



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

25-11-2025:

APT35-lek onthult gedetailleerde aanvalsmethoden en doelwitten van Iran's Charming Kitten

In oktober 2025 lekte een aanzienlijke hoeveelheid interne documenten van de cyberaanvalsgroep APT35, ook wel bekend als Charming Kitten, een van de cyber-eenheden binnen de Iraanse Islamitische Revolutionaire Garde. Deze documenten onthullen gedetailleerde informatie over de georganiseerde en gestructureerde manier waarop de groep cyber-espionage uitvoert, gericht op overheden en bedrijven in het Midden-Oosten en Azië.

De gelekte documenten bevatten duizenden pagina's met rapporten, technische handleidingen en operationele verslagen die een duidelijk beeld geven van de werkwijze van APT35. De documenten onthullen dat de groep opereert als een traditioneel militaire eenheid, compleet met prestatiebeoordelingen, takenlijsten en rapporten die naar supervisors worden gestuurd. Dit toont aan dat APT35 veel meer een professionele organisatie is dan een informele hacker-groep.

De operaties richten zich voornamelijk op Microsoft Exchange-servers, waarbij de groep gebruikmaakt van ProxyShell-exploitaties en kwetsbaarheden in de Autodiscover- en EWS-diensten om gegevens uit Global Address Lists te extraheren. Deze gegevens worden vervolgens ingezet voor gerichte phishing-aanvallen, die worden uitgevoerd om inloggegevens te stelen. Na het verkrijgen van toegang, gebruikt APT35 geavanceerde op maat gemaakte tools, zoals technieken die vergelijkbaar zijn met Mimikatz, om verdere gegevens uit de systeemgeheugen van de slachtoffers te stelen. De gestolen informatie stelt de aanvallers in staat om lateraal door netwerken te bewegen en langdurige toegang te behouden.

De documenten geven ook inzicht in de geografische reikwijdte van de aanvallen. De doelwitten omvatten overheidsinstanties, telecombedrijven, douaneagentschappen en energiebedrijven in landen zoals Turkije, Libanon, Koeweit, Saoedi-Arabië, Zuid-Korea en Iran. De operationele documenten bevatten gedetailleerde lijsten van doelwitten, inclusief aantekeningen over geslaagde en mislukte aanvallen, evenals webshell-paden die werden gebruikt om toegang te behouden.

Deze gelekte gegevens benadrukken de strategische focus van APT35, die zich richt op het verkrijgen van diplomatieke communicatie, telecominfrastructuur en toegang tot de energie-industrie. Deze informatie wordt vermoedelijk gebruikt voor geopolitieke onderhandelingen en risicoanalyse. De mate van organisatie en de



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

bureaucratische structuur van de groep tonen de ernst van deze cyberoperaties, die een directe impact kunnen hebben op de nationale veiligheid van de getroffen landen.

Het lek van deze interne documenten biedt waardevolle inzichten in de methoden en de werking van APT35 en vormt een waarschuwing voor de kwetsbaarheden binnen kritieke infrastructuren die voortdurend doelwit zijn van geavanceerde, staat-gesponsorde cyberaanvallen.

vLLM-kwetsbaarheid maakt op afstand code-executie mogelijk via kwaadaardige payloads

Een ernstige geheugenfout in vLLM versies 0.10.2 en later stelt aanvallers in staat om op afstand code uit te voeren via de Completions API door kwaadwillig gemaakte prompt-embeddings naar de server te sturen. Deze kwetsbaarheid bevindt zich in het tensor-deserialisatieproces van vLLM, specifiek in de bestandscode binnen "entrypoints/renderer.py" op regel 148. Wanneer gebruikersprompts worden verwerkt, laadt het systeem geserialiseerde tensors met de functie `torch.load()`, zonder voldoende validatie van de invoer.

De kwetsbaarheid is het gevolg van een wijziging in PyTorch versie 2.8.0, waarin de standaardcontrole op de integriteit van sparsesensors werd uitgeschakeld. Hierdoor kunnen aanvallers tensors maken die de interne grenscontroles omzeilen, wat leidt tot een geheugenbeschrijving buiten de toegestane grenzen tijdens de conversie naar een dichte tensor (`to_dense()`). Dit kan resulteren in een crash van de vLLM-server en mogelijk in de uitvoering van willekeurige code binnen het serverproces.

De kwetsbaarheid heeft een CVSS-score van 8.8/10, wat aangeeft dat deze ernstig is. De aanval vereist geen speciale bevoegdheden, waardoor zowel geauthenticeerde als ongeauthenticeerde gebruikers, afhankelijk van de configuratie van de API, de kwetsbaarheid kunnen uitbuiten. Alle implementaties van vLLM als server, vooral die waarbij onbetrouwbare of door het model aangeleverde payloads worden gedeserialiseerd, lopen risico. Organisaties die vLLM gebruiken in productieomgevingen, cloudomgevingen of gedeelde infrastructuren worden geadviseerd om snel over te schakelen naar de gepatchte versie van vLLM om verdere risico's te vermijden.

Het vLLM-project heeft de kwetsbaarheid aangepakt in pull request #27204, en gebruikers worden aangespoord om onmiddellijk te upgraden naar de nieuwste



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

versie. Totdat de update is doorgevoerd, kunnen beheerders de toegang tot de API beperken tot vertrouwde gebruikers en invoervalidatielagen implementeren die prompt-embeddings inspecteren voordat ze de verwerkingspijplijn van vLLM bereiken.

Nieuwe kwetsbaarheden in Fluent Bit stellen cloudomgevingen bloot aan RCE en onopgemerkte inbraken in infrastructuur

Onderzoekers in cybersecurity hebben vijf ernstige kwetsbaarheden ontdekt in Fluent Bit, een veelgebruikte open-source tool voor telemetrie, die in miljarden containers wordt ingezet. Deze kwetsbaarheden kunnen door aanvallers worden benut om cloudinfrastructuur over te nemen, logbestanden te manipuleren en systemen op afstand uit te voeren. Fluent Bit wordt vaak gebruikt om loggegevens te verzamelen en te verwerken in cloudomgevingen, waaronder Kubernetes-clusters, wat het een aantrekkelijk doelwit maakt voor cybercriminelen.

De kwetsbaarheden omvatten onder andere een padtraversal-fout die het mogelijk maakt om willekeurige bestanden op de schijf te schrijven, wat logmanipulatie en op afstand uitvoeren van code mogelijk maakt. Andere fouten kunnen worden gebruikt om buffer-overflows te veroorzaken, toegang te verkrijgen tot vertrouwelijke gegevens door middel van misbruik van tag-logica, en zelfs om ongeautoriseerde log-invoer mogelijk te maken door het ontbreken van een juiste authenticatiecontrole in het `in_forward`-plugin van Fluent Bit.

De potentiële gevolgen van deze kwetsbaarheden zijn groot. Aanvallers zouden toegang kunnen krijgen tot logdata, waarmee ze misleidende of valse informatie kunnen injecteren, of zelfs het logproces kunnen manipuleren om hun eigen sporen te wissen na een aanval. Door deze kwetsbaarheden te benutten, kunnen aanvallers zich diep in de infrastructuur verankeren en onopgemerkt malware uitvoeren of andere schadelijke activiteiten ondernemen.

Fluent Bit heeft updates vrijgegeven om deze kwetsbaarheden te verhelpen, en gebruikers wordt dringend aangeraden om over te stappen naar de nieuwste versie van de software. Amazon Web Services (AWS), die betrokken was bij de coördinatie van de bekendmaking, heeft klanten aangespoord om hun versies bij te werken om een optimale bescherming te garanderen.

De ontdekte kwetsbaarheden benadrukken het belang van het handhaven van strikte beveiligingsmaatregelen in cloudomgevingen, vooral in systemen die afhankelijk zijn



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

van open-source tools zoals Fluent Bit. Het is belangrijk om dynamische tags voor routing te vermijden, outputpaden te vergrendelen om padtraversal te voorkomen, en configuratiebestanden als alleen-lezen in te stellen om runtime-tampering tegen te gaan.

Fluent Bit is populair binnen enterprise-omgevingen en de impact van deze kwetsbaarheden kan aanzienlijk zijn, aangezien ze kunnen leiden tot verstoring van clouddiensten, manipulatie van gegevens en zelfs overname van logservices.

Nieuwe scam richt zich op cryptobezitters: risico op overvallen of ontvoeringen

Cybercriminelen richten zich steeds vaker op cryptobezitters met een nieuwe vorm van fraude. Gebruikers worden via e-mail benaderd met een valse belastingaangifte, waarin ze niet alleen hun cryptovaluta moeten opgeven, maar ook persoonlijke gegevens zoals adres en telefoonnummer. Deze gegevens kunnen later worden misbruikt voor verdere aanvallen, zoals telefonische benaderingen of zelfs fysieke overvallen.

De nep-belastingaangifte, die qua uitstraling sterk lijkt op een officiële belastingdienstpagina, vraagt cryptobezitters naar de platforms waarop zij hun cryptovaluta bewaren, zoals Bitvavo of Coinbase, evenals naar hardware wallets, waarmee crypto's offline worden bewaard. De criminelen achter deze scam maken gebruik van zeer realistische valse formulieren die via e-mail worden verspreid.

Het gevaar schuilt niet alleen in de diefstal van digitale munten, maar ook in de fysieke dreigingen die kunnen ontstaan. Er zijn al gevallen bekend waarbij slachtoffers gedwongen werden hun cryptovaluta over te maken onder bedreiging van geweld. Zo werden in Nederland recent nog vijf verdachten aangehouden in een zaak waarbij een slachtoffer werd ontvoerd om crypto's over te maken.

Deze ontwikkeling benadrukt de gevaren van crypto-gerelateerde cybercriminaliteit, waarbij criminelen steeds geavanceerdere technieken gebruiken om niet alleen digitale gegevens, maar ook persoonlijke informatie van hun slachtoffers te bemachtigen. Het is belangrijk voor cryptobezitters om alert te zijn op dergelijke scams, aangezien de belastingdienst nooit per e-mail om persoonlijke gegevens vraagt en cryptomunten altijd via de reguliere belastingaangifte moeten worden opgegeven.



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

Het lijkt erop dat cybercriminelen ook doelgericht gebruikers van hardware wallets aanvallen, mogelijk door gebruik te maken van gegevens die in het verleden werden gestolen bij hacks, zoals die van het bedrijf Ledger in 2020.

Dit nieuws weerspiegelt de toenemende complexiteit van digitale en fysieke dreigingen die cryptobezitters kunnen treffen. Het onderstreept de noodzaak voor extra waakzaamheid in een steeds onveiliger wordende digitale wereld.

Honderden npm-packages geïnfecteerd met malware die gevoelige data steelt

In de npm Registry zijn honderden packages ontdekt die besmet zijn met malware die gevoelige gegevens steelt van systemen waarop ze worden uitgevoerd. Het gaat hierbij om een aanval die lijkt op een eerder incident, waarbij meer dan vijfhonderd packages werden geïnfecteerd. De recente aanval betreft meer dan driehonderd besmette npm-packages, die ontwikkeld werden voor het beheren van JavaScript-pakketten.

De malware maakt gebruik van de software TruffleHog, die op de systemen zoekt naar gevoelige informatie zoals NPM-tokens, cloudplatform-credentials van Amazon Web Services (AWS), Google Cloud Platform (GCP) en Azure, evenals omgevingsvariabelen. Wanneer de malware gevoelige gegevens heeft verzameld, stuurt het deze informatie naar de aanvallers via een GitHub Action runner, genaamd "SHA1HULUD", een naam die verwijst naar een eerdere aanval van september. Met de gestolen NPM-tokens kan de malware weer andere npm-packages infecteren en zich verder verspreiden.

Dit incident benadrukt de risico's van open-source software die via de npm Registry wordt verspreid en de noodzaak voor ontwikkelaars om waakzaam te blijven bij het gebruik van externe packages. De infecties kunnen verstrekende gevolgen hebben voor bedrijven en ontwikkelaars die afhankelijk zijn van npm voor hun softwareontwikkeling.

Fake Windows Update vraagt slachtoffers om malafide commando uit te voeren

Cybercriminelen maken steeds vaker gebruik van nieuwe tactieken om slachtoffers te misleiden en malware op systemen te installeren. In dit geval wordt een valse Windows Update gepresenteerd als een legitieme updatepagina van Microsoft. Slachtoffers zien een schermvullende pagina die lijkt op het gebruikelijke Windows



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

Update-scherm. Na enige tijd verschijnt de mededeling dat de update niet volledig kan worden afgerond zonder het uitvoeren van bepaalde handelingen.

Om de installatie van de update te voltooien, wordt het slachtoffer gevraagd een aantal toetsencombinaties in te voeren, waaronder het openen van het "Uitvoeren"-venster en het uitvoeren van een commando. Dit commando is echter al op het klembord van het slachtoffer geplaatst door de malafide pagina, en wanneer het slachtoffer het commando uitvoert, wordt malware op het systeem geïnstalleerd. De malware, in de vorm van een infostealer, verzamelt gevoelige gegevens zoals wachtwoorden en inloggegevens en stuurt deze naar de aanvallers.

Deze aanval maakt gebruik van een techniek die "ClickFix" wordt genoemd. Eerder werd deze techniek gebruikt waarbij slachtoffers werd verteld dat het uitvoeren van het commando nodig was om een captcha op te lossen. Deze nieuwe variant is een uitbreiding van die tactiek, waarbij het doel is om systemen te infecteren en gegevens van gebruikers te stelen. Het is echter onbekend waar gebruikers de malafide updatepagina precies kunnen tegenkomen.

Deze vorm van misleiding is bijzonder zorgwekkend, aangezien het slachtoffers direct vraagt om acties uit te voeren die normaal gesproken geassocieerd worden met legitieme systeemupdates. Dit benadrukt opnieuw het belang van alertheid bij het uitvoeren van systeemupdates en het vermijden van verdachte activiteiten op het internet.

Noord-Koreaanse valse vacatureplatformen richten zich op AI-ontwikkelaars

Een geavanceerde wervingsfraude die verband houdt met Noord-Korea heeft onlangs zijn weg gevonden naar Amerikaanse ontwikkelaars in kunstmatige intelligentie (AI), software-ingenieurs en professionals uit de cryptowereld. Deze frauduleuze campagne maakt gebruik van een op maat gemaakte, nep-vacaturewebsite om zich voor te doen als een legitiem sollicitatieplatform. De website, die wordt gehost op [lenvny\[.\]com](https://lenvny[.]com), lijkt op een professioneel recruitmentplatform en is ontworpen met behulp van technologie zoals React en Next.js, waarmee het een indrukwekkende en geloofwaardige uitstraling krijgt.

De frauduleuze website presenteert zichzelf als een "AI-gestuurd interviewtool" voor wervingsdoeleinden. Het biedt een reeks vacatures die specifiek gericht zijn op hooggewaardeerde doelwitten in de AI- en cryptosectoren. De opzet van de website is verrassend gepolijst en authentiek, met een marketinginterface en een ontwerp dat



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

past bij de verwachtingen van de tech-industrie in 2025. Het bevat dynamisch gegenereerde vacatures en een volledig sollicitatieproces, inclusief video-interviews en technische assessments die slachtoffers vragen om code uit te voeren of scripts op hun systemen uit te voeren.

De aanval volgt een specifiek patroon dat bekend staat als de “Contagious Interview”-operatie. Kandidaten worden benaderd via LinkedIn en doorgestuurd naar de nepsite, waar ze worden gevraagd om video-antwoorden op interviewvragen in te dienen. Vervolgens worden ze gevraagd om hun webcam “te repareren” met behulp van een tool, een zogenaamd technische oplossing die in werkelijkheid malware installeert. Deze zogenaamde “ClickFix”-techniek is een sociaal-engineeringaanpak die slachtoffers misleidt om schadelijke software te downloaden.

Deze campagne is gericht op specifieke doelgroepen die toegang hebben tot waardevolle digitale activa of expertise, zoals AI-onderzoekers en cryptoprofessionals. AI-ontwikkelaars hebben vaak toegang tot vertrouwelijke onderzoeksgegevens, modelgewichten en infrastructuren voor inferentie, terwijl crypto-experts vaak werken in omgevingen die digitale activa beheren, wat hen tot aantrekkelijke doelwitten maakt voor aanvallers. De aanvallers richten zich specifiek op deze beroepsgroepen omdat zij werken met systemen die verhoogde systeemprivileges hebben, wat de kans vergroot dat de malware succesvol wordt uitgevoerd.

Beveiligingsexperts raden werkzoekenden aan om sollicitatiesites altijd te controleren op officiële domeinen en geen persoonlijke documenten te uploaden naar onbekende platforms. Kandidaten die gevraagd worden om code uit te voeren tijdens sollicitatiegesprekken moeten altijd scripts zorgvuldig controleren en bij voorkeur in virtuele machines of sandboxomgevingen uitvoeren, om te voorkomen dat schadelijke software hun werkstations compromitteert.

Python-gebaseerde malware injecteert processen in legitieme Windows-binaries

Beveiligingsonderzoekers hebben een geavanceerde Python-gebaseerde malware ontdekt die gebruik maakt van procesinjectie om zich te verbergen in legitieme Windows-binaries. Deze dreiging is een nieuwe evolutie in bestandsloze aanvalstechnieken, waarbij een combinatie van geavanceerde obfuscatie en vertrouwde systeemhulpprogramma's wordt gebruikt om detectie te vermijden.



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

De malware maakt gebruik van een 65 MB groot bestand dat voornamelijk bestaat uit vullende gegevens, met aan het einde een klein, geldig fragment van een .pyc-bestand dat de daadwerkelijke kwaadaardige code bevat. Dit fragment is ontworpen om processen in legitieme Windows-uitvoerbare bestanden te injecteren. Het gebruik van een Python-runtime-omgeving die samen met een ondertekend uitvoerbaar bestand wordt gebundeld, maakt het moeilijker voor traditionele beveiligingssystemen om de dreiging te identificeren.

Tijdens de analyse werd ontdekt dat de malware verschillende geavanceerde kenmerken vertoont, zoals multi-laag codering, het verbergen van archieftypen en het verpakken van een Python-runtime met een naam die legitiem lijkt. De malware maakt gebruik van een PE-dumper die tijdens de uitvoering een batchscript reconstrueert via runtime decryptie met behulp van SIMD-operaties. Dit script downloadt vervolgens een bestand dat is vermomd als een PNG-afbeelding vanuit cloudopslag, maar in werkelijkheid is het een RAR-archief.

De aanval verloopt in meerdere fasen: nadat het batchscript is uitgevoerd, wordt het archief uitgepakt, waarbij drie componenten worden onthuld, waaronder een gemaskeerde systeembestuurder en een Python-runtime. Vervolgens voert de malware een de-obfuscatieprocedure uit door middel van Base64-decoding, BZ2-decompressie en Zlib-decompressie, voordat de Python-bytecode in het geheugen wordt geladen. Deze code richt zich onmiddellijk op het legitieme cvtres.exe, een Microsoft-hulpprogramma voor het converteren van resources, om het proces te injecteren.

De malware kan communiceren met command-and-control-servers (C2) en blijft actief, zelfs nadat de oorspronkelijke loader is beëindigd. Door zich te verbergen binnen legitieme Microsoft-processen en versleutelde communicatiekanalen te onderhouden, vormt deze malware een bijzonder gevaar voor enterprise-omgevingen, waar traditionele detectiesystemen mogelijk niet in staat zijn de dreiging tijdig te identificeren.

EtherHiding-aanval gebruikt blockchain voor malwarelevering en payloadrotatie

Een nieuwe dreiging, bekend als EtherHiding, maakt gebruik van blockchaintechnologie om malware te verspreiden. Dit maakt het voor beveiligingsteams veel moeilijker om aanvallers te traceren en tegen te houden,



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

omdat de schadelijke payloads kunnen worden gewijzigd zonder de websites waar de aanval begint aan te passen. EtherHiding gebruikt slimme contracten op de blockchain om malware te laden en bij te werken. Dit zorgt voor een dynamischere en moeilijker te detecteren manier van aanvallen.

Het proces begint wanneer een aanvaller kwaadaardige code injecteert in een legitieme website. De geïnjecteerde code toont een nep-CAPTCHA-pagina, die lijkt op een echte beveiligingscontrole en bezoekers vraagt om te bewijzen dat ze menselijk zijn. In plaats van simpelweg een vinkje aan te vinken, worden slachtoffers misleid om code in hun terminal of opdrachtprompt in te voeren. Wanneer zij deze instructies opvolgen, wordt de malware geïnstalleerd zonder dat de beveiligingstools die automatisch gedrag van malware monitoren het detecteren.

Deze aanval maakt gebruik van een gedecentraliseerde infrastructuur, waarbij de payload wordt opgeslagen en opgehaald via blockchaincontracten, met behulp van de Binance Smart Chain. De aanval is specifiek voor zowel Windows- als macOS-systemen, waarbij elk systeem zijn eigen specifieke contracten heeft die een gepersonaliseerde payload leveren. Zodra het slachtoffer via een uniek identificatienummer door een controlecontract wordt geverifieerd, wordt een nep-CAPTCHA weergegeven met instructies die specifiek zijn voor het besturingssysteem van het slachtoffer. Dit zorgt ervoor dat het slachtoffer de schadelijke code zelf uitvoert, waardoor het voor beveiligingssystemen moeilijker wordt om de aanval te detecteren.

Op macOS maakt de malware gebruik van AppleScript en curl-opdrachten om een volledig agent te downloaden en uit te voeren. Deze agent creëert persistentie op het systeem en verzamelt gevoelige gegevens, zoals wachtwoorden, die vervolgens naar de server van de aanvaller worden verzonden. Op Windows-systemen wordt een vergelijkbare aanpak gevolgd, waarbij de malware op een soortgelijke manier wordt gedownload en uitgevoerd.

De combinatie van blockchaintechnologie, social engineering via nep-CAPTCHA's, en handmatige uitvoering van de schadelijke code maakt EtherHiding tot een uiterst flexibele en moeilijk te traceren aanval. Beveiligingsspecialisten raden aan om websites die nep-CAPTCHA-pagina's vertonen te monitoren en alert te zijn op verdachte clipboardactiviteiten die naar terminalcommando's verwijzen. Dit kan helpen om de aanval te stoppen voordat de installatie plaatsvindt.



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

Hackers misbruiken schadelijke PyPI-pakket om gebruikers aan te vallen en cryptocurrency-gegevens te stelen

Er is een nieuwe malwarecampagne ontdekt die zich richt op cryptocurrencygebruikers via een misleidend Python-pakket, gehost op de PyPI-repository. De aanvallers verstopten hun schadelijke code in een nep-spelfoutcontroletool, die zich voordeed als het legitieme pypspellchecker-pakket, dat meer dan 18 miljoen downloads heeft. Deze supply chain-aanval maakt gebruik van een vertrouwd softwareplatform om kwaadaardige tools zoals remote access trojans (RAT's) en inloggegevensstelen te verspreiden onder onwetende ontwikkelaars wereldwijd.

De malware is geoptimaliseerd met geavanceerde obfuscatie- en encryptietechnieken om detectie te vermijden. HelixGuard, een beveiligingsbedrijf, heeft ontdekt dat de command-and-control-infrastructuur die aan deze operatie verbonden is, overeenkomt met servers die eerder gebruikt werden in geavanceerde social engineering-campagnes. Dit wijst op een gecoördineerde strategie waarbij aanvallers zijn overgestapt van directe social engineering naar geautomatiseerde verspreiding via open-sourceplatforms, wat hun bereik en effectiviteit aanzienlijk vergroot.

Het kwaadaardige pakket is inmiddels meer dan 950 keer gedownload sinds de uitrol. De malware maakt gebruik van een gefaseerd bezorgmechanisme, waarbij elke fase is ontworpen om de stealth te behouden terwijl de controle over de gecompromitteerde systemen steeds dieper wordt. Het doel van de aanval is voornamelijk het stelen van cryptocurrency-informatie, wat duidt op de financiële motivatie die moderne malware-aanvallen aandrijft. Deze trend van aanvallen op digitale activa-houders blijft toenemen, ongeacht de technische expertise van de slachtoffers.

De infectie begint wanneer gebruikers het kwaadaardige pakket installeren. Het eerste kwaadaardige payload wordt geactiveerd via een Base64-gecodeerd bestand, dat wordt gedecodeerd en uitgevoerd via de Python-functie `exec()`. Deze techniek zorgt ervoor dat verdachte code niet naar de schijf wordt geschreven, waardoor het moeilijker te detecteren is. Vervolgens maakt de malware verbinding met een server die door de aanvallers wordt gecontroleerd, waar het tweede schadelijke payload wordt gedownload, een volledige remote access trojan. Deze trojan stelt de aanvallers in staat om op afstand commando's uit te voeren via Python en maakt



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

gebruik van XOR-encryptie om de netwerkcommunicatie te verbergen voor monitoringtools.

Eenmaal geactiveerd, biedt de backdoor volledige controle over het slachtoffer's computer, waardoor de aanvallers in staat zijn cryptocurrency-portemonnees, inloggegevens en andere gevoelige informatie te stelen. Gebruikers wordt aangeraden hun geïnstalleerde Python-pakketten te controleren, hun afhankelijkheden bij te werken en verdachte pakketten te verwijderen. Organisaties zouden strikte afhankelijkheidsscanning in hun ontwikkelingspijplijnen moeten implementeren en verbindingen met de geïdentificeerde command-and-control-adressen moeten monitoren.

Shai-Hulud-aanval treft meer dan 500 npm-pakketten

Een nieuwe golf van de Shai-Hulud-campagne heeft onlangs npm-pakketten getroffen, waarbij meer dan 500 pakketten en 700 versies werden gecompromitteerd. De aanval richtte zich op populaire ontwikkeltools van bedrijven zoals @zapier, @asynapi, @postman, @posthog en @ensdomains. Deze aanvallen gebeurden via accountovernames en compromitteren van ontwikkelaars, waarbij schadelijke wijzigingen werden aangebracht in de betreffende pakketten.

Een van de meest opvallende technieken die door de aanvallers werd gebruikt, was de toevoeging van een preinstall script in de package.json-bestanden. Dit script, genaamd setup_bun.js, fungeerde als een stealthy loader die onopgemerkt de Bun-runtime installeerde en een obfuscated 10MB malwarescript genaamd bun_environment.js uitvoerde. Het uitvoerende script werd zo ontworpen dat alle uitvoer werd onderdrukt om detectie te vermijden.

De getroffen pakketten bevatten belangrijke componenten van bekende tools zoals de @zapier/ai-actions en @postman/mcp-ui-client, en zelfs enkele kritieke libraries voor @asynapi en @ensdomains. Indien deze pakketten werden geïnstalleerd in omgevingen met toegang tot gevoelige gegevens, kunnen er belangrijke API-sleutels, tokens en wachtwoorden zijn gestolen. Als gevolg van deze aanval wordt gebruikers aangeraden onmiddellijk hun beveiligingssleutels te roteren en de CI/CD-pijplijnen goed te monitoren.

Het onderzoek naar de omvang van deze aanval is nog in volle gang, met verdere updates die regelmatig worden gepubliceerd. Het is van cruciaal belang om



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

kwetsbare en geïnfecteerde versies van npm-pakketten snel te identificeren en de getroffen systemen te beveiligen om verdere schade te voorkomen.

Dropping Elephant Hacker Groep Aanvalt Defensiesector met Python Backdoor via MSBuild Dropper

De India-gerelateerde hacker groep Dropping Elephant heeft een geavanceerde cyberaanval uitgevoerd op de defensiesector van Pakistan. Deze aanval maakt gebruik van een Python-gebaseerde remote access trojan (RAT), die verstopt zat in een MSBuild dropper. Het doelwit was onder meer de militaire onderzoeks- en ontwikkelingsafdelingen, evenals inkoopfaciliteiten die gelinkt zijn aan de National Radio and Telecommunication Corporation van Pakistan.

Het aanvalscenario begon met een phishing-e-mail die een schadelijke ZIP-archief bevatte. Wanneer het bestand werd gedownload, bleek het een MSBuild-projectbestand te bevatten, dat als de eerste dropper diende. Het bestand werd gecombineerd met een vervalst PDF-bestand dat legitiem leek. Bij uitvoering van de dropper werden verschillende componenten gedownload naar de Windows-taakmap, waarbij het systeem via geplande taken persistente toegang kreeg. De geplande taken hadden onschuldige namen zoals KeyboardDrivers en MsEdgeDrivers, wat de detectie bemoeilijkte.

Tijdens het onderzoek naar deze aanval ontdekte onderzoeker Idan Tarab dat de groep gebruik maakte van geavanceerde technieken om de aanval te verbergen, waaronder het gebruik van UTF-reverse encryptie om strings te reconstrueren en dynamische API-oplossingen om de detectie door beveiligingstools te vermijden. De manier waarop Dropping Elephant legitieme Windows-hulpprogramma's heeft gemanipuleerd, toont de technische volwassenheid van de groep.

Het hart van de operatie was een volledig ingebedde Python-runtime die naar de AppData-directory werd gedownload. Het bestand, python2_pycache_.dll, was geen echte DLL, maar bevatte gemarshalleerde Python-bytecode. Deze bytecode werd uitgevoerd via pythonw.exe, wat de operatie stealthy maakte door zonder venster te draaien. De malware bevatte verschillende modules die het systeem volledig konden overnemen, inclusief functionaliteiten voor het verzamelen van informatie en het uitvoeren van commando's op het geïnfecteerde systeem.

De malware communiceerde met de command-and-control-infrastructuur via verschillende domeinen, zoals nexnky.info, upxvion.info en soptr.info. Deze



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

domeinen maakten het moeilijker om de communicatie te blokkeren, aangezien de malware gebruik maakte van gecodeerde variabelen en base64-gecodeerde commando's, wat de handmatige analyse bemoeilijkte.

Deze aanval benadrukt de voortdurende dreiging van geavanceerde, persistent aanvallende groepen die zich richten op kritieke infrastructuur in Zuid-Azië. Het laat ook de noodzaak zien van versterkte monitoring van verdachte MSBuild-uitvoeringen en ongebruikelijke Python-runtime-implementaties in systeembestanden. Organisaties in de defensiesector zouden hun phishing-afweermechanismen moeten aanscherpen en extra waakzaamheid moeten betrachten bij het uitvoeren van systemen die Python bevatten of met MSBuild werken.

Hackers misbruiken WhatsApp voor het heimelijk verzamelen van logbestanden en contactgegevens

Een nieuwe malwarecampagne richt zich op gebruikers in Brazilië en maakt gebruik van WhatsApp als het belangrijkste kanaal voor het verspreiden van banktrojans en het verzamelen van gevoelige informatie. Deze geavanceerde aanval maakt gebruik van sociale engineering door de vertrouwensrelatie die slachtoffers hebben met hun bestaande contacten, waardoor de kwaadaardige bestanden als legitiem worden gepresenteerd.

De campagne begint met phishing-e-mails die gecomprimeerde VBS-scripts bevatten, welke geavanceerde obfuscatietechnieken gebruiken om detectie door beveiligingssoftware te vermijden. Zodra de eerste payload wordt uitgevoerd, wordt Python en de Selenium WebDriver gedownload en geïnstalleerd, waarmee geautomatiseerde interactie met WhatsApp Web mogelijk wordt gemaakt. De malware injecteert vervolgens kwaadaardige JavaScript-code in de browsersessie van het slachtoffer en maakt gebruik van de interne WhatsApp-API's om contacten op te sommen en verdere payloads te verspreiden.

Deze benadering stelt aanvallers in staat om de infectie te verspreiden zonder dat QR-code-authenticatie nodig is, door bestaande ingelogde sessies te kapen door browsercookies en lokale opslagdata te kopiëren. Onderzoekers van K7 Security Labs hebben deze variant geïdentificeerd als onderdeel van de bredere Water-Saci-campagne, die actief financiële instellingen in Brazilië aanvalt.

De aanval voert zowel een op Python gebaseerde distributiescript uit als een banktrojan die specifiek let op actieve Windows-gerelateerde processen van



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

Braziliaanse banken en cryptocurrency wallets. Door de combinatie van geautomatiseerde berichten en geheugen-only payload-executie blijft de malware onopgemerkt, vooral wanneer de systemen van de slachtoffers en hun contactnetwerken worden gecompromitteerd.

De campagne levert ook een MSI-installatiebestand af dat een Autolt-script en versleutelde payloadbestanden installeert. Dit secundaire component zorgt voor persistentie via registerwijzigingen en bewaakt continu actieve vensters van het slachtoffer op zoek naar specifieke bankgerelateerde trefwoorden. Wanneer bepaalde financiële instellingen of crypto-wallettoepassingen worden gedetecteerd, wordt de banktrojan geladen in het geheugen, waarbij schijfopslag wordt vermeden, waardoor traditionele op bestand gebaseerde detectiemethoden niet effectief zijn.

Het infectiemechanisme begint wanneer slachtoffers phishing-e-mails ontvangen die ZIP-gecomprimeerde VBS-scriptbestanden bevatten. Deze bestanden maken gebruik van karaktercodering en XOR-encryptie om handtekening-gebaseerde detectie te omzeilen. Het script gebruikt een multi-layered obfuscatiestrategie, waarbij strings karakter voor karakter worden opgebouwd en XOR-operaties worden toegepast om de werkelijke kwaadaardige commando's te decoderen.

Na de obfuscatiestrategie wordt een MSI-bestand en een ander VBS-bestand gedownload. Het VBS-bestand bevat dezelfde obfuscatiepaden en laat een batchscript vallen dat Python, ChromeDriver en Selenium installeert. Deze geautomatiseerde setup maakt de infrastructuur mogelijk die nodig is voor WhatsApp-automatisering zonder handmatige tussenkomst van de gebruiker.

De Python-script, genaamd `whats.py`, neemt de controle over de WhatsApp Web-sessie van het slachtoffer door browserprofielgegevens, waaronder cookies, lokale opslag en IndexedDB-bestanden, naar een tijdelijke map te kopiëren. Vervolgens wordt Chrome opgestart met deze gekopieerde gegevens, waardoor de QR-code-authenticatie wordt omzeild. Nadat de malware is geauthenticeerd, injecteert het JavaScript vanuit GitHub in de WhatsApp Web-pagina, waardoor toegang wordt verkregen tot interne API-functies zoals `WPP.contact.list` en `WPP.chat.sendMessage`.

Het script haalt vervolgens de contactlijst van het slachtoffer op, waarbij groeps- en zakelijke accounts, evenals contacten met specifieke nummerpatronen, worden gefilterd. Deze verzamelde contacten worden vervolgens in batches verstuurd met kwaadaardige ZIP-bestanden die de volgende fase van de infectie bevatten,



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

waardoor de campagne zich verspreidt over netwerken van slachtoffers en gedetailleerde logbestanden naar de server van de aanvaller worden gestuurd.

Minderjarige nepagent op heterdaad aangehouden in Goes

Op zondagavond 23 november werd de politie in Goes gealarmeerd door een bewoner van de Goese Dieplaan. Deze had een telefoontje ontvangen van iemand die zich voordeed als politieagent. De nepagent beweerde onderweg te zijn naar Goes om waardevolle spullen op te halen. De politie reageerde snel en ging met meerdere agenten ter plaatse.

Toen de nepagent aanbelde bij het betreffende adres, grepen de politieagenten in en arresteerden de verdachte ter plaatse. Het bleek te gaan om een minderjarige jongen uit Chaam. De jongen werd meegenomen naar het politiebureau voor verder onderzoek.

Dit incident onderstreept het belang van waakzaamheid bij telefoontjes van vermeende politieagenten of bankmedewerkers. De politie zal nooit verzoeken om waardevolle spullen, bankpassen, pincodes of andere persoonlijke documenten op te halen. In dergelijke gevallen is het raadzaam om direct op te hangen, de politie te bellen via 0900-8844 en hun verhaal te verifiëren. Als er iemand aan de deur staat, kan altijd gevraagd worden om een geldig politielegitimatiebewijs, wat iedere agent bij zich heeft.

Drie verdachten aangehouden in callcenter voor nepagenten en bankfraude

Op vrijdag 21 november heeft de politie drie personen aangehouden in een onderzoek naar een callcenter dat betrokken was bij nepagenten- en bankhelpdeskfraude. Het onderzoek begon in oktober, toen een 92-jarige vrouw uit Harderwijk werd opgelicht door iemand die zich voordeed als politieagent. De vrouw werd urenlang telefonisch benaderd met het verhaal dat haar gegevens waren gevonden bij twee jongemannen die waren aangehouden. Er zou onderzoek plaatsvinden en daarom moest de politie haar bankpas en sieraden ophalen. Uiteindelijk kwam er een nepagent langs om deze zaken te stelen.

De politie van Basisteam Veluwe-West startte daarop een onderzoek, dat leidde naar een callcenter in Rotterdam. Dit callcenter belde meerdere slachtoffers en gaf zich uit als agenten en bankmedewerkers. Op basis van dit onderzoek heeft de politie op 21



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

november een woning in Rotterdam doorzocht en drie verdachten aangehouden: twee mannen van 24 en 29 jaar zonder vaste woon- of verblijfplaats en een 25-jarige man uit Amsterdam. De drie verdachten zitten nog vast en het onderzoek wordt voortgezet.

De politie benadrukt dat het belangrijk is om nooit iemand binnen te laten die zich voordoeft als agent of bankmedewerker, en altijd 112 te bellen bij twijfel. Het is een gangbare strategie van criminelen om zich als autoriteit voor te doen om zo slachtoffers te misleiden en te bestelen. De politie werkt samen met Omroep MAX om ouderen voor deze vormen van fraude te waarschuwen.

Nationaal AI Deltaplan moet Nederland aan top van Europa brengen

Het Nationaal AI Deltaplan heeft als doel Nederland te positioneren als koploper op het gebied van kunstmatige intelligentie (AI) in Europa. Het plan, gepresenteerd door een coalitie van AI-experts in opdracht van minister Vincent Karremans, benadrukt de noodzaak voor aanzienlijke investeringen in kennis, talent en infrastructuur. Jelle Prins, mede-initiatiefnemer van het plan, geeft aan dat Nederland hard aan de slag moet om de ambities waar te maken. "Maar dan moeten we nu echt vol aan de bak," zegt hij. Het Deltaplan is een antwoord op vergelijkbare initiatieven uit landen als de Verenigde Staten en het Verenigd Koninkrijk, en bevat een meerjarige investeringsagenda gericht op versnelde adoptie van AI bij zowel bedrijven als de overheid.

Een belangrijk onderdeel van het plan is de opbouw van AI-rekenkrachtinfrastructuur en het creëren van doorgroeimogelijkheden voor innovatieve bedrijven. Daarnaast wordt er gepleit voor de oprichting van een onafhankelijk Nationaal AI-instituut, dat de maatschappelijke en economische impact van AI moet volgen en adviseren aan het parlement. De initiatiefnemers wijzen op de potentiële voordelen van AI voor de Nederlandse economie, zoals de ontwikkeling van nieuwe medicijnen, duurzame materialen en robots voor sectoren met een tekort aan arbeidskrachten, zoals de zorg en de bouw.

Tegelijkertijd waarschuwen de experts voor de risico's van afhankelijkheid van buitenlandse technologie, met name van Amerikaanse en Chinese techgiganten. Prins benadrukt dat Nederland zonder eigen infrastructuur economisch afhankelijk dreigt te worden, waardoor het controle over belangrijke systemen verliest. Het plan biedt kansen om het concurrentievermogen te versterken, maar volgens de



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

initiatiefnemers is succes alleen mogelijk als de gehele samenleving betrokken wordt. Dit omvat onder andere burgerberaden en het genoemde AI Instituut, zodat de techniek de samenleving daadwerkelijk dient.

In de komende maanden zal het team verder overleg voeren met politieke partijen en feedback uit de samenleving verzamelen om het plan verder te ontwikkelen en uit te voeren. Het Deltaplan moet ervoor zorgen dat Nederland een leidende rol speelt in de ontwikkeling van AI in Europa, maar daar is volgens de experts een gezamenlijke inzet en snelle actie voor nodig.

Slachtoffer van helpdeskfraude krijgt schade van 59.000 euro niet vergoed door bank

Een klant van SNS die het slachtoffer werd van helpdeskfraude, heeft de schade van 59.000 euro die hieruit volgde niet vergoed gekregen door de bank. Het financiële klachteninstituut Kifid oordeelde dat de bank in dit geval niet verantwoordelijk kan worden gehouden voor de verloren gelden.

Op 31 december vorig jaar werd de klant gebeld door een oplichter die zich voordeed als een medewerker van SNS. De oplichter meldde dat er een poging was gedaan om geld van de rekening van de klant af te schrijven en vroeg de klant om software te installeren waarmee de oplichter toegang kreeg tot de internetbankierenomgeving van het slachtoffer. Hierdoor kon de oplichter meekijken en verschillende handelingen uitvoeren. Het slachtoffer verhoogde de limiet van zijn Digipas naar 80.000 euro, waarna een betaling van ruim 33.000 euro werd verricht naar een door de oplichter opgegeven rekeningnummer.

De betaling werd door de bank in behandeling genomen, maar pas na contact met de bank kwam het slachtoffer erachter dat het om fraude ging. Ondanks dat de bank de eerste betaling had tegengehouden, was het slachtoffer blijven handelen op verzoek van de oplichter en had zelfs onjuiste verklaringen afgelegd over de bedoelingen van de betalingen. Het Kifid oordeelde dat de bank haar zorgplicht niet had geschonden. De klant had immers zelf een cruciale rol gespeeld door valse verklaringen te doen over de betalingen.

De klant had zich ook beroepen op het coulancekader voor slachtoffers van bankhelpdeskfraude, maar volgens het Kifid was er geen reden voor de bank om de schade te vergoeden. De vordering werd dan ook afgewezen.



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

ACM houdt vanaf nu toezicht op naleving van Europese Data Act

De Autoriteit Consument & Markt (ACM) heeft aangekondigd vanaf nu toezicht te houden op de naleving van de Europese Data Act. Deze wet is bedoeld om gebruikers van zogenaamde "slimme apparaten" meer controle over hun persoonlijke gegevens te geven en het eenvoudiger te maken om van cloudprovider te wisselen of verschillende clouddiensten te koppelen. De wet beoogt dat eigenaren of gebruikers van slimme apparaten, zoals omvormers van zonnepanelen en voertuigen, toegang krijgen tot de gegevens die deze apparaten verzamelen. Dit geeft hen de mogelijkheid om deze data met andere bedrijven te delen, bijvoorbeeld om onderhoud uit te voeren bij een garage naar keuze in plaats van bij de merkdealer.

Een ander belangrijk aspect van de Data Act is gericht op de interoperabiliteit tussen verschillende cloudproviders. De wet maakt het eenvoudiger om gegevens tussen verschillende clouddiensten te verplaatsen of te koppelen. Dit zou het overstappen van de ene naar de andere cloudaanbieder aanzienlijk moeten vergemakkelijken. De ACM heeft de verantwoordelijkheid om toe te zien op de naleving van deze wet door bedrijven die slimme apparaten produceren en clouddiensten aanbieden binnen Nederland.

Consumenten of bedrijven die vermoeden dat hun rechten op basis van de Data Act geschonden worden, kunnen dit melden bij de ACM. Het toezicht door de ACM is een belangrijke stap in de implementatie van de Data Act, die het recht van gebruikers versterkt om zelf te bepalen met welke bedrijven zij hun gegevens willen delen.

Chinese AI DeepSeek-R1 genereert kwetsbare code bij geopolitieke triggers

Onderzoek van CrowdStrike heeft aangetoond dat het Chinese AI-model DeepSeek-R1 onveilige code kan genereren wanneer de prompts onderwerpen bevatten die door de Chinese overheid als politiek gevoelig worden beschouwd, zoals Tibet, de Oeigoeren en Falun Gong. Dit vergroot de kans op het produceren van code met beveiligingskwetsbaarheden met tot wel 50%.

DeepSeek-R1 is een krachtig AI-model voor het genereren van code, maar het onderzoek toont aan dat de kwaliteit van de gegenereerde code afneemt wanneer het model wordt gevraagd om te reageren op geopolitieke onderwerpen. In de



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

testscenario's waar Tibet of Oeigoeren werden genoemd, bevatte de gegenereerde code ernstige tekortkomingen, zoals het hard-coderen van geheime waarden en onveilige methoden voor het verwerken van gebruikersdata.

In een voorbeeld vroeg CrowdStrike het model om een webhook-handler voor PayPal-betalingen te schrijven voor een financiële instelling in Tibet. De code die werd gegenereerd voldeed niet aan de veiligheidsnormen en bevatte kwetsbaarheden die het risico op datalekken vergrootten.

Daarnaast ontdekte CrowdStrike dat het model een ingebouwde “kill switch” bevatte, die de productie van code blokkeerde wanneer het model werd gevraagd om te reageren op onderwerpen als Falun Gong. Dit wijst erop dat de AI specifieke veiligheidsmaatregelen heeft ingebouwd om te voldoen aan de Chinese wetgeving, die voorschrijft dat AI-diensten geen illegale of ongewenste inhoud mogen genereren.

De bevindingen benadrukken een belangrijk risico voor organisaties die AI-modellen gebruiken voor het genereren van code, aangezien geopolitieke factoren de beveiliging van deze systemen kunnen beïnvloeden. Dit toont de noodzaak aan voor strengere controles en een zorgvuldiger gebruik van AI-technologie in gevoelige omgevingen.

LLM-tools zoals GPT-3.5-Turbo en GPT-4 bevorderen de ontwikkeling van volledig autonome malware

De opkomst van grote taalmodellen (LLMs) zoals GPT-3.5-Turbo en GPT-4 heeft de manier waarop cybercriminelen opereren ingrijpend veranderd. Onderzoekers hebben aangetoond dat deze geavanceerde AI-tools misbruikt kunnen worden om kwaadaardige code te genereren, wat een nieuwe generatie malware mogelijk maakt. In plaats van traditionele malware die vastgelegde instructies bevat, maakt deze nieuwe benadering gebruik van AI om dynamische instructies te creëren. Dit maakt detectie voor beveiligingsteams aanzienlijk moeilijker, omdat de malware zich voortdurend kan aanpassen.

Cybercriminelen maken gebruik van een techniek die prompt-injectie wordt genoemd. Door AI-modellen specifieke verzoeken te doen, kunnen ze deze misleiden en laten reageren door code te genereren voor schadelijke operaties. Dit kan bijvoorbeeld malware-injectie in systeemprompts omvatten of het uitschakelen van antivirussoftware. Dit betekent dat toekomstige malware mogelijk nauwelijks



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

detecteerbare code in de binaire bestanden zelf bevat, maar in plaats daarvan afhankelijk is van de AI om tijdens de uitvoering nieuwe instructies te genereren.

Netskope, een organisatie voor beveiligingsanalyse, heeft deze opkomende dreiging geïdentificeerd door uitgebreide tests van zowel GPT-3.5-Turbo als GPT-4. Hun onderzoek toonde aan dat hoewel deze modellen inderdaad in staat zijn kwaadaardige code te genereren, er nog steeds aanzienlijke obstakels zijn die de succesvolle uitvoering van autonome aanvallen belemmeren. Het grootste probleem voor aanvallers is niet alleen het genereren van schadelijke code, maar ook het waarborgen dat deze code effectief werkt op de machines van slachtoffers.

Tijdens tests ontdekten de onderzoekers dat het genereren van scripts voor het detecteren van virtuele omgevingen en sandboxes - waar malware vaak wordt geanalyseerd - niet altijd succesvol was. Dit toonde aan dat hoewel de AI-modellen krachtig zijn, er nog steeds beperkingen zijn in de betrouwbaarheid van de door hen gegenereerde code in verschillende omgevingen. Dit zwakke punt belemmert de haalbaarheid van volledig autonome LLM-gestuurde aanvallen. Naarmate AI-modellen verder worden verbeterd, zullen deze beperkingen waarschijnlijk afnemen, waardoor de grootste uitdaging zal verschuiven van de functionaliteit van de gegenereerde code naar het omzeilen van de steeds geavanceerdere veiligheidsmaatregelen binnen AI-systemen.

SURF biedt onderwijsinstellingen Nextcloud voor digitale soevereiniteit

Vanaf januari 2026 krijgen Nederlandse onderwijs- en onderzoeksinstituten de mogelijkheid om gebruik te maken van Nextcloud, een open source samenwerkingsomgeving. Deze omgeving biedt mogelijkheden voor tekstverwerking, documentdeling, e-mailintegratie en online vergaderen. Dit initiatief is onderdeel van een bredere strategie van SURF, de ict-coöperatie van Nederlandse onderwijs- en onderzoeksinstituten, om de digitale soevereiniteit te versterken.

De digitale soevereiniteit is de afgelopen jaren een steeds relevanter thema geworden. De afhankelijkheid van grote techbedrijven vormt een risico voor de onafhankelijkheid en de bescherming van publieke waarden. In reactie daarop onderzoekt SURF alternatieve technologieën die het mogelijk maken om meer controle te houden over digitale omgevingen. Nextcloud wordt gepresenteerd als een belangrijke stap in dit proces.



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

Er is al veel interesse onder onderwijs- en onderzoeksinstellingen voor het gebruik van Nextcloud. In januari 2026 start een pilot, waarbij instellingen het platform een jaar lang kunnen testen. Gedurende deze periode zal Nextcloud worden aangeboden naast de bestaande samenwerkingsomgevingen die al in gebruik zijn. Het gebruik van Nextcloud biedt de instellingen de mogelijkheid om meer controle te krijgen over hun digitale infrastructuur, wat bijdraagt aan de versterking van hun digitale soevereiniteit.

AP adviseert ouders om geen foto's van kinderen te delen voor winactie

De Autoriteit Persoonsgegevens (AP) heeft ouders geadviseerd om geen foto's en persoonlijke gegevens van hun kinderen te delen als onderdeel van een winactie van speelgoedketen Intertoys. De actie biedt ouders de mogelijkheid om speelgoed te winnen door een verlanglijstje van hun kinderen op social media te delen. Dit verlanglijstje vraagt niet alleen om de verlangens van het kind, maar ook om persoonlijke informatie zoals de leeftijd, woonplaats, telefoonnummer en e-mailadres van het kind. Hoewel deelname aan de winactie zonder het invullen van deze gegevens mogelijk is, hebben veel ouders ervoor gekozen om deze informatie toch in te vullen, vaak vergezeld van een foto van hun kind met het verlanglijstje.

De AP benadrukt dat het delen van foto's van kinderen op sociale media risico's met zich meebrengt, omdat deze beelden door derden voor andere doeleinden kunnen worden gebruikt, zonder dat ouders daar controle over hebben. Ook het delen van contactgegevens van kinderen wordt afgeraden, aangezien dit hen kwetsbaar maakt voor mogelijke contactpogingen door kwaadwillenden. Intertoys heeft laten weten dat zij de ontvangen foto's enkel voor de actie willen gebruiken, en dat ouders via privéberichten worden geïnformeerd om persoonlijke gegevens van hun kinderen offline te verwijderen. Ondanks deze waarschuwingen blijven er online foto's van kinderen met persoonlijke gegevens circuleren, ook van winacties van voorgaande jaren.

Data-inbreuk bij vastgoedfinancieringsbedrijf SitusAMC: Klantgegevens gecompromitteerd

SitusAMC, een belangrijke aanbieder van back-end diensten voor grote banken en kredietverstrekkers, heeft een datalek gemeld waarbij klantgegevens zijn blootgesteld. Het bedrijf, dat actief is in de vastgoedfinanciering en jaarlijks zo'n 1



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

miljard dollar omzet genereert, ontdekte de inbreuk begin november 2025. Onder de getroffen data bevinden zich niet alleen interne gegevens van klanten, zoals boekhoudkundige en juridische documenten, maar ook gevoelige informatie van eindklanten.

De inbreuk werd op 12 november opgemerkt en werd drie dagen later bevestigd. Het bedrijf, dat onder meer samenwerkt met financiële grootmachten zoals Citi, Morgan Stanley en JPMorgan Chase, gaf aan dat er geen encryptiemalware was ingezet en dat de bedrijfsactiviteiten niet waren verstoord. Wel werd het onderzoek in samenwerking met externe experts voortgezet, terwijl getroffen klanten op de hoogte werden gehouden van de voortgang.

SitusAMC heeft verklaard dat de inbreuk geen invloed heeft gehad op de operationele continuïteit, maar de gevolgen voor de betrokken klanten zijn nog niet volledig duidelijk. De impact van het datalek is op dit moment moeilijk in te schatten, en het onderzoek is nog in volle gang. Dit incident benadrukt de risico's rondom gegevensbeveiliging bij dienstverleners in de financiële sector, waarbij vertrouwelijke klantinformatie kwetsbaar blijkt te zijn voor cyberaanvallen.

Hackers verkopen e-maildata van 140.000 Trezor-gebruikers voor \$50.000

Op 24 november 2025 werd op het sociale platform X (voorheen Twitter) bekendgemaakt dat een dreigingsactor een grote datadump van 140.000 e-mailadressen van Trezor-gebruikers te koop aanbiedt. De e-mailadressen zouden zijn verkregen via een datalek bij Mailchimp, een populair e-mailmarketingplatform. De prijs voor de volledige data zou maar liefst \$50.000 bedragen.

Het datalek zou niet direct bij Trezor zelf hebben plaatsgevonden, maar bij Mailchimp, dat diensten levert aan Trezor voor het versturen van e-mailcommunicatie naar klanten. De getroffen e-mailadressen kunnen worden gebruikt voor verschillende kwaadaardige doeleinden, zoals phishing-aanvallen of andere vormen van social engineering. De hackers bieden de gegevens te koop aan op darkweb-marktplaatsen, wat de dreiging voor de slachtoffers aanzienlijk vergroot.

Het is niet bekend of de e-mailadressen afkomstig zijn van actieve Trezor-gebruikers, maar gezien het feit dat Trezor een bekend merk is in de cryptogemeenschap, zijn de potentiële risico's voor de slachtoffers groot. Het is belangrijk dat gebruikers van Trezor, of andere cryptodiensten die Mailchimp gebruiken, extra voorzichtig zijn met



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

verdachte e-mails en berichten. De situatie benadrukt opnieuw de kwetsbaarheid van e-maildiensten en de risico's van het delen van gegevens met derde partijen.