

BOLETÍN INFORMATIVO N° 200



SEPTIEMBRE 2025



Investigadores descubren malware MalTerminal con tecnología GPT-4 que crea ransomware y shell inverso

Investigadores en ciberseguridad han descubierto lo que dicen es el ejemplo más antiguo conocido hasta la fecha de un malware que incorpora capacidades de modelo de lenguaje grande (LLM).

17.500 dominios de phishing atacan a 316 marcas en 74 países en el aumento global de PhaaS

Las ofertas de phishing como servicio (PhaaS) conocidas como Lighthouse y Lucid se han vinculado a más de 17.500 dominios de phishing dirigidos a 316 marcas de 74 países.



Hackers de la RPDC usan ClickFix para distribuir malware BeaverTail en estafas de empleos criptográficos

Se ha observado que actores de amenazas vinculados a la República Popular Democrática de Corea (también conocida como RPDC o Corea del Norte) utilizan señuelos estilo ClickFix para distribuir un malware conocido llamado BeaverTail e InvisibleFerret.



UNC1549 hackea 34 dispositivos en 11 empresas de telecomunicaciones mediante ofertas de empleo en LinkedIn y malware MINIBIKE.

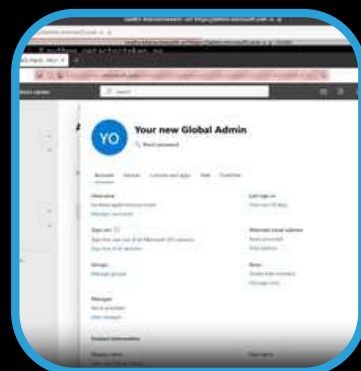
Un grupo de ciberespionaje con vínculos con Irán, conocido como UNC1549, ha sido atribuido a una nueva campaña dirigida a empresas de telecomunicaciones europeas, infiltrándose con éxito en 34 dispositivos de 11 organizaciones como parte de una actividad con temática de reclutamiento en LinkedIn.

INCIDENTES DE SISTEMAS



Microsoft corrige un fallo crítico de Entra ID que permite la suplantación de administrador global en todos los inquilinos.

Una falla crítica en la validación de token en Microsoft Entra ID (anteriormente Azure Active Directory) podría haber permitido a los atacantes suplantar a cualquier usuario, incluidos los administradores globales, en cualquier inquilino.



Google corrige el fallo de día cero CVE-2025-10585 en Chrome, ya que el exploit V8 activo amenaza a millones de personas.

Google lanzó el miércoles actualizaciones de seguridad para el navegador web Chrome para abordar cuatro vulnerabilidades, incluida una que, según dijo, ha sido explotada.

RECOMENDACIONES

LECTURA DE SEGURIDAD



Los LLM pueden impulsar las decisiones de ciberseguridad, pero no son para todos

Los LLM están pasando rápidamente de la experimentación al uso diario en ciberseguridad. Los equipos están empezando a utilizarlos para analizar la inteligencia de amenazas , guiar la respuesta a incidentes y ayudar a los analistas a gestionar el trabajo repetitivo

La nueva técnica de Rowhammer contra DDR5 logra la escalada de privilegios

Al aplicar ingeniería inversa a los mecanismos Target Row Refresh (TRR) que hasta ahora han protegido a la RAM DDR5 contra cambios de bits, los investigadores de seguridad pudieron derrotar a todos los DIMM en sus pruebas.



La IA necesita ética para evitar daños en el mundo real

En este video de Help Net Security, Brittany Allen, arquitecta sénior de confianza y seguridad en Sift , explora cómo el auge de los agentes de IA está creando nuevos riesgos de fraude. Explica cómo estos agentes, aunque diseñados para asistir a los usuarios, pueden ayudar involuntariamente a los estafadores al realizar tareas sin reconocer intenciones maliciosas.



NOTICIAS DE NUESTROS PARTNERS



Cerrando la brecha de seguridad de la IA agente: Por qué la protección de la identidad debe evolucionar ahora

La IA de Agentic está aumentando los riesgos para la seguridad de la identidad. Descubre cómo adoptarla sin abrir nuevas vías a los atacantes.



El marco de evaluación cibernética v4.0 eleva el nivel: 6 preguntas que todo equipo de seguridad debería plantearse sobre su postura de seguridad

Una guía práctica para las actualizaciones clave de detección y respuesta en CAF v4.0, incluida la detección basada en anomalías, la búsqueda de amenazas dirigida por máquinas y los requisitos de postura de seguridad proactiva.



IBM y BharatGen colaboran para acelerar la adopción de IA en India gracias a los grandes modelos lingüísticos de Indic

IBM y BharatGen anunciaron hoy una colaboración estratégica para impulsar la adopción de la Inteligencia Artificial (IA) en India, impulsada por los Modelos Multimodales y Grandes Lenguajes (LLM) soberanos de BharatGen, adaptados al singular panorama lingüístico y cultural de la India.



BENCHMARKING EN CIBERSEGURIDAD

Nos es muy grato informarle que, como parte de la estrategia de difusión de los servicios de Riesgo y Seguridad de Mainsoft, hemos creado una serie de benchmarkings sin costo para empresas seleccionadas.

MIDE LA SEGURIDAD DE
TU ORGANIZACIÓN SIN
COSTO

[HAZ CLICK](#)

EVENTOS CERCANOS DE NUESTROS PARTNERS

DARKTRACE

[Más Información](#)

BeyondTrust

[Más Información](#)

IBM
Gold Partner

[Más Información](#)

