

Alerta ID:	100
Fecha del reporte:	14/05/2026
Entidad:	Todas las entidades del ecosistema digital
Título:	Modelo de IA “Mythos” demuestra capacidades avanzadas en descubrimiento automatizado de vulnerabilidades
Herramienta de detección	N/A
Activo involucrado:	Aplicaciones empresariales, código fuente, sistemas operativos, navegadores, pipelines DevSecOps, infraestructura cloud, software open source y entornos de desarrollo
Tipo de alerta:	Alerta
Nivel de riesgo:	Alto

Objetivo:

Informar a las organizaciones del sector salud en Colombia sobre las nuevas capacidades demostradas por el modelo de inteligencia artificial “Mythos”, desarrollado por Anthropic, el cual ha mostrado resultados significativamente superiores en descubrimiento automatizado de vulnerabilidades dentro de código fuente, aplicaciones web, firmware y software nativo.

El propósito de esta alerta es que las entidades evalúen el impacto que la aceleración del descubrimiento automatizado de vulnerabilidades puede generar sobre:

- gestión de vulnerabilidades;
- tiempos de exposición;
- procesos DevSecOps;
- seguridad de aplicaciones;
- monitoreo de infraestructura;
- y estrategias de defensa organizacional.



Asimismo, se busca generar conciencia sobre el cambio progresivo del panorama de amenazas derivado del uso de modelos de IA avanzados capaces de:

- analizar grandes volúmenes de código;
- identificar vulnerabilidades complejas;
- correlacionar fallas;
- construir cadenas de explotación;
- realizar ingeniería inversa;
- y automatizar parcialmente actividades ofensivas.

El incidente evidencia una evolución importante en la industria de ciberseguridad, donde las capacidades previamente limitadas a investigadores altamente especializados comienzan a ser potenciadas mediante inteligencia artificial avanzada. En el contexto del sector salud, este escenario representa un riesgo significativo debido a la alta dependencia tecnológica existente sobre:

- plataformas clínicas;
- aplicaciones web;
- sistemas heredados;
- software médico;
- servicios cloud;
- dispositivos conectados;
- y múltiples componentes open source.

Descripción:

Investigadores independientes y firmas especializadas en ciberseguridad realizaron múltiples pruebas técnicas sobre “Mythos Preview”, el nuevo modelo de inteligencia artificial desarrollado por Anthropic, identificando capacidades avanzadas para descubrimiento



automatizado de vulnerabilidades sobre aplicaciones web, software nativo, firmware y código fuente complejo.

De acuerdo con el análisis publicado por SecurityWeek, la firma XBOW ejecutó evaluaciones comparativas entre Mythos y otros modelos de inteligencia artificial orientados a investigación ofensiva y análisis de seguridad. Los resultados mostraron que Mythos logró desempeños considerablemente superiores en tareas relacionadas con auditoría automatizada de código y descubrimiento de vulnerabilidades complejas.

Las pruebas evidenciaron capacidades particularmente efectivas para:

- análisis de código fuente;
- descubrimiento de vulnerabilidades nativas;
- análisis de firmware;
- ingeniería inversa;
- evaluación de aplicaciones web;
- y razonamiento contextual sobre sistemas complejos.
-

XBOW indicó que Mythos mostró mejores resultados cuando operaba sobre entornos reales y combinaba análisis dinámico con acceso directo al código fuente de las aplicaciones evaluadas. Los investigadores también señalaron que el modelo fue capaz de correlacionar múltiples condiciones vulnerables y construir escenarios ofensivos complejos de manera mucho más eficiente que modelos anteriores.

Anthropic había informado previamente que Mythos logró identificar vulnerabilidades históricas dentro de:

- Linux;
- OpenBSD;
- FFmpeg;
- Otros componentes ampliamente utilizados dentro de ecosistemas empresariales y open source.



Algunas pruebas demostraron incluso capacidades para construir cadenas de explotación y razonar sobre posibles rutas de escalamiento de privilegios dentro de sistemas reales.

Investigaciones adicionales realizadas por Mozilla señalaron que modelos avanzados de IA, incluyendo Mythos, contribuyeron a acelerar procesos de identificación y corrección de vulnerabilidades dentro de Firefox durante 2026. Este escenario evidencia una evolución importante en las capacidades ofensivas y defensivas asociadas al uso de inteligencia artificial dentro de la industria de ciberseguridad.

No obstante, las evaluaciones independientes también identificaron limitaciones importantes relacionadas con:

- validación práctica de exploits;
- interpretación contextual;
- priorización de hallazgos;
- razonamiento operativo;
- y diferenciación entre vulnerabilidades teóricas y escenarios realmente explotables.

XBOW concluyó que Mythos aún presenta dificultades para determinar correctamente el impacto operacional real de ciertos hallazgos y puede sobreestimar la criticidad de algunas vulnerabilidades bajo determinados contextos.

El principal riesgo asociado a este tipo de tecnologías corresponde a la reducción progresiva del tiempo entre:

- descubrimiento de vulnerabilidades;
- análisis técnico;
- generación de exploits;
- y explotación activa.



A medida que los modelos de IA continúan evolucionando, las organizaciones podrían enfrentar ventanas de exposición considerablemente más reducidas para aplicar parches, corregir configuraciones inseguras y ejecutar procesos de mitigación.

En el contexto del sector salud, este escenario resulta especialmente relevante debido a la alta dependencia tecnológica existente sobre:

- plataformas clínicas;
- aplicaciones web institucionales;
- dispositivos médicos conectados;
- sistemas heredados;
- infraestructura cloud;
- servicios SaaS;
- y componentes open source.

La automatización progresiva del descubrimiento de vulnerabilidades podría incrementar significativamente el riesgo de explotación acelerada sobre infraestructuras críticas y servicios institucionales expuestos.

Consecuencias de la explotación

La evolución de modelos de IA especializados en descubrimiento automatizado de vulnerabilidades puede generar impactos significativos sobre la confidencialidad, integridad y disponibilidad de los sistemas institucionales.

Uno de los principales riesgos asociados corresponde a la reducción acelerada del tiempo disponible para corregir vulnerabilidades antes de que sean explotadas activamente. Históricamente, las organizaciones contaban con ventanas de tiempo relativamente amplias entre divulgación, análisis técnico y explotación masiva. Sin embargo, el uso de inteligencia artificial avanzada podría reducir significativamente estos tiempos mediante automatización parcial de actividades ofensivas.



Las capacidades demostradas por modelos como Mythos podrían facilitar:

- análisis automatizado de código;
- correlación de vulnerabilidades;
- descubrimiento de fallas complejas;
- generación de pruebas de concepto;
- y construcción de cadenas de explotación.

Otro aspecto crítico corresponde al incremento potencial del volumen de vulnerabilidades descubiertas dentro de software empresarial, sistemas operativos y plataformas ampliamente utilizadas. Esto podría generar presión significativa sobre equipos de seguridad y procesos de remediación.

En entornos del sector salud, una explotación acelerada podría comprometer:

- plataformas clínicas;
- aplicaciones hospitalarias;
- servicios web expuestos;
- infraestructura cloud;
- dispositivos médicos conectados;
- y servicios institucionales críticos.

Diversos especialistas también alertaron sobre el riesgo de que actores maliciosos utilicen capacidades similares para automatizar reconocimiento ofensivo, acelerar explotación y optimizar campañas de ataque dirigidas.



Modo de explotación del ataque



Fase 1 – Recolección y análisis automatizado de código

El proceso inicia mediante el análisis automatizado de grandes volúmenes de código fuente utilizando modelos avanzados de IA.

Los sistemas pueden evaluar:

- repositorios completos;
- firmware;
- aplicaciones web;
- librerías open source;
- y componentes empresariales.

Mythos demostró capacidades avanzadas en auditoría automatizada de código y razonamiento contextual sobre vulnerabilidades complejas.

Fase 2 – Identificación de patrones vulnerables

Posteriormente, la IA identifica errores lógicos, validaciones inseguras y condiciones vulnerables dentro de aplicaciones complejas.

Las pruebas realizadas demostraron capacidades relevantes para:

- descubrimiento de vulnerabilidades nativas;
- análisis de firmware;
- evaluación de aplicaciones web;
- ingeniería inversa;
- y correlación de fallas complejas.



Durante esta etapa, el modelo puede construir hipótesis ofensivas y detectar posibles rutas de explotación dentro de sistemas reales.

Fase 3 – Construcción de cadenas de explotación

Los modelos avanzados pueden correlacionar múltiples vulnerabilidades para construir escenarios ofensivos más complejos.

Anthropic indicó previamente que Mythos logró encadenar vulnerabilidades dentro del kernel Linux para alcanzar escenarios de escalamiento de privilegios.

Esto podría facilitar:

- automatización de razonamiento ofensivo;
 - identificación de rutas explotables;
 - correlación de debilidades;
 - y generación acelerada de escenarios de ataque.
-

Fase 4 – Generación y validación de exploits

Posteriormente, el modelo puede asistir parcialmente en tareas relacionadas con:

- generación de pruebas de concepto;
- validación técnica;
- análisis dinámico;
- y optimización de explotación.

No obstante, los análisis independientes señalaron que Mythos aún presenta limitaciones importantes relacionadas con validación práctica y contextualización operacional de ciertos hallazgos.



Fase 5 – Explotación acelerada y expansión del compromiso

Finalmente, una explotación exitosa podría facilitar:

- acceso no autorizado;
- movimiento lateral;
- robo de información;
- escalamiento de privilegios;
- compromiso cloud;
- y automatización ofensiva.

El principal riesgo asociado a este escenario corresponde a la reducción progresiva del tiempo entre descubrimiento y explotación efectiva de vulnerabilidades. Esto podría afectar significativamente la capacidad de respuesta de organizaciones con procesos lentos de parchado o gestión de vulnerabilidades.

Vulnerabilidades asociadas



Riesgos derivados de automatización ofensiva basada en IA

- **Descripción:** Los modelos avanzados de inteligencia artificial están comenzando a automatizar parcialmente tareas tradicionalmente ejecutadas por investigadores especializados en seguridad ofensiva. Esto incluye análisis automatizado de código fuente, identificación de errores lógicos, correlación de vulnerabilidades y construcción de posibles rutas de explotación dentro de aplicaciones complejas. Estas capacidades permiten acelerar procesos relacionados con:
 - descubrimiento de vulnerabilidades;
 - análisis ofensivo;
 - generación de pruebas de concepto;



- y razonamiento sobre escenarios explotables.
 - **Impacto:** La automatización progresiva de estas actividades podría reducir significativamente el tiempo requerido para desarrollar ataques sofisticados, incrementando el riesgo de explotación temprana sobre sistemas expuestos, aplicaciones institucionales y plataformas críticas.
-

Incremento de exposición sobre software heredado y componentes open source

- **Descripción:** Las capacidades de análisis automatizado demostradas por Mythos permiten identificar vulnerabilidades complejas dentro de sistemas heredados, librerías open source, firmware y componentes ampliamente utilizados que históricamente podían permanecer sin detección durante largos periodos de tiempo.

Este escenario resulta especialmente relevante para organizaciones que dependen de:

- software legacy;
 - plataformas sin soporte actualizado;
 - componentes open source;
 - y aplicaciones desarrolladas sobre arquitecturas antiguas.
 - **Impacto:** La explotación de vulnerabilidades previamente desconocidas podría afectar servicios críticos institucionales y aumentar significativamente el riesgo de compromiso operacional dentro de infraestructuras corporativas y entornos del sector salud.
-

Riesgos asociados a reducción de tiempos de explotación

- **Descripción:** La evolución de herramientas de IA orientadas a descubrimiento automatizado de vulnerabilidades podría reducir considerablemente el tiempo existente entre:



- identificación de fallas;
- análisis técnico;
- generación de exploits;
- y explotación activa.

A medida que estas capacidades evolucionan, los actores maliciosos podrían acelerar progresivamente sus procesos ofensivos y optimizar campañas dirigidas.

- **Impacto:** Las organizaciones podrían enfrentar ventanas de exposición considerablemente más reducidas para ejecutar procesos de parchado, remediación y mitigación, aumentando la probabilidad de explotación antes de implementar controles correctivos.

Exposición derivada de análisis automatizado de código fuente

- **Descripción:** Los modelos avanzados de IA poseen capacidades para analizar grandes volúmenes de código fuente y correlacionar múltiples condiciones vulnerables dentro de aplicaciones empresariales y plataformas complejas.

El riesgo asociado incrementa especialmente en escenarios donde existe exposición de:

- repositorios internos;
 - código fuente corporativo;
 - pipelines CI/CD;
 - firmware;
 - y componentes de software propietarios.
- **Impacto:** La exposición o robo de repositorios podría facilitar procesos acelerados de análisis ofensivo, descubrimiento de vulnerabilidades y construcción automatizada de escenarios de explotación dirigidos contra aplicaciones institucionales.
-



Riesgos asociados a correlación automatizada de vulnerabilidades

- **Descripción:** Las nuevas capacidades de razonamiento contextual demostradas por modelos como Mythos permiten correlacionar múltiples vulnerabilidades individuales para construir cadenas de explotación más complejas y efectivas.

Esto podría facilitar escenarios relacionados con:

- escalamiento de privilegios;
 - evasión de controles;
 - movimiento lateral;
 - y explotación encadenada de fallas aparentemente aisladas.
- **Impacto:** Los atacantes podrían aprovechar múltiples debilidades menores para desarrollar escenarios avanzados de compromiso progresivo sobre infraestructura crítica y plataformas institucionales.

Exposición sobre entornos DevSecOps y pipelines automatizados

- **Descripción:** Los entornos de desarrollo modernos contienen repositorios, secretos, pipelines CI/CD y herramientas automatizadas que pueden ser analizadas rápidamente mediante capacidades avanzadas de inteligencia artificial.

Estos entornos suelen mantener acceso privilegiado hacia:

- infraestructura cloud;
- servicios SaaS;
- repositorios institucionales;
- plataformas de despliegue;



- y sistemas críticos organizacionales.
- **Impacto:** Una explotación exitosa sobre componentes DevSecOps podría facilitar acceso a secretos cloud, credenciales privilegiadas, procesos automatizados de despliegue y múltiples servicios integrados dentro de la infraestructura corporativa.

Recomendaciones de prevención

Las organizaciones del sector salud deben fortalecer urgentemente sus estrategias de gestión de vulnerabilidades y resiliencia frente a amenazas potenciadas por inteligencia artificial.

La prevención debe enfocarse especialmente en reducción de superficie de ataque, aceleración de parchado y fortalecimiento de capacidades DevSecOps.

Implementar gestión acelerada de vulnerabilidades

Los tiempos de exposición podrían reducirse significativamente debido al uso de IA ofensiva.

Controles recomendados:

- Priorizar parchado de activos críticos.
- Reducir tiempos de remediación.
- Implementar gestión continua de vulnerabilidades.
- Automatizar validaciones de exposición.
- Monitorear explotación activa.
- Aplicar parchado basado en riesgo.
- Validar vulnerabilidades explotables.



Fortalecer seguridad DevSecOps y protección de código

Los repositorios y pipelines representan objetivos prioritarios frente a análisis automatizado ofensivo.

Controles recomendados:

- Restringir acceso a código fuente.
- Proteger pipelines CI/CD.
- Implementar análisis SAST/DAST.
- Validar dependencias open source.
- Implementar control de secretos.
- Segmentar entornos de desarrollo.
- Auditar accesos administrativos.

Reducir superficie de exposición institucional

La reducción de superficie expuesta continúa siendo crítica frente a automatización ofensiva basada en IA.

Controles recomendados:

- Eliminar servicios innecesarios.
- Restringir exposición a Internet.
- Segmentar infraestructura crítica.
- Implementar Zero Trust.
- Aplicar mínimos privilegios.
- Fortalecer controles de acceso.
- Monitorear activos expuestos.



Fortalecer monitoreo y capacidades defensivas basadas en IA

Las organizaciones deben evolucionar progresivamente sus capacidades defensivas.

Controles recomendados:

- Implementar SIEM avanzado.
- Integrar capacidades UEBA.
- Implementar XDR/EDR.
- Automatizar correlación de eventos.
- Detectar actividad anómala.
- Implementar análisis contextual.
- Fortalecer inteligencia de amenazas.

Recomendaciones de detección y mitigación

Las organizaciones deben fortalecer capacidades de monitoreo orientadas a detectar explotación acelerada de vulnerabilidades y automatización ofensiva.

Monitorear explotación temprana de vulnerabilidades

Controles recomendados:

- Detectar escaneo automatizado.
- Monitorear actividad anómala.
- Correlacionar eventos ofensivos.
- Integrar inteligencia de amenazas.
- Detectar explotación de zero-days.
- Monitorear intentos de escalamiento.
- Analizar actividad privilegiada.



Implementar monitoreo sobre activos críticos y expuestos

Los activos expuestos continúan siendo objetivos prioritarios frente a análisis ofensivo automatizado.

Controles recomendados:

- Monitorear aplicaciones críticas.
- Detectar accesos sospechosos.
- Analizar tráfico perimetral.
- Monitorear APIs y servicios web.
- Detectar explotación automatizada.
- Correlacionar eventos cloud.
- Monitorear actividad VPN.

Fortalecer detección sobre entornos DevSecOps

Los entornos de desarrollo representan objetivos estratégicos frente a automatización ofensiva basada en IA.

Controles recomendados:

- Monitorear repositorios.
- Detectar cambios anómalos.
- Auditar pipelines.
- Monitorear acceso a secretos.
- Detectar ejecución sospechosa.
- Correlacionar eventos DevOps.



- Validar integridad de artefactos.

Recomendaciones de respuesta



En caso de sospecha de explotación acelerada o descubrimiento automatizado de vulnerabilidades dentro del entorno institucional, las organizaciones deben activar inmediatamente sus procedimientos de respuesta a incidentes.

La respuesta debe enfocarse especialmente en contención rápida, protección de credenciales y validación de integridad sobre activos críticos.

Contener sistemas potencialmente comprometidos

Controles recomendados:

- Aislar sistemas afectados.
- Restringir accesos comprometidos.
- Revocar sesiones activas.
- Bloquear actividad sospechosa.
- Segmentar entornos críticos.
- Limitar privilegios temporalmente.
- Monitorear movimiento lateral.

Ejecutar análisis forense y validación de integridad

Las organizaciones deben identificar vulnerabilidades explotadas, actividad privilegiada y posibles escenarios de persistencia.



Controles recomendados:

- Revisar logs de explotación.
- Analizar actividad privilegiada.
- Validar integridad de aplicaciones.
- Revisar configuraciones críticas.
- Identificar vulnerabilidades explotadas.
- Ejecutar análisis forense.
- Correlacionar eventos recientes.

Proteger credenciales y accesos institucionales

Las credenciales privilegiadas representan uno de los principales objetivos dentro de escenarios de explotación automatizada.

Controles recomendados:

- Rotar credenciales comprometidas.
- Revocar tokens activos.
- Implementar MFA reforzado.
- Auditar accesos privilegiados.
- Monitorear reutilización de credenciales.
- Revisar secretos cloud.
- Fortalecer controles PAM.

Restaurar operación segura y fortalecer controles posteriores

Una vez contenido el incidente, las organizaciones deben fortalecer capacidades defensivas y reducir exposición residual.



Controles recomendados:

- Aplicar parches pendientes.
- Restaurar configuraciones seguras.
- Actualizar controles de monitoreo.
- Reforzar playbooks SOC/CSIRT.
- Implementar lecciones aprendidas.
- Capacitar equipos técnicos.
- Validar exposición residual.

Conclusiones



Las capacidades demostradas por Mythos confirman una evolución significativa en el uso de inteligencia artificial aplicada al descubrimiento automatizado de vulnerabilidades.

Aunque las evaluaciones independientes identificaron limitaciones relacionadas con razonamiento contextual y validación práctica, Mythos mostró capacidades avanzadas en auditoría automatizada de código, análisis de firmware, descubrimiento de vulnerabilidades y evaluación de aplicaciones complejas.

El principal riesgo asociado a este escenario corresponde a la reducción progresiva del tiempo entre descubrimiento y explotación de vulnerabilidades, lo cual podría impactar significativamente las estrategias tradicionales de gestión de vulnerabilidades y remediación.

Las organizaciones del sector salud deben fortalecer urgentemente controles relacionados con:

- gestión acelerada de vulnerabilidades;
- DevSecOps;
- monitoreo continuo;
- reducción de superficie de ataque;
- y protección de infraestructura crítica.



El crecimiento acelerado de capacidades ofensivas potenciadas por IA continuará transformando el panorama de amenazas y requerirá estrategias defensivas más ágiles, automatizadas y resilientes.

Fuentes:

- https://www.securityweek.com/mythos-proves-potent-in-vulnerability-discovery-less-convincing-elsewhere/?utm_source=chatgpt.com
- https://www.anthropic.com/?utm_source=chatgpt.com
- https://www.techradar.com/pro/security/mozilla-says-anthropics-mythos-preview-and-other-ai-models-helped-it-identify-and-ship-423-firefox-security-bug-fixes-in-just-one-month?utm_source=chatgpt.com
- https://it.slashdot.org/story/26/04/07/2115208/anthropic-unveils-claude-mythos-powerful-ai-with-major-cyber-implications?utm_source=chatgpt.com

En caso de dudas, inquietudes o requerir asistencia adicional relacionada con esta alerta de seguridad, puede comunicarse directamente con el **CSIRT Salud** a través de las líneas telefónicas **(+57) 316 893 1490 - 318 155 3570** o mediante el correo electrónico **csirtsalud@minsalud.gov.co**. Nuestro equipo está disponible para brindar el acompañamiento necesario.

