

Guía de Seguridad de las TIC

CCN-STIC 887F

Guía de respuesta a incidentes de seguridad en AWS



Marzo 2024



Catálogo de Publicaciones de la Administración General del Estado
<https://cpage.mpr.gob.es>

cpage.mpr.gob.es

Edita:



Pº de la Castellana 109, 28046 Madrid
© Centro Criptológico Nacional, 2024
NIPO: 083-24-124-5

Fecha de Edición: marzo de 2024

LIMITACIÓN DE RESPONSABILIDAD

El presente documento se proporciona de acuerdo con los términos en él recogidos, rechazando expresamente cualquier tipo de garantía implícita que se pueda encontrar relacionada. En ningún caso, el Centro Criptológico Nacional puede ser considerado responsable del daño directo, indirecto, fortuito o extraordinario derivado de la utilización de la información y software que se indican incluso cuando se advierta de tal posibilidad.

AVISO LEGAL

Quedan rigurosamente prohibidas, sin la autorización escrita del Centro Criptológico Nacional, bajo las sanciones establecidas en las leyes, la reproducción parcial o total de este documento por cualquier medio o procedimiento, comprendidos la reprografía y el tratamiento informático, y la distribución de ejemplares del mismo mediante alquiler o préstamo públicos.

ÍNDICE

1. GUÍA DE RESPUESTA A INCIDENTES DE SEGURIDAD EN AWS	4
1.1 DESCRIPCIÓN DEL USO DE ESTA GUÍA.....	4
1.3 ANTES DE EMPEZAR	6
2. RESPUESTA A INCIDENTES EN AWS	10
2.1 PREPARACIÓN.....	10
2.2 DETECCIÓN, ANÁLISIS E INVESTIGACIÓN	17
2.3 CONTENCIÓN, MITIGACIÓN Y RECUPERACIÓN	24
2.4 ACTIVIDAD POST-CIBERINCIDENTE	32
3. TECNOLOGÍAS Y SERVICIOS DE REFERENCIA PARA LA RESPUESTA A INCIDENTES EN AWS	34
3.1 PREPARACIÓN.....	34
3.2 DETECCIÓN, ANÁLISIS E INVESTIGACIÓN	38
3.3 CONTENCIÓN, MITIGACIÓN Y RECUPERACIÓN	43
4. RECURSOS ADICIONALES	48

1. GUÍA DE RESPUESTA A INCIDENTES DE SEGURIDAD EN AWS

1.1. Descripción del uso de esta guía

El propósito de esta Guía es ayudar a las capacidades de respuesta a ciberincidentes y su adecuado tratamiento, en entornos de nube pública de Amazon Web Services (AWS) al Sector Público, a los sistemas de información de las entidades del sector privado que se encuentren bajo el ámbito de aplicación del Esquema Nacional de Seguridad y a los operadores de servicios esenciales y proveedores de servicios digitales de acuerdo con el Real Decreto-ley 12/2018, de 7 de septiembre, de seguridad de las redes y sistemas de información.

Esta guía debe utilizarse como complemento a la guía **CCN-STIC 887A – Guía de Configuración Segura para AWS**. Y deberá tenerse especialmente en cuenta lo dispuesto en las guías de referencia **CCN-STIC 817 Esquema Nacional de Seguridad - Gestión de ciberincidentes** y la **Guía Nacional de Notificación y Gestión de Ciberincidentes**.

Del mismo modo, en el ámbito de la respuesta a ciberincidentes, constituyen el marco normativo de referencia los **artículos 33 y 34** del Esquema Nacional de Seguridad y la **Instrucción nacional de notificación y gestión de ciberincidentes** recogida en el Anexo del Real Decreto 43/2021, de 26 de enero, por el que se desarrolla el Real Decreto-ley 12/2018, de 7 de septiembre, de seguridad de las redes y sistemas de información.

En cuanto al público al que se dirige, esta guía está especialmente enfocada para:

- Equipos de Respuesta a Ciberincidentes internos a las organizaciones.
- Responsables de Seguridad de la Información (de obligada designación tanto para los operadores de servicios esenciales, de acuerdo con lo previsto en el artículo 16.3 del Real Decreto Ley 12/2018, como para las entidades bajo el ámbito de aplicación del ENS, de acuerdo con lo previsto en el artículo 13.2.c de la propia norma).
- Responsables de Sistemas de Información.
- En general, gestores de programas de ciberseguridad, administradores de red y sistemas y administradores y responsables de nube.

Por último, esta guía se estructura en:

- El presente apartado, introductorio.
- El apartado 2. Tecnologías y servicios de referencia para la respuesta a incidentes en AWS: se exponen las principales tecnologías que intervienen en la gestión de incidentes en AWS.
- El apartado 3. Respuesta a incidentes en AWS: aporta una visión general de los fundamentos de la respuesta a incidentes de seguridad en el entorno de la nube de AWS.
- Por último, en el apartado 4, recursos adicionales: se facilitan una serie

de recursos de código abierto desarrollados por el equipo de respuesta a incidentes de clientes de AWS.

Los apartados 2 y 3 se presentan en línea con las fases definidas en la guía **CCN-STIC 817 Esquema Nacional de Seguridad - Gestión de ciberincidentes** y en la **Guía nacional de notificación y gestión de ciberincidentes**, a saber:

- Preparación. Se trata de una fase inicial en la que el ciberincidente todavía no ha ocurrido y la organización enfoca sus esfuerzos en prepararse para cualquier suceso que pueda ocurrir.
- Detección, análisis e identificación. El objetivo de esta fase es tener la capacidad de identificar o detectar, con la menor dilación posible, cualquier ciberincidente que pueda sufrir la organización, desencadenando los procesos de notificación correspondientes.
- Contención. En el momento en que se ha identificado un ciberincidente, la máxima prioridad es contener el impacto de éste en la organización, de forma que se pueda evitar lo antes posible la propagación a otros sistemas o redes, minimizando o evitando un impacto mayor y la extracción de información fuera de la organización a través de la separación de equipos de la red afectada, inhabilitando cuentas comprometidas o cambiando contraseñas.
- Mitigación. Las medidas de mitigación dependerán del tipo de ciberincidente, ya que en algunos casos será necesario contar con apoyo de proveedores de servicios, y en otros casos puede suponer incluso el borrado completo de los sistemas afectados y la recuperación desde una copia de seguridad.
- Recuperación. La finalidad de la fase de recuperación consiste en devolver el nivel de operación a su estado normal y que las áreas de negocio afectadas puedan retomar su actividad. Es importante no precipitarse en la puesta en producción de sistemas que se han visto implicados en ciberincidentes.
- Actividad post-ciberincidente. Una vez que el ciberincidente está controlado y la actividad ha vuelto a la normalidad, es momento de llevar a cabo un proceso al que no se le suele dar toda la importancia que merece: las lecciones aprendidas. Además de las reflexiones sobre lo sucedido con el fin de aprender y evitar que una situación similar se repita, conviene realizar un informe del ciberincidente que detalle la causa y coste de este, así como las medidas que la organización debe tomar para prevenir futuros ciberincidentes similares.

Los asuntos relacionados con la seguridad y la conformidad son una responsabilidad compartida entre AWS, la cual puede ser consultada en la guía CCN-STIC 887A Guía de Configuración Segura para AWS en la sección 1.3.

1.2. Antes de empezar

Para comenzar, es importante comprender las particularidades de las operaciones de seguridad y de la respuesta ante ciberincidentes en la nube respecto a los ciberincidentes ocurridos en entornos tradicionales.

1.2.1. Principios y objetivos de la respuesta a incidentes en AWS

Aunque los objetivos y los principios de seguridad (como el de mínimo privilegio o el modelo de defensa en profundidad) son, a grandes rasgos, compartidos en entornos nube como en entornos on-premise, es recomendable tener en cuenta los siguientes principios, relevantes para la respuesta a ciberincidentes en la nube de AWS:

- Establecer objetivos de respuesta. Algunos objetivos comunes incluyen contener y mitigar el problema, recuperar recursos afectados, preservar los datos para análisis forenses, etc.
- Responder utilizando la nube. Implementando patrones de respuesta dentro de la nube.
- Saber lo que tiene y lo que necesita. Conservar los registros, recursos, instantáneas y otras pruebas copiándolos y almacenándolos en una cuenta centralizada en la nube dedicada a la respuesta. Utilizar etiquetas, metadatos y mecanismos que apliquen políticas de retención.
- Utilizar mecanismos de redistribución. Si una anomalía de seguridad puede atribuirse a un error de configuración, la solución puede ser tan sencilla como eliminar la variación redistribuyendo los recursos con la configuración adecuada.
- Automatizar en la medida de lo posible. A medida que surjan problemas o se repitan incidentes, se deben crear mecanismos para clasificar y responder de forma programática a los sucesos comunes.
- Elegir soluciones escalables. Implementando mecanismos de detección y respuesta escalables en todos sus entornos para reducir eficazmente el tiempo entre la detección y la respuesta.
- Aprender y mejorar el proceso. Se debe ser proactivo a la hora de identificar lagunas en sus procesos, herramientas o personal, y poner en marcha un plan para solucionarlas.

1.2.2. Dominios de incidentes de seguridad en la nube

Existen tres dominios dentro de la responsabilidad del cliente en los que pueden producirse incidentes de seguridad: servicio, infraestructura y aplicación:

- **Dominio de servicio:** Un incidente en el dominio de servicio es aquel al que responde exclusivamente con los mecanismos de la API de AWS o que tiene causas raíz asociadas con la configuración o permisos de recursos, y podría tener un registro de logs orientado a dicho servicio. Por ejemplo, podrían afectar a la cuenta de AWS, a los permisos IAM, a la facturación o a los metadatos de recursos.
- **Dominio de infraestructura.** Los incidentes en este dominio incluyen datos o actividad relacionada con la red, como procesos e instancias de Amazon EC2, tráfico dentro de la VPC y otras áreas como contenedores. La respuesta en este dominio a menudo implica la adquisición de datos relacionados con el incidente para su análisis forense, como las capturas de tráfico de red, los bloqueos dentro de un volumen EBS o la memoria volátil de una instancia.
- **Dominio de aplicación.** En este dominio se consideran los actores que podrían actuar contra cuentas, recursos o datos de AWS. En este ámbito conviene disponer de un marco de riesgos para determinar los riesgos específicos a los que se expone la organización y actuar en consecuencia.

1.2.3. Diferencias clave en la respuesta a incidentes en AWS

A continuación, se exponen las diferencias clave de la respuesta a incidentes en la nube de AWS frente a los entornos on-premise.

Diferencia 1. La seguridad como responsabilidad compartida

Tal y como ya se ha adelantado, el modelo de responsabilidad compartida alivia parte de la carga operativa del cliente porque AWS opera, administra y controla los componentes desde el sistema operativo del host y la capa de virtualización hasta la seguridad física de las instalaciones en las que opera el servicio.

A medida que cambia la responsabilidad compartida en la nube, también cambian las opciones de respuesta ante incidentes. Planificar y comprender estas compensaciones y ajustarlas a sus necesidades de gobierno es un paso crucial en la respuesta a incidentes.

Además de la relación directa que se tiene con AWS, puede haber otras entidades que tengan responsabilidades en el modelo de responsabilidad particular. Por ejemplo, se puede tener unidades organizativas internas que asuman la responsabilidad de algunos aspectos de las operaciones. También se puede tener socios u otras partes que desarrollen, administren u operen parte de la tecnología en la nube.

Con carácter general, son responsabilidad del cliente:

- Todas las actividades que ocurran en la cuenta, incluyendo el acceso no

autorizado.

- Configurar y utilizar correctamente los servicios de AWS.
- Mantener actualizado el correo electrónico de la cuenta raíz de AWS para las notificaciones.
- Tomar las medidas adecuadas para asegurar, proteger y hacer copias de seguridad de la cuenta y su contenido.
- No revelar las credenciales ni las claves de acceso a terceros no autorizados.

Algunas de las causas más comunes de la ocurrencia de incidentes de seguridad son:

- Dependencia de versiones obsoletas de software en las aplicaciones.
- Escasez de actuaciones en materia de gestión de vulnerabilidades.
- Selección de configuraciones inseguras en los recursos de AWS.
- Gestión insuficiente de la seguridad del software de la aplicación.
- Respuesta ineficaz a los hallazgos de Amazon GuardDuty y otros controles de detección.
- Acciones involuntarias de revelación de credenciales y secretos de seguridad.

Diferencia 2. Dominio de servicio en la nube

Debido a las diferencias en la responsabilidad de seguridad que existen en los servicios en la nube, se introdujo un nuevo dominio para los incidentes de seguridad: el dominio de servicio, que ha sido expuesto en el apartado 1.2.2 Dominios de incidentes de seguridad en la nube. El dominio de servicio abarca la cuenta de AWS, los permisos de Amazon IAM, los metadatos de recursos, la facturación y otras áreas. Este dominio es diferente para la respuesta a incidentes debido a cómo se responde. La respuesta dentro del dominio de servicio se realiza normalmente mediante la revisión y la emisión de llamadas a la API, en lugar de la respuesta tradicional basada en el host y la red. En el dominio de servicio, no se interactúa con el sistema operativo de un servicio afectado.

Diferencia 3. APIs para el aprovisionamiento de la infraestructura

Otra diferencia proviene de la característica de la nube de autoservicio bajo demanda. Los clientes interactúan con la nube de AWS utilizando una API RESTful a través de endpoints públicos y privados disponibles en muchas ubicaciones geográficas de todo el mundo. Los clientes pueden acceder a estas API con credenciales de AWS. A diferencia del control de acceso on-premise, estas credenciales no están necesariamente vinculadas a una red o a un dominio de Microsoft Active Directory. En su lugar, las credenciales están asociadas a una entidad de seguridad de IAM dentro de una cuenta de AWS. Se puede acceder a estos endpoints de API fuera de la red corporativa, lo que será importante comprender cuando se responda a un incidente en el que las credenciales se utilicen fuera de la red o geografía previstas.

Debido a la naturaleza basada en API de AWS, una fuente de registro importante es AWS CloudTrail, que rastrea las llamadas a la API de administración realizadas en las cuentas de AWS y donde se puede encontrar información sobre las llamadas a la API.

Gracias a esta naturaleza basada en las APIs, con AWS las organizaciones pueden aprovechar plenamente las ventajas de la adopción de nube a través de una estrategia operativa fundamentada en la automatización. La infraestructura como código (IaC) es un patrón de entornos automatizados altamente eficientes en los que los servicios de AWS se implementan, configuran, reconfiguran y destruyen mediante código facilitado por servicios nativos como AWS CloudFormation o soluciones de terceros. Esto empuja a que la implementación de la respuesta a incidentes sea altamente automatizada, lo cual es deseable para evitar errores humanos.

Diferencia 4. Naturaleza dinámica de la nube

La nube es dinámica; permite crear y eliminar recursos rápidamente. Con el autoescalado, los recursos se pueden aumentar y reducir en función del aumento del tráfico. Con una infraestructura efímera y cambios rápidos, puede que un recurso que se esté investigando ya no exista o haya sido modificado. Comprender la naturaleza efímera de los recursos de AWS y cómo puede realizar un seguimiento de la creación y eliminación de los recursos de AWS será importante para el análisis de incidentes. Puede utilizar AWS Config para realizar un seguimiento del historial de configuración de sus recursos de AWS.

La naturaleza dinámica de la nube posee un gran potencial para la respuesta a incidentes de seguridad. A diferencia de los sistemas tradicionales, la capacidad para crear, modificar o replicar la infraestructura rápidamente permite la adaptación de la aplicación a los eventos de seguridad que puedan ocurrir. Por ejemplo, ante la indisponibilidad de una ubicación, es posible replicar la infraestructura fácilmente en una ubicación diferente.

Diferencia 5. Acceso a los datos

Dentro de un marco de análisis forense en el que se requiere un acceso privilegiado a los datos, el acceso a los datos es diferente debido al modelo de responsabilidad. El cliente podrá acceder a los datos y recursos de los que ostente responsabilidad, pero no podrá acceder físicamente a los datos y recursos que están bajo la responsabilidad del proveedor. Por ejemplo, el cliente podrá acceder a los datos de una instancia EC2 y correr las herramientas de análisis forense que considere necesarias, pero no podrá acceder a los datos de la tecnología subyacente que soporta el servicio EC2. En cualquier caso, AWS pone a disposición del cliente servicios para ayudar en el análisis forense. Con carácter general, los datos se recopilan a través de llamadas a la API. Se deberá practicar y entender cómo realizar la recopilación de datos a través de API para estar preparado y verificar el almacenamiento adecuado para una recopilación y acceso efectivos. Además, se deberá establecer mecanismos para garantizar la seguridad en la cadena de

custodia de los datos forenses.

2. RESPUESTA A INCIDENTES EN AWS

A continuación, se presentan una visión general de los fundamentos de la respuesta a incidentes de seguridad dentro del entorno de la nube de Amazon Web Services (AWS).

2.1. Preparación

Es importante hacer hincapié en la preparación, no sólo para establecer una capacidad de respuesta a incidentes de modo que la organización esté preparada para responder a ellos, sino también para prevenirlos asegurándose de que los sistemas, redes y aplicaciones son lo suficientemente seguros. Más allá de disponer de una línea base de configuración de seguridad (ver la **CCN-STIC 887A – Guía de Configuración Segura para AWS**), es recomendable que la preparación en AWS se realice en torno a los ámbitos de personal, procesos y tecnologías. Cada uno de estos ámbitos es igualmente importante para la respuesta eficaz a los ciberincidentes, sin que ningún plan o procedimiento de gestión de incidentes se considere completo si no contempla los tres.

2.1.1. Personas

Para responder a los incidentes de seguridad se necesita identificar a las partes interesadas que apoyarían la respuesta.

Definición de roles y responsabilidades

Debido a que no podemos predecir ni codificar todas las posibles direcciones que tomará un incidente de seguridad, se recomienda otorgar más peso a la automatización para tareas simples y repetitivas y dejar que las personas tomen las decisiones difíciles. La gestión de eventos de seguridad poco claros requiere de una buena definición de roles y de comunicación entre diferentes responsables. Hay que tener bien definidos los roles y responsabilidades y si debe participar alguna tercera parte. La creación de una matriz de asignación de responsables, aprobadores, consultados e informados (RACI) para un incidente puede parecer un trámite burocrático, pero permite una identificación rápida y directa, y describe claramente el liderazgo en diferentes etapas y ámbitos de la gestión del incidente.

En este sentido, los marcos de seguridad de referencia, como el ENS (artículo 13.3) establecen la distinción de los roles del responsable de la seguridad y del responsable del sistema, así como la independencia jerárquica entre ambos. En caso de que ambas funciones recaigan en la misma persona por ausencia justificada de recursos en la organización, se deberían aplicar medidas compensatorias para garantizar la distinción de responsabilidades.

Formación del personal de respuesta

La formación del personal de respuesta a ciberincidentes sobre las tecnologías que

utiliza la organización será crucial para que puedan responder adecuadamente a un incidente de seguridad.

Además de los conceptos tradicionales de respuesta a ciberincidentes, también es importante que comprendan los servicios de AWS y su entorno. Existen varios mecanismos tradicionales para formar al personal, como la formación online y presencial. También es importante considerar la realización de simulacros como mecanismo de formación.

Por una parte, AWS ofrece [talleres de seguridad online](#) en los que puede obtener experiencias prácticas con los servicios de seguridad y monitorización de AWS. AWS también ofrece una serie de opciones de formación y vías de aprendizaje a través de formación digital, formación presencial, socios y certificaciones. Para obtener más información, consulte [Formación y certificación de AWS](#).

Equipos de respuesta y soporte de AWS

AWS Support

AWS Support ofrece una gama de planes que proporcionan acceso a herramientas y experiencia que respaldan la salud operativa de las soluciones de AWS. Si se necesita soporte técnico y más recursos para planificar, implementar y optimizar el entorno de AWS, se puede seleccionar el plan de soporte que mejor se adapte al caso de uso de AWS en la organización.

El Centro de soporte de la consola de administración de AWS es el punto de contacto central para obtener soporte para problemas que afecten a sus recursos de AWS. El acceso a AWS Support está controlado por IAM, requiriéndose una política específica para que el usuario IAM pueda acceder a AWS Support.

Tal y como se expone en el apartado Gestión diaria [op.ext.2] de la guía **CCN-STIC 887A Guía de configuración segura para AWS**, es recomendable la asignación a algún usuario del rol iam_support_role para poder llevar a cabo la gestión e los incidentes con el soporte técnico de AWS.

Equipo de respuesta a incidentes de clientes de AWS (CIRT)

El Equipo de Respuesta a Incidentes de Clientes de AWS (CIRT) es un equipo global especializado de AWS que proporciona soporte a los clientes durante eventos de seguridad activos en el lado del cliente del Modelo de Responsabilidad Compartida de AWS.

Los clientes de AWS pueden ponerse en contacto con AWS CIRT a través de un caso de AWS Support.

Soporte de respuesta DDoS

AWS ofrece AWS Shield, un servicio administrado de protección contra denegación de servicio distribuido (DDoS) que protege las aplicaciones web que se ejecutan en AWS.

AWS Shield en su modalidad Advanced ofrece un nivel de protección superior contra ataques dirigidos y también le proporciona acceso las 24 horas del día, los 7 días de la semana al equipo de respuesta de AWS Shield Response Team (SRT), ya sea de forma proactiva cuando el equipo detecta un evento DDoS o bajo solicitud del cliente.

Para más información, se puede consultar el apartado Protección frente a la denegación de servicio [mp.s.4] de la guía **CCN-STIC 887A Guía de configuración segura para AWS**.

2.1.2. Procedimientos

Cuando se produce un incidente de seguridad, unos pasos y flujos de trabajo claros ayudarán a la organización a responder a tiempo. Si bien es posible que la organización ya disponga de procedimientos de respuesta a incidentes, es importante actualizar, iterar y probarlos con regularidad, especialmente si se deben incorporar nuevas particularidades sobre la gestión de incidentes en la nube.

Desarrollar y probar un plan de respuesta a incidentes

El plan o procedimiento de respuesta a incidentes es el primer documento a desarrollar para la respuesta a incidentes. Es un documento de alto nivel que normalmente suele incluir las siguientes secciones:

- Visión general del equipo de respuesta, definiendo los objetivos y funciones del equipo de respuesta a incidentes.
- Funciones y responsabilidades: enumera las partes interesadas en la respuesta a incidentes y detalla sus funciones.
- Plan de comunicación: detalla la información de contacto y cómo se comunicará durante un incidente.
- Fases de la respuesta a incidente y acciones de alto nivel a tomar dentro de cada fase.
- Definición de la gravedad y priorización de incidentes: detalla cómo clasificar la gravedad de los incidentes, cómo priorizarlos y cómo afectan las definiciones de gravedad a los procedimientos de escalado y notificación. En este sentido, se recomienda seguir los criterios de determinación del nivel de peligrosidad e impacto que se definen en la Guía y en la Instrucción Nacional de Notificación y Gestión de Incidentes

Documentar y centralizar diagramas de arquitectura

La comprensión de cómo están estructurados los sistemas y las redes es clave no sólo para la respuesta a incidentes, sino también para verificar la coherencia entre las aplicaciones.

Se debe desarrollar documentación y repositorios internos que detallen elementos como los siguientes:

- Estructura de cuentas de AWS (número de cuentas, cómo están organizadas, quiénes son los responsables, etc.).
- Patrones de servicios de AWS (qué servicios se usan, qué servicios son los más usados, etc.).
- Patrones de arquitectura.
- Patrones de autenticación en AWS (cómo se autentican los desarrolladores, se usan roles o usuarios IAM, está la autenticación centralizada en un proveedor de identidades externo, etc.).
- Patrones de autorización en AWS.
- Registro y monitorización (qué fuentes de registro se usan y dónde se guardan, habilitación de GuardDuty, agregación de eventos en un SIEM, etc.).
- Topología de red (cuántos dispositivos hay conectados a la red, cómo se conecta la red con AWS en el plano físico y lógico, etc.).
- Infraestructura externa.

Desarrollo de guías de respuesta a incidentes

Una parte clave de la preparación de sus procesos de respuesta a incidentes es el desarrollo de guías. Las guías de respuesta a incidentes proporcionan una serie de directrices y pasos concretos a seguir cuando se produce un incidente de seguridad. Contar con una estructura y unos pasos claros simplifica la respuesta y reduce la probabilidad de que se produzcan errores humanos.

Las guías deben crearse para escenarios de incidentes como incidentes previstos (como DoS/DDoS, ransomware o compromiso de credenciales) y descubrimientos y alertas de seguridad como los hallazgos de GuardDuty.

El siguiente [enlace](#) proporciona algunos ejemplos de guías de respuesta a incidentes.

Ejecutar simulacros periódicos

Para una adecuada respuesta a incidentes en la organización es importante revisar continuamente las capacidades de respuesta ante incidentes. Los simulacros son un método que puede utilizarse para realizar esta evaluación. Los simulacros utilizan escenarios de eventos de seguridad del mundo real diseñados para imitar las tácticas, técnicas y procedimientos de un actor de amenazas y permiten a una organización ejercitar y evaluar sus capacidades de respuesta a incidentes respondiendo a estos eventos cibernéticos simulados como podrían ocurrir en la realidad, permitiendo de una manera estructurada, practicar las guías y procedimientos de respuesta ante incidentes en escenarios realistas.

Entre los tipos de simulacros se encuentran:

- Ejercicios tipo tabletop: Consiste en sesiones basadas en el debate en las que participan las distintas partes interesadas en la respuesta a incidentes para practicar las funciones y responsabilidades y utilizar las herramientas de comunicación y las guías de respuesta a incidentes establecidas. Este ejercicio de simulación se centra en los procesos, las personas y la colaboración.
- Ejercicios de Purple Team: Los ejercicios del Purple Team aumentan el nivel de colaboración entre el personal de respuesta a incidentes (Blue Team) y los actores de amenazas simuladas (Red Team). El Equipo Azul se compone generalmente de miembros del Centro de Operaciones de Seguridad (SOC), pero también puede incluir otras partes interesadas que participarían durante un evento real. El Equipo Rojo se compone generalmente de un equipo de pruebas de penetración o de partes interesadas clave, formadas en seguridad ofensiva.
- Ejercicios de Red Team: Durante un ejercicio del Equipo Rojo, la ofensiva (Equipo Rojo) lleva a cabo un simulacro para lograr un objetivo o conjunto de objetivos de un alcance predeterminado. Los defensores (Equipo Azul) no tendrán necesariamente conocimiento del alcance y duración del ejercicio, lo que proporciona una evaluación más realista de cómo responderían a un incidente real.

2.1.3. Tecnología

Desarrollar una estructura de cuentas de AWS

AWS Organizations ayuda a administrar y gobernar de forma centralizada un entorno de AWS a medida que crece y escala los recursos de AWS. Una organización de AWS consolida sus cuentas de AWS para que pueda administrarlas como una sola unidad. Puede utilizar unidades organizativas (OU) para agrupar cuentas y administrarlas como una sola unidad.

Para la respuesta a incidentes, es útil tener una estructura de cuenta de AWS que admita las funciones de respuesta a incidentes, que incluye una OU de seguridad y una OU forense.

En la guía **CCN-STIC 887D Guía de configuración segura para entornos multi-cuenta en AWS** se profundiza en el funcionamiento de AWS Organizations y se proponen arquitecturas de unidades organizativas y cuentas recomendadas.

Desarrollar y aplicar una estrategia de etiquetado

Las etiquetas, que permiten asignar metadatos a los recursos de AWS, constan de una clave y un valor definidos por el usuario. Se pueden crear etiquetas para clasificar los recursos por finalidad, propietario, entorno, tipo de datos procesados y otros criterios.

Contar con una estrategia de etiquetado coherente puede acelerar los tiempos de

respuesta al permitir identificar y discernir rápidamente información contextual sobre un recurso de AWS. Las etiquetas también pueden servir como mecanismo para iniciar automatizaciones de respuesta.

Para obtener más información sobre qué etiquetar, puede consultar la documentación sobre el etiquetado de recursos de AWS. En primer lugar, se deberán definir las etiquetas que se desean implementar en la organización. Después, implementar y aplicar esta estrategia de etiquetado.

En el blog de AWS [Implement AWS resource tagging strategy using AWS Tag Policies and Service Control Policies \(SCPs\)](#) se puede encontrar información detallada sobre la implementación y la aplicación de una estrategia de etiquetado.

Actualizar la información de contacto de la cuenta de AWS

Para cada una de las cuentas de AWS, es importante disponer de información de contacto precisa y actualizada para que las partes interesadas correctas reciban notificaciones importantes de AWS sobre temas como la seguridad, la facturación y las operaciones. Para cada cuenta de AWS, se tiene un contacto principal y contactos alternativos para seguridad, facturación y operaciones.

Es una práctica recomendada utilizar una lista de distribución de correo electrónico como información de contacto para eliminar las dependencias hacia una sola persona.

Preparar el acceso a las cuentas de AWS

Durante un incidente los diferentes equipos que se encarguen de responder se deberán tener accesos adecuados a cada función para realizar las diferentes tareas que se requieran.

Será necesario identificar y debatir la estrategia de cuentas de AWS y la estrategia de identidades en la nube con los arquitectos de la nube de la organización para comprender qué métodos de autenticación y autorización están configurados, por ejemplo:

- Federación: un usuario asume un rol de IAM en una cuenta de AWS de un proveedor de identidades.
- Acceso entre cuentas: un usuario asume un rol de IAM en varias cuentas de AWS.
- Autenticación: un usuario se autentica como un usuario de AWS IAM creado dentro de una sola cuenta de AWS.

Comprender el panorama de las amenazas

Para la respuesta a incidentes, un modelo de amenazas puede ayudar a identificar los vectores de ataque que un actor de amenazas podría haber utilizado durante un incidente. Comprender de qué se está defendiendo será crucial para responder a tiempo.

La inteligencia sobre amenazas cibernéticas son los datos y el análisis de la

intención, la oportunidad y la capacidad de un actor de amenazas. La obtención y el uso de inteligencia sobre amenazas es útil para detectar un incidente en una fase temprana y para comprender mejor el comportamiento del actor de la amenaza. La inteligencia sobre ciberamenazas incluye indicadores estáticos como direcciones IP o hashes de archivos de malware. También incluye información de alto nivel, como patrones de comportamiento e intenciones. Puede recopilar información sobre amenazas de varios proveedores de ciberseguridad y de repositorios de código abierto.

Para integrar y maximizar la inteligencia sobre amenazas para el entorno de AWS, se puede utilizar algunas capacidades listas para usar e integrar propias listas de inteligencia sobre amenazas. Amazon GuardDuty utiliza fuentes de inteligencia sobre amenazas internas de AWS y de terceros.

Seleccionar y configurar los logs para análisis y alertas

Durante una investigación de seguridad, es necesario poder revisar los registros pertinentes para registrar y comprender todo el alcance y la cronología del incidente. Los registros también son necesarios para generar alertas que indiquen que se han producido determinadas acciones de interés. Es fundamental seleccionar, habilitar, almacenar y configurar los mecanismos de consulta y recuperación, así como configurar las alertas.

- Seleccionar y habilitar las fuentes de logs. Es preciso definir cuáles van a ser las fuentes de logs relevantes para las cargas de trabajo de la cuenta de AWS. Para ello se debe elegir entre las herramientas de monitorización definidas a lo largo del apartado 2 de este documento, así como en la guía **CCN-STIC 887G Guía de configuración segura para monitorización y gestión en AWS**.
- Seleccionar el almacenamiento de registros. La elección del almacenamiento de logs suele estar relacionada con la herramienta de consulta que utilice, las capacidades de retención, la familiaridad y el coste. Cuando se habilite los logs de servicio de AWS, se debe proporcionar una instalación de almacenamiento; normalmente un bucket de Amazon S3 o un grupo de logs de CloudWatch.
- Identificar un periodo de retención de logs adecuado. Cuando se utilice un bucket de S3 o un grupo de registros de CloudWatch para almacenar registros, se debe establecer ciclos de vida adecuados para cada fuente de registro con el fin de optimizar los costes de almacenamiento y recuperación. Por lo general, se dispone de entre 3 y 12 meses de registros disponibles para su consulta, con una retención de hasta siete años.
- Seleccionar y aplicar mecanismos de consulta de registros. En AWS, los principales servicios que se pueden utilizar para consultar logs son CloudWatch Logs Insights para datos almacenados en grupos de logs de CloudWatch, Amazon Athena y Amazon OpenSearch Service para datos almacenados en Amazon S3 y CloudTrail Lake y Security Lake para la consulta de los registros procedentes de múltiples fuentes. También se

pueden utilizar herramientas de consulta de terceros, como un sistema de información de seguridad y gestión de eventos (SIEM).

- El proceso de selección de una herramienta de consulta de logs debe tener en cuenta los aspectos relacionados con las personas, los procesos y la tecnología de las operaciones de seguridad. Seleccionar una herramienta que cumpla los requisitos operativos, empresariales y de seguridad, y que sea accesible y fácil de mantener a largo plazo. Se debe tener en cuenta que las herramientas de consulta de registros funcionan de forma óptima cuando el número de registros que hay que escanear se mantiene dentro de los límites de la herramienta.
- AWS proporciona alertas de forma nativa a través de servicios de seguridad, como Amazon GuardDuty, AWS Security Hub y AWS Config. También se pueden utilizar motores de generación de alertas personalizados para alertas de seguridad no cubiertas por estos servicios o para alertas específicas relevantes.

Desarrollar capacidades forenses

Los conceptos de la investigación forense tradicional en las instalaciones se aplican a AWS. Una vez que se tenga el entorno y la estructura de la cuenta de AWS configurados para la investigación forense, se querrá definir las tecnologías necesarias para llevar a cabo eficazmente metodologías forenses sólidas en las cuatro fases:

- **Recopilación:** Recopilación de logs de AWS relevantes, como AWS CloudTrail, AWS Config, VPC Flow Logs y logs a nivel de host. Recopilar instantáneas, backups y volcados de memoria de los recursos de AWS afectados.
- **Examen:** Examinar los datos recopilados extrayendo y evaluando la información relevante.
- **Análisis:** Analizar los datos recopilados para comprender el incidente y extraer conclusiones.
- **Elaboración de informes:** Presentar la información resultante de la fase de análisis.

En este sentido, la creación de copias de seguridad de sistemas y bases de datos clave es fundamental para fines forenses. En AWS puede tomar instantáneas de varios recursos. Las instantáneas le proporcionan copias de seguridad puntuales de esos recursos. Existen varios servicios de AWS que pueden ayudar a realizar backups y recuperaciones. Para más información sobre la realización de copias de seguridad se puede consultar el apartado Copias de seguridad [mp.info.6] de la Guía **CCN-STIC 887A Guía de Configuración Segura para AWS**.

2.2. Detección, análisis e investigación

La adecuada implantación de las medidas de seguridad antedichas ayudará a detectar los posibles incidentes de seguridad de la información. La detección

consiste en la identificación de una amenaza que ha penetrado en el sistema de información o de una anomalía que requieran de un análisis en profundidad. El análisis consiste en la confirmación de si el evento detectado constituye un incidente de seguridad y la identificación de la naturaleza y la gravedad del mismo. Como último aspecto de esta fase, la investigación consiste en la averiguación de detalle del incidente, identificando aspectos como el vector de entrada o el alcance del mismo.

2.2.1. Detección

Una alerta es el componente principal de la fase de detección. Genera un hallazgo para iniciar el proceso de respuesta a incidentes basándose en la actividad de interés de la cuenta de AWS.

Si bien la precisión de las alertas es un reto; no siempre es posible determinar con total certeza si se ha producido un incidente, si está en curso o si se producirá en el futuro. No obstante, se recomienda seguir las siguientes prácticas a la hora de preparar la detección de incidentes en el entorno AWS.

Establecer las fuentes de alertas

Se recomienda considerar el uso de las siguientes fuentes para definir alertas:

- Hallazgos de servicios de AWS como GuardDuty, Security Hub, Inspector o Config.
- Logs. Los logs de servicios, infraestructura y aplicaciones de AWS almacenados en buckets de S3 y grupos de logs de CloudWatch pueden analizarse y correlacionarse para generar alertas.
- Actividad de facturación. Los cambios repentinos en la actividad de facturación pueden indicar incidentes de seguridad.
- Inteligencia sobre ciberamenazas. Al suscribirse a una fuente de inteligencia sobre ciberamenazas de terceros, se puede correlacionar esa información con otras herramientas de registro y monitorización para identificar posibles indicadores de eventos.
- Herramientas de partners. Para la respuesta ante incidentes, los productos de socios con detección y respuesta de endpoints (EDR) o SIEM pueden ayudar a respaldar los objetivos de respuesta ante incidentes.
- Seguridad de AWS. AWS Support podría ponerse en contacto con los clientes si se identifican actividades abusivas o maliciosas.
- Contacto único. Dado que pueden ser los clientes, desarrolladores u otro personal de la organización quienes adviertan algo inusual, es importante disponer de un método conocido y bien publicitado para ponerse en contacto con el equipo de seguridad. Las opciones más populares son los sistemas de tickets, las direcciones de correo electrónico de contacto y los formularios web. Si la organización trabaja con el público en general, es posible que también se necesite un mecanismo de contacto de seguridad

de cara al público.

La detección como parte de la ingeniería de control de seguridad

Los mecanismos de detección forman parte integrante del desarrollo de los controles de seguridad. A medida que se definen los controles directivos y preventivos, deben construirse controles de detección y respuesta relacionados.

Por ejemplo, una organización establece un control relacionado con el usuario raíz de una cuenta de AWS, que sólo debe utilizarse para actividades específicas y muy bien definidas. Se asocia con un control preventivo implementado mediante el uso de una política de control de servicios (SCP) de la organización de AWS. Si se produce una actividad de usuario root superior a la línea de base esperada, un control detectivo implementado con una regla EventBridge y un tema SNS alertará al centro de operaciones de seguridad (SOC). El control de respuesta implica que el SOC seleccione la guía de respuesta adecuada, se debe realizar un análisis y trabajar hasta que se resuelva el incidente.

Los controles de seguridad se definen mejor mediante el modelado de amenazas de las cargas de trabajo que se ejecutan en AWS. La criticidad de los controles de detección se establecerá observando el análisis de impacto empresarial (BIA) para la carga de trabajo concreta. Las alertas generadas por los controles de detección no se gestionan a medida que llegan, sino en función de su criticidad inicial, que se ajustará durante el análisis. La criticidad inicial es una ayuda para la priorización; el contexto en el que se produjo la alerta determinará su verdadera criticidad.

Por ejemplo, una organización utiliza Amazon GuardDuty como componente del control de detecciones utilizado para instancias EC2 que forman parte de una carga de trabajo. Se genera el hallazgo `Impact:EC2/SuspiciousDomainRequest.Reputation`, que le informa de que la instancia de Amazon EC2 incluida en su carga de trabajo está consultando un nombre de dominio sospechoso de ser malicioso. Esta alerta está configurada por defecto como de gravedad baja y, a medida que avanza la fase de análisis, se determina que un actor no autorizado ha desplegado varios cientos de instancias EC2 de tipo `p4d.24xlarge`, lo que ha aumentado significativamente el coste operativo de la organización. En este punto, el equipo de respuesta a incidentes toma la decisión de ajustar la criticidad de esta alerta a alta, aumentando el sentido de urgencia y acelerando las acciones posteriores. Puede observarse que la severidad del hallazgo de GuardDuty no puede ser cambiada, más bien, la alerta de la organización basada en el hallazgo tendrá que ajustar su criticidad.

La implementación de controles de detección

Es importante comprender cómo se implementan los controles de detección porque ayudan a determinar cómo se utilizará la alerta para el suceso concreto. Hay dos implementaciones principales de estos controles:

- Detección basada en el comportamiento. La detección basada en el comportamiento utiliza modelos matemáticos de aprendizaje automático

(ML) o inteligencia artificial (IA). La detección se realiza por inferencia; por tanto, la alerta puede no reflejar necesariamente un suceso real.

- Detección basada en reglas. Es determinista; los clientes pueden establecer los parámetros exactos de la actividad sobre la que se debe alertar.

Las implementaciones modernas de sistemas de detección, como los sistemas de detección de intrusos (IDS), suelen incluir ambos mecanismos. A continuación, se exponen algunos ejemplos de detecciones basadas en reglas y en comportamiento con GuardDuty:

- El hallazgo `Exfiltration:IAMUser/AnomalousBehavior`, informa que “se observó una solicitud API anómala en tu cuenta”. En este caso, el modelo ML evalúa todas las peticiones API de la cuenta e identifica eventos anómalos que están asociados a técnicas utilizadas por adversarios, lo que indica que este hallazgo está basado en el comportamiento.
- En el hallazgo `Impact:S3/MaliciousIPCaller`, GuardDuty está analizando llamadas API desde el servicio Amazon S3 en CloudTrail, comparando el elemento de registro `SourceIPAddress` con una tabla de direcciones IP públicas que incluye feeds de inteligencia de amenazas. Una vez que encuentra una coincidencia directa con una entrada, genera el hallazgo, es decir, se trata de un hallazgo basado en reglas.

Se recomienda implementar una combinación de alertas basadas tanto en comportamiento como en reglas, ya que no siempre es posible implementar alertas basadas en reglas para cada actividad dentro de su modo de amenaza.

Detección basada en personas

La otra fuente importante de detección procede de personas internas o externas a la organización. Las personas internas son normalmente los empleados y los contratistas, y las externas son entidades como los investigadores de seguridad, las fuerzas de seguridad, las noticias y los medios sociales.

Aunque la detección basada en la tecnología puede configurarse sistemáticamente, la detección basada en las personas se presenta de diversas formas, como correos electrónicos, tickets, correo, publicaciones en las noticias, llamadas telefónicas e interacciones en persona. Se puede esperar que las notificaciones de detección basadas en la tecnología se entreguen casi en tiempo real, pero no hay expectativas de plazos para la detección basada en las personas. Es imperativo que la cultura de seguridad incorpore, facilite y potencie los mecanismos de detección basados en las personas para un enfoque de seguridad de defensa en profundidad.

2.2.2. Análisis

Los registros, las capacidades de consulta y la inteligencia sobre amenazas son algunos de los componentes de apoyo necesarios para la fase de análisis. Muchos de los registros utilizados para la detección también se utilizan para el análisis y

requerirán la incorporación y configuración de herramientas de consulta.

Validación, alcance y evaluación de impacto de la alerta

Durante la fase de análisis, se realiza un análisis exhaustivo de los registros con el objetivo de validar las alertas, definir el alcance y evaluar el impacto del posible compromiso. La validación de la alerta es el punto de entrada de la fase de análisis. El personal de respuesta a incidentes debe buscar entradas de registro de varias fuentes y ponerse en contacto directo con los propietarios de la carga de trabajo afectada.

La determinación del alcance es el siguiente paso, en el que se inventariarán todos los recursos implicados y se ajustará la criticidad de la alerta después de que las partes interesadas acuerden que es poco probable que se trate de un falso positivo.

Por último, el análisis de impacto detalla la interrupción real del negocio.

Una vez identificados los componentes de la carga de trabajo afectados, los resultados del alcance pueden correlacionarse con el objetivo de punto de recuperación (RPO) y el objetivo de tiempo de recuperación (RTO) de la carga de trabajo relacionada, ajustando la criticidad de la alerta, lo que iniciará la asignación de recursos y toda la actividad que tenga lugar a continuación. No todos los incidentes interrumpirán directamente las operaciones de una carga de trabajo que soporta un proceso de negocio. Incidentes como la revelación de datos confidenciales, el robo de propiedad intelectual o el secuestro de recursos (como en la minería de criptomonedas) pueden no detener o debilitar un proceso de negocio inmediatamente, pero pueden tener consecuencias más adelante.

Enriquecer los registros y hallazgos de seguridad

Una vez establecidas las fuentes de registros de eventos y hallazgos de seguridad, se pueden enriquecer con inteligencia sobre amenazas y contexto organizativo.

Los servicios de inteligencia sobre amenazas se utilizan para asignar reputación y atribuir propiedad a direcciones IP públicas, nombres de dominio y hashes de archivos. Estas herramientas están disponibles como servicios de pago y gratuitos.

Las organizaciones que adoptan Amazon Athena como herramienta de consulta de logs obtienen la ventaja de los trabajos de AWS Glue para cargar información de inteligencia sobre amenazas como tablas. Las tablas de inteligencia sobre amenazas se pueden utilizar en consultas SQL para correlacionar elementos de log como direcciones IP y nombres de dominio, lo que proporciona una vista enriquecida de los datos que se van a analizar.

AWS no proporciona inteligencia sobre amenazas directamente a los clientes, pero servicios como Amazon GuardDuty hacen uso de la inteligencia sobre amenazas para el enriquecimiento y la generación de hallazgos. También puede cargar listas de amenazas personalizadas en GuardDuty basadas en su propia inteligencia sobre amenazas.

Por otra parte, los hallazgos se pueden enriquecer con automatización. La automatización es una parte integral de la gobernanza de la nube de AWS. Se puede utilizar a lo largo de las distintas fases del ciclo de vida de respuesta a incidentes.

La fase de análisis puede aprovechar el mecanismo de detección y reenviar el cuerpo de la alerta a un motor capaz de consultar los registros y enriquecer los observables para contextualizar el suceso.

El cuerpo de la alerta, en su forma fundamental, se compone de un recurso y una identidad. Por ejemplo, podría implementar una automatización para consultar CloudTrail en busca de la actividad de la API de AWS realizada por la identidad o el recurso del cuerpo de alerta en el momento de la alerta, proporcionando información adicional como EventSource, EventName, SourceIPAddress y UserAgent de la actividad de la API identificada. Al realizar estas consultas de forma automatizada, el equipo de respuesta puede ahorrar tiempo durante el triaje y obtener contexto adicional para ayudar a tomar decisiones mejor informadas.

La publicación del blog [How to enrich AWS Security Hub findings with account metadata](#) contiene un ejemplo de cómo utilizar la automatización para enriquecer los hallazgos de seguridad y simplificar el análisis.

Recolectar y analizar la evidencia forense

El proceso forense tiene las siguientes características fundamentales:

- Coherente: Sigue exactamente los pasos documentados, sin desviaciones.
- Repetible: Produce exactamente los mismos resultados cuando se repite con el mismo artefacto.
- Habitual: Está públicamente documentado y ampliamente adoptado.

Es importante mantener la cadena de custodia de los artefactos recopilados durante la respuesta a incidentes. Utilizar la automatización y hacer que se genere documentación automática de esta recopilación puede ayudar, además de almacenar los artefactos en repositorios de sólo lectura. El análisis sólo debe realizarse en réplicas exactas de los artefactos recopilados para mantener la integridad.

Con estas características en mente, y basándose en las alertas relevantes y la evaluación del impacto y el alcance, se necesitará recopilar los datos que serán relevantes para la investigación y el análisis posteriores.

Los logs del plano de servicio/control pueden recopilarse para su análisis local o, idealmente, consultarse directamente mediante servicios nativos de AWS (cuando proceda). Los datos (incluidos los metadatos) se pueden consultar directamente para obtener información relevante o para adquirir los objetos de origen; por ejemplo, utilizar la CLI de AWS para adquirir los metadatos de objetos y buckets de Amazon S3 y adquirir directamente los objetos de origen. Los recursos deben

recopilarse de manera coherente con el tipo de recurso y el método de análisis previsto.

Por ejemplo, las bases de datos pueden recopilarse creando una copia o instantánea del sistema que ejecuta la base de datos, creando una copia o instantánea de toda la base de datos o consultando y extrayendo determinados datos y registros de la base de datos relevantes para la investigación.

En el caso de las instancias de Amazon EC2, existe un conjunto específico de datos que deben recopilarse y un orden específico de recopilación que debe realizarse para adquirir y conservar la mayor cantidad de datos para el análisis y la investigación.

En concreto, el orden de respuesta para adquirir y conservar la mayor cantidad de datos de una instancia de Amazon EC2 es el siguiente:

1. Adquirir metadatos de instancia: Por ejemplo, ID de instancia, tipo, dirección IP, ID de VPC/subred, región, ID de imagen de máquina de Amazon (AMI), grupos de seguridad adjuntos, hora de lanzamiento, etc.
2. Habilitar protecciones y etiquetas de instancia como la protección de terminación, configurando el comportamiento de apagado para que se detenga (si está configurado para terminar), deshabilitando los atributos Delete on Termination para los volúmenes EBS adjuntos y aplicando las etiquetas adecuadas tanto para la denotación visual como para su uso en posibles automatizaciones de respuesta (por ejemplo, al aplicar una etiqueta con el nombre Status y el valor Quarantine, realizar la adquisición forense de datos y aísla la instancia).
3. Adquirir disco (instantáneas EBS): Cada instantánea contiene la información que se necesita para restaurar sus datos (desde el momento en que se tomó la instantánea) en un nuevo volumen EBS.
4. Adquirir memoria: Dado que las snapshots de EBS solo capturan los datos que se han escrito en su volumen de Amazon EBS, lo que podría excluir los datos almacenados o almacenados en caché en la memoria por sus aplicaciones o sistema operativo, es imprescindible adquirir una imagen de la memoria del sistema utilizando una herramienta comercial o de código abierto de terceros adecuada para adquirir los datos disponibles del sistema.
5. (Opcional) Realizar una recopilación de datos específica (disco/memoria/registros) a través de una respuesta en vivo en el sistema sólo si el disco o la memoria no se pueden adquirir de otra manera, o si existe una razón comercial u operativa válida. Hacer esto modificará valiosos datos y artefactos del sistema.
6. Retirar la instancia del servicio: Retirar la instancia de los grupos de autoescalado, cancelar el registro de la instancia en los balanceadores de carga y ajustar o aplicar un perfil de instancia predefinido con permisos minimizados o sin permisos.
7. Aislar o contener la instancia: Compruebar que la instancia está aislada de

forma efectiva de otros sistemas y recursos del entorno finalizando e impidiendo las conexiones actuales y futuras a y desde la instancia.

8. En función de la situación y los objetivos, seleccionar una de las siguientes opciones:
 - a. Desmantelar y apagar el sistema (recomendado). Apagar el sistema una vez obtenidas las pruebas disponibles para verificar la mitigación más eficaz contra un posible impacto futuro de la instancia en el entorno.
 - b. Continuar ejecutando la instancia dentro de un entorno aislado instrumentado para su monitorización. Aunque no se recomienda como enfoque estándar, si una situación amerita la observación continua de la instancia (como cuando se necesitan datos o indicadores adicionales para realizar una investigación y un análisis exhaustivos de la instancia), se puede considerar apagar la instancia, crear una AMI de la instancia y volver a lanzar la instancia en su cuenta forense dedicada dentro de un entorno sandbox que esté preparado para estar completamente aislado y configurado con instrumentación para facilitar la supervisión continua de la instancia (por ejemplo, registros de flujo de VPC o VPC Traffic Mirroring).

Bitácora del incidente

Se deben recoger documentar formalmente todas las acciones realizadas, los análisis efectuados y la información identificada durante el análisis para que pueda ser utilizada en las fases posteriores y, en última instancia, en un informe final. Las anotaciones deben ser sucintas y precisas, y confirmar que se incluye la información pertinente para verificar la comprensión efectiva del incidente y mantener una cronología exacta. También son útiles cuando se involucra a personas ajenas al equipo principal de respuesta al incidente.

2.3. Contención, mitigación y recuperación

La fase de contención, mitigación y recuperación tiene como primer objetivo - atendiendo a la peligrosidad del incidente- mitigar su impacto. Es decir, reducir la afectación a los sistemas de información como consecuencia del incidente e impedir su propagación. En segundo lugar, el objetivo es eliminar por completo los restos o la persistencia del incidente en los sistemas de información y, por último, conseguir que el sistema vuelva a su modo de funcionamiento normal.

2.3.1. Contención

La contención, en lo que a respuesta a incidentes respecta, se refiere al proceso o a la aplicación de una estrategia que actúa para minimizar el alcance del incidente y a contener sus efectos.

Una estrategia de contención depende de una miríada de factores y puede ser diferente de una organización a otra en términos de aplicación de tácticas de

contención, tiempo y propósito. Algunos de los criterios para determinar la estrategia de contención adecuada son el potencial daño y alcance a recursos, las necesidades de preservación de prueba, los requisitos de disponibilidad del servicio, el tiempo y los recursos disponibles para aplicar la estrategia, el nivel de eficacia (contención total o parcial) o la duración de la solución temporal.

Pese a ello, en lo que respecta a los servicios en AWS, los pasos fundamentales de contención se pueden reducir a:

1. Contención en la fuente: Utilizar el filtrado y el enrutamiento para impedir el acceso desde una fuente determinada.
2. Contención técnica y de acceso: Eliminar el acceso para impedir el acceso no autorizado a los recursos afectados.
3. Contención en el destino: Utilizar filtrado y enrutamiento para impedir el acceso a un recurso de destino.

Contención en la fuente

La contención en la fuente o en el origen es el uso y la aplicación de filtrado o enrutamiento dentro de un entorno para impedir el acceso a los recursos desde una dirección IP de origen o un rango de red específicos. Aquí se destacan ejemplos de contención de origen mediante servicios de AWS:

- Grupos de seguridad: crear y aplicar grupos de seguridad de aislamiento a instancias de Amazon EC2 o eliminar reglas de un grupo de seguridad existente puede ayudar a contener el tráfico no autorizado a una instancia de Amazon EC2 o un recurso de AWS. Es importante tener en cuenta que las conexiones rastreadas existentes no se cerrarán como resultado de cambiar los grupos de seguridad – sólo el tráfico futuro será bloqueado efectivamente por el nuevo grupo de seguridad.
- Políticas: Las políticas de los buckets de Amazon S3 pueden configurarse para bloquear o permitir el tráfico procedente de una dirección IP, un rango de red o un endpoint de VPC. Las políticas crean la capacidad de bloquear direcciones sospechosas y el acceso al bucket de Amazon S3
- AWS WAF: Las listas de control de acceso web (ACL web) pueden configurarse en AWS WAF para proporcionar un control detallado de las solicitudes web a las que responden los recursos. Puede añadir una dirección IP o un rango de red a un conjunto de IP configurado en AWS WAF y aplicar condiciones de coincidencia, como bloqueo, al conjunto de IP. Esto bloqueará las solicitudes web a un recurso si la dirección IP o los rangos de red del tráfico de origen coinciden con los configurados en las reglas del conjunto de IP.

Un ejemplo de contención de origen se puede ver en el siguiente diagrama con un analista de respuesta a incidentes modificando un grupo de seguridad de una instancia de Amazon EC2 con el fin de restringir nuevas conexiones sólo a ciertas direcciones IP. Como se indica en el punto de los grupos de seguridad, las

conexiones rastreadas existentes no se cerrarán como resultado del cambio de los grupos de seguridad.

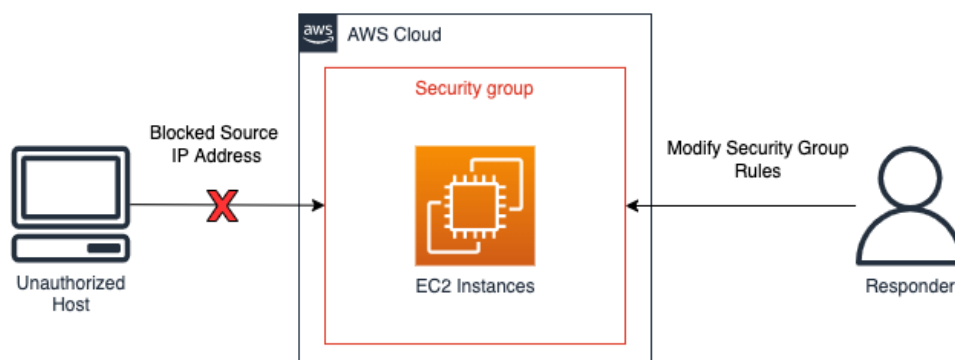


Fig. 2 – Bloqueo de las peticiones procedentes desde una IP a través de la modificación del grupo de seguridad de la instancia

Contención técnica y de acceso

Este tipo de contención se centra en evitar el uso no autorizado de un recurso limitando el acceso de las funciones y las entidades de IAM con acceso al recurso. Esto incluye la restricción de los permisos de las entidades, la revocación temporal de credenciales de seguridad, etc. Aquí se destacan ejemplos de técnica y contención de acceso utilizando los servicios de AWS:

- Restricción de permisos: Los permisos asignados a una entidad de IAM deben seguir el principio de privilegios mínimos. Sin embargo, durante un incidente de seguridad activo, es posible que necesite restringir aún más el acceso a un recurso específico por parte de una entidad IAM. En este caso, es posible contener el acceso a un recurso eliminando los permisos de la entidad con el servicio IAM, bien sea desde la consola de administración, la CLI de AWS o un SDK de AWS.
- Revocar claves: Las claves de acceso de IAM son utilizadas por las entidades de IAM para acceder o administrar recursos. Para contener el acceso de una entidad de IAM de la que una clave de acceso ha sido comprometida, la clave de acceso puede ser desactivada o eliminada. En este sentido, es importante tener en cuenta que:
 - Una clave de acceso puede ser reactivada después de haber sido desactivada,
 - Una clave de acceso no es recuperable una vez que ha sido eliminada,
 - Una entidad de seguridad IAM puede tener hasta dos claves de acceso en un momento dado (si bien, debería evitarse),
 - Los usuarios o aplicaciones que utilicen la clave de acceso perderán el acceso una vez que la clave se haya desactivado o eliminado.

- Revocar credenciales de seguridad temporales: Las credenciales de seguridad temporales pueden ser empleadas por una organización para controlar el acceso a los recursos de AWS. Las credenciales temporales suelen ser utilizadas por roles de IAM y no tienen que ser rotadas o revocadas explícitamente porque tienen un tiempo de vida limitado. En casos donde ocurre un incidente de seguridad que involucra una credencial de seguridad temporal antes de su expiración, se podrían alterar los permisos efectivos de dicha credencial de seguridad utilizando el servicio IAM.

Para los incidentes de seguridad en los que la clave de acceso de IAM de un usuario se ve comprometida por un usuario no autorizado que creó credenciales de seguridad temporales, las credenciales de seguridad temporales se pueden revocar utilizando dos métodos:

- Adjuntando una política en línea al usuario de IAM que impida el acceso en función de la hora de emisión del token de seguridad,
 - Eliminando y volviendo a crear el usuario de IAM con las claves de acceso comprometidas.
- AWS WAF – Ciertas técnicas empleadas por usuarios no autorizados incluyen patrones de tráfico malicioso comunes, como solicitudes que contienen inyección SQL y cross-site scripting (XSS). AWS WAF puede configurarse para que coincida y deniegue el tráfico que emplea estas técnicas utilizando las declaraciones de reglas integradas de AWS WAF.

Un ejemplo de contención técnica y de acceso puede verse en el siguiente diagrama, con un usuario del equipo de respuesta a incidentes rotando claves de acceso o eliminando una política de IAM para evitar que un usuario de IAM acceda a un bucket de Amazon S3.

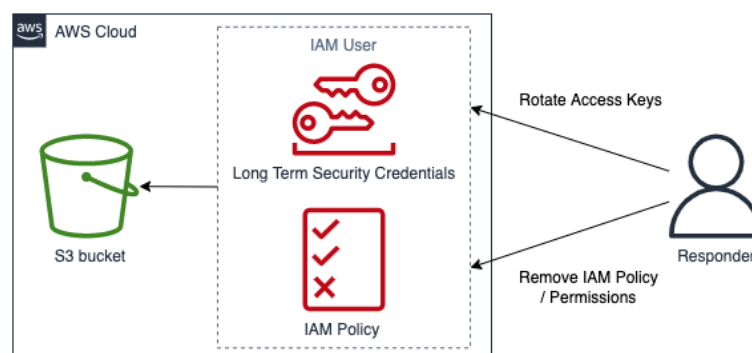


Fig. 3 – Rotación de claves de acceso y eliminación de permisos para evitar que un usuario IAM acceda a un bucket de S3

Contención en el destino

La contención en el destino es la aplicación de filtrado o enrutamiento dentro de un entorno para impedir el acceso a un host o recurso objetivo. Entre los ejemplos

de contención de destino que utilizan los servicios de AWS se incluyen:

- **ACL de red:** Las ACL de red (NACL) que se configuran en subredes que contienen recursos de AWS pueden tener reglas de denegación añadidas. Estas reglas de denegación se pueden aplicar para impedir el acceso a un recurso de AWS en particular; sin embargo, la aplicación de una NACL afectará a todos los recursos de la subred, no sólo a los recursos a los que se está accediendo sin autorización. Las reglas listadas dentro de una NACL se procesan en orden descendente, por lo que la primera regla en una NACL existente debería configurarse para denegar el tráfico no autorizado al recurso y subred objetivo. Como alternativa, se puede crear una NACL completamente nueva con una única regla de denegación para el tráfico entrante y saliente y asociarla a la subred que contiene el recurso objetivo para impedir el acceso a la subred utilizando la nueva NACL.
- **Apagado:** Apagar completamente un recurso puede ser efectivo para contener los efectos del uso no autorizado. El cierre de un recurso también impedirá el acceso legítimo para las necesidades del negocio y evitará que se obtengan datos forenses volátiles, por lo que esta debe ser una decisión intencionada y debe ser validada contra las políticas de seguridad de una organización.
- **VPC de aislamiento:** las VPC de aislamiento se pueden utilizar para proporcionar una contención eficaz de los recursos al tiempo que se proporciona acceso al tráfico legítimo como soluciones antivirus o EDR que requieren acceso a Internet o a una consola de gestión externa. Las VPC de aislamiento pueden preconfigurarse antes de un evento de seguridad para permitir direcciones IP y puertos válidos, y los recursos objetivo pueden trasladarse inmediatamente a esta VPC de aislamiento durante un evento de seguridad activo para contener el recurso al tiempo que se permite que el tráfico legítimo sea enviado y recibido por el recurso objetivo durante las fases posteriores de respuesta al incidente. Un aspecto importante del uso de una VPC de aislamiento es que los recursos, como las instancias EC2, deben apagarse y volver a iniciarse en la nueva VPC de aislamiento antes de su uso. Para efectuar esta operación, [puede seguir los pasos descritos en este documento](#).
- **Grupos de autoescalado y balanceadores de carga:** Los recursos de AWS adjuntos a grupos de autoescalado y balanceadores de carga deben separarse y darse de baja como parte de los procedimientos de contención en el destino. La desvinculación y anulación del registro de los recursos de AWS se puede realizar mediante la consola de administración de AWS, la CLI de AWS y el SDK de AWS.

En el siguiente diagrama se muestra un ejemplo de contención en el destino con un miembro del equipo de respuesta a incidentes que añade una NACL a una subred para bloquear una solicitud de conexión de red procedente de un host no autorizado.

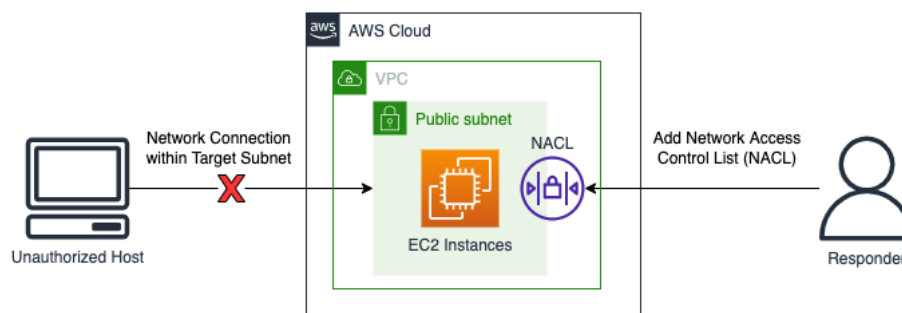


Fig. 4 – Adición de una lista de control de acceso a la subred para el bloqueo de una solicitud procedente de un host no autorizado.

2.3.2. Mitigación

La mitigación, en relación con la respuesta a incidentes de seguridad, es la eliminación de recursos sospechosos o no autorizados en un esfuerzo por devolver la cuenta a un estado seguro conocido. La estrategia de mitigación depende de múltiples factores, en función de los requisitos de negocio de la organización. Algunos pasos de alto nivel para la mitigación son:

1. Identificar y solucionar todas las vulnerabilidades que fueron explotadas.
2. Eliminar el malware, elementos inapropiados y otros componentes.
3. Si se descubren más hosts afectados (por ejemplo, nuevas infecciones de malware), repetir los pasos de detección y análisis para identificar todos los demás hosts afectados y, a continuación, contener y mitigar el incidente para éstos.

Para los recursos de AWS, estos pasos se pueden refinar aún más a través de los eventos detectados y analizados con los registros disponibles o herramientas automatizadas como CloudWatch Logs y Amazon GuardDuty. Esos eventos deben ser la base para determinar qué remediaciones deben efectuarse para restaurar adecuadamente el entorno a un estado seguro conocido.

El primer paso de la mitigación es determinar qué recursos se han visto afectados dentro de la cuenta de AWS. Esto se consigue mediante el análisis de sus fuentes de datos de registro, recursos y herramientas automatizadas disponibles. A través de estos datos, se deberá:

- Identificar las acciones no autorizadas realizadas por las identidades de IAM en la cuenta.
- Identificar accesos o cambios no autorizados en su cuenta.
- Identificar la creación de recursos o usuarios de IAM no autorizados.

- Identificar sistemas o recursos con cambios no autorizados.

Una vez identificada la lista de recursos, se debe evaluar cada uno de ellos para determinar el impacto en la empresa si se elimina o restaura el recurso. Por ejemplo, si un servidor web aloja su aplicación empresarial y su eliminación causaría tiempo de inactividad, entonces se debería considerar la recuperación del recurso a partir de copias de seguridad seguras verificadas o relanzar el sistema desde una AMI limpia antes de eliminar el servidor afectado.

Una vez que haya concluido el análisis de impacto en el negocio, entonces, utilizando los eventos del análisis de registro, se debe ir a las cuentas y realizar las remediaciones apropiadas, tales como:

- Rotación o eliminación de claves (este paso elimina la capacidad del actor de continuar realizando actividades dentro de la cuenta).
- Rotación de las credenciales de usuario IAM potencialmente no autorizadas.
- Eliminar recursos no reconocidos o no autorizados. En este sentido, es importante tener en cuenta la posibilidad de realizar copias de seguridad de dichos recursos por motivos normativos, de conformidad o legales. Por ejemplo, podría conservarse una instancia de Amazon EC2 creando una instantánea de Amazon EBS antes de eliminarla.

En el caso de las infecciones de malware, es posible que sea necesario ponerse en contacto con un socio de AWS u otro proveedor. Si se utiliza el módulo GuardDuty Malware para Amazon EBS, es posible que disponga de recomendaciones para los hallazgos proporcionados.

Una vez que se hayan erradicado los recursos afectados identificados, se recomienda realizar una revisión de seguridad de toda la cuenta. Esto se puede hacer utilizando las reglas de AWS Config, utilizando soluciones de código abierto como Prowler y ScoutSuite, o a través de otros proveedores. También se debería considerar la posibilidad de realizar escaneos de vulnerabilidades en sus recursos de cara al público (Internet) para evaluar el riesgo residual.

La mitigación es un paso del proceso de respuesta a incidentes y puede ser manual o automatizada, dependiendo del incidente y de los recursos afectados. La estrategia global debe alinearse con las políticas de seguridad y las necesidades empresariales de la organización, y verificar que se mitigan los efectos negativos a medida que se eliminan los recursos o configuraciones inadecuados.

2.3.3. Recuperación

La recuperación es el proceso de restaurar los sistemas a un estado seguro conocido, validar que las copias de seguridad son seguras o no se han visto afectadas por el incidente antes de la restauración, realizar pruebas para verificar que los sistemas funcionan correctamente después de la restauración y abordar las vulnerabilidades asociadas al incidente de seguridad.

El orden de la recuperación depende de los requisitos de la organización. Si bien, como parte del proceso de recuperación, debe realizar un análisis del impacto en el negocio para determinar, como mínimo:

- Las prioridades empresariales o de dependencia,
- El plan de restauración,
- Autenticación y autorización.

Algunos de los pasos de alto nivel para la recuperación de sistemas son:

- Restaurar las copias de seguridad limpias.
 - a. Se debe comprobar que las copias de seguridad se evalúan antes de ser restauradas para asegurarse de que la infección no está presente y evitar el resurgimiento del incidente de seguridad. Por ello, las copias de seguridad deben evaluarse periódicamente como parte de las pruebas de recuperación de desastres para verificar que el mecanismo de copia de seguridad funciona correctamente y que la integridad de los datos cumple los objetivos del punto de recuperación.
 - b. Si es posible, se deben utilizar copias de seguridad anteriores a la primera marca de tiempo del evento identificada como parte del análisis de la causa raíz.
- Reconstruir los sistemas desde cero, incluida la redistribución desde una fuente de confianza utilizando la automatización, incluso en una nueva cuenta de AWS no afectada por el incidente.
- Sustitución de los archivos comprometidos por versiones limpias. Se debe estar absolutamente seguro de que los archivos que se están recuperando son seguros y no están comprometidos por el incidente.
- Instalación de parches.
- Cambio de contraseñas. Esto incluye:
 - Las contraseñas de las entidades IAM que podrían haber sido objeto de compromiso. Se recomienda utilizar roles y federación como parte de la estrategia basada en el principio de mínimo privilegio.
 - Reforzar la seguridad del perímetro de la red (conjuntos de reglas de cortafuegos, listas de control de acceso al enrutador de frontera).

Una vez recuperados los recursos, es importante recoger las lecciones aprendidas para actualizar las políticas, procedimientos y guías de respuesta a incidentes.

En resumen, es imperativo implantar un proceso de recuperación que facilite la vuelta a las operaciones seguras conocidas. La recuperación puede llevar mucho tiempo y requiere una estrecha vinculación con las estrategias de contención para equilibrar el impacto empresarial con el riesgo de reinfección. Los procedimientos de recuperación deben incluir pasos para restaurar los recursos y servicios, las entidades de IAM y realizar una revisión de seguridad de la cuenta para evaluar el

riesgo residual

2.4. Actividad post-ciberincidente

En esta fase se debe documentar la causa originaria y el coste del ciberincidente, especialmente en términos de compromiso de información o de impacto en los servicios prestados, así como las medidas que la organización ha tomado durante la resolución del mismo y las iniciativas que se deberán llevar a cabo para prevenir futuros ciberincidentes de naturaleza similar.

El panorama de las ciberamenazas cambia constantemente y es importante ser igual de dinámico en la capacidad de la organización para proteger eficazmente sus entornos. La clave para la mejora continua es iterar las actuaciones realizadas durante los incidentes y las simulaciones con el fin de mejorar las capacidades para detectar, responder e investigar eficazmente posibles incidentes de seguridad, reduciendo las posibles vulnerabilidades, el tiempo de respuesta y el retorno a operaciones seguras. Los siguientes mecanismos pueden ayudar a verificar que la organización permanece preparada con las últimas capacidades y conocimientos para responder eficazmente, sea cual sea la situación.

Establecer un marco de trabajo para aprender de los incidentes

Implantar un marco de trabajo y una metodología de aprendizaje de las lecciones aprendidas, no sólo ayudará a mejorar las capacidades de respuesta ante incidentes, sino que también ayudará a evitar que el incidente se repita. Aprender de cada incidente puede ayudar a evitar que se repitan los mismos errores, exposiciones o configuraciones erróneas, no sólo mejorando la postura de seguridad, sino también minimizando el tiempo perdido en situaciones evitables.

Establecer métricas para el éxito

Las métricas son necesarias para medir, evaluar y mejorar eficazmente las capacidades de respuesta ante incidentes. Sin métricas, no hay referencia con la que medir con precisión o incluso identificar lo bien que su organización está (o no está) funcionando. Existen algunas métricas comunes a la respuesta ante incidentes que son un buen punto de partida para una organización que busca establecer expectativas y referencias para trabajar hacia la excelencia operativa. Por ejemplo:

- **Tiempo medio de detección:** Es el tiempo medio que se tarda en descubrir un posible incidente de seguridad. Concretamente, es el tiempo transcurrido entre la aparición del primer indicador de peligro y la identificación o alerta inicial.
- **Tiempo medio de reconocimiento:** Es el tiempo medio que se tarda en reconocer y priorizar un posible incidente de seguridad. Concretamente, es el tiempo que transcurre entre la generación de una alerta y el momento en que un miembro del SOC o el equipo de respuesta a incidentes identifica y prioriza la alerta para su tratamiento.

- Tiempo medio de respuesta. Es el tiempo medio que se tarda en iniciar la respuesta inicial a un posible incidente de seguridad. Concretamente, es el tiempo transcurrido entre la alerta inicial o el descubrimiento de un posible incidente de seguridad y las primeras acciones emprendidas para responder.
- Tiempo medio de contención. Es el tiempo medio que se tarda en contener un posible incidente de seguridad. Concretamente, es el tiempo transcurrido entre la alerta inicial o el descubrimiento de un posible incidente de seguridad y la finalización de las acciones de respuesta que impiden eficazmente que el atacante cause más daños.
- Tiempo medio de recuperación. Es el tiempo medio que se tarda en volver a operar con total seguridad tras un posible incidente de seguridad. Concretamente, es el tiempo transcurrido entre la alerta inicial o el descubrimiento de un posible incidente de seguridad y el momento en que la organización vuelve a funcionar con normalidad y seguridad sin verse afectada por el incidente.
- Tiempo de permanencia del atacante. Es el tiempo medio que un usuario no autorizado tiene acceso a un sistema o entorno. Es similar al tiempo medio de contención, salvo que el marco temporal comienza con el momento inicial en que el atacante obtuvo acceso al sistema o entorno, que puede ser anterior a la alerta o descubrimiento inicial.

Usar indicadores de compromiso

Un indicador de compromiso (IoC) es un artefacto observado en una red, sistema o entorno que puede identificar una actividad maliciosa o un incidente de seguridad. Los IoCs pueden existir en una variedad de formas, incluyendo direcciones IP, dominios, artefactos a nivel de red como *flags* TCP o cargas útiles; artefactos a nivel de sistema o host como ejecutables, nombres de archivos y hashes, entradas de archivos de registro y más. También pueden ser una combinación de elementos o actividades, como la existencia de elementos o artefactos específicos en un sistema, acciones realizadas en cierto orden o actividad de red que pueden indicar una amenaza, ataque o metodología de ataque específicos.

A medida se que trabaje para mejorar iterativamente el programa de respuesta a incidentes, se debe implementar un marco para recopilar, gestionar y utilizar los IoC como mecanismo para construir y mejorar continuamente las detecciones y alertas y mejorar la velocidad y eficacia de las investigaciones. Se puede empezar por incorporar la recopilación y gestión de IoCs en las fases de análisis e investigación de los procesos de respuesta a incidentes. Al identificar, recopilar y almacenar los IoCs de forma proactiva como parte estándar del proceso, se puede crear un repositorio de datos (como parte de un programa de inteligencia sobre amenazas más completo) que, a su vez, se puede utilizar para mejorar las detecciones y alertas existentes, crear detecciones y alertas adicionales, identificar dónde y cuándo se vio un artefacto anteriormente, crear y consultar

documentación sobre cómo se realizaron anteriormente las investigaciones que incluían IoCs coincidentes, etc.

2.4.1. Continuar la formación y la concienciación

La formación y la concienciación son esfuerzos en evolución y continuos que deben perseguirse y mantenerse con determinación. Existen diversos mecanismos para verificar que el equipo de la organización mantiene una concienciación, unos conocimientos y unas capacidades acordes con la evolución de la tecnología y el panorama de las ciberamenazas.

Un mecanismo es emplear la formación continua como parte estándar de los objetivos y operaciones de los equipos. El personal de respuesta a incidentes y las partes interesadas deben estar capacitados eficazmente en la detección, respuesta e investigación de incidentes en AWS. La formación debe ser continua para verificar que el equipo mantiene el conocimiento de los últimos avances tecnológicos, actualizaciones y mejoras que se pueden aprovechar para mejorar la eficacia y la eficiencia de la respuesta, así como adiciones o actualizaciones que se pueden aprovechar para mejorar la investigación y el análisis.

Otro mecanismo consiste en verificar que los simulacros se realizan con regularidad (por ejemplo, trimestralmente) y se centran en resultados específicos para la organización. Aunque la realización de ejercicios de simulación es una forma excelente de generar una línea de base inicial, las pruebas continuas son clave para las mejoras sostenidas y para mantener un reflejo actualizado y preciso del estado actual de las operaciones de respuesta a incidentes. La realización de pruebas con las situaciones de seguridad más recientes y críticas y las capacidades de respuesta más importantes o novedosas, y la incorporación de las lecciones aprendidas a la formación, las operaciones y los procesos/procedimientos verificarán que se es capaz de mejorar continuamente los procesos de respuesta y el programa de seguridad en su conjunto.

3. TECNOLOGÍAS Y SERVICIOS DE REFERENCIA PARA LA RESPUESTA A INCIDENTES EN AWS

En los siguientes apartados se aportan algunas recomendaciones y se definen los principales servicios y tecnologías de AWS que influyen en las diferentes fases de la gestión de los ciberincidentes expuestas en el apartado anterior, si bien, dependiendo de su uso, algunos servicios se pueden emplear en varias de las fases del ciberincidente.

3.1. Preparación

Más allá de las tecnologías y servicios de AWS que aquí se exponen para atender la fase de preparación a los incidentes, es importante tener en cuenta que esta fase está principalmente compuesta por las actuaciones señaladas en el apartado 2.1 de este documento, es decir, la formación de las personas que deben responder a los incidentes, la elaboración de procedimientos de respuesta y la

preparación de la tecnología.

IAM

IAM (Identity and Access Management) es el principal servicio para la gestión de identidades y accesos en AWS. Con IAM se puede especificar quién o qué puede acceder a los servicios y recursos en AWS, administrar de forma centralizada los permisos específicos y analizar el acceso para perfeccionar los permisos en todo AWS. Es, por tanto, uno de los servicios más importantes de cara a la preparación del entorno de AWS para la prevención de incidentes utilizando una estrategia basada en el principio de mínimo privilegio.

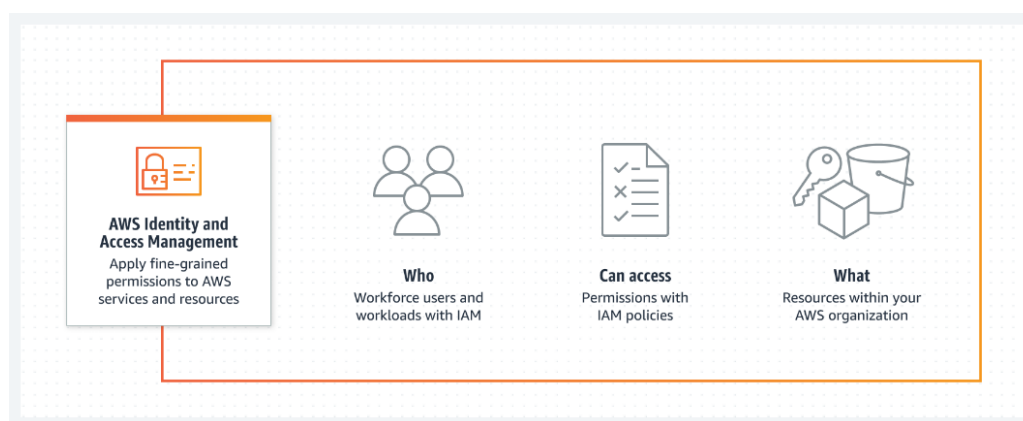


Fig. 5 – Diagrama de síntesis del funcionamiento de IAM

Las configuraciones y las prácticas recomendadas para IAM se desarrollan con detalle en el apartado Control de Acceso [op.acc] de la guía **CCN-STIC 887A – Guía de Configuración Segura para AWS**, y en el documento de [prácticas recomendadas](#) del fabricante.

AWS Security Hub

AWS Security Hub proporciona una visión completa del estado de seguridad en AWS y ayuda a verificar el entorno con los estándares y las prácticas de seguridad recomendadas. AWS Security Hub recopila datos de todas las cuentas de AWS, de los servicios y de los productos de terceros compatibles, ayudando a la organización a analizar sus tendencias de seguridad y a identificar los problemas de seguridad de mayor prioridad, todo ello desde un único lugar. Los resultados se resumen visualmente en paneles integrados con gráficos y tablas que permiten su procesamiento.

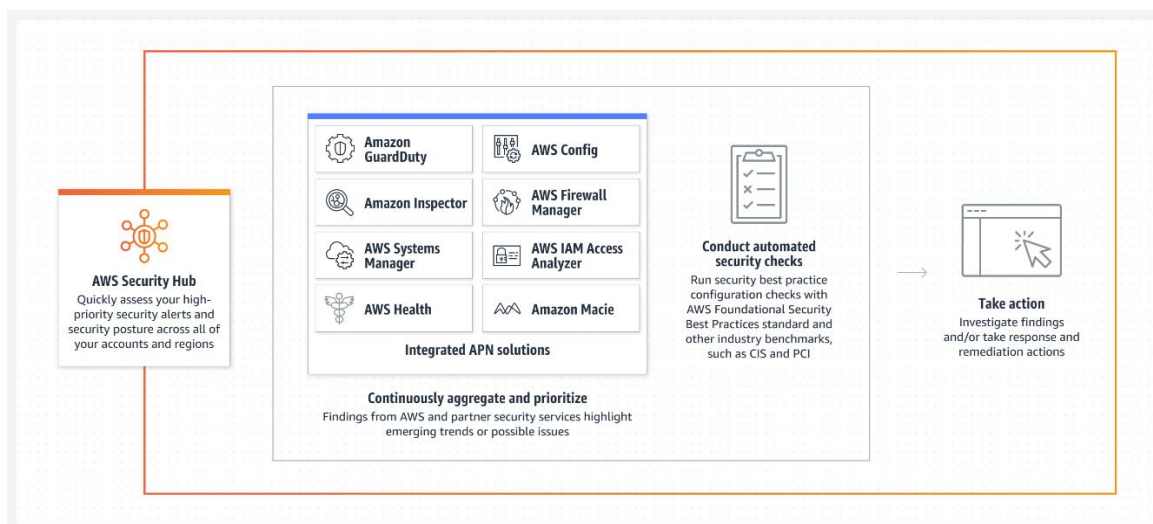


Fig. 6 – Diagrama de funcionamiento de AWS Security Hub

Para más información sobre el servicio Security Hub, se puede consultar la documentación del fabricante.

AWS Config Rules

Las AWS Config Rules representan los ajustes de configuración ideal de los recursos y permiten evaluar su adecuación. AWS Config proporciona reglas predefinidas y personalizables denominadas reglas administradas para comenzar de una manera adecuada.

Aunque AWS Config realiza un seguimiento continuo de los cambios de configuración que se producen entre los recursos, también comprueba si estos cambios no cumplen con las condiciones de las reglas. Si un recurso no cumple con la regla, AWS Config marca el recurso y la regla como *no conforme*.

Por ejemplo, cuando se crea un volumen de EC2, AWS Config puede evaluar el volumen con respecto a una regla que requiere que los volúmenes se cifren. Si el volumen no está cifrado, AWS Config marca el volumen y la regla como no conforme. AWS Config también puede comprobar todos los recursos para determinar si cumplen los requisitos de toda la cuenta. Por ejemplo, AWS Config puede comprobar si el número de volúmenes de EC2 en una cuenta permanece dentro de un total deseado, o si una cuenta utiliza AWS CloudTrail para iniciar sesión.

También se pueden crear reglas personalizadas para evaluar recursos adicionales que AWS Config aún no registra.



Fig. 7 – Diagrama de funcionamiento de AWS Config

Para más información sobre AWS Config Rules, se puede consultar la [documentación del fabricante](#).

AWS Fault Injector Simulator

AWS Fault Injector Simulator es un servicio gestionado que permite realizar experimentos de inyección de errores en las cargas de trabajo de AWS. Estos experimentos estresan una aplicación al crear eventos disruptivos para que se pueda observar cómo responde la aplicación y utilizar la información obtenida para mejorar el rendimiento y la resiliencia. Con AWS Fault Injector Simulator se puede configurar y ejecutar experimentos que ayudan a crear las condiciones reales necesarias para descubrir problemas en las aplicaciones, que de otro modo, podrían ser difíciles de encontrar.

Para ejecutar experimentos, primero se debe crear una plantilla de experimento, que es el modelo del experimento. Contiene las acciones, los objetivos y las condiciones de finalización del experimento.

NOTA: Es importante señalar que AWS Fault Injector Simulator lleva a cabo acciones reales sobre los recursos de AWS, por lo que antes de empezar a ejecutar experimentos en producción se debería completar una fase de planificación y ejecución de experimentos en un entorno de preproducción. Además, se recomienda limitar estos ejercicios a equipos de seguridad con experiencia y conocimientos profundos de AWS.

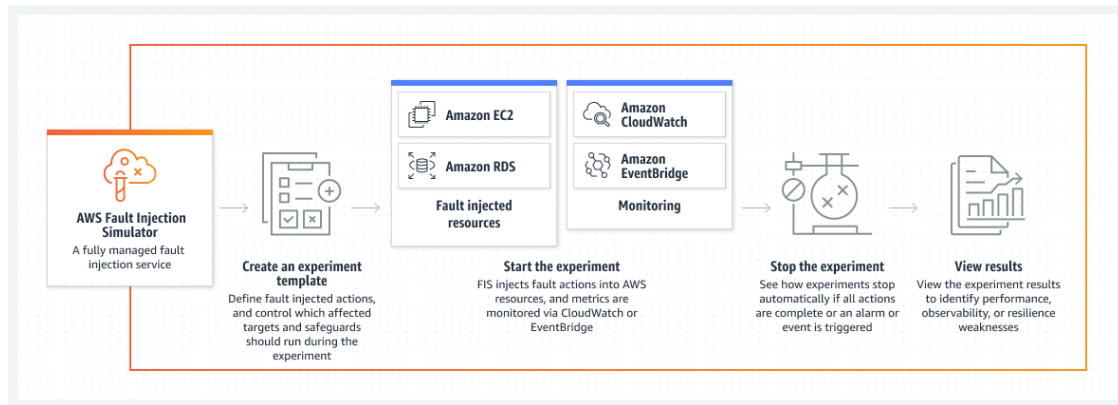


Fig. 8 – Diagrama de funcionamiento de AWS Fault Injector Simulator

Para más información sobre AWS Fault Injection Simulator, se puede consultar la [documentación del fabricante](#).

3.2. Detección, análisis e investigación

Amazon GuardDuty

Amazon GuardDuty es un servicio de detección de amenazas inteligente que supervisa de manera continua las cargas de trabajo y cuentas de AWS para detectar actividades maliciosas y envía hallazgos de seguridad para su visibilidad y resolución.

GuardDuty se encarga de analizar los eventos en las cuentas de AWS a partir de eventos de administración de AWS CloudTrail, eventos de datos de S3 de AWS CloudTrail, registros de flujo de Amazon VPC (datos del tráfico de red) y registros de DNS.

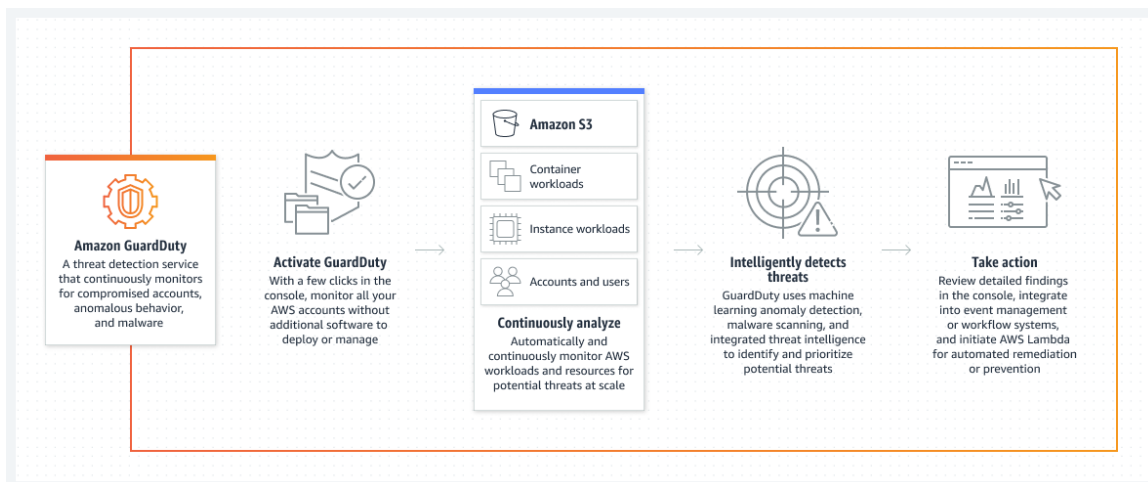


Fig. 9 – Funcionamiento de Amazon GuardDuty

Para más información sobre GuardDuty y su activación se puede consultar la [guía del fabricante](#) y la guía **CCN-STIC 887G Guía de configuración segura para monitorización y gestión en AWS**.

Amazon CloudWatch

Amazon CloudWatch es un servicio de monitorización y administración que suministra datos e información procesable para aplicaciones y recursos de infraestructura locales, híbridos y de AWS. Con CloudWatch, se puede recopilar y hacer un seguimiento de métricas en una sola plataforma, lo que permite monitorear aplicaciones y sistemas individuales aislados (servidor, red, base de datos, etc.).

La página de inicio de CloudWatch muestra automáticamente las métricas disponibles por defecto sobre todos los servicios de AWS que se utilicen. También se pueden crear paneles personalizados para mostrar colecciones de métricas que se elijan.

Asimismo, se pueden crear alarmas que vigilen métricas y enviar notificaciones o realizar cambios automáticamente en los recursos que se están monitorizando cuando se supera determinado umbral.

En síntesis, las funciones de CloudWatch son la recopilación de datos, la monitorización, las alarmas y la detección de anomalías.

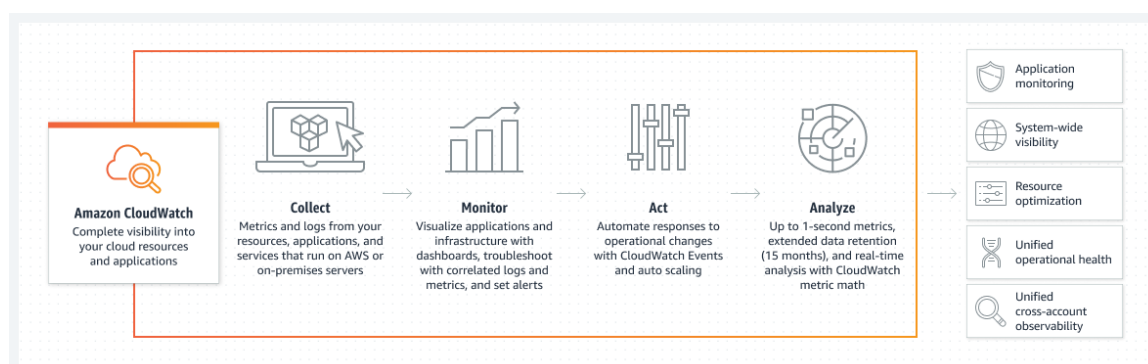


Fig. 10 – Diagrama de funcionamiento de CloudWatch

Para más información sobre CloudWatch se puede consultar la [guía del fabricante](#) y la guía **CCN-STIC 887G Guía de configuración segura para monitorización y gestión en AWS**.

AWS CloudTrail

AWS CloudTrail es un servicio de registro que facilita la gobernanza, conformidad y las auditorías de operaciones y riesgos en la cuenta de AWS. Las acciones que realiza un usuario, rol o servicio de AWS se registran como eventos en CloudTrail. Los eventos registrados incluyen las acciones llevadas a cabo en la consola de administración, los SDK, las herramientas de línea de comandos y otros servicios de AWS.

El historial de eventos simplifica el análisis de seguridad, el seguimiento de

cambios de recursos y la resolución de problemas. Además, se puede usar CloudTrail para detectar actividad inusual en las cuentas de AWS. Estas funciones ayudan a simplificar los análisis operativos y la solución de problemas.

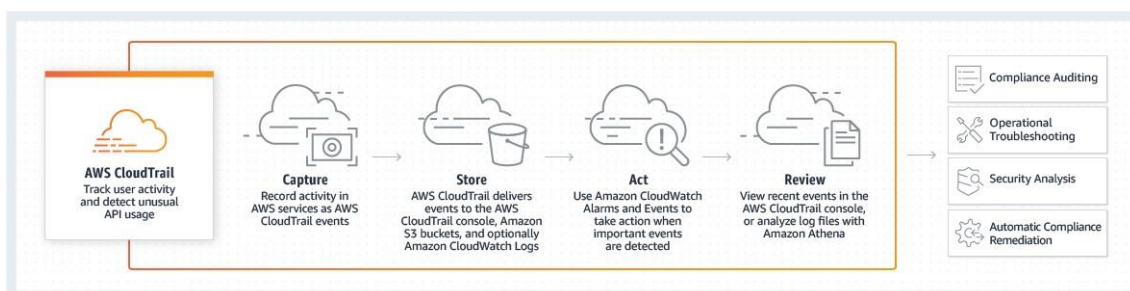


Fig. 11 – Diagrama de funcionamiento de CloudTrail

Por otro lado, AWS CloudTrail Lake permite ejecutar consultas basadas en SQL sobre los eventos, que se agregan en almacenes de datos de eventos. Este tipo de consultas permite ofrecer una visión más detallada y personalizable de los eventos que las simples búsquedas de claves y valores en el historial de eventos, dado que se pueden ejecutar complejas consultas en SQL con múltiples campos de búsqueda y en todas las regiones simultáneamente.

Para más información sobre CloudTrail se puede consultar la [guía del fabricante](#) y la guía **CCN-STIC 887G Guía de configuración segura para monitorización y gestión en AWS**.

Amazon Security Lake

Amazon Security Lake es un servicio de *data lake* de datos de seguridad totalmente gestionado. Se puede utilizar para centralizar automáticamente los datos de seguridad de los entornos de AWS, proveedores de SaaS, sistemas on-premise, etc. Estos *data lakes* están respaldados por Amazon S3 y en todo caso el cliente conserva la propiedad de los datos. Security Lake automatiza la recopilación de datos de eventos y logs relacionados con la seguridad de los servicios integrados de AWS y de terceros. También ayuda a la administración del ciclo de vida de los datos con configuraciones de retención y replicación personalizables.

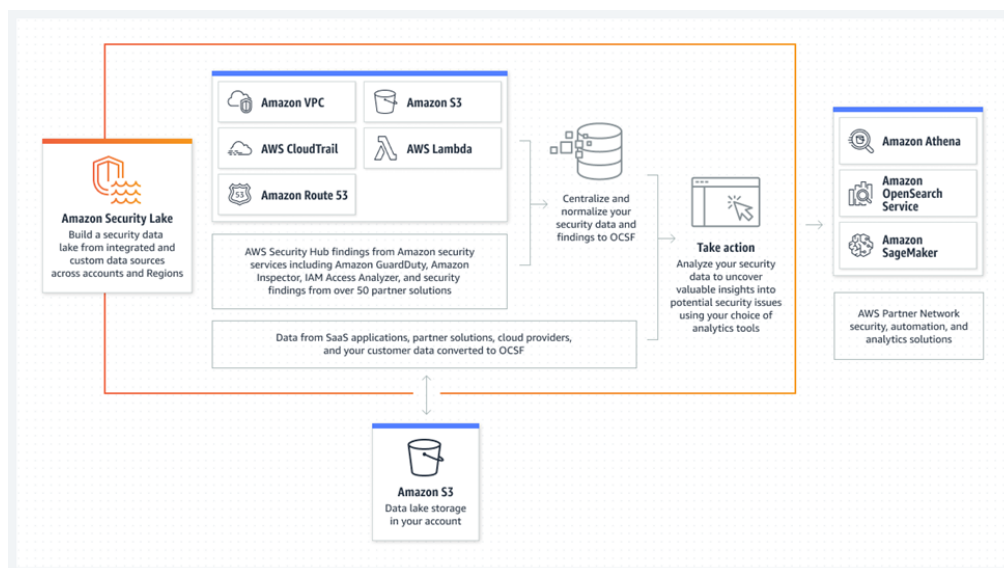


Fig. 12 – Diagrama de funcionamiento de Security Lake

Amazon Detective

Amazon Detective es una herramienta que ayuda al análisis, investigación e identificación rápida de la causa raíz de los hallazgos de seguridad o actividades sospechosas. Amazon Detective recopila automáticamente datos de registro de los recursos de AWS como los intentos de inicio de sesión, las llamadas a la API y el tráfico de red de AWS CloudTrail o los registros de flujo de Amazon VPC para, a continuación, utilizar el aprendizaje automático, el análisis estadístico y la teoría de grafos para generar visualizaciones que ayudan a llevar a cabo investigaciones de seguridad rápidas y eficientes.

Detective permite:

- Investigar rápidamente cualquier actividad que esté fuera de la norma.
- Identificar patrones que pueden indicar un problema de seguridad.
- Comprender todos los recursos afectados por un hallazgo.

Para usar Amazon Detective se debe habilitar primero Amazon GuardDuty.

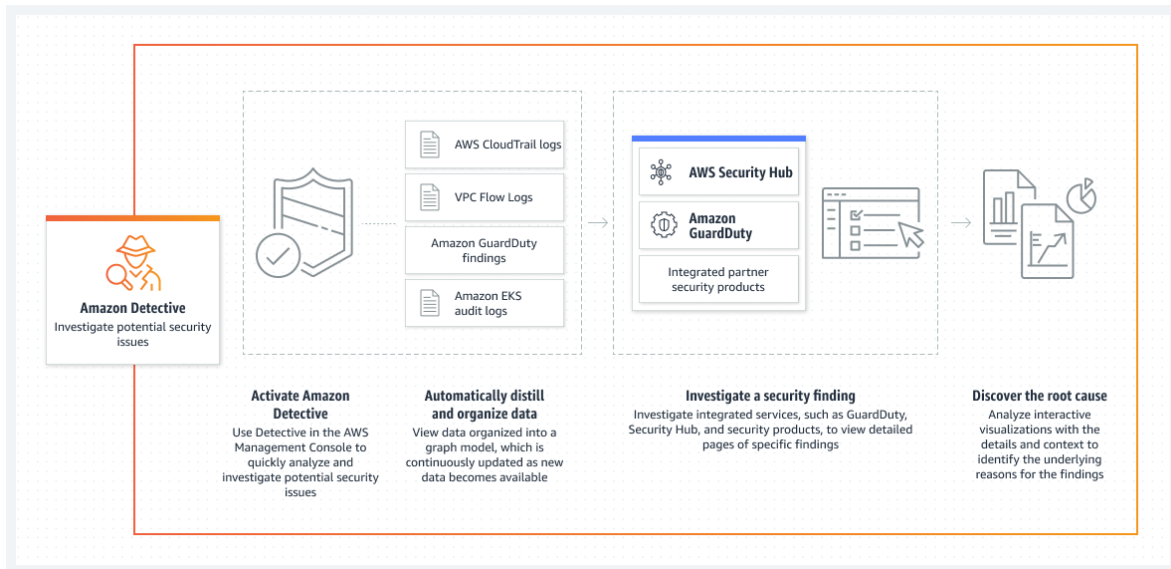


Fig. 13 – Diagrama de funcionamiento de Amazon Detective

Para más información sobre Amazon Detective se puede consultar la [guía del fabricante](#).

AWS Systems Manager

AWS Systems Manager permite centralizar datos operacionales de diferentes servicios de AWS y automatizar tareas en recursos de AWS. Permite la creación de grupos lógicos de recursos, tales como aplicaciones, capas diferentes de una pila de aplicaciones o entornos de producción y de desarrollo. Con Systems Manager, se puede seleccionar un grupo de recursos y ver su actividad reciente, sus cambios en la configuración de recursos, sus notificaciones relacionadas, sus alertas operativas, su inventario de software y su estado de conformidad con parches. Systems Manager proporciona un lugar centralizado para ver y administrar recursos de AWS. Esto permite tener un control y una visibilidad completos de las operaciones.

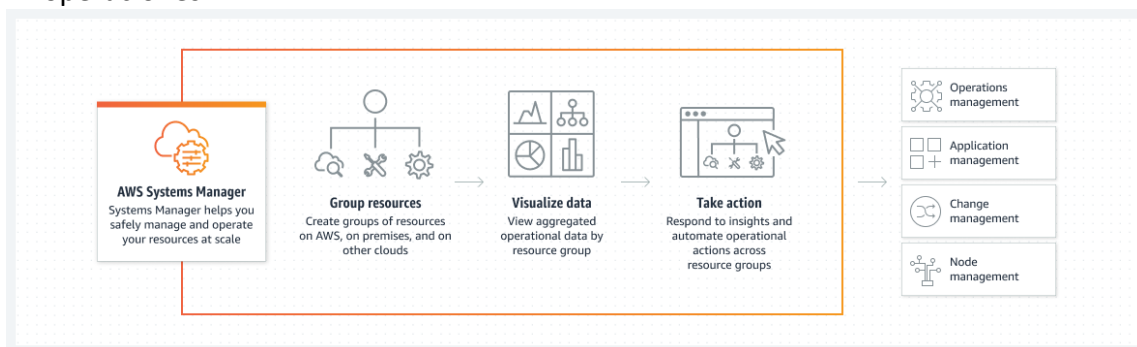


Fig. 14 – Diagrama de funcionamiento de AWS Systems Manager

Se puede consultar más información sobre Systems Manager en la [documentación del fabricante](#).

3.3. Contención, mitigación y recuperación

AWS Lambda

AWS Lambda es un servicio que permite ejecutar código sin aprovisionar ni administrar servidores. Lambda ejecuta el código en una infraestructura de computación de alta disponibilidad y realiza todas las tareas de administración de los recursos de computación, incluido el mantenimiento del servidor y del sistema operativo, el aprovisionamiento de capacidad y el escalado automático.

Con Lambda se puede ejecutar código para prácticamente cualquier tipo de aplicación o servicio de backend. Asimismo, se puede utilizar para automatizar la respuesta a incidentes. Únicamente es necesario suministrar el código en alguno de los [lenguajes que admite Lambda](#).

El código se organiza en funciones de Lambda. Las funciones se ejecutan solo cuando es necesario y se escala de manera automática, desde unas pocas solicitudes por día hasta miles por segundo, pagándose únicamente por el tiempo que se consuma y sin que se apliquen cargos cuando el código no se está ejecutando.

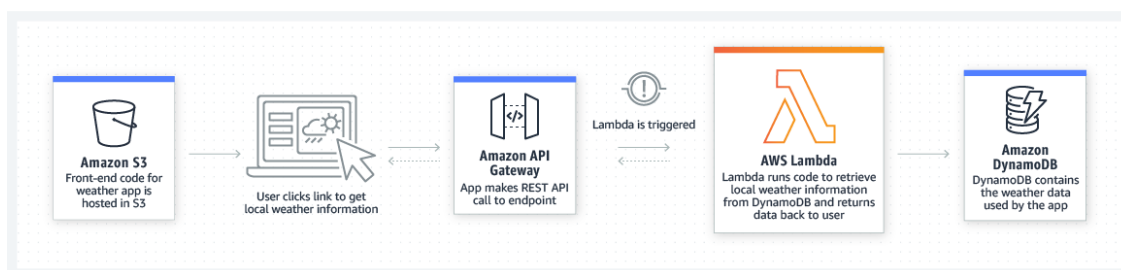


Fig. 15 – Ejemplo de implementación de AWS Lambda en una aplicación web

Se puede consultar más información sobre AWS Lambda en la [documentación del fabricante](#).

AWS Network Firewall

AWS Network Firewall es un servicio administrado que facilita la implementación de protecciones de red básicas para todas sus Amazon Virtual Private Clouds (VPC). El motor de AWS Network Firewall permite definir reglas de firewall que brindan un control minucioso sobre el tráfico de la red. AWS Network Firewall trabaja junto a AWS Firewall Manager de modo que permite crear políticas basadas en reglas de AWS Network Firewall y luego aplicar de manera centralizada estas políticas en sus cuentas y VPC.

El sistema de prevención de intrusiones (IPS) de AWS Network Firewall proporciona una inspección activa del flujo de tráfico para que se pueda identificar y bloquear las vulnerabilidades mediante la detección basada en firmas. AWS Network Firewall también ofrece filtrado web que puede detener el tráfico a URL incorrectas conocidas y monitorear nombres de dominio completamente autorizados.

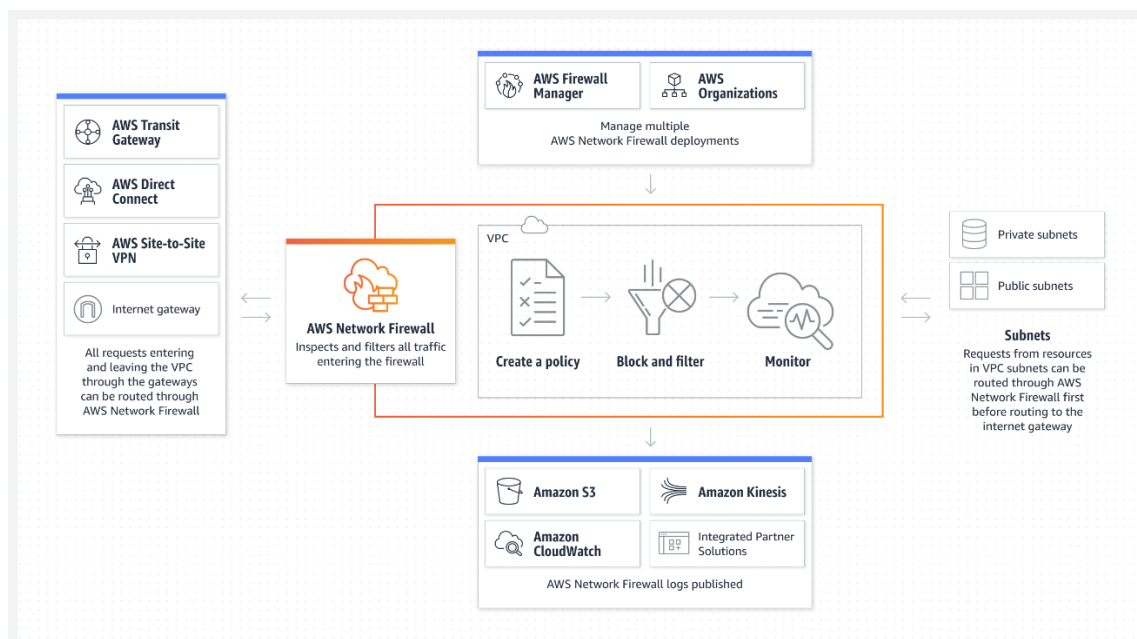


Fig. 16 – Diagrama de funcionamiento e integración con otros servicios de Network Firewall

Se puede encontrar la implementación de esta tecnología en la siguiente [guía del fabricante](#).

Amazon EventBridge

Amazon EventBridge proporciona un flujo de eventos de sistema casi en tiempo real que describen cambios en los recursos de Amazon Web Services (AWS). Mediante reglas sencillas que se puede configurar rápidamente, se puede asignar los eventos y dirigirlos a uno o más flujos o funciones de destino. EventBridge conoce los cambios operativos a medida que se producen. EventBridge responde a estos cambios operativos y toma medidas correctoras según sea necesario, enviando mensajes para responder al entorno, activando funciones, realizando cambios y captando información de estado.

También es posible utilizar EventBridge para programar acciones automatizadas que se disparen automáticamente en determinados momentos a través de expresiones cron o de frecuencia.

Entre otros, se pueden configurar como destinos para EventBridge: Instancias EC2, funciones de AWS Lambda, grupos de registro de Amazon CloudWatch Logs, Automatizaciones de Systems Manager, plantillas de evaluación de Amazon Inspector, Temas de Amazon SNS, colas de Amazon SQS, etc.

Se puede consultar más información sobre Amazon EventBridge en la [documentación del fabricante](#).

Amazon Simple Notification Service

Amazon Simple Notification Service (Amazon SNS) es un servicio administrado con el que se ofrece la entrega de los mensajes de los publicadores a los suscriptores.

Los publicadores se comunican de forma asíncrona con los suscriptores mediante el envío de mensajes a un *tema*, que es un punto de acceso lógico y un canal de comunicación. Los clientes pueden suscribirse al tema de SNS y recibir mensajes publicados mediante un tipo de punto de enlace compatible, como Amazon Kinesis Data Firehose, Amazon SQS, AWS Lambda, HTTP, correo electrónico, notificaciones push móviles y mensajes de texto móviles (SMS).

Amazon SNS permite el envío de mensajes de aplicación a aplicación (A2A), o de aplicación a persona (A2P). Entre otros casos de uso, Amazon SNS se puede utilizar para desencadenar notificaciones a través de correo electrónico o SMS cuando se producen eventos como cambios en grupos de EC2 AutoScaling, archivos nuevos cargados en buckets de S3 o superación de umbrales de métricas de CloudWatch.

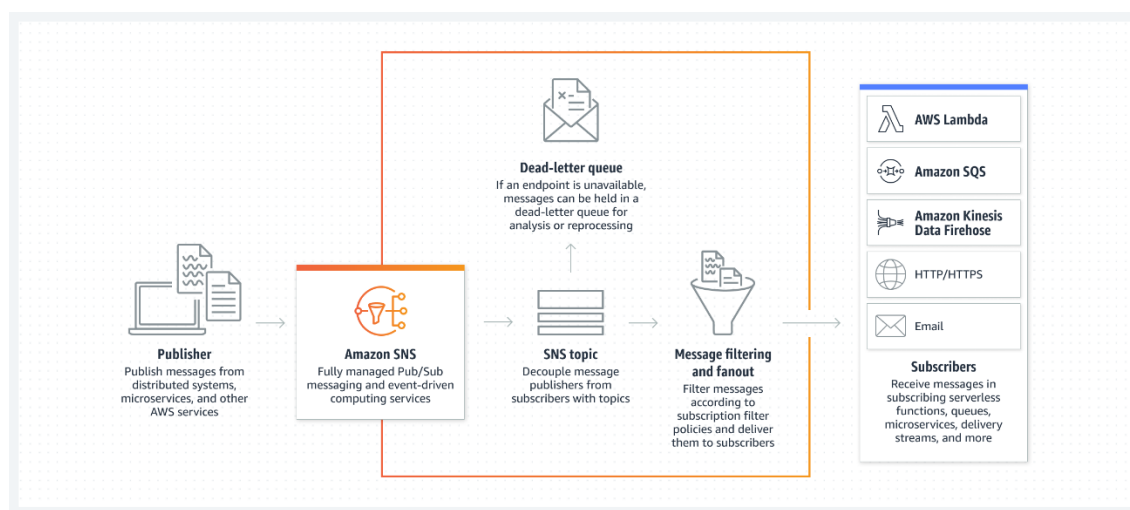


Fig. 17 – Diagrama de funcionamiento de SNS

Se puede consultar más información sobre Amazon Simple Notification Service en la [documentación del fabricante](#).

AWS Step Functions

AWS Step Functions es un servicio de orquestación sin servidor que le permite integrarse con las funciones de AWS Lambda y otros servicios de AWS para crear aplicaciones. A través de la consola gráfica de Step Functions, puede ver el flujo de trabajo de su aplicación como una serie de pasos basados en eventos.

Con los controles integrados de Step Functions, puede examinar el estado de cada paso del flujo de trabajo para asegurarse de que la aplicación se ejecuta en orden y según lo esperado. Según su caso de uso, puede hacer que Step Functions llame servicios de AWS para realizar tareas. Puede crear flujos de trabajo que procesen y publiquen modelos de aprendizaje automático.

Una de las funciones para las que se puede utilizar AWS Step Functions es la automatización de seguridad. Por ejemplo:

- Se puede automatizar la respuesta a un escenario en el que se haya

expuesto un usuario y una clave de acceso.

- Corregir automáticamente las respuestas a los incidentes de seguridad de acuerdo con las acciones definidas.
- Advertir a los empleados de la recepción de correos electrónicos de suplantación de identidad a los pocos segundos de recibirlos.

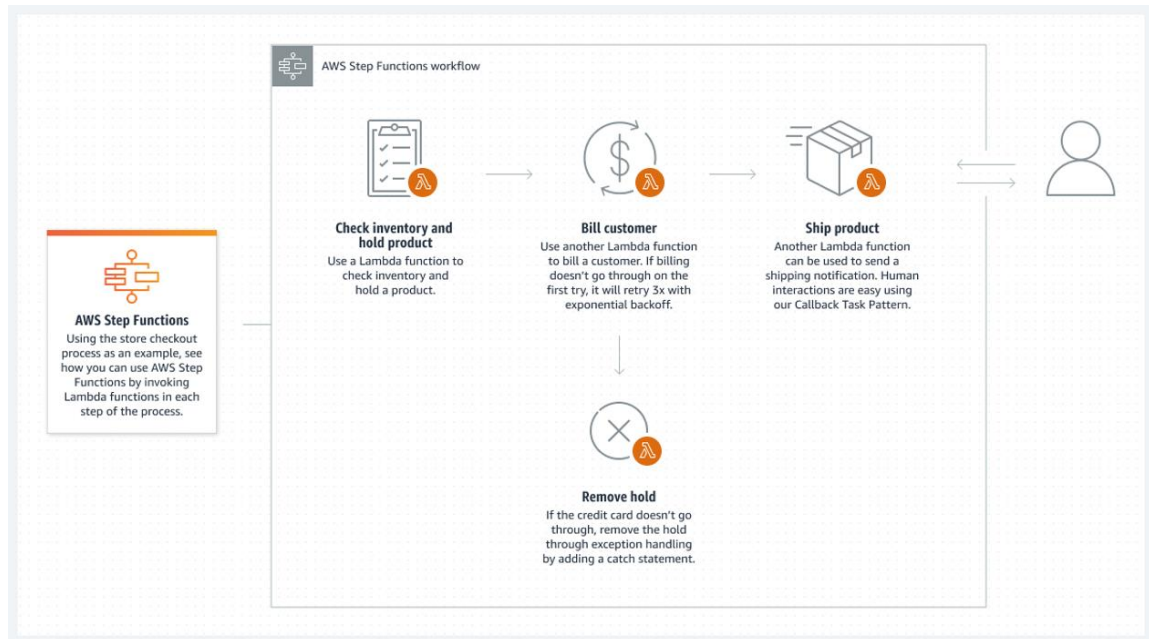


Fig. 18 – Diagrama de funcionamiento de AWS Step Functions

Se puede consultar más información sobre AWS Step Functions en la [documentación del fabricante](#).

AWS Systems Manager – Incident Manager

AWS Systems Manager Incident Manager es una consola de administración de incidentes diseñada para ayudar a los usuarios a mitigar y recuperarse de incidentes que afectan a las aplicaciones en AWS. Un incidente es cualquier interrupción no planificada o reducción de la calidad de los servicios.

Incident Manager reduce el tiempo de resolución de incidentes notificando a los encargados de responder sobre el impacto, proporcionando herramientas de colaboración para hacer que los servicios se vuelvan a poner en funcionamiento. Para lograr el objetivo principal de reducir el tiempo de resolución de incidentes, Incident Manager automatiza los planes de respuesta y permite el escalado del equipo de respuesta.

Entre las características de Incident Manager se incluyen:

- Planes de respuesta: crear y automatizar planes de respuesta iniciados por las alarmas de CloudWatch o los eventos de EventBridge.

- Automatización de la guía del incidente: definir las guías de ejecución mediante Systems Manager Automation para automatizar la respuesta y proporcionar pasos detallados al equipo de respuesta.
- Interacción y escalamiento: conectar automáticamente a las personas correctas para cada incidente único. Interactuar a través de diferentes métodos de contacto y escalar a través del equipo para que garanticen la visibilidad y la participación activa durante los incidentes.
- Colaboración activa: el equipo de respuesta a incidentes podrá responder activamente a los incidentes mediante la integración con el AWS Chatbot.
- Seguimiento de incidentes: revisar los detalles del incidente para obtener información durante un incidente. Crear y corregir elementos de seguimiento mientras sigue las guías de incidente.

Se puede encontrar más información sobre AWS Systems Manager – Incident Manager en [la documentación del fabricante](#).

AWS Backup

AWS Backup es un servicio totalmente administrado que facilita la centralización y la automatización de la protección de datos en todos los servicios de AWS, en la nube y en las instalaciones del cliente. Si bien este servicio se describe a grandes rasgos en el apartado Copias de Seguridad [mp.info.6] de la guía **CCN-STIC 887A – Guía de Configuración Segura para AWS**, es importante tener aquellos aspectos relacionados especialmente con la fase de la recuperación de los incidentes de seguridad.

En primer lugar, los listados de las copias de seguridad pueden ser consultados mediante programación o mediante la consola de AWS Backup, y estas consultas se pueden realizar por recurso protegido o por almacén de copias.

Asimismo, la restauración de las copias de seguridad también puede realizarse desde la consola o mediante programación, a través de la operación de API `StartRestoreJob`. Para cada restauración se crea un trabajo de restauración con un identificador de trabajo único. Es importante tener en cuenta que, al utilizar AWS Backup para restaurar una copia de seguridad, se crea un nuevo recurso en función de la copia de seguridad que se restaure. Esto sirve para evitar que la actividad de restauración destruya los recursos existentes.

AWS Backup se puede utilizar para la restauración de diferentes recursos de AWS, como pueden ser datos de S3, máquinas virtuales, sistemas de archivos, volúmenes de Amazon EBS, bases de datos, instancias S3, etc.

Para más información sobre el servicio AWS Backup, puede consultar este [enlace](#).

AWS Elastic Disaster Recovery

AWS es un servicio para la recuperación de desastres en AWS. Este servicio ofrece

capacidades de recuperación y se administra desde la consola de administración de AWS. Con AWS Elastic Disaster Recovery se puede recuperar las aplicaciones en AWS desde la infraestructura en la nube, la infraestructura física y servicios de virtualización como VMware, vSphere o Microsoft Hyper-V. Además, se puede utilizar Amazon Elastic Disaster Recovery para recuperar instancias de Amazon EC2 en una región diferente o para recuperar todas las aplicaciones y bases de datos que se ejecuten en versiones compatibles de con los sistemas operativos Windows y Linux.

Para configurar una replicación, es preciso instalar un agente en los servidores a replicar y añadirlos en la consola de AWS Elastic Disaster Recovery. Los datos se replicarán en una subred de preproducción en la cuenta de AWS, en la región que se seleccione. Si se necesita recuperar aplicaciones, se pueden lanzar instancias de recuperación en AWS en cuestión de minutos, utilizando el estado más actualizado del servidor a recuperar o un punto anterior en el tiempo. Después de que las aplicaciones se ejecuten en AWS, se puede optar por mantenerlas o por iniciar la replicación de datos de vuelta al sitio principal cuando se resuelvan los problemas.

Se pueden realizar pruebas no disruptivas para confirmar que la implementación se ha completado. Durante el funcionamiento normal, se debe mantener la monitorización de la replicación y realizar periódicamente simulacros no disruptivos de recuperación y vuelta atrás.

Para más información sobre el servicio AWS Elastic Disaster Recovery, puede consultar este [enlace](#).

4. RECURSOS ADICIONALES

A continuación, se proporcionan un conjunto de herramientas para la gestión de incidentes desarrolladas por diferentes equipos de AWS:

- [Virtual SOC with MITRE Attack integration with AWS Security Hub example](#). Proyecto que amplía las capacidades de Security Hub para mapear los hallazgos que detecta en la infraestructura de AWS con información del marco MITRE Attack.
- [Automated Incident Response and Forensics Framework](#). Marco de trabajo en AWS para el despliegue procesos automatizados de respuesta a incidentes y de análisis forense.
- [AWS Incident Response Playbooks Workshop](#). Proyecto que crea un entorno *sandbox* en una cuenta de AWS para facilitar el desarrollo de guías de respuesta a incidentes que mejoren la capacidad de respuesta a dichos incidentes.
- [AWS Twitch “Saferoom”](#) – Participación de los clientes en redes sociales.
- [AWS Cloud Saga](#) – Herramienta para que los clientes prueben los controles

y alertas de seguridad dentro de su entorno AWS utilizando alertas basadas en eventos de seguridad vistos por el CIRT de AWS.

- [IR Response Playbooks](#) – Colección de archivos que se proporciona como marco de ejemplo para la creación de guías de seguridad para la defensa en escenarios de ataque en AWS.
- [Blog del CIRT de AWS disponible públicamente](#).
- [Assisted Log Enabler](#) – Ayuda a los clientes a habilitar los registros.
- [Athenea Bootstrap](#) – Método de instalación y configuración rápida de Athenea; crea un entorno de análisis totalmente configurado.
- [Security Workshops](#) – Aprende a instalar y utilizar herramientas de código abierto como Assisted Log Enabler, CloudSaga y Security Analytics Bootstrap.

