



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

19-11-2025:

Wereldwijd websites uit de lucht door storing bij Cloudflare

Op 18 november 2025 heeft een technische storing bij netwerkbedrijf Cloudflare wereldwijd geleid tot onbereikbaarheid van verschillende websites. Cloudflare, een van de grootste bedrijven op het gebied van netwerkbeveiliging, speelt een cruciale rol in het beschermen van miljoenen websites tegen digitale aanvallen. De oorzaak van de storing is nog onbekend, maar Cloudflare heeft aangegeven op de hoogte te zijn van het probleem en is bezig met onderzoek. In Nederland hebben onder andere de websites van NH Nieuws, AT5 en het Financieele Dagblad te maken gehad met downtime. Ook de websites van politieke partijen zoals VVD en D66, evenals de stichting Alzheimer Nederland, waren niet bereikbaar. Het platform Allestoringen.nl heeft inmiddels meer dan 2000 meldingen ontvangen van gebruikers die moeite hadden met het laden van sites. De storing heeft wereldwijd gevolgen, aangezien vele andere populaire sites ook werden getroffen. Cloudflare heeft bevestigd dat het probleem actief wordt onderzocht en er nog geen exacte oplossing is gepresenteerd. De impact van de storing blijft groot, aangezien veel bedrijven en organisaties wereldwijd afhankelijk zijn van de diensten van Cloudflare voor hun digitale aanwezigheid. De incidenten benadrukken de kwetsbaarheid van de infrastructuur die het internet ondersteunt.

Storing bij Cloudflare raakt wereldwijd websites en apps

Op 18 november 2025 heeft een grote storing bij Cloudflare wereldwijd geleid tot verstoringen van vele websites en apps, waaronder populaire platforms zoals X en ChatGPT. De storing begon rond 12:30 uur Nederlandse tijd en resulteerde in wijdverspreide 500-foutmeldingen bij gebruikers van Cloudflare-diensten. Het probleem betrof een storing in het Global Network van Cloudflare, die meerdere klanten beïnvloedde en zorgde voor onbereikbaarheid van diverse websites. Bovendien waren de Cloudflare Dashboard en de Cloudflare API tijdelijk niet beschikbaar.

Als gevolg van de storing werden duizenden meldingen ingediend via DownDetector.com, die de omvang van het probleem aantoonde. Diverse websites die gebruik maken van Cloudflare voor content delivery network (CDN) services en DDoS-bescherming waren niet bereikbaar of traag in het laden. Op fora zoals Hacker



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

News en Reddit kwamen honderden reacties binnen van verontruste gebruikers die de gevolgen van de storing ondervonden.

Cloudflare heeft de storing inmiddels verholpen en meldde dat het probleem werd geïdentificeerd en opgelost. Details over de oorzaak van de storing zijn voornamelijk niet beschikbaar. De impact van de storing heeft echter opnieuw het belang van robuuste infrastructuur voor de online dienstverlening benadrukt.

Ziekenhuizen in Eindhoven en Veldhoven schorten afspraken op door netwerkstoring

Het Máxima Medisch Centrum (MMC), met vestigingen in Eindhoven en Veldhoven, heeft vanwege een ernstige netwerkstoring alle geplande afspraken, behandelingen en operaties voor dinsdag 18 november moeten annuleren. Het ziekenhuis heeft aangegeven dat door de storing toegang tot elektronische patiëntendossiers onmogelijk is, wat ernstige gevolgen heeft voor de zorgverlening. Hoewel de storingen begonnen zijn met telefonieproblemen, heeft het MMC benadrukt dat er geen sprake is van een hack, maar van een technisch probleem binnen het netwerk. Het ziekenhuis is momenteel alleen bereikbaar voor spoedeisende vragen en roept patiënten op geen niet-spoedeisende vragen te stellen om de beschikbare middelen niet verder te belasten. De storingen hebben invloed op de toegang tot verschillende systemen, waardoor het ziekenhuis tijdelijk niet in staat is om zorg te leveren zoals gebruikelijk. Medewerkers kunnen bijvoorbeeld geen gegevens in patiëntendossiers raadplegen, wat de zorg voor patiënten bemoeilijkt. Het MMC benadrukt dat de situatie nauwlettend wordt gevolgd en dat spoedeisende hulp bij andere locaties kan worden geboden als dat nodig blijkt.

Storing bij Odido en Ziggo verholpen na grote uitval in Nederland

Op 18 november 2025 kampten meerdere Nederlandse providers, waaronder Odido en Ziggo, met een grote storing die duizenden klanten in het hele land trof. Klanten meldden op verschillende platformen, zoals Allestoringen.nl, dat ze geen internetverbinding konden maken en problemen ondervonden bij het tv-kijken. De storing begon in de vroege ochtenduren en bereikte zijn piek rond 9.30 uur, met meldingen die opliepen tot bijna tienduizend voor Odido en enkele duizenden voor Ziggo. De oorzaak van de storing bleek een stroomstoring in een datacenter te zijn, wat leidde tot de tijdelijke uitval van de diensten van deze providers. Zowel Odido als Ziggo gaven aan dat de problemen werden veroorzaakt door deze onvoorziene



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

technische storing, die het netwerk ernstig verstoorde. Technici van de providers werden snel ingezet om de situatie te verhelpen, en na enkele uren waren de meeste diensten weer operationeel. De storing had niet alleen impact op de telefonie- en internetdiensten van Odido en Ziggo, maar ook op andere diensten zoals internetbankieren. ASN Bank ondervond op hetzelfde moment ook technische problemen door de landelijke netwerkuitval, hoewel klanten geen problemen ervaarden met pinbetalingen. De banken gaven aan dat het probleem was opgelost zodra de verbindingen van de providers waren hersteld. Dit incident benadrukt de kwetsbaarheid van digitale diensten, die steeds meer afhankelijk zijn van betrouwbare datacenters en netwerkverbindingen. De bedrijven werken nu aan maatregelen om de kans op soortgelijke storingen in de toekomst te verkleinen. De oorzaak van de stroomstoring wordt verder onderzocht om herhaling te voorkomen.

Cloudflare wijt wereldwijde storing aan latente bug

Op 18 november 2025 kwam het wereldwijd tot een verstoring van verschillende internetdiensten door Cloudflare, een belangrijke speler in het beheer van webverkeer en bescherming tegen aanvallen. Volgens Cloudflare-CTO Dane Knecht was de oorzaak van de storing een "latente bug" die tot problemen leidde bij de botmitigatie-oplossing van het bedrijf. De bug crashte na een routine configuratie-aanpassing, wat resulteerde in verstoringen in de diensten van Cloudflare, waaronder traag ladende of niet werkende websites en applicaties.

De storing begon rond 12:30 uur Nederlandse tijd en trof een breed scala aan bedrijven, organisaties en websites die afhankelijk zijn van Cloudflare. In een verklaring op X (voorheen Twitter) benadrukte Knecht dat het niet om een cyberaanval ging. Ondanks de transparantie over de oorzaak, gaf hij aan dat de impact en de tijd die het kostte om het probleem op te lossen onacceptabel waren. Cloudflare kondigde maatregelen aan om herhaling van dergelijke incidenten in de toekomst te voorkomen.

EU werkt aan digitale soevereiniteit om technologische afhankelijkheid te verminderen

De Europese Unie kampt met de toenemende afhankelijkheid van technologie van buitenlandse partijen, met name uit de Verenigde Staten en China. Dit creëert veiligheidsrisico's voor kritieke infrastructuur en Europese digitale soevereiniteit. Om



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

deze afhankelijkheid te verkleinen, neemt de EU serieuze stappen door de oprichting van strategische initiatieven en samenwerkingen te bevorderen. Een recente bijeenkomst in Berlijn bracht vertegenwoordigers van EU-lidstaten, het bedrijfsleven en academische instellingen samen om de technologische autonomie van Europa te versterken. Onderwerpen als de inzet van opensource software, digitale identificatie, en de oprichting van 'EU Inc.' werden besproken om bedrijven binnen de EU te helpen gemakkelijker over landsgrenzen heen te opereren zonder extra kosten. Deze initiatieven kunnen de veerkracht van de Europese digitale infrastructuur tegen cyberdreigingen verbeteren. Daarnaast investeert het moederbedrijf van supermarktketen Lidl 11 miljard euro in een datacenter in Berlijn om de rekenkracht te vergroten en toekomstige ontwikkelingen op het gebied van kunstmatige intelligentie beter te ondersteunen, zonder afhankelijk te zijn van Amerikaanse of Chinese technologie. Deze stappen zijn een reactie op de groeiende geopolitieke druk, waarbij de invloed van buitenlandse techbedrijven op de Europese cyberspace steeds groter wordt. De EU wil haar digitale autonomie vergroten om de veiligheid van haar digitale netwerken en kritieke infrastructuren te waarborgen. Hoewel veel van deze initiatieven nog in een beginstadium verkeren, is het duidelijk dat de EU nu serieus werk maakt van het versterken van haar digitale soevereiniteit. UNC1549 hackers gebruiken aangepaste tools voor aanvallen op lucht- en defensiesector

Sinds medio 2024 voert de Iraans gesteunde hacker groep UNC1549 gerichte aanvallen uit op organisaties in de luchtvaart-, luchtvaart- en defensiesector wereldwijd. De groep maakt gebruik van een geavanceerde benadering die zowel phishing-aanvallen als de uitbuiting van vertrouwde connecties tussen de primaire doelwitten en hun derde leveranciers combineert. Deze strategie blijkt bijzonder effectief tegen goed beveiligde organisaties zoals defensiecontractanten, die vaak hun leveranciers als zwakkere schakels in het netwerk gebruiken, wat hen tot de eerste doelwitten maakt voor inbreuken. UNC1549 heeft zijn methoden in de loop van de tijd verfijnd en toont een aanzienlijke mate van operationele en tactische complexiteit. De aanvallers gebruiken zorgvuldig samengestelde phishing e-mails, die specifiek gericht zijn op de rollen van de slachtoffers, om aanvankelijk toegang te verkrijgen tot netwerken. Zodra de hackers toegang hebben verkregen, maken ze gebruik van creatieve technieken voor laterale beweging, waaronder het stelen van broncode van slachtoffers om spear-phishing campagnes op te zetten met lookalike-domeinen die beveiligingsproxies kunnen omzeilen. Daarnaast misbruiken ze interne ticketingsystemen om inloggegevens van nietsvermoedende medewerkers te verkrijgen. Google Cloud-analisten hebben vastgesteld dat UNC1549 aangepaste tools inzet die speciaal zijn ontworpen om detectie te omzeilen en forensisch



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

onderzoek te bemoeilijken. Elk exploitatiepayload dat tijdens de onderzoeken werd geïdentificeerd, had een unieke hash, zelfs wanneer meerdere exemplaren van dezelfde achterdeurvariant binnen een enkel slachtoffer netwerk werden aangetroffen. Deze mate van aanpassing benadrukt de aanzienlijke middelen en de toewijding van de groep aan operationele beveiliging. Een van de meest technisch significante aspecten van de operaties van UNC1549 betreft hun gebruik van search order hijacking voor malwarepersistentie. Deze techniek houdt in dat kwaadaardige DLL-bestanden in legitieme software-installatiemapjes worden geplaatst, zodat aanvallers blijvende uitvoering kunnen verkrijgen wanneer beheerders of gebruikers de legitieme software uitvoeren.

De groep heeft met succes deze kwetsbaarheid geëxploiteerd in veelgebruikte bedrijfsoplossingen zoals FortiGate, VMWare, Citrix, Microsoft en NVIDIA uitvoerbare bestanden. Bij het verkrijgen van initiële toegang installeren de aanvallers opzettelijk legitieme software om misbruik te maken van deze zoekvolgorde-hijacking mogelijkheid. De TWOSTROKE backdoor, een aangepaste C++-achterdeur, is een voorbeeld van deze technische verfijning. Het communiceert via SSL-versleutelde TCP-verbindingen op poort 443, wat het moeilijk maakt om het van legitiem verkeer te onderscheiden. TWOSTROKE genereert een unieke slachtoffer-ID door de volledig gekwalificeerde DNS-computernaam op te halen en deze via XOR-encryptie om te zetten naar een bot-ID.

UNC1549's campagne is gericht op langdurige persistente toegang en anticipeert op de reacties van onderzoekers. Ze implementeren strategisch achterdeuren die maandenlang inactief blijven, waarbij ze alleen worden geactiveerd nadat slachtoffers pogingen ondernemen om de situatie te verhelpen. Dit, gecombineerd met het uitgebreide gebruik van reverse SSH-shells en domeinen die de industrieën van de slachtoffers nabootsen, creëert een complexe omgeving voor verdedigers om zich tegen te beschermen.

Lazarus APT-groep onthult nieuwe ScoringMathTea RAT voor remote command execution

De Lazarus APT-groep, gelieerd aan de Noord-Koreaanse overheid, heeft een nieuwe versie van hun Remote Access Trojan (RAT) geïntroduceerd, genaamd ScoringMathTea. Dit malwareprogramma maakt deel uit van een bredere cyberaanval, bekend als Operation DreamJob, en richt zich met name op bedrijven die technologieën voor onbemande luchtvaartuigen (UAV) leveren aan Oekraïne. Het



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

doel van deze aanvallen is het stelen van gevoelige productiekennis en intellectueel eigendom. ScoringMathTea is geavanceerd opgebouwd en kan op verschillende manieren schade aanrichten. Het biedt aanvallers volledige controle over geïnfecteerde systemen, waarmee ze op afstand opdrachten kunnen uitvoeren en malware kunnen laden zonder sporen achter te laten. Het malwareprogramma maakt gebruik van meerdere technische mechanismen, zoals dynamische API-resolutie en encryptie via meerdere lagen, om detectie te omzeilen. Door gebruik te maken van technieken zoals polymorfe vervorming en een gepersonaliseerd substitutiecijfer wordt het extreem moeilijk om de code te analyseren.

Een van de meest geavanceerde kenmerken van ScoringMathTea is de mogelijkheid om plugins in het geheugen te laden en uit te voeren. Dit maakt het voor aanvallers mogelijk om hun code volledig in het geheugen te draaien, zonder dat er bestanden op de harde schijf van de geïnfecteerde machine worden achtergelaten. Dit biedt een langdurige toegang tot de systemen van de slachtoffers, wat het malware bijzonder gevaarlijk maakt voor bedrijven die kwetsbaar zijn voor deze aanvallen. Deze nieuwe vorm van RAT benadrukt de evolutie van Lazarus' aanvallen, waarbij steeds geavanceerdere technologieën worden ingezet om vertrouwelijke informatie te stelen en de werking van doelwitten te verstoren. Beveiligingsexperts waarschuwen dat deze malware een aanzienlijke bedreiging vormt, niet alleen vanwege de geavanceerde technieken, maar ook vanwege het specifieke doelwit van de aanval: bedrijven die kritieke technologie leveren aan Oekraïne in het kader van de voortdurende conflicten.

MI5 waarschuwt Britse parlementariërs voor Chinese spionage via LinkedIn

De Britse veiligheidsdienst MI5 heeft een belangrijke waarschuwing afgegeven met betrekking tot Chinese spionage-inspanningen via LinkedIn. Twee personen van de Chinese inlichtingendienst zouden actief parlementsleden en andere belangrijke afgevaardigden benaderen via de netwerkwebsite LinkedIn. Deze nep-headhunters zijn bezig met het verkrijgen van vertrouwelijke en gevoelige informatie door contact te leggen met individuen in politieke en economische sectoren, waaronder medewerkers van denktanks.

MI5 waarschuwt de betrokkenen voor aanbiedingen die ogenschijnlijk onschuldige voordelen beloven, zoals volledig betaalde reizen naar China, maar in werkelijkheid gericht zijn op het verkrijgen van insider-informatie. Naast het aanbieden van reizen, worden er ook betalingen gedaan via contant geld of cryptovaluta als beloning voor



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

de verstrekte informatie. De veiligheidsdienst heeft bevestigd dat de LinkedIn-accounts van de betreffende individuen onder beheer staan van de Chinese veiligheidsdienst. Dit incident komt op het moment dat er al eerder zorgen waren over Chinese spionage via LinkedIn, waarbij vrouwelijke agenten werden ingezet om contact te leggen met duizenden Britten die werkzaam zijn in gevoelige sectoren. De MI5 benadrukt dat deze pogingen gericht zijn op het verzamelen van waardevolle informatie die kan worden ingezet voor politieke en economische invloed. De Britse overheid heeft dan ook herhaaldelijk gewaarschuwd voor het risico van spionage via sociale netwerken en de rol van digitale platforms in moderne spionageactiviteiten.

Google brengt updates uit voor actief aangevallen kwetsbaarheid in Chrome

Google heeft recent beveiligingsupdates vrijgegeven om een kwetsbaarheid in de JavaScript-engine V8 van de Chrome-browser te verhelpen. Het probleem, aangeduid als CVE-2025-13223, wordt momenteel actief misbruikt door cybercriminelen. De kwetsbaarheid bevindt zich in V8, die wordt gebruikt door Chrome en andere browsers voor het uitvoeren van JavaScript-code. Google heeft de kwetsbaarheid als 'high' geclassificeerd, wat betekent dat aanvallers mogelijk code kunnen uitvoeren binnen de context van de browser. Dit kan leiden tot het uitlekken of manipuleren van gevoelige gegevens, zoals inloggegevens of persoonlijke informatie.

De kwetsbaarheid werd op 12 november gemeld door de Google Threat Analysis Group, die verantwoordelijk is voor het detecteren van aanvallen door overheden en andere gerichte aanvallers. Hoewel Google geen specifieke details heeft gegeven over de waargenomen aanvallen, is de impact aanzienlijk. Gebruikers kunnen gevoelige informatie verliezen of last krijgen van andere schadelijke gevolgen. Het lek kan niet op zichzelf een systeem volledig overnemen, maar zou, in combinatie met andere kwetsbaarheden, wel kunnen worden gebruikt voor verdere aanvallen. Om het probleem te verhelpen, heeft Google de update naar versie 142.0.7444.175/176 van Chrome uitgerold voor Windows en macOS, en versie 142.0.7444.175 voor Linux. De update wordt doorgaans automatisch geïnstalleerd, maar gebruikers kunnen handmatig controleren of zij de nieuwste versie hebben geïnstalleerd. Voor Microsoft Edge, dat is gebaseerd op dezelfde Chromium-engine als Chrome, is er op dit moment nog geen update beschikbaar. Eerdere waarschuwingen van Google voor vergelijkbare kwetsbaarheden in juni en



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

september van dit jaar benadrukken de voortdurende dreiging van aanvallen op de V8-engine.

DoorDash kwetsbaarheid in e-mailspoofing vormt risico voor phishingaanvallen

Een kwetsbaarheid in het systeem van DoorDash maakte het mogelijk voor kwaadwillenden om e-mails te versturen die volledig op officiële DoorDash-correspondentie leken. Deze kwetsbaarheid stelde aanvallers in staat om via het DoorDash for Business-platform e-mails te genereren die uit de naam van het bedrijf werden verzonden, zonder dat ze gemarkeerd werden als verdachte berichten. Dit creëerde een ideale mogelijkheid voor phishingcampagnes, waarin aanvallers zich voordeden als DoorDash om vertrouwelijke informatie van slachtoffers te verkrijgen. De kwetsbaarheid werd ontdekt door een beveiligingsonderzoeker die aantoonde dat een invoerveld voor 'budgetnaam' op het platform onvoldoende werd gecontroleerd. Hierdoor konden aanvallers kwaadaardige HTML-code in de e-mail plaatsen, wat leidde tot de mogelijkheid om onveilige links of bijlagen te verspreiden. Het systeem van DoorDash stuurde de e-mail vervolgens door zonder enige waarschuwing, waardoor het leek alsof het een betrouwbare communicatie was van het bedrijf. Na de ontdekking van de kwetsbaarheid werd deze snel gepatcht, maar de communicatie tussen de onderzoeker en DoorDash leidde tot een controversiële situatie. De onderzoeker klaagde dat het bedrijf de kwetsbaarheid te lang had genegeerd, wat mogelijk risico's voor de digitale veiligheid heeft vergroot. De kwestie legt de spanningen bloot over de ethiek van kwetsbaarheidsrapportage, waarbij bedrijven en beveiligingsonderzoekers soms verschillende verwachtingen hebben over de omgang met ontdekte problemen.

Microsoft Azure getroffen door 15 Tbps DDoS-aanval via 500.000 IP-adressen

Op 17 november 2025 meldde Microsoft dat zijn Azure-netwerk was getroffen door een enorme DDoS-aanval van maar liefst 15,72 terabit per seconde (Tbps), uitgevoerd door het Aisuru-botnet. De aanval werd uitgevoerd vanaf meer dan 500.000 IP-adressen en maakte gebruik van extreem hoge UDP-stromen. Deze stroom was gericht op een specifiek publiek IP-adres in Australië en bereikte bijna 3,64 miljard pakketten per seconde (bps).

Het Aisuru-botnet is een Turbo Mirai-achtig botnet dat bekend staat om zijn recordbrekende DDoS-aanvallen. Het botnet maakt gebruik van gecompromitteerde



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

IoT-apparaten, zoals routers en camera's, voornamelijk in de Verenigde Staten en andere landen. De recente aanval werd gekarakteriseerd door plotselinge UDP-bursts met minimale bronspoofing en het gebruik van willekeurige bronpoorten, wat de opsporing vereenvoudigde en de handhaving door providers vergemakkelijkte. Eerder dit jaar, in september 2025, werd hetzelfde botnet verantwoordelijk gehouden voor een DDoS-aanval van 22,2 Tbps, die een recordaantal van 10,6 miljard pakketten per seconde bereikte. Ondanks de enorme kracht van deze aanvallen, kon Azure de aanval snel neutraliseren, zonder langdurige gevolgen voor de service. Aisuru's operaties omvatten niet alleen DDoS-aanvallen, maar ook pogingen om de DNS-diensten van Cloudflare te verstoren door via kwaadaardige query-verzoeken de populariteit van zijn eigen domeinen te verhogen. Dit gedrag leidde tot een verstoring van het domeinrangschikkingssysteem van Cloudflare, dat daarop besloot verdachte domeinen te verwijderen om verdere schade te voorkomen.

Microsoft en Cloudflare blijven waakzaam tegenover deze nieuwe generatie botnets, die de grens van traditionele DDoS-aanvallen verleggen en nieuwe uitdagingen presenteren voor het beschermen van kritieke infrastructuren.

0-Day kwetsbaarheid in Microsoft-ondertekende driver mogelijk voor EDR/XDR bypass

Op 17 november 2025 werd bekend dat er een 0-day kwetsbaarheid is ontdekt in een door Microsoft ondertekende driver, die mogelijk wordt misbruikt om Endpoint Detection and Response (EDR) en Extended Detection and Response (XDR) systemen te omzeilen. Deze driver zou op verschillende darkwebmarkten beschikbaar zijn, wat een verhoogd risico vormt voor organisaties die afhankelijk zijn van deze beveiligingssystemen. De kwetsbaarheid stelt aanvallers in staat om de beveiligingsmechanismen van EDR- en XDR-oplossingen te omzeilen, waardoor ze onopgemerkt toegang kunnen krijgen tot netwerken. Het gebruik van een legitieme, door Microsoft goedgekeurde driver maakt het moeilijker voor traditionele beveiligingssystemen om de aanval te detecteren, wat de dreiging vergroot. Omdat het gaat om een 0-day kwetsbaarheid, is er op dit moment geen patch beschikbaar om het probleem te verhelpen. Dit onderstreept het belang voor organisaties om proactief te blijven in hun beveiligingsmaatregelen en alert te zijn op potentiële misbruikscenario's van dergelijke kwetsbaarheden.



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

Pig-butcheringsfraude groeit door gebruik van AI-assistenten

Pig-butcheringsfraude, een van de meest schadelijke vormen van cybercriminaliteit, heeft zich op wereldwijde schaal uitgebreid en veroorzaakt jaarlijks miljarden dollars aan schade. Deze langdurige beleggingsfraudes draaien om het opbouwen van vertrouwen via emotionele manipulatie en valse handelsplatformen, waarna de slachtoffers hun levensspaargeld verliezen. Tegenwoordig opereren de fraudeurs op industriële schaal, waarbij criminele netwerken gebruik maken van geavanceerde technologieën om hun bereik te vergroten en het succes van hun operaties te verbeteren. Een opvallend kenmerk van moderne pig-butcheringsoperaties is het gebruik van kunstmatige intelligentie (AI) om geloofwaardige valse identiteiten te creëren en om meerdere gesprekken met slachtoffers gelijktijdig te onderhouden. Scammers maken gebruik van door AI gegenereerde foto's om overtuigende online profielen te maken, waardoor het voor slachtoffers vrijwel onmogelijk wordt om nep-profielen te herkennen. De AI-tools die worden ingezet, kunnen realistische afbeeldingen van niet-bestaande personen creëren, compleet met verschillende poses en achtergronden. Dit helpt de fraudeurs om betrouwbaar over te komen op datingapps en sociale-mediaplatforms. De technologie stelt hen in staat om snel en op grote schaal geloofwaardige profielen te maken die een basis visuele inspectie door potentiële slachtoffers doorstaan.

Volgens onderzoekers van Cyfirma vertrouwen criminele netwerken steeds meer op AI-ondersteunde identiteitsfabricage en geautomatiseerde berichtgeneratie om hun operaties op te schalen. Deze technologie maakt het mogelijk om grote aantallen slachtoffers gelijktijdig te benaderen, terwijl de operationele continuïteit wordt gewaarborgd, zelfs wanneer individuele accounts of domeinen geblokkeerd worden.

De technische infrastructuur die deze fraude ondersteunt, omvat geavanceerde back-end systemen die gebruik maken van klantrelatiebeheertools (CRM) om het gedrag van slachtoffers te volgen en waardevolle doelwitten te identificeren. Deze systemen voeren automatisering uit voor het onboarden van nieuwe slachtoffers, het afhandelen van de eerste gesprekken en het genereren van geloofwaardige financiële rapportages op valse handelsplatformen. Wanneer slachtoffers proberen geld op te nemen, worden geautomatiseerde barrières zoals verificatiekosten of belastingbetalingen geactiveerd, waardoor de slachtoffers nog meer geld kwijtraken voordat de fraude uiteindelijk in elkaar stort.



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

Yurei ransomware: encrypteren van bestanden, operatie-model en datatransfer methoden ontrafeld

Yurei ransomware, een nieuwe dreiging op het gebied van ransomware, werd voor het eerst openbaar geïdentificeerd in september 2025. Deze Go-gebaseerde malware volgt een typisch ransomware-operatiemodel door bedrijfsnetwerken binnen te dringen, belangrijke gegevens te versleutelen, back-ups te verwijderen en losgeld te eisen voor de gestolen informatie. De groep achter de aanval opereert via een gespecialiseerde dark web-site waar ze slachtoffers benaderen en de betalingsvoorwaarden bespreken op basis van de financiële situatie van het doelwit. Bekende slachtoffers van de Yurei-aanvallen bevinden zich in Sri Lanka en Nigeria, met voornamelijk getroffen sectoren zoals transport en logistiek, IT-software, marketing en reclame, en de voedsel- en drankindustrie. In tegenstelling tot veel moderne ransomware-operaties is er geen duidelijk bewijs dat Yurei samenwerkt met andere cybercriminaliteitsgroepen of gebruikmaakt van Ransomware-as-a-Service (RaaS)-modellen. De dreigingsactoren stellen de losgelddbetalingen vast op basis van een case-by-case analyse van de financiële situatie van het slachtoffer, hoewel de exacte bedragen nog niet openbaar zijn gemaakt. Onderzoekers van ASEC hebben ontdekt dat Yurei ransomware zich onderscheidt door een geavanceerde encryptiemethode. De malware maakt gebruik van het ChaCha20-Poly1305-algoritme voor bestandversleuteling, waarbij een 32-byte sleutel en een 24-byte nonce als willekeurige waarden worden gegenereerd. Deze sleutels worden vervolgens beschermd door de secp256k1-ECIES-methode met een ingebedde publieke sleutel, zodat alleen de dreigingsactor met de overeenkomstige privésleutel de bestanden kan ontsleutelen. De encryptie begint met Yurei die het geïnfecteerde systeem scant om alle beschikbare schijven en mogelijke versleutelingdoelen te identificeren. Kritieke systeemdirectory's zoals Windows, System32 en Program Files worden doelbewust uitgesloten om volledige systeemstoringen te voorkomen. Bestanden met extensies zoals .sys, .exe, .dll en .Yurei (de eigen versleutelde bestandsmarker) worden ook overgeslagen om herhaalde versleuteling te voorkomen. De versleuteling van bestanden gebeurt in blokken van 64 KB, waarbij de versleutelde sleutel en nonce aan het begin van elk bestand worden opgeslagen met de "||"-scheidingsteken. De ransom note, opgeslagen als "_README_Yurei.txt", dreigt met het verwijderen van de ontsleutelsleutel en het lekken van gestolen gegevens, waaronder databases, financiële documenten en persoonlijke informatie, op het dark web als het slachtoffer niet binnen vijf dagen reageert.



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

Payroll Pirates: netwerk van criminele groepen kapen salaris systemen

Een nieuw en geavanceerd cybercrimineel netwerk, bekend als Payroll Pirates, heeft sinds medio 2023 wereldwijd doelwitten aangevallen. Het netwerk richt zich voornamelijk op salarisadministraties, kredietunies en handelsplatformen, met als doel het stelen van inloggegevens van werknemers en het omleiden van salarisbetalingen naar de bankrekeningen van de aanvallers. De criminele groep maakt gebruik van malvertising, waarbij valse advertenties verschijnen op zoekmachines die gebruikers naar phishingwebsites lokken.

Wanneer werknemers hun inloggegevens invoeren op deze nep-websites, worden de gegevens direct naar de aanvallers gestuurd. Door deze tactiek heeft het netwerk inmiddels meer dan 200 platforms aangevallen en meer dan 500.000 gebruikers getroffen. De aanvallen begonnen met Google Ads, die valse links naar payroll websites promootten. Gebruikers die hun bedrijfspagina zochten, werden via gesponsorde advertenties naar phishingwebsites geleid die identiek waren aan de legitieme inlogpagina's. In juni 2024 werd het netwerk opnieuw actief, met verbeterde tools die in staat waren om twee-factor-authenticatie te omzeilen. Dit werd mogelijk gemaakt door real-time interactie via Telegram-bots, die slachtoffers vroegen om verificatiecodes zodra ze hun wachtwoord hadden ingevoerd. De aanvallers gebruikten hierbij verborgen PHP-scripts die de gestolen gegevens zonder detectie naar de aanvallers stuurden. Het gevaar van deze operaties ligt in de snelheid en verfijning waarmee de aanvallers werken. Zodra een slachtoffer zijn inloggegevens invoert, worden deze onmiddellijk naar de criminelen gestuurd, die via een Telegram-bot in real-time verder communiceren met de slachtoffers om extra informatie te verkrijgen, zoals beveiligingscodes en antwoorden op beveiligingsvragen. Dit maakt het voor de slachtoffers bijna onmogelijk om te herkennen dat ze het slachtoffer zijn van een aanval voordat het te laat is.

Xanthorox: Een AI-tool voor het genereren van kwaadaardige code op de donkere webmarkten

Xanthorox is een nieuwe, verontrustende AI-tool die snel verspreid is binnen de darknet gemeenschappen en wordt gebruikt voor het creëren van schadelijke code op aanvraag. Dit platform, dat wordt vergeleken met reguliere chatbots zoals ChatGPT, mist de veiligheidsbeperkingen die doorgaans dergelijke systemen beschermen. Het werd in oktober 2024 voor het eerst geïntroduceerd op een privé Telegram-kanaal en breidde zich al snel uit naar diverse darknet-forums, waar het



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

sindsdien aan populariteit wint. Xanthorox is niet zomaar een tool voor ethisch hacken, zoals de maker beweert, maar heeft veel geavanceerdere capaciteiten die het gevaarlijk maken voor de cybersecuritygemeenschap. De tool kan eenvoudig door gebruikers worden aangestuurd met tekstinstructies om malware en ransomware te genereren. Dit maakt het mogelijk voor zelfs technisch minder onderlegde individuen om geavanceerde schadelijke software te ontwikkelen. Een bijzonder zorgwekkende functie van Xanthorox is de "Agentex"-versie, waarbij gebruikers eenvoudig opdrachten kunnen geven zoals "geef me ransomware die dit doet", waarna de tool automatisch werkende uitvoerbare code genereert op basis van deze input.

Het platform biedt een maandelijks abonnement voor \$300 en een jaarlijks abonnement voor \$2.500, met betalingen die uitsluitend via cryptocurrency worden uitgevoerd, wat het moeilijk maakt om de identiteit van de gebruikers en ontwikkelaars te traceren. De tool is volledig zelfvoorzienend en draait op specifieke servers, wat het moeilijk maakt voor onderzoekers om het systeem te stoppen of te controleren. Onderzoekers van Trend Micro ontdekten dat Xanthorox goed gecommuniceerde, functionele code produceert die klaar is voor gebruik of als basis kan dienen voor complexere aanvallen. Het gegenereerde malware kan verschillende programmeertalen gebruiken, zoals C/C++ en JavaScript, en heeft de capaciteit om encryptie en obfuscatie toe te passen. Dit maakt het moeilijk voor beveiligingssystemen om de code te detecteren en te neutraliseren. Ondanks deze geavanceerde functionaliteit, heeft de tool beperkingen. Zo kan Xanthorox niet in real-time gegevens van het internet of het dark web ophalen, wat zijn nut voor verkenning of gegevensverzameling beperkt. Deze ontdekking benadrukt de potentiële dreiging die Xanthorox vormt voor zowel individuen als organisaties. Hoewel de maker van de tool beweert dat het bedoeld is voor penetratietesten, blijkt uit het gebruik en de functionaliteit van het platform dat het voornamelijk wordt ingezet voor het ontwikkelen van schadelijke software. De impact van deze tool op de bredere cyberbeveiligingsgemeenschap is groot, aangezien het de toegang tot geavanceerde aanvalsmethoden vergemakkelijkt voor kwaadwillende actoren.

Dreiging door Lynx ransomware via gecompromitteerde RDP-inloggegevens na verwijdering van serverback-ups

Lynx ransomware vormt een steeds grotere bedreiging voor bedrijfsnetwerken, waarbij een toenemende focus ligt op zowel datadiefstal als de vernietiging van



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

infrastructuur. De aanvalsmethoden die worden ingezet zijn steeds geavanceerder, waarbij aanvallers zich niet alleen richten op het versleutelen van bestanden, maar ook zorgvuldig de netwerkomgeving analyseren en kwetsbaarheden benutten voordat de ransomware daadwerkelijk wordt ingezet. Een recente aanvalsgolf heeft aangetoond hoe aanvallers misbruik maken van gecompromitteerde Remote Desktop Protocol (RDP)-inloggegevens. Deze gegevens zijn vaak het resultaat van malware-infecties, datalekken of het werk van zogenaamde initial access brokers. Wat deze aanvallen onderscheidt van andere is de uitgebreide voorbereidingstijd die de aanvallers nemen voordat ze de ransomware daadwerkelijk verspreiden. In plaats van direct over te gaan tot encryptie van systemen, besteden de aanvallers dagen aan het onderzoeken van het netwerk, het creëren van backdoors en het identificeren van waardevolle doelwitten. De aanvallers maken gebruik van een methode waarbij ze na de initiële toegang, die vaak plaatsvindt via een geldige RDP-verbinding, in korte tijd zich lateraal binnen het netwerk bewegen. Dit gebeurt bijvoorbeeld via een ander gecompromitteerd administrateuraccount, waarmee ze de controle over de netwerkstructuur kunnen verkrijgen. De aanvallers installeerden ook remote access software, zoals AnyDesk, om hun aanwezigheid te behouden, zelfs wanneer oorspronkelijke toegangspunten worden ontdekt. Een zorgwekkend aspect van deze aanval is de vernietiging van back-ups voordat de ransomware wordt uitgerold. Nadat de aanvallers gegevens hadden verzameld en voorbereid voor de dubbele extortie, verwijderden ze systematisch alle back-upbestanden. Dit vergroot de impact van de ransomware aanzienlijk, omdat de getroffen organisaties niet meer kunnen terugvallen op hun back-upsystemen om de versleutelde gegevens te herstellen. Het proces van gegevensdiefstal en de daaropvolgende vernietiging van back-ups verhoogt de druk op slachtoffers om het geëiste losgeld te betalen, aangezien alternatieve herstelopties niet meer beschikbaar zijn. De tijdlijn van de aanval toont een gedegen planning van de aanvallers, waarbij het aanvallen van back-ups en het voorbereiden van de ransomware-infectie ruim een week in beslag nam. De uiteindelijke encryptie van systemen en servers vond pas plaats nadat het netwerk grondig was geanalyseerd, wat de effectiviteit van de aanval vergrootte. Deze campagne benadrukt de noodzaak voor organisaties om niet alleen te focussen op het beschermen van netwerken tegen directe aanvallen, maar ook om robuuste back-upstrategieën te implementeren die bestand zijn tegen ransomware-aanvallen.



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

Remcos RAT C2-activiteit in kaart gebracht, met de gebruikte poorten voor communicatie

Remcos, een commercieel verkrijgbare remote access tool (RAT), die oorspronkelijk werd gepromoot als administratieve software, wordt momenteel actief misbruikt door cybercriminelen. Deze malware stelt aanvallers in staat om op afstand commando's uit te voeren, bestanden te stelen, schermen vast te leggen, toetsaanslagen te loggen en gebruikersreferenties te verzamelen via command-and-control (C2)-servers, die gebruik maken van HTTP- of HTTPS-protocollen. Ondanks dat het als legitieme software wordt gepositioneerd met zowel gratis als betaalde versies, worden ongeautoriseerde exemplaren op grote schaal gebruikt voor datadiefstal en ongeoorloofde toegang tot systemen. De verspreiding van Remcos vindt voornamelijk plaats via e-mailcampagnes die kwaadaardige bijlagen bevatten, evenals via bestanden die worden gehost op gecompromitteerde websites. Daarnaast gebruiken aanvallers gespecialiseerde loaders zoals GuLoader en Reverse Loader om Remcos als een tweede payload te leveren, waarmee ze initiële detectiesystemen kunnen omzeilen. Zodra de malware is geïnstalleerd, vestigt deze persistentie en houdt continu contact met de C2-infrastructuur, waardoor een betrouwbare achterdeur wordt gecreëerd voor doorlopende aanvallen. Beveiligingsanalisten van Censys hebben tussen 14 oktober en 14 november 2025 meer dan 150 actieve Remcos C2-servers wereldwijd gedetecteerd, wat duidt op de wijdverspreide adoptie van dit gereedschap door aanvallers. De meeste van deze servers opereerden op poort 2404, de standaardpoort voor Remcos, maar er werd ook activiteit waargenomen op andere poorten zoals 5000, 5060, 5061, 8268 en 8808, wat aangeeft dat operators flexibel zijn in hun implementatiestrategieën.

De communicatie tussen de malware en de C2-servers vindt doorgaans plaats via HTTP en HTTPS op voorspelbare poorten. Het netwerkverkeer bevat vaak gecodeerde POST-verzoeken en ongewone TLS-configuraties die karakteristieke patronen creëren. De aanvallers hergebruiken vaak certificaten over meerdere servers, maken gebruik van sjablonen en gebruiken goedkope hostingproviders in landen zoals de Verenigde Staten, Nederland en Duitsland. Deze infrastructuur maakt het voor netwerkbeheerders mogelijk om communicatie te detecteren en te blokkeren. De persistentie van de malware wordt gehandhaafd via geplande taken en registerinstellingen, wat de aanvallers in staat stelt om toegang te behouden, zelfs na een herstart van het systeem. Deze combinatie van commando-uitvoering, bestandsoverdracht en robuuste persistentie maakt Remcos bijzonder gevaarlijk voor



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

organisaties met zwakke beveiligingsmaatregelen, waardoor het essentieel is om netwerkmonitoring en endpointdetectie onmiddellijk te implementeren.

Onderzoek onthult rol Tuoni C2 bij mislukte poging tot cyberinbraak in vastgoedsector

Cybersecurityonderzoekers hebben gedetailleerde informatie gedeeld over een mislukte cyberaanval op een vastgoedbedrijf in de Verenigde Staten, waarbij het Tuoni C2-framework werd gebruikt. Dit relatief nieuwe command-and-control (C2) systeem, dat sinds begin 2024 beschikbaar is, werd in de aanval ingezet om stealthy, in-memory payloads te leveren. Tuoni is oorspronkelijk bedoeld voor beveiligingsprofessionals voor penetratietests en red teaming, en is beschikbaar in een gratis versie via GitHub. De aanval vond plaats in oktober 2025 en werd vermoedelijk opgestart door middel van sociale manipulatie, waarbij de aanvallers zich via Microsoft Teams voordeden als vertrouwde leveranciers of collega's om toegang te krijgen. Het slachtoffer werd misleid om een PowerShell-opdracht uit te voeren, die vervolgens een tweede PowerShell-script ophaalde van een externe server. Dit script gebruikte steganografische technieken om een payload te verbergen binnen een bitmapafbeelding (BMP). Het uiteindelijke doel van de ingebedde payload was het uitvoeren van shellcode in het geheugen van het systeem. De uitvoering van de payload resulteerde in het activeren van een agent, de "TuoniAgent.dll", die verbinding maakte met de C2-server en remote control mogelijk maakte. Het framework zelf is complex, maar de aanvallers gebruikten een geavanceerd maar traditionele C2-aanvalsmethode. Onderzoekers vermoeden dat AI-tools een rol speelden in de ontwikkeling van de aanvallen, gezien de gescrite commentaren en de modulaire structuur van de eerste loader. Hoewel de aanval uiteindelijk niet succesvol was, toont het incident aan hoe geavanceerde red teaming-tools voor kwaadaardige doeleinden kunnen worden misbruikt. Eerder dit jaar werd in een ander onderzoek ook het gebruik van AI-gebaseerde hulpmiddelen voor het versnellen van kwetsbaarheidsuitbuitingen besproken. De voortdurende dreiging van dergelijke geavanceerde aanvallen benadrukt de noodzaak van voortdurende waakzaamheid in de cybersecuritysector.

Politie houdt twee mannen aan voor grootschalige creditcardfraude in Amsterdam



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

De politie heeft op 18 november twee mannen uit Amsterdam aangehouden op verdenking van grootschalige creditcardfraude. De verdachten zouden tientallen slachtoffers hebben gemaakt, die gezamenlijk voor tienduizenden euro's werden bestolen. Het fraudeursduo zou phishingmails hebben verzonden die afkomstig leken van creditcardmaatschappij ICS. Deze e-mails bevatten links naar phishingsites, waar slachtoffers hun gegevens invulden. De daders gebruikten de verkregen informatie vervolgens voor frauduleuze transacties. Tijdens de arrestatie van de verdachten werd onder meer designer kleding en tientallen gegevensdragers in beslag genomen. De politie werkte samen met ICS in dit onderzoek en roept slachtoffers van phishing en creditcardfraude op om aangifte te doen. Het is van belang dat slachtoffers zich niet schamen om aangifte te doen, aangezien dit kan bijdragen aan het onderzoek en mogelijk tot een hogere straf leidt.

Kamervragen over Amerikaanse overname van cloudbedrijf Solvinity

Op 7 november 2025 werd bekend dat het Nederlandse cloudbedrijf Solvinity werd overgenomen door het Amerikaanse techbedrijf Kyndryl. Deze overname heeft geleid tot bezorgdheid binnen de Nederlandse overheid, die gebruik maakt van de diensten van Solvinity, zoals DigiD en MijnOverheid. Daarnaast draaien ook het beveiligde communicatiesysteem van het ministerie van Justitie en Veiligheid en andere kritieke overheidsdiensten op de infrastructuur van Solvinity. De overname heeft ophef veroorzaakt, vooral gezien de Amerikaanse wetgeving die mogelijk invloed heeft op de privacy en vertrouwelijkheid van gegevens die door de overheid worden beheerd. Solvinity was tot nu toe een vertrouwde partij voor de Nederlandse overheid, maar de nieuwe Amerikaanse eigenaar zou onder andere moeten voldoen aan wetgeving zoals de CLOUD Act, de Foreign Intelligence Surveillance Act (FISA) en Executive Order 12333. Dit zou kunnen betekenen dat Amerikaanse autoriteiten toegang krijgen tot gevoelige overheidsinformatie. In reactie op de overname heeft Kamerlid Kathmann (GroenLinks/PvdA) Kamervragen gesteld aan minister Karremans van Economische Zaken. Zij vroeg onder andere of de Nederlandse overheid geprobeerd heeft de overname door een buitenlandse partij te voorkomen en of er een strategie is voor het geval de vertrouwelijkheid van overheidsinformatie in gevaar komt. Ook werd er gevraagd of er mogelijkheden zijn om contracten met Solvinity te beëindigen als blijkt dat de veiligheid van kritieke overheidsinformatie niet gegarandeerd kan worden onder de nieuwe eigenaar. De minister heeft drie weken de tijd om antwoord te geven op de vragen van de Kamer. De zorgen rondom de overname wijzen op de



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

groeijende bezorgdheid over de impact van buitenlandse overnames op de privacy en veiligheid van nationale gegevensinfrastructuren.

Van Marum: jaarlijks minimaal 1 miljard euro extra nodig voor digitalisering

Demissionair staatssecretaris Van Marum voor Digitalisering heeft aangegeven dat er tot 2030 jaarlijks minstens 1 miljard euro extra moet worden geïnvesteerd in digitalisering. In een brief aan de Tweede Kamer benadrukt hij dat digitalisering niet langer een keuze is, maar een noodzaak. Volgens Van Marum liggen kritieke systemen stil bij uitval, wat grote gevolgen zou hebben voor de samenleving. Hij wijst erop dat de beveiliging van de digitale infrastructuur essentieel is en dat de huidige geopolitieke situatie de druk vergroot om de digitale weerbaarheid te versterken. Het kabinet lanceerde in juli 2025 de Nederlandse Digitaliseringsstrategie (NDS), een gezamenlijk initiatief van de overheid en verschillende publieke diensten. De strategie richt zich op zes prioriteiten, waaronder cloudtechnologie, kunstmatige intelligentie, de versterking van digitale weerbaarheid, en het versterken van het digitale vakmanschap van ambtenaren. De financiering voor deze plannen wordt omschreven in een investeringsagenda, waaruit blijkt dat er jaarlijks minimaal 1 miljard euro extra nodig is om de gestelde doelen te behalen. Het geld zal onder andere worden gebruikt voor de versterking van de digitale weerbaarheid van gemeentelijke systemen, de uitbreiding van detectie- en responscapaciteit, en voor de uitvoering van de Cyberveiligheidswet. De investering moet bijdragen aan het creëren van een weerbaar digitaal ecosysteem en het vergroten van de digitale autonomie. Daarnaast is er aandacht voor de oprichting van een federatief Security Operations Center, dat moet helpen bij het monitoren van cyberdreigingen. Van Marum benadrukt dat de investeringsagenda voorlopig een eerste indicatie is van de kosten en dat een gedetailleerder rapport in de eerste helft van 2026 zal volgen.

Help4U: Europees platform biedt steun voor jongeren die slachtoffer zijn van online seksueel misbruik

Help4U is een nieuw digitaal platform, ontwikkeld door Europol en CENTRIC, dat jongeren ondersteunt die te maken hebben met seksueel misbruik of andere vormen van online schade. Het platform is ontworpen om eenvoudig, privé en toegankelijk te zijn, zodat kinderen en tieners betrouwbare informatie kunnen vinden, hun rechten kunnen begrijpen en in contact kunnen komen met mensen die hen kunnen helpen.



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

Help4U biedt praktische en begrijpelijke begeleiding voor jongeren onder de 18 jaar, evenals informatie voor ouders, leraren en andere professionals die hen ondersteunen. Het platform richt zich op het versterken van de veiligheid van jongeren zowel online als offline, door hen duidelijk, betrouwbare informatie en toegang tot hulp te bieden op het moment dat deze het hardst nodig is. Het platform is ontstaan uit een samenwerkingsproject tussen België, Duitsland, Ierland, Nederland en Slovenië, en is inmiddels uitgebreid naar andere Europese landen, waaronder Bulgarije, Cyprus, Kroatië, Griekenland, Hongarije, Italië, Portugal, Spanje en Roemenië. Help4U is ontwikkeld in nauwe samenwerking met experts uit verschillende disciplines, zoals psychologie, onderwijs, IT, gegevensbescherming, wetshandhaving en de academische wereld. Door de gezamenlijke inzet van nationale initiatieven onder één platform te bundelen, versterkt Help4U de Europese reactie op online seksueel misbruik en zorgt het ervoor dat steun de juiste slachtoffers bereikt. Jongeren die online ongewenste ervaringen hebben, kunnen Help4U gebruiken om antwoorden te vinden, hulp in te schakelen en weer controle over hun situatie te krijgen. Dit initiatief is een belangrijk voorbeeld van hoe Europese samenwerking, technologie en co-creatie samenkomen om kinderen te beschermen en de collectieve reactie van Europa tegen kindermisbruik te versterken.

Google-topman waarschuwt voor blind vertrouwen in AI

Sundar Pichai, CEO van Google, heeft gewaarschuwd voor het onkritisch vertrouwen in kunstmatige intelligentie (AI) en het gebruik van AI in kritieke systemen. In een interview met de BBC benadrukte Pichai dat zelfs de nieuwste AI-technologieën, waaronder Google's eigen Gemini-model, vatbaar zijn voor fouten en onnauwkeurigheden. Dit benadrukt de risico's van het inzetten van AI in systemen die direct invloed kunnen hebben op de cyberbeveiliging en informatiebeveiliging van organisaties. Pichai erkent dat bij de ontwikkeling van AI de druk om snel te innoveren vaak ten koste gaat van het grondig inbouwen van beschermingsmaatregelen tegen schadelijke effecten. Dit geldt voor zowel commerciële toepassingen van AI, zoals de integratie van Gemini in de Google-zoekmachine, als voor bredere toepassingen in de cybersecurity. Ondanks de kracht van AI blijft het essentieel om kritisch te blijven en te zorgen voor de juiste toezicht- en controlemechanismen. Google heeft de afgelopen jaren miljarden geïnvesteerd in AI, maar Pichai waarschuwt voor de mogelijke gevaren van een 'AI-zeepbel'. De enorme investeringen in de sector, die vaak gepaard gaan met onrealistische verwachtingen, kunnen uiteindelijk leiden tot teleurstellingen. Cybersecurity-experts



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

moeten zich bewust zijn van de kwetsbaarheden die AI met zich meebrengt, niet alleen op het gebied van fouten in de technologie, maar ook in de manier waarop AI wordt ingezet voor aanvallen of misbruik van systemen.

Sadistische sekte ronselt jongeren via Roblox: ernstige gevolgen voor Nederlandse kinderen

Een groeiend probleem binnen online gaming is de ontdekking van een internationale sekte, genaamd "The Cult of Spawn", die jongeren ronselt via het populaire platform Roblox. Het onderzoek van EenVandaag heeft blootgelegd dat deze sekte actief jongeren in Nederland en andere landen manipuleert en aanzet tot zelfverminking, het delen van naaktbeelden en zelfs zelfmoord. De sekte maakt gebruik van het populaire horrorspel Forsaken om jonge spelers in gesloten, vaak privé online groepen te lokken. Een van de belangrijkste zorgen is dat leden van deze groep jongeren manipuleren door gebruik te maken van dreigementen, chantage en psychologische druk. De politie heeft bevestigd dat deze sekte ook actief is in Nederland, wat de ernst van de situatie benadrukt. Er is grote bezorgdheid over de effectiviteit van de bestaande mechanismen om dergelijke groeperingen op online platformen te stoppen. Het slachtoffer van deze groeperingen kan leiden tot tragische gevolgen, zoals de zelfmoorddreigingen die naar voren kwamen in het onderzoek, waaronder de ervaring van een Nederlandse moeder wiens dochter door sekteleiden werd aangespoord om haar leven te beëindigen. De aanwezigheid van dergelijke sekten is al langere tijd een probleem op platforms als Roblox en Minecraft, waar ook geweld en seksueel misbruik steeds vaker voorkomen. Dit heeft ertoe geleid dat het kabinet nu onderzoekt of de maatregelen van deze platformen voldoende zijn om jongeren te beschermen tegen dergelijke dreigingen. De complexiteit van het probleem wordt versterkt door de anonimiteit die online platforms bieden, waardoor het voor de politie en technologische bedrijven moeilijk is om effectief op te treden tegen dergelijke misbruiknetwerken.

AI-teddybeer van de markt gehaald na ongepaste antwoorden

Een teddybeer met ingebouwde AI-chatfunctie is van de markt gehaald nadat onderzoek aantoonde dat de knuffel ongepaste en gevaarlijke antwoorden gaf. De beer, genaamd Kumma en geproduceerd door de Singaporese speelgoedfabrikant FoloToy, bevatte een microfoon en luidspreker zodat kinderen konden praten met het



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

speelgoed. De beer maakte gebruik van AI-technologie, waaronder chatbots zoals ChatGPT, maar zonder adequate veiligheidsmaatregelen. Dit leidde tot situaties waarin de teddybeer gevaarlijke of ongepaste antwoorden gaf op vragen van kinderen. Onderzoekers ontdekten dat de beer, wanneer een onderzoeker vroeg waar messen in huis te vinden waren, antwoord gaf dat messen meestal op een veilige plek bewaard worden, zoals in een keukenla of messenblok. Hoewel de beer adviseerde een volwassene om hulp te vragen, was de reactie ongepast. Eveneens gaf de teddybeer gedetailleerde antwoorden op vragen over seksuele voorkeuren en technieken, waarbij het speelgoed zelfs instructies gaf over bondage en andere seksuele handelingen.

Na deze bevindingen werd de toegang tot het AI-model van ChatGPT geblokkeerd en werd de verkoop van de teddybeer tijdelijk stopgezet. FoloToy verklaarde dat een uitgebreid veiligheidsonderzoek zou worden uitgevoerd. De Public Interest Research Group, een Amerikaanse consumentenorganisatie, wijst op de gevaren van speelgoed met AI en noemt het de opkomst van een experiment dat wordt uitgevoerd op kinderen. De organisatie benadrukt dat chatbots zoals ChatGPT niet bedoeld zijn voor kinderen, maar door dit soort producten alsnog toegankelijk worden gemaakt.

Roblox voert verplichte leeftijdsverificatie in voor 2 miljoen Nederlandse kinderen

Roblox, het populaire gamingplatform, heeft aangekondigd dat het vanaf december 2025 verplicht wordt voor Nederlandse gebruikers onder de 18 jaar om een leeftijdsverificatie te doorlopen. Dit wordt gedaan door middel van een selfie die via een algoritme wordt geanalyseerd om de leeftijd van de gebruiker te schatten. Het doel van deze maatregel is om de communicatie tussen volwassenen en kinderen beter te controleren en de risico's van online misbruik te verkleinen. De wijziging wordt wereldwijd pas in 2026 geïmplementeerd, maar Nederland en andere landen zoals Australië en Nieuw-Zeeland vormen de voorlopers. Deze stap volgt op groeiende bezorgdheid over de veiligheid op Roblox. Recent onderzoek heeft aangetoond dat de platformen worden misbruikt door criminele groeperingen, zoals de sekte The Cult of Spawn, die jongeren via het platform werven en onder druk zetten om deel te nemen aan schadelijke activiteiten. De groep maakt gebruik van het populaire horrospeel Forsaken om jongeren te manipuleren en hen naar afgesloten chatgroepen te lokken, waar ze worden bedreigd en onder druk gezet om zelfverminking en zelfmoord te overwegen. Deze problematiek is niet nieuw; ook



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

andere platformen zoals Minecraft werden eerder al in verband gebracht met soortgelijke risico's voor kinderen.

Roblox, dat wereldwijd meer dan 150 miljoen dagelijkse actieve gebruikers heeft, zegt dat de leeftijdsverificatie zal worden geïmplementeerd met een stapsgewijze aanpak om te testen hoe het systeem in verschillende regio's werkt. De verwachting is dat de veranderingen moeten bijdragen aan het versterken van de veiligheid op het platform, waar techbedrijven en wetshandhavers zich steeds vaker geconfronteerd zien met de uitdagingen van digitale misbruikpraktijken en moeilijkheden in het blokkeren van criminele netwerken.

AI-chatbots geven misleidend financieel advies, onderzoek wijst op risico's

Uit recent Brits onderzoek blijkt dat populaire AI-chatbots, zoals ChatGPT en Microsoft Copilot, misleidend en onjuist financieel advies geven. Dit is zorgwekkend gezien het groeiende gebruik van AI in uiteenlopende digitale diensten, waaronder financiële en juridische ondersteuning. Het onderzoek van de Britse consumentenorganisatie Which? toonde aan dat de onderzochte chatbots in sommige gevallen foutieve belastingadviezen gaven en gebruikers aanmoedigden om belangrijke financiële limieten te overschrijden, wat kan leiden tot juridische complicaties. Bijvoorbeeld, ChatGPT adviseerde ten onrechte om de jaarlijkse stortingslimieten voor belastingvrije spaarrekeningen te overschrijden en raadde onterecht aan om een reisverzekering af te sluiten voor EU-reizen, terwijl dat niet verplicht is. Daarnaast werd ook vastgesteld dat de AI-tools van Meta en Google onnauwkeurige informatie gaven over het claimen van compensatie voor vertraagde vluchten en het inhouden van betalingen aan aannemers, wat juridische risico's met zich meebrengt. Het onderzoek benadrukt het belang van kritische beoordeling van AI-adviezen, vooral in het kader van financiële planning en juridische zaken. Omdat dergelijke systemen geen bescherming bieden onder bestaande financiële regelingen, lopen gebruikers het risico om verkeerd geadviseerd te worden, wat kan leiden tot ernstige consequenties. Dit wijst op een bredere uitdaging voor de digitale veiligheid, waarbij AI-systemen soms onbedoeld schadelijke gevolgen kunnen hebben voor gebruikers die hun vertrouwen stellen in deze technologieën zonder menselijke supervisie. Het onderzoek benadrukt dat de afhankelijkheid van AI in kritieke domeinen zoals financieel advies en contracten zorgvuldig moet worden geëvalueerd, gezien de risico's van onnauwkeurigheid en de beperkte juridische bescherming voor gebruikers.



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

Onderzoek naar meldpunt voor burgemeesters om online content te verwijderen

Er wordt momenteel onderzoek gedaan naar de mogelijkheid om een landelijk meldpunt in te richten waarmee burgemeesters online content kunnen laten verwijderen die leidt tot openbare ordeverstoringen. Het meldpunt zou bedoeld zijn om snel op te treden tegen online oproepen tot geweld of andere verstoringen die via sociale media of andere platforms worden verspreid. Dit onderzoek wordt uitgevoerd door het Centrum voor Criminaliteitspreventie en Veiligheid (CCV), met de verwachting dat de resultaten begin 2026 bekend zullen zijn. Minister Van Oosten van Justitie en Veiligheid benadrukte dat er de afgelopen jaren steeds meer aandacht is gekomen voor het aanpakken van online aangejaagde openbare ordeverstoringen. Het voorstel roept echter juridische vragen op, aangezien het de grenzen van de vrijheid van meningsuiting raakt, een fundamenteel recht in Nederland en Europa. Er moet een zorgvuldige afweging worden gemaakt tussen het beschermen van de openbare orde en het waarborgen van digitale rechten. Een relevante gebeurtenis die dit debat aanwakkerde, was de uitspraak van de Raad van State in juni 2025, waarbij werd bepaald dat de burgemeester van Utrecht onterecht een 'online gebiedsverbod' had opgelegd aan een jongen die via Telegram opriep tot verstoringen. Deze uitspraak maakt duidelijk dat het reguleren van online uitingen door lokale overheden juridische en technische complicaties met zich meebrengt. De mogelijke oprichting van een meldpunt moet ook in aanmerking nemen hoe platforms snel en effectief reageren op verzoeken tot verwijdering van content, zonder inbreuk te maken op privacy en andere digitale rechten van gebruikers. Dit onderzoek biedt mogelijk nieuwe inzichten in de rol van lokale autoriteiten bij het beheer van online dreigingen, en kan ook van invloed zijn op hoe cybersecurity-maatregelen in de toekomst zullen worden toegepast in relatie tot openbare orde.

ACM en Europese Commissie onderzoeken de impact van clouddiensten op de digitale markt

De Autoriteit Consument & Markt (ACM) werkt samen met de Europese Commissie aan een diepgaand onderzoek naar de invloed van grote clouddienstaanbieders, zoals Microsoft en Amazon, op de concurrentie en digitale veiligheid. Er wordt onderzocht of deze bedrijven als 'gatekeeper' moeten worden aangemerkt onder de Digital Markets Act (DMA), wat zou betekenen dat ze strengere regels moeten volgen om eerlijke concurrentie te waarborgen. Dit is belangrijk, omdat dergelijke bedrijven,



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

door hun omvang en marktdominantie, invloed kunnen hebben op de toegang tot en controle over data, wat gevolgen heeft voor de privacy en beveiliging van gebruikers.

Naast de mogelijke herkwalficatie van Microsoft Azure en Amazon Web Services als 'core platform services', wordt er ook gekeken naar de effectiviteit van de DMA in het voorkomen van oneerlijke praktijken, zoals lock-in effecten die bedrijven en gebruikers zouden kunnen belemmeren in hun keuze voor clouddiensten. Ook de toegang tot gegevens voor zakelijke gebruikers wordt geëvalueerd, wat van belang is voor bedrijven die afhankelijk zijn van clouddiensten voor hun digitale bedrijfsvoering. De resultaten van dit onderzoek kunnen belangrijke implicaties hebben voor de manier waarop bedrijven hun digitale infrastructuur beheren en beveiligen. Het wordt verwacht dat de bevindingen in 2026 beschikbaar zullen zijn.

Franse overheidsdienst Pajemploi getroffen door datalek, 1,2 miljoen mensen mogelijk getroffen

De Franse overheidsdienst Pajemploi, die verantwoordelijk is voor het beheer van sociale zekerheid voor ouders en thuiszorgverleners, heeft onlangs melding gemaakt van een datalek. Het incident heeft mogelijk persoonlijke gegevens blootgelegd van 1,2 miljoen mensen die gebruikmaken van de dienst, waaronder zorgverleners die voor particuliere werkgevers werken. De gegevens die mogelijk zijn buitgemaakt, bevatten onder andere volledige namen, geboortedata, adressen, sofinummers, en de naam van de bankinstelling van de betrokkene. Volgens Pajemploi waren er geen aanwijzingen dat bankrekeningnummers, e-mailadressen, telefoonnummers of wachtwoorden zijn gestolen. De cyberaanval werd op 14 november 2025 gedetecteerd en prompt werd actie ondernomen om de aanval te stoppen en de systemen te beschermen. Het Franse overheidsorgaan heeft de autoriteiten, waaronder de Franse toezichthouder voor gegevensbescherming (CNIL) en de Nationale Agentschap voor de Veiligheid van Informatie Systemen (ANSSI), ingelicht. Ondanks de omvang van het datalek, heeft het incident geen invloed gehad op de werking van de diensten van Pajemploi, zoals het verwerken van aangiften en het betalen van salarissen. Pajemploi heeft aangegeven dat de getroffen personen persoonlijk op de hoogte zullen worden gebracht en raadt hen aan extra waakzaam te zijn voor mogelijke fraude via e-mail, sms of telefoon, nu hun gestolen gegevens mogelijk voor phishing-aanvallen kunnen worden gebruikt. Het lek komt kort na andere incidenten in Frankrijk, waaronder een datalek bij de werkloosheidsdienst



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

Pôle Emploi, wat de bezorgdheid over de bescherming van persoonlijke gegevens binnen overheidsdiensten vergroot.

Oplichting rond visumaanvraag of aanvraag ETA

Er zijn meldingen van oplichting met betrekking tot visumaanvragen en de aanvraag voor een Electronic Travel Authorisation (ETA). Slachtoffers van deze fraude melden dat zij terechtkomen bij tussenpersonen die hogere bedragen vragen voor de visum- of ETA-aanvraag dan wanneer dit via de officiële overheidswebsite zou gebeuren. In sommige gevallen wordt ook gevraagd om persoonlijke gegevens, zoals een identiteitsbewijs, via e-mail door te geven. In bepaalde gevallen ontvangt de melder een visum dat geldig blijkt te zijn, terwijl het voor andere melders onduidelijk is wat er na de betaling gebeurt. De meeste meldingen hebben betrekking op de aanvraag van een ETA, aangezien dit document vanaf 2 april 2025 verplicht is voor iedereen die voor minder dan zes maanden naar het Verenigd Koninkrijk reist. Het is belangrijk om goed te controleren of de aanvraag via de officiële kanalen wordt gedaan, bijvoorbeeld via de website Nederland Wereldwijd, waar de officiële aanvraagformulieren voor visums en ETA's te vinden zijn. Indien er twijfel bestaat over een betaling of de gegevens die zijn verstrekt, wordt aangeraden contact op te nemen met de fraudehulpdesk voor persoonlijk advies.

Van Weel benadrukt urgentie rondom AI in buitenlands beleid

De demissionaire minister van Buitenlandse Zaken, David van Weel, heeft op de AI Summit Brainport 2025 zijn bezorgdheid geuit over de groeiende invloed van kunstmatige intelligentie (AI) op zowel de nationale als internationale veiligheid. Volgens Van Weel is het essentieel dat AI een integraal onderdeel wordt van het buitenlands beleid van Nederland en Europa. Hij verwees naar uitspraken van de Russische president Poetin in 2017, waarin hij stelde dat werelddominantie uiteindelijk afhangt van de beheersing van AI. Hoewel Van Weel het niet eens is met Poetin's opvatting, erkent hij dat deze uitspraak deels correct is. Van Weel benadrukte de noodzaak voor effectieve wet- en regelgeving rondom AI op internationaal niveau. Nederland, samen met de Europese Unie, de NAVO en de Verenigde Naties, wil zich inzetten voor het ontwikkelen van een toekomstbestendig beleid om de wereldwijde impact van AI te reguleren. In dit verband wordt gesproken over het belang van AI in verschillende sectoren, waaronder landbouw en



Cybercrimeinfo - cyberdreigingsanalyse (Openbare versie)

cybersecurity. De minister wijst ook op de ethische verschillen tussen westerse waarden en die van andere landen, zoals China, dat AI gebruikt voor massale surveillantie van zijn bevolking. Volgens Van Weel is de invloed van AI in verschillende domeinen zichtbaar, van het slagveld in Oekraïne tot de verspreiding van desinformatie en de gevolgen voor de samenleving. Hij roept op om AI in samenhang te bekijken en te zorgen dat het beleid zowel de nationale belangen beschermt als de wereldwijde stabiliteit waarborgt.