

Informe de amenazas Gen H1 2026

La confianza se ha convertido en una nueva superficie de ataque.

Fuente: Gen Digital - H1 2026 Threat Report

Revisión: 24/07/2026

Durante el primer semestre de 2026, los atacantes han desplazado sus técnicas hacia entornos que usuarios y empresas ya consideran legítimos: reservas, mensajería, soporte técnico, anuncios, pagos, identidad, herramientas de desarrollo y agentes de IA.

Resumen ejecutivo

La confianza se ha convertido en superficie de ataque

Gen identifica un patrón claro: muchas amenazas efectivas ya no dependen de malware evidente ni de engaños burdos. Funcionan porque aparecen dentro de contextos cotidianos en los que usuarios y empresas ya confían.

- Las estafas se integran en reservas, procesos de compra, páginas de soporte y anuncios.
- La toma de cuentas puede apoyarse en funciones normales, como dispositivos vinculados en aplicaciones de mensajería.
- Los atacantes aprovechan marcas conocidas, flujos de pago, sistemas de actualización, herramientas de desarrollo y permisos delegados.
- La IA introduce un nuevo nivel de riesgo cuando los agentes pueden navegar, leer archivos, instalar software o ejecutar acciones con permisos concedidos por el usuario.

20,3 M

ataques de falsas páginas de soporte técnico bloqueados

1,9 B

intentos de rastreo bloqueados durante H1 2026

3,3 M

alertas de brecha con fuente atribuida enviadas a usuarios Norton y LifeLock

Indicadores principales

Cifras destacadas del semestre

Indicador	Dato comunicado por Gen	Lectura práctica
Estafas de tiendas falsas	114,2 millones de ataques bloqueados, +109%	Los comercios falsos escalan con rapidez mediante plantillas reutilizables, anuncios y dominios desechables.
Suplantación de organismos públicos	+387%	Los delincuentes explotan la confianza en instituciones públicas para obtener dinero o datos personales.
Suplantación familiar	más de +454%	La ingeniería social se apoya en relaciones de confianza y canales cotidianos.
Soporte técnico falso	20,3 millones de ataques bloqueados	El objetivo sigue siendo conseguir acceso remoto, pagos o información financiera.
Publicidad fraudulenta	más de 304 millones de impresiones detectadas en UE y Reino Unido en menos de un mes	La publicidad online permite que el fraude llegue al usuario en entornos aparentemente legítimos.
Rastreo online	aprox. 1,9 mil millones de intentos bloqueados	La privacidad también se erosiona por rastreo persistente y perfilado continuo.
Web skimming	1 millón de ataques bloqueados, +212%	Los pagos online siguen siendo un punto crítico porque el usuario espera introducir datos de tarjeta.

Amenazas y fraude

Cuando el ataque parece parte del proceso normal

Muchas amenazas actuales no buscan parecer sospechosas. Se integran en procesos normales y aprovechan el momento en el que la víctima espera recibir una solicitud, un aviso, una llamada o una validación.

Escenario	Cómo se aprovecha la confianza	Riesgo
Reservas y hoteles	Mensajes o solicitudes que encajan con una reserva real o un flujo de pago esperado.	Pago fraudulento, robo de datos y validaciones falsas.
Mensajería	Abuso de funciones legítimas como vincular un dispositivo o aprobar una sesión.	Acceso persistente a conversaciones y posibilidad de suplantar al usuario.
Soporte técnico	Páginas o alertas que simulan errores del sistema y servicios de ayuda.	Instalación de herramientas remotas, pagos indebidos o robo de información.
Tiendas online falsas	Anuncios y dominios que llevan a tiendas convincentes con procesos de checkout falsos.	Pérdida económica y exposición de datos de pago.

Implicación para empresas

- La formación debe centrarse en contexto, no solo en reconocer enlaces extraños.
- Los controles de navegador, reputación web, protección antiphishing y seguridad en pagos cobran más importancia.
- Los usuarios necesitan validaciones sencillas antes de autorizar sesiones, pagos, accesos remotos o cambios de cuenta.

Identidad y privacidad

Datos expuestos y acceso persistente

El riesgo de identidad en H1 2026 no se limita a una contraseña robada. Datos expuestos, sesiones de confianza y contexto real pueden reutilizarse para que el atacante parezca autorizado.

Área	Señal destacada	Interpretación
Brechas de datos	18.618 eventos de brecha identificados, +94,5% frente a H2 2025	Más información personal disponible para phishing, fraude y toma de cuentas.
Alertas de brecha	3,3 millones de alertas con fuente atribuida, +628,1%	Aumenta la presión sobre usuarios y organizaciones para responder a filtraciones.
Registros con email	más de 15,7 millones de registros con direcciones de correo	El email sigue siendo un identificador clave para campañas dirigidas.
Actividad bancaria	alertas de actividad en cuentas bancarias +734%	Mayor exposición de señales financieras y validaciones asociadas a identidad.
Rastreo	310,8 millones de intentos bloqueados al mes de media	La privacidad se degrada también por seguimiento continuo, no solo por brechas visibles.

Medidas recomendables

- Activar alertas de brechas y supervisión de identidad donde esté disponible.
- Usar gestor de contraseñas y credenciales únicas por servicio.
- Revisar dispositivos vinculados en servicios de mensajería, correo y plataformas críticas.
- Reducir el rastreo del navegador y bloquear publicidad o páginas maliciosas cuando sea posible.

IA y herramientas

Agentes de IA: un nuevo frente de seguridad

El riesgo no está solo en textos falsos, voces sintéticas o imágenes manipuladas. Aparece cuando un agente puede actuar con permisos del usuario.

Comportamiento de alto riesgo	Por qué importa
Ejecutar comandos peligrosos del sistema	Puede permitir cambios no autorizados o compromiso del equipo.
Abrir canales remotos de comando	Facilita control externo del sistema.
Descargar y ejecutar código de internet	Introduce riesgo de malware o scripts no verificados.
Leer archivos de credenciales	Puede exponer contraseñas, tokens o claves.
Crear acceso remoto persistente	Permite que el atacante mantenga control después del incidente inicial.
Intentar anular instrucciones del agente	Indica manipulación del comportamiento esperado del sistema de IA.

Criterio práctico

La seguridad debe actuar antes de que el usuario o el agente concedan confianza: antes de ejecutar herramientas, instalar paquetes, abrir sesiones, leer archivos sensibles o conectar servicios.

Recomendaciones DEISA

Controles prioritarios para clientes

Prioridad	Recomendación	Objetivo
1	Protección web, antiphishing y reputación de URLs	Bloquear tiendas falsas, soporte técnico fraudulento, anuncios maliciosos y páginas de pago falsas.
2	Protección de identidad y alertas de brechas	Detectar exposición de datos y acelerar la respuesta del usuario.
3	Gestor de contraseñas y MFA	Reducir el impacto de credenciales filtradas y toma de cuentas.
4	Revisión de sesiones y dispositivos vinculados	Evitar accesos persistentes en mensajería, correo y servicios cloud.
5	Control de permisos y ejecución en herramientas de IA	Evitar acciones automáticas peligrosas con permisos concedidos.
6	Formación orientada a contexto	Preparar al usuario frente a ataques que parecen procesos normales.

Lectura comercial

El informe refuerza la necesidad de soluciones que combinen protección del dispositivo, navegador, identidad, privacidad y asistencia frente a estafas. La amenaza ya no está solo en el archivo malicioso: también aparece en el momento en que el usuario confía en un proceso aparentemente legítimo.

Fuente

Origen de la información

Elemento	Detalle
Informe original	Gen H1 2026 Threat Report
URL	https://www.gendigital.com/blog/insights/reports/threat-report-h1-2026
Comunicado Gen	Gen Half-Year Threat Report: Attackers are Moving Closer to the Systems People Trust
Fecha del comunicado	15/07/2026

Para consultar el informe completo, revisar la publicación oficial de Gen Digital.