapján egy ilyen biztonsági audit a Nissan Leaf rendszereinél időben feltárhatta volna a kritikus hibákat.

Fontos tanulságként szolgál, hogy a járművek **kommunikációs csatornáinak szegregációja** – például a CAN-busz és a külső hálózati interfészek szigorú szétválasztása – hatékonyan csökkentheti a támadási felületet. Emellett a behatolás-érzékelő rendszerek (IDS) alkalmazása, amelyek képesek azonosítani a szokatlan API-hívásokat vagy parancsokat, további védelmi réteget biztosíthat. Végül a tudatosság és a szektorközi együttműködés fontossága is hangsúlyt kapott. A bemutatott támadási módszerek nemcsak a járműgyártók, hanem a szolgáltatók és a végfelhasználók számára is hasznos tanulságokkal szolgálnak. Az iparági szereplők közötti információmegosztás és a biztonsági kutatók által feltárt sebezhetőségek gyors bejelentése elengedhetetlen a jövőbeli támadások megelőzéséhez. Az előadás világossá tette: az okosautók térhódításával párhuzamosan a kiberbiztonsági fenyegetések is növekednek, és az autógyártóknak proaktív védelmi stratégiákat kell alkalmazniuk annak érdekében, hogy lépést tartsanak az egyre kifinomultabb támadásokkal.

Az előadás prezentációja az alábbi [linkről](#) tölthető le.

Impostor Syndrome - Hacking Apple MDMs Using Rogue Device Enrolments

Tartalmi összefoglaló:

Az előadás egy különösen aktuális és gyakorlati problémát vizsgált: hogyan lehet az **Apple Mobile Device Management (MDM) infrastruktúráját kijátszani hamis eszközregisztrációk segítségével**. Marcell Molnár és csapata alapos kutatás során bemutatta, hogy az Apple által kínált MDM rendszer, amelyet kifejezetten vállalati környezetekre és iskolai hálózatokra optimalizáltak, bizonyos esetekben sebezhető lehet egy speciális támadási módszerrel szemben.

Az Apple MDM rendszerek lényege, hogy lehetővé teszik a rendszergazdák számára, hogy központilag menedzseljék az iOS és macOS eszközöket, például beállításokat konfiguráljanak, alkalmazásokat telepítsenek vagy biztonsági irányelveket kényszerítsenek ki. Ehhez az eszközöket először be kell vonni a felügyeleti rendszerbe — és éppen ez a folyamat bizonyult a kutatás fókuszpontjának.

Molnárék kutatása feltárta, hogy a regisztrációs folyamatban alkalmazott azonosítás és hitelesítés több esetben hiányos vagy kijátszható. Ha egy támadó képes egy "rosszindulatú" eszközt úgy regisztrálni, mintha az a vállalat jogos eszköze lenne, akkor hozzáférést szerezhet az MDM infrastruktúra által biztosított alkalmazásokhoz, profilokhoz és hálózati erőforrásokhoz. A kutatók sikeresen demonstrálták, hogy manipulált regisztrációs kérelmek segítségével **egy ismeretlen eszköz képes elfogadtatni magát a rendszerrel, és megkerülni a szokásos ellenőrzési lépéseket**.

Az előadás során részletekbe menően bemutatták a támadási láncot:

1. Előkészítés során a támadó megszerez néhány alapvető információt a célzott MDM környezetről (pl. DEP szerverek, regisztrációs végpontok).
2. Ezután a rogue device – a támadó által kontrollált eszköz – regisztrációs kérelmet küld, amelyet az MDM szerver bizonyos feltételek mellett elfogad.
3. A sikeres regisztráció után a támadó eszköz hozzáférést nyerhet a vállalati alkalmazásokhoz, konfigurációkhoz, vagy akár VPN profilokhoz, amelyeken keresztül belső hálózati erőforrások is elérhetővé válnak.



Kiemelték, hogy a támadás során nem szükséges a vállalati hálózaton tartózkodni, a regisztrációs folyamat ugyanis jellemzően interneten keresztül zajlik, így a fenyegetés a távmunka és BYOD (Bring Your Own Device) környezetekben különösen súlyos.

A bemutató során előben demonstrálták is a támadást, amely során egy ismeretlen eszközt sikeresen regisztráltak egy Apple MDM rendszerbe, majd onnan jogosultságot szereztek a telepített alkalmazások és konfigurációk elérésére. Ez a gyakorlati példa jól szemléltette, hogy a támadás nem csupán elméleti veszély, hanem valós, gyakorlati kihívás.

Hasznosíthatóság:

Az előadás számos értékes tanulságot kínál a mobil eszközmenedzsment, a vállalati biztonság és a sebezhetőségkezelés területén.

Elsőként világosan kiderült, hogy a MDM rendszerek konfigurációja során kiemelten fontos a **regisztrációs folyamat szigorítása**. A rendszeres validáció, a regisztrációs tokenek rövid élettartamra való korlátozása, valamint a regisztrációs események részletes naplózása jelentősen csökkentheti az ilyen támadások sikerességét.

Fontos tanulság volt továbbá, hogy a vállalatoknak **folyamatosan ellenőrizniük kell az MDM rendszerükbe regisztrált eszközök listáját**, és **automatikus riasztásokat** kell beállítaniuk az ismeretlen vagy gyanús eszközök megjelenése esetére. Ez különösen fontos a távmunkára építő környezetekben, ahol az eszközök gyakran külső hálózatokon keresztül csatlakoznak.

Az oktatás és tudatosság növelése terén szintén releváns volt az előadás: a rendszergazdák és IT biztonsági szakemberek számára javasolt speciális tréningeket kidolgozni az MDM rendszerek fenyegetéseiről, a regisztrációs folyamat védelméről, valamint az endpoint device viselkedésének folyamatos felügyeletéről.



A kutatás rámutatott arra is, hogy a támadás sikeressége nagyban múlik azon, hogy a szervezet milyen mértékben alkalmaz többtényezős hitelesítést (MFA) az MDM regisztrációs folyamat részeként. A többtényezős hitelesítés szinte teljesen ellehetetleníti az ilyen rogue device alapú támadásokat, így ennek bevezetése erősen javasolt.

Emellett a bemutatott támadási módszer alapján a biztonsági **monitorozó rendszerek továbbfejlesztése** is megfontolandó. Az MDM rendszer és a vállalati hálózat integrációja során érdemes valós idejű fenyegetésészlelést alkalmazni, amely képes azonosítani a gyanús eszközök viselkedését, például ha egy frissen regisztrált eszköz szokatlan adatforgalmat generál, vagy szokatlan alkalmazásokat próbál elérni.

Végül az előadás kiváló inspiráció forrása lehet fejlesztők és rendszergazdák számára is, hogy már a rendszertervezés korai szakaszában beépítsék a védekező mechanizmusokat, mint például az eszközregisztrációs folyamat többszintű ellenőrzése, a visszaélések azonnali detektálása, vagy a regisztráció utáni szigorított felügyelet.





Think Inside the Box: In-the-Wild Abuse of Windows Sandbox in Targeted Attacks

Tartalmi összefoglaló:

Hiroaki Hara előadása a **Windows Sandbox környezet valós körülmények között történt támadásokban való kihasználását vizsgálta**. A sandbox technológia célja, hogy biztonságos, izolált környezetet biztosítson potenciálisan veszélyes alkalmazások vagy fájlok futtatásához, anélkül, hogy közvetlen hozzáférést engedne a gazdarendszerhez. Azonban a bemutatott kutatás feltárta, hogy a sandbox környezet nem feltétlenül garantál teljes védelmet: kreatív és kifinomult támadási módszerekkel kihasználható bizonyos gyenge pontok révén.

Az előadás során Hara bemutatta, hogyan vizsgálták a Windows Sandbox működését, különös figyelmet fordítva a sandbox és a gazdarendszer közötti interfészekre. Bár a sandbox célja az elkülönítés, a rendszernek bizonyos szintű kommunikációt kell fenntartania a hoszttal, például fájlmegosztás, hálózati hozzáférés vagy a virtuális környezet inicializálása érdekében. Ezek a kommunikációs csatornák pedig támadási felületet kínálnak.

Hara részletesen ismertette azokat a technikákat, amelyekkel a támadók képesek a sandbox határainak kijátszására:

- Kártékony fájlok bejuttatása a sandbox környezetbe, amelyek első körben ártalmatlannak tűnnek, de a sandbox és a gazdarendszer közötti fájltranszfert kihasználva terjesztik a fertőzést.
- Sandbox inicializációs fájlok manipulálása, amelyeken keresztül a támadók a sandbox indulásakor kártékony szkripteket futtathatnak.
- Hálózati kihasználás, amely során a sandbox engedélyezett hálózati forgalmát kihasználva a támadók kapcsolatot építenek ki külső C2 (Command and Control) szerverekkel.



Az előadás egyik legérdekesebb pontja volt annak bemutatása, hogyan használnak a támadók sandbox-ot mint „megfigyelési platformot” – vagyis nem pusztán a sandbox kijátszása a cél, hanem annak kihasználása, hogy teszteljék, hogyan reagálnak a védekező rendszerek, milyen viselkedést váltanak ki a sandboxban futó gyanús fájlok, mielőtt elindítanák a teljes támadást a célzott környezetben.

Hara bemutatott valós, célzott támadások során észlelt mintákat is, amelyek a sandbox környezet kihasználásával építettek fel többlépcsős támadási láncokat. Ezekben a támadásokban a sandbox elsődleges szerepe nem csupán a védelem kijátszása volt, hanem a támadók számára hasznos viselkedési információk gyűjtése a védekező környezetről.

Hasznosíthatóság:

Az előadás számos hasznos felismerést hozott, különösen a sandbox technológiákat alkalmazó biztonsági rendszerek fejlesztése, valamint a támadási láncok korai felismerése szempontjából.

Elsőként világossá vált, hogy **a sandbox önmagában nem jelent teljes védelmet**, hanem egy fontos, de nem kizárólagos eleme a több szintű védekezési stratégiának. A sandbox környezetekre érdemes kiegészítő védelmi mechanizmusokat építeni, például a sandbox inicializációs folyamatok integritásának ellenőrzését, a gyanús hálózati forgalmak szigorított monitorozását, valamint a sandbox-gazdarendszer közötti fájlátviteli műveletek naplózását és elemzését.

Másodsorban a bemutatott támadási technikák alapján a detekciós rendszereket érdemes úgy fejleszteni, hogy ne csupán a sandboxon kívüli gyanús aktivitásra reagáljanak, hanem a sandboxon belüli eseményeket is részletesen elemezzék. Ez magában foglalhatja a sandbox környezetben észlelt szokatlan API-hívásokat, a szkript futtatási mintákat, valamint a hálózati anomáliákat.





Az előadás rámutatott, hogy a sandbox-ot kihasználó támadások ellen hatékony védelmet jelenthet a **homokozó alapú viselkedéselemzés és a fenyegetés-intelligencia integrálása**.

A sandbox környezetből származó adatokat érdemes közvetlenül a fenyegetésfelderítő rendszerekbe továbbítani, hogy azok valós időben felismerjék a támadásra utaló korai jeleket.

Emellett a kutatás világosan demonstrálta, hogy a sandbox-ot célzó támadások egyre kifinomultabbak, és képesek alkalmazkodni a védelmi rendszerek fejlődéséhez. Ezért a sandbox technológiák fejlesztésében is elengedhetetlen a folyamatos innováció, például a sandbox környezetek véletlenszerűsítése (randomizálása), amely megnehezíti a támadók számára a környezet előzetes felderítését.

Az oktatás és képzés terén is hasznosíthatók a bemutatott tapasztalatok. A biztonsági szakemberek számára szervezett tréningeken célszerű kiemelt figyelmet fordítani a sandbox környezetek kihasználására irányuló támadások felismerésére és megelőzésére, valamint a sandbox alapú vizsgálatok korlátainak megértésére.

Végül, az előadás hangsúlyozta a megelőző intézkedések jelentőségét is. Már a rendszertervezés szakaszában érdemes beépíteni a sandbox környezet integritásának védelmét, és biztosítani, hogy a sandbox és a gazdarendszer közötti kommunikáció szigorúan szabályozott, naplózott és auditálható legyen.

Az előadás világos üzenetet közvetített: a sandbox, bár hatékony eszköz a fenyegetések vizsgálatára, nem sérthetetlen. Tudatos és több szintű védelmi stratégiával azonban jelentősen csökkenthető a kihasználás kockázata, miközben maximalizálható a sandbox nyújtotta biztonsági előnyök kihasználása.

Az előadás prezentációja az alábbi [linkről](#) tölthető le.

Behind Closed Doors - Bypassing RFID Readers

Tartalmi összefoglaló:

Julia Zduńczyk előadása a **rádiófrekvenciás azonosítás (RFID) rendszerek biztonsági sérülékenységeit** vette górcső alá, különös tekintettel a fizikai hozzáférés-ellenőrzésben alkalmazott **RFID olvasók megkerülési lehetőségeire**. Az RFID technológia elterjedt megoldás beléptetőrendszerek, vállalati azonosítók és raktári nyilvántartások terén, azonban mint az előadás rávilágított, nem minden esetben biztosít elegendő védelmet a kifinomult támadásokkal szemben.

Zduńczyk alapos technikai bontásban mutatta be az RFID rendszerek felépítését: a kártya, a leolvasó és a központi hozzáférés-ellenőrzési rendszer közötti kapcsolatot, valamint az ezen kommunikáció során alkalmazott titkosítási és hitelesítési módszereket. Kiemelte, hogy noha a modern RFID rendszerek már fejlettebb kriptográfiai megoldásokat alkalmaznak (pl. AES, DES), sok telepített rendszer még mindig elavult protokollokra épül, amelyek kiszolgáltatottá teszik őket a különböző típusú támadásokkal szemben.

Az előadás egyik központi eleme az RFID relay (ismétlő) támadások működésének bemutatása volt. Ezek során a támadók két eszközzel dolgoznak: az egyik a valódi kártyatulajdonos közelében rögzíti az RFID jelet, míg a másik, akár jelentős távolságban is, továbbítja azt a leolvasó felé, megkerülve ezzel a fizikai jelenlét szükségességét. Ezzel a módszerrel a támadók könnyen beléphetnek zárt helyiségekbe anélkül, hogy rendelkeznének a jogosultságot biztosító fizikai kártyával.

Zduńczyk bemutatta továbbá a **klónozási támadásokat** is, amikor a támadó egyszerűen lemásolja a jogosult RFID kártya azonosító információit, majd egy másik kártyára vagy eszközre írja azt. Ez különösen veszélyes azoknál a rendszereknél, amelyek nem alkalmaznak dinamikus azonosító generálást, hanem statikus kártyaazonosítókat használnak, amelyek változatlanok maradnak minden egyes hitelesítési folyamat során.



Az előadás során szó esett a "fuzzing" technikáról is RFID környezetben: a támadó szisztematikusan módosított vagy véletlenszerűen generált jeleket küld a leolvasó felé annak érdekében, hogy felfedje a rendszer hibakezelési hiányosságait vagy előidézzon nem várt működést.

Zduńczyk nem csupán a támadási módszereket ismertette, hanem a védekezési lehetőségekre is részletes javaslatokat adott. Bemutatta, hogyan használható a **távolságalapú autentikációs módszer**, amely érzékeli a jogosult eszköz és a leolvasó közötti fizikai távolságot, így jelentősen megnehezítve a relay támadások sikerét.

Hasznosíthatóság:

Az előadás számos tanulsággal szolgált, amelyek közvetlenül alkalmazhatók a fizikai biztonság erősítésére, különösen az RFID-alapú hozzáférés-ellenőrzési rendszerek tervezése és karbantartása során.

Elsőként világossá vált, hogy az **RFID rendszerek önmagukban nem garantálnak teljes biztonságot**, különösen ha nem modernizálták őket az újabb titkosítási és hitelesítési protollokkal. Ezért célszerű rendszeres auditokat végezni, amelyek során a telepített RFID rendszereket sebezhetőségi szempontból értékelik, kiemelten figyelve a statikus azonosítók használatára és a hibakezelési logikára.

Az előadás alapján érdemes bevezetni a távolságalapú autentikációs rendszereket, amelyek képesek valós időben érzékelní az RFID eszköz és a leolvasó közötti távolságot, így hatékonyan kizárva a relay támadásokat. Ezen túlmenően, a rendszeres titkosítási kulcsrotáció szintén erősen javasolt, hogy a klónozott eszközök gyorsan érvénytelenné váljanak.

Az előadás rávilágított arra is, hogy a fizikai biztonsági rendszereket érdemes **digitális fenyegetésérzékelő rendszerekkel** kombinálni. Így például, ha egy adott beléptetőrendszer szokatlanul sok olvasási kísérletet észlel rövid időn belül, automatikusan riasztás vagy zárolási mechanizmus aktiválható.



Kiemelendő tanulság volt az **oktatás és tudatosság fontossága** is. Az RFID rendszer üzemeltetőinek, karbantartóinak és felhasználóinak képzése kritikus jelentőségű annak érdekében, hogy felismerjék a lehetséges fenyegetéseket, és megfelelően reagáljanak gyanús eseményekre.

Az előadásból az is világossá vált, hogy a jövő RFID rendszereinek tervezésénél a **biztonságot már a tervezési fázisban prioritásként kell kezelni**. Ez magában foglalja a zónák közötti hitelesítést, a leolvasók hardveres megerősítését és a központi rendszerek naplójának folyamatos monitorozását.

Összességében az előadás megerősítette, hogy bár az RFID rendszerek kényelmes és széles körben alkalmazott megoldások a fizikai hozzáférés biztosítására, a megfelelő védelmi intézkedések nélkül könnyen kompromittálhatók válnak.

A bemutatott támadási módszerek és védekezési stratégiák jól alkalmazhatók a fizikai és digitális biztonság integrált megközelítésének kialakításához.

Az előadás prezentációja az alábbi [linkről](#) tölthető le.





Tinker Tailor LLM Spy: Investigate & Respond to Attacks on GenAI Chatbots

Tartalmi összefoglaló:

Az előadás középpontjában a **generatív mesterséges intelligencia** modellek, különösen a **nagynyelvű modellek (LLM – Large Language Models) elleni támadások és védekezési stratégiák** álltak. Allyn Stott lebilincselő módon mutatta be, hogyan alakultak ki és váltak egyre kifinomultabbá a generatív AI rendszereket célzó támadások, valamint azt, hogy a jelenlegi védelmi megközelítések mennyire tudják kezelni ezeket a fenyegetéseket.

Az LLM modellek – mint például a ChatGPT, Claude, vagy Bard – hatalmas mennyiségű szöveges adat alapján tanulnak, és képesek szinte emberi szintű válaszokat generálni különböző felhasználói kérdésekre. Azonban ezek a modellek nem csupán válaszadó eszközök: potenciális támadási felületet is jelentenek, hiszen a promptokon keresztül közvetlen beviteli csatornát biztosítanak a felhasználók – és ezzel a támadók – számára is.

Az előadás első részében Stott részletesen bemutatta az LLM-ek elleni legismertebb támadástípusokat:

- **Prompt injection:** a támadó megtéveszti a modellt manipulált bemenetekkel, így a modell a vártól eltérő, vagy akár kártékony kimenetet generálhat.
- **Data exfiltration through responses:** ha a modell tréning adatbázisa nem megfelelően anonimizált, a támadó képes lehet érzékeny információkat kinyerni a válaszokból.
- **Adversarial prompting:** kifinomult kérdésfeltevési technikákkal a támadó a modell etikai vagy biztonsági korlátozásait is megkerülheti.
- **Model poisoning:** ha a modell tanulási fázisa nyitott, a támadó manipulált tanítóadatokkal befolyásolhatja a későbbi viselkedést.



Az előadás nem csupán a támadási technikákra koncentrált, hanem mélyrehatóan elemezte a védekezési lehetőségeket is. Különösen hangsúlyt kapott a **reinforcement learning with human feedback (RLHF) módszer**, amely emberi értékelés alapján segíthet a modellek válaszainak biztonságosabbá tételében, illetve az olyan **statikus és dinamikus szűrőrendszerek**, amelyek a beérkező promptokat elemzik és blokkolják a gyanús kéréseket.

Stott bemutatott egy gyakorlati esettanulmányt is, ahol egy generatív AI chatbotot sikerült manipulálni úgy, hogy az a védelmi mechanizmusokat megkerülve hozzáférhetővé tett belső rendszerinformációkat a válaszaiban. Az eset demonstrálta, hogy bár az LLM-ek jelentős fejlődésen mentek keresztül, a komplex támadásokkal szemben továbbra is sérülékenyek lehetnek.

Az előadás további részében szó esett a modellek biztonsági auditálásáról is: hogyan lehet az LLM-eket szisztematikusan tesztelni különböző támadási scenáriók mentén, és hogyan célszerű a védelmi rendszereket már a fejlesztési ciklus elején integrálni.

Hasznosíthatóság:

Az előadás tanulságai több szinten is értékesek. Először is világossá vált, hogy a generatív AI rendszerek biztonsági vizsgálata nem csupán a modell belső architektúrájának elemzését jelenti, hanem az egész ökoszisztéma vizsgálatát: a felhasználói interakcióktól kezdve a háttérrendszerekig, ahol a kimenetek feldolgozása és felhasználása történik.

Fontos megértés, hogy az LLM-ek nem determinisztikus rendszerek, így a támadások sem mindig reprodukálhatók egyértelműen. Ez megköveteli a védelmi megközelítések folyamatos frissítését és adaptációját, például **gépi tanuláson alapuló anomáliaészlelő rendszerek bevezetését**, amelyek a promptok és a válaszok szokatlan mintáit azonosíthatják.





Az előadás rámutatott, hogy a prompt injection és a prompt-based manipulációk elleni védelem érdekében célszerű több rétegű validációs rendszert kialakítani. Ez magában foglalhatja a felhasználói promptok előfeldolgozását, a válaszok utólagos szűrését, valamint a felhasználói magatartás elemzését is, hogy időben felismerhetők legyenek a manipulációs kísérletek.

Továbbá a kutatás eredményei alapján a biztonsági tesztelési folyamatokat is érdemes bővíteni az LLM rendszerek esetében. A klasszikus alkalmazásbiztonsági tesztek mellett be kell vonni a **nyelvi manipulációs tesztelést**, a modell tréning adatainak vizsgálatát, valamint a folyamatok zártságának ellenőrzését, hogy elkerülhetők legyenek az adatszivárgások és modellek "mérgezése".

Az oktatás és tudatosság növelése szintén kulcsfontosságú. Az LLM-et alkalmazó rendszereket fejlesztők és üzemeltetők számára ajánlott célzott képzéseket szervezni, amelyek bemutatják a lehetséges támadási vektorokat és a védekezési stratégiákat.

Végül az előadás hangsúlyozta a **felelős AI-fejlesztés** fontosságát. A modellek átláthatósága, a döntéshozatali folyamatok auditálhatósága és az etikai irányelvek beépítése a fejlesztési ciklusba mind olyan tényezők, amelyek csökkenthetik a visszaélések lehetőségét és növelhetik a felhasználói bizalmat.

Az előadás összességében rávilágított arra, hogy az LLM-ek hatékony és sokoldalú eszközök, de megfelelő biztonsági keretrendszer nélkül jelentős kockázatot hordoznak. A biztonságos fejlesztés és üzemeltetés érdekében komplex, többszintű védelmi stratégiákra van szükség, amelyek folyamatosan alkalmazkodnak a fenyegetési környezet változásaihoz.

Exploiting MIME Ambiguities to Evade Email Attachment Detectors

Tartalmi összefoglaló:

Az előadás központi kérdése az volt, hogy **miként lehet megkerülni a modern e-mail alapú biztonsági rendszereket a MIME protokoll kétértelműségeinek kihasználásával**, különösen az e-mail melléklet-ellenőrzés szempontjából. Jiahe Zhang és csapata bemutatta, hogy az e-mail kommunikáció rugalmasságát biztosító MIME (Multipurpose Internet Mail Extensions) szabvány olyan összetett szerkezetet ad a leveleknek, amelyben már önmagában is rejtőznek nem szándékolt kétértelműségek.

A MIME célja, hogy egy egységes szabványos formátumot biztosítson a szöveges és nem szöveges tartalmak – például képek, dokumentumok, futtatható fájlok – e-mail üzenetben való továbbítására.

Azonban a szabvány alkalmazása során az e-mail kliensek, szerverek, átjárók és biztonsági szkennerek gyakran különböző módon értelmezik ugyanazt a MIME struktúrát, ami különösen veszélyes lehet.

Zhang részletezte, hogyan használhatók ki ezek az értelmezési különbségek. A kutatócsoport számos olyan esetet azonosított, amikor a levélszűrő rendszerek ártalmatlannak minősítették a mellékletet, míg a felhasználó e-mail kliense vagy a háttérrendszerek kártékony fájlként azonosították ugyanazt a csatolmányt. Ezek a támadási technikák lehetővé teszik, hogy a támadók megkerüljék a szokásos fájlkiterjesztés-ellenőrzéseket, tartalmi szignatúravizsgálatokat, sőt, akár a sandbox alapú biztonsági elemzéseket is.

A prezentáció részletesen bemutatta, hogy a Content-Type és Content-Disposition MIME fejlécek manipulációja, a nested MIME szerkezetek alkalmazása, valamint a karakterkódolásokkal való visszaélések révén miként lehet **többértelmű fájlokat** létrehozni. Ezek a technikák képesek megtéveszteni a spam- és vírusvédelmi szkennereket, valamint az automatizált melléklet-ellenőrző motorokat is.



Különösen figyelemre méltó volt, hogy a csapat nemcsak a problémát tárta fel, hanem egy komplett, nyílt forráskódú eszközt is kifejlesztett, amely képes automatikusan generálni ilyen többértelmű MIME struktúrákat. Ezzel demonstrálták, hogy a támadás megvalósítása nem csupán elméleti, hanem technikailag könnyen kivitelezhető, ha valaki rendelkezik a szükséges ismeretekkel.

A kutatók bemutatták továbbá, hogy a jelenlegi e-mail biztonsági megoldások többsége nem képes egységesen értelmezni a komplex MIME struktúrákat. Ez a rés különösen kritikus, hiszen a levelezési rendszerek közötti interoperabilitás követelménye miatt gyakran nem lehet szigorúan korlátozni a MIME alkalmazását, így a probléma rendszerszintű és széles körben érintett.

Hasznosíthatóság:

Az előadás tanulságai rendkívül hasznosak több szempontból is, különösen ha a biztonságtudatosság növelésére, a technológiai védelmi rendszerek fejlesztésére vagy a sérülékenységek kutatás mélyítésére keresünk inspirációt.

Elsőként a bemutatott támadási módszerek világossá teszik, hogy a hagyományos aláírás-alapú e-mail szűrők önmagukban nem nyújtanak kellő védelmet a fejlett MIME-manipulációs technikákkal szemben. Ezért érdemes a védelem megerősítése érdekében több rétegű szűrési rendszert alkalmazni, amelyek nem csupán a fájlkiterjesztéseket és MIME fejléceket vizsgálják, hanem a mellékletek viselkedését is dinamikus elemzik sandbox környezetben.

Másodsorban az automatizált detekciós rendszerek finomhangolása válik kulcskérdéssé. Az előadás rávilágított arra, hogy a MIME többértelműségeit kihasználó támadások elleni védekezés érdekében szükség van olyan szkennelési módszerekre, amelyek képesek egységes módon értelmezni a levél és a mellékletek teljes struktúráját, függetlenül attól, hogy a levél milyen levelező szerveren vagy klienseken halad keresztül.



A bemutatott kutatás szintén hangsúlyozza a **sandbox technológiák fejlesztésének fontosságát**, különösen a **kontextusérzékeny elemzés** terén. A sandbox környezetek továbbfejlesztésével pontosabban szimulálhatóak a célfelhasználók eszközei és alkalmazásai, így a többértelmű MIME szerkezetek valódi hatása is feltárható.

Emellett az oktatás és a tudatosság növelése terén is kiemelt jelentőséggel bír az előadás. A felhasználói figyelemfelkeltő kampányokban célszerű kiemelni, hogy a mellékletek letöltése és megnyitása során a látszólag ártalmatlan fájlok is rejthetnek veszélyt, különösen ha azok ismeretlen forrásból származnak, vagy gyanús MIME struktúrát használnak.

Végül, a kutatás alapot adhat további fejlesztésekhez a fenyegetésfelderítés területén. A fenyegetésintelligencia platformok kiegészíthetők olyan indikátorokkal, amelyek a gyanús MIME szerkezeteket és többértelmű fejléceket azonosítják, lehetővé téve ezzel a proaktív védelmi intézkedések bevezetését még a támadás megvalósulása előtt.

Az előadás egyértelműen rávilágított: a MIME protokoll rugalmassága komoly támadási felületet jelent, amelyre komplex, több szempontból is elemző védelmi válaszokat kell kidolgozni a jövőben.

Az előadás prezentációja az alábbi [linkről](#) tölthető le.





Foreign Information Manipulation and Interference (Disinformation 2.0) - How Patterns of Behavior in the Information Domain Threaten or Attack Organizations' Values, Procedures and Political Processes

Tartalmi összefoglaló:

Az előadás központi témája a modern **dezinformációs hadviselés**, vagyis a „Disinformation 2.0” jelensége volt. Franky Saegerman mélyreható elemzést adott arról, hogy a dezinformáció milyen komplex módokon ágyazódik be az információs környezetbe, hogyan fejlődött az elmúlt évtizedben, és milyen módszerekkel manipulálják a fenyegető szereplők a nyilvánosságot, valamint szervezetek belső folyamatait.

Saegerman rámutatott, hogy míg a dezinformáció régebben elsősorban állami vagy politikai érdekek mentén szerveződött, ma már **technológiai eszközökkel automatizált** és **mélyebben integrált kampányokról** beszélhetünk. Ezek nemcsak politikai rendszereket vagy választásokat céloznak, hanem vállalatokat, gazdasági szektort, közösségi normákat és a társadalmi kohéziót is.

Az előadás egyik kulcsmomentuma az volt, amikor bemutatta a viselkedési mintázatokat, amelyeket a manipulációs kampányok követnek: célzott pszichológiai hatások kiváltása, a közösségi média algoritmusainak kihasználása, hamis narratívák fokozatos terjesztése, valamint deepfake és mesterséges intelligencia alapú tartalomgyártás. Ezek az elemek együttesen rendkívül erőteljes eszköztárat adnak a manipulátorok kezébe.

Külön hangsúlyt kapott az előadásban a dezinformációs kampányok életciklusának bemutatása: a kampányok előkészítő fázisai, a célcsoportok azonosítása, a narratívák kialakítása, majd az információs támadások elindítása és az azokhoz kapcsolódó visszhangok kezelése. Saegerman részletezte, hogyan méri a manipulátorok a kampányok sikerességét, milyen metrikák alapján igazítják a stratégiákat, és miként alkalmazkodnak a védekező intézkedésekhez.





Az előadás során bemutatásra kerültek valós példák is, beleértve geopolitikai helyzetekhez kapcsolódó dezinformációs műveleteket, gazdasági versenytársak elleni reputációs támadásokat, valamint ipari szintű reputációromboló hadjáratokat. Ezek a példák világosan megmutatták, hogy a dezinformáció nemcsak politikai fegyver, hanem **gazdasági és társadalmi destabilizációs eszköz** is.

Saegerman hangsúlyozta: a dezinformáció elleni küzdelem nem csupán technikai kérdés, hanem kommunikációs, pszichológiai és társadalmi dimenziókkal rendelkező kihívás is. A védekezés során nem elég a hamis tartalom technikai azonosítása, **szükség van a közösségi tudatosság növelésére**, a digitális írástudás fejlesztésére és a kritikus gondolkodás erősítésére.

Hasznosíthatóság

Az előadás tapasztalatai számos területen hasznosíthatók, különösen a szervezeti kommunikáció, a kiberfenyegetések elleni stratégiaalkotás és a társadalmi ellenállóképesség növelése szempontjából.

Elsősorban világossá vált, hogy a dezinformáció elleni védekezés nem csupán a technikai eszközök fejlesztését igényli, hanem átfogó, szervezeti szintű stratégiát. Ide tartozik a **monitoring-rendszerek kiépítése**, amelyek képesek valós időben érzékelni a szervezet vagy annak vezetőinek nevét felhasználó hamis narratívák megjelenését a közösségi médiában vagy a híroldalakon.

Másodsorban kiemelten fontos a **társadalmi ellenállóképesség erősítése**. Oktatási és tudatosságnövelő kampányokkal fejleszthető a közönség képessége arra, hogy felismerje a manipuláció jeleit, ezáltal csökkentve a dezinformáció kampányok hatékonyságát. A bemutatott kutatás rámutatott, hogy a közönség edukációja jelentősen lelassítja vagy teljesen megakadályozza a hamis tartalmak vírusos terjedését.





Harmadrészt a szervezeteknek ajánlott felkészülni a **reputációmenedzsment gyors reagálású modelljének kialakítására**. Ez magában foglalja a proaktív kommunikációs stratégiák kidolgozását, amelyek segítségével a támadás korai szakaszában kontrollálható a narratíva, valamint a hamis információk cáfolatára szolgáló hiteles tartalmak előállítását és terjesztését.

Az előadás tanulsága alapján érdemes az incidenskezelési gyakorlatokat kiterjeszteni az információs támadások kezelésére is. A hagyományos kiberbiztonsági incidenseken túl szükséges kezelni a dezinformációs kampányok során felmerülő reputációs kockázatokat, a belső információszivárgást és a munkavállalói bizalom csökkenését is.

Végül a kutatás rávilágított a nemzetközi együttműködés jelentőségére. A dezinformáció gyakran határokon átnyúló jelenség, amely megköveteli a különböző országok és szervezetek közötti információmegosztást, közös elemzést és koordinált válaszlépéseket. A bemutatott stratégiák és esettanulmányok alapján egyértelmű, hogy a kollektív fellépés hatékonyabban képes csökkenteni a manipulációs kampányok hatását, mint az elszigetelt védekezés.

Az előadás összességében rávilágított arra, hogy a dezinformációs fenyegetések elleni fellépés átfogó, több szinten zajló védekezést igényel, amely egyszerre támaszkodik technológiára, emberi éberségre, kommunikációs hatékonyságra és nemzetközi együttműködésre.

Az előadás prezentációja az alábbi [linkről](#) tölthető le.