



LEY DE SERVICIOS DIGITALES

Informe de Transparencia

para el período de informe que finaliza en diciembre de 2025

Informe publicado: 17 de febrero de 2026

1. Introducción

Este informe de transparencia se ha elaborado de conformidad con el artículo 15 del Reglamento (UE) 2022/2065, la Ley de Servicios Digitales (DSA). Describe nuestras prácticas de moderación de contenido y nuestras decisiones de cumplimiento en relación con áreas de productos específicas y tiene como objetivo proporcionar información clara, accesible y completa sobre cómo gestionamos y moderamos el contenido en nuestra plataforma.

Este informe cubre específicamente las acciones y procedimientos de moderación aplicados a los siguientes productos: las plataformas en línea y los servicios digitales de Awaze Group que facilitan el listado, descubrimiento y reserva de alojamiento vacacional a corto plazo, incluido contenido generado por el usuario, como listados de propiedades, descripciones, imágenes, reseñas y comunicaciones relacionadas.

Este informe es parte de nuestro compromiso continuo con la transparencia, la rendición de cuentas y el cumplimiento de las obligaciones establecidas en la DSA.

Nombre del proveedor de servicios	Este informe cubre las siguientes entidades legales del Grupo Awaze: Novasol, Fincallorca, Ardennes Étape y SandyBlue - Colectivamente el “Grupo Awaze”
Fecha de publicación del informe	17 de febrero de 2026
Fecha de publicación del último informe anterior	17 de febrero de 2025
Fecha de inicio del período del informe	1 de enero de 2025
Fecha de finalización del período del informe	31 de diciembre de 2025

2. Órdenes recibidas de las autoridades de los Estados miembros de la UE

De conformidad con los artículos 9 y 10 de la Ley de Servicios Digitales (LSD), esta sección proporciona información sobre las órdenes recibidas de las autoridades competentes de los Estados miembros de la UE durante el período del informe. Estas órdenes se refieren a contenidos ilícitos y solicitudes de información.

2.1. Órdenes de actuación contra contenidos ilícitos por parte de las autoridades de los Estados miembros

Tipo de contenido ilegal	Número de pedidos recibidos	Estado miembro que emite la orden	Tiempo medio para confirmar la recepción	Tiempo medio para hacer efectiva la orden
Bienestar animal	0	N / A	N / A	N / A
Infracciones de información al consumidor	0	N / A	N / A	N / A
Violencia cibernética	0	N / A	N / A	N / A
Violaciones de la protección de datos y la privacidad	0	N / A	N / A	N / A
Discurso ilegal o dañino	0	N / A	N / A	N / A
Infracciones de propiedad intelectual	0	N / A	N / A	N / A
Efectos negativos sobre el discurso cívico o las elecciones	0	N / A	N / A	N / A
Protección de menores	0	N / A	N / A	N / A
Riesgo para la seguridad pública	0	N / A	N / A	N / A
Estafas/fraudes	0	N / A	N / A	N / A
autolesión	0	N / A	N / A	N / A
Productos inseguros, no conformes o prohibidos	0	N / A	N / A	N / A
Violencia	0	N / A	N / A	N / A
Tipo de contenido ilegal no especificado por la autoridad	0	N / A	N / A	N / A
Todos los demás tipos	0	N / A	N / A	N / A
Total:	0	N / A	N / A	N / A

2.2 Orden de facilitar información sobre los destinatarios del servicio por parte de las autoridades del Estado miembro

Motivo del informe/Tipo de contenido ilegal	Número de avisos recibidos	Número de avisos recibidos por los señalizadores de confianza	Número de acciones realizadas a raíz de los avisos	No. procesados únicamente por medios automatizados medio	Tiempo medio para actuar
Bienestar animal	0	N / A	N / A	N / A	N / A
Infracciones de información al consumidor	0	N / A	N / A	N / A	N / A
Violaciones de la protección de datos y la privacidad	0	N / A	N / A	N / A	N / A
Discurso ilegal o dañino 0		N / A	N / A	N / A	N / A
Infracciones de propiedad intelectual	0	N / A	N / A	N / A	N / A
Efectos negativos sobre el discurso cívico o las elecciones	0	N / A	N / A	N / A	N / A
Protección de menores	0	N / A	N / A	N / A	N / A
Riesgo para la seguridad pública 0		N / A	N / A	N / A	N / A
Estafas/fraudes	0	N / A	N / A	N / A	N / A
autolesión	0	N / A	N / A	N / A	N / A
Productos inseguros, no conformes o prohibidos	0	N / A	N / A	N / A	N / A
Violencia	0	N / A	N / A	N / A	N / A
Tipo de contenido ilegal no especificado por la autoridad	0	N / A	N / A	N / A	N / A
Todos los demás tipos	0	N / A	N / A	N / A	N / A
Total:	0	N / A	N / A	N / A	N / A

3. Informes/avisos de usuarios

Esta sección describe la cantidad y la naturaleza de las denuncias presentadas por usuarios, otras personas y entidades sobre contenido que consideran ilegal o que infringe los términos y condiciones de nuestra plataforma. Además, describe cómo gestionamos las medidas de moderación de contenido en respuesta a las denuncias de los usuarios.

Informes recibidos de los usuarios

Motivo del informe/Tipo de contenido ilegal	Número de avisos recibidos	Número de avisos recibidos por los señalizadores de confianza	Número de acciones realizadas a raíz de los avisos	No. procesados únicamente por medios automatizados medio	Tiempo medio para actuar
Bienestar animal	0	N / A	N / A	N / A	N / A
Infracciones de información al consumidor	0	N / A	N / A	N / A	N / A
Violaciones de la protección de datos y la privacidad	0	N / A	N / A	N / A	N / A
Discurso ilegal o dañino 0		N / A	N / A	N / A	N / A
Infracciones de propiedad intelectual	0	N / A	N / A	N / A	N / A
Efectos negativos sobre el discurso cívico o las elecciones	0	N / A	N / A	N / A	N / A
Protección de menores	0	N / A	N / A	N / A	N / A
Riesgo para la seguridad pública 0		N / A	N / A	N / A	N / A
Estafas/fraudes	0	N / A	N / A	N / A	N / A
autolesión	0	N / A	N / A	N / A	N / A
Productos inseguros, no conformes o prohibidos	0	N / A	N / A	N / A	N / A
Violencia	0	N / A	N / A	N / A	N / A
Tipo de contenido ilegal no especificado por la autoridad	0	N / A	N / A	N / A	N / A
Todos los demás tipos	0	N / A	N / A	N / A	N / A
Total:	0	N / A	N / A	N / A	N / A

4. Moderación de contenido realizada por iniciativa propia de Awaze

En Awaze nos comprometemos a mantener un entorno seguro y libre de abusos tanto para nuestros clientes como para sus usuarios finales. Como proveedor de una plataforma de atención al cliente, nuestra plataforma permite a las empresas interactuar con los usuarios mediante mensajería, correo electrónico e integraciones, todo lo cual conlleva un riesgo de abuso si no se protege adecuadamente.

Utilizamos un enfoque estratificado para la moderación de contenido, que combina sistemas de detección automatizados, herramientas de administración interna y procesos de revisión manual. Esta sección ofrece un resumen de las acciones de moderación de contenido que realizamos por iniciativa propia, sin obligación legal ni notificación a terceros.

La moderación, realizada por iniciativa propia, incluye tanto la detección proactiva mediante sistemas automatizados como la revisión manual por parte de moderadores de contenido. Nos comprometemos a garantizar que todo el personal involucrado en la moderación de contenido cuente con las habilidades, los conocimientos y los recursos necesarios para desempeñar sus responsabilidades de forma justa, precisa y conforme a las leyes y políticas internas aplicables.

Moderación de contenido por iniciativa propia

Tipo de contenido ilegal u otra violación de la AUP	Número de artículos moderados	Número de elementos detectados utilizando únicamente métodos automatizados medio	Tipo de restricción aplicada
Estafas/fraudes	Correos electrónicos entrantes filtrados: 2.691.775 (anualizados; incluye 52.046 detecciones de suplantación de identidad) Enlaces maliciosos encontrados: 1135 clics en URL no seguras detectados (correo electrónico) + 98 262 eventos de protección DNS bloqueados (web, incluidos 7165 de phishing) Cargas maliciosas encontradas: 1010 detecciones de malware entrante (correo electrónico)	Correos entrantes filtrados: 2.691.775 Enlaces maliciosos encontrados: 1.135 (correo electrónico) + 98.262 (web) Se encontraron 1010 cargas maliciosas	Restricción de visibilidad: Los correos electrónicos están en cuarentena/ Bloqueado; las solicitudes web están bloqueadas por el filtrado DNS o la política de firewall. Eliminación de contenido: Los correos electrónicos entrantes pueden ser rechazados (no entregados) si coinciden con la categoría de spam. Controles de malware/suplantación de identidad. Suspensión de cuentas: No se utilizan en estos controles (se gestionan mediante procesos independientes de RR. HH./TI cuando es necesario).
Otro tipo de violaciones a los términos y condiciones de la plataforma	Total de quejas de spam: 96.665 (Rechazo de spam: puerta de enlace de correo electrónico segura; estimación anualizada de estado estable)	Quejas de spam automatizadas: 96.665 (detección automática de spam en la puerta de enlace de correo electrónico segura)	Suspender permisos de la plataforma: No aplicable. Se trata de rechazos de la puerta de enlace de correo electrónico, no de medidas de cumplimiento de la plataforma.
Total:	191.072	197.072	N / A

4. Descripción cualitativa de los medios automatizados

Awaze utiliza controles de seguridad automatizados para reducir estafas, phishing, suplantación de identidad y malware en el acceso al correo electrónico y a la web.

Protecciones de correo electrónico: Utilizamos una pasarela de correo electrónico segura (Mimecast) para inspeccionar el correo electrónico entrante antes de su entrega. La pasarela utiliza comprobaciones automatizadas (reputación, autenticación, Análisis de contenido, escaneo de malware y detección de suplantación de identidad para rechazar o poner en cuarentena mensajes asociados con estafas, fraudes, suplantación de identidad o malware. También utilizamos controles de autenticación de correo electrónico en todos nuestros dominios (SPF, DKIM y DMARC). Estos controles ayudan a los sistemas receptores a verificar que los mensajes que afirman provenir de dominios de Awaze estén autorizados y no hayan sido alterados, y reducen los intentos de suplantación de identidad y suplantación de dominio.

Protecciones web: Utilizamos un servicio de seguridad SD-WAN administrado (Cato) que aplica filtrado DNS, políticas de firewall y prevención de intrusiones. Esto bloquea el acceso a dominios maliciosos conocidos, infraestructura de phishing y destinos de comando y control, y bloquea patrones de tráfico de alto riesgo en el borde de la red.

Nota del informe: Las cifras anuales se estiman utilizando las ventanas de informes de los proveedores (Mimecast agosto-diciembre de 2025; Cato 12 de octubre-31 de diciembre de 2025) y suponen un estado estable sin cambios en las políticas o el volumen.

Fines precisos

Las herramientas automatizadas tienen como objetivo proteger la actividad de mensajería entrante y saliente, prevenir el abuso, preservar la reputación del remitente y garantizar el cumplimiento de las directrices y la legislación (por ejemplo, las directrices de envío de correo electrónico y la conformidad con la ley CAN-SPAM). Sus objetivos específicos incluyen:

- Detectar actividades y patrones sospechosos durante el registro;
- Evaluar el contenido en el momento de su creación y acceso a través de enlaces, cargas y otros contenidos generados por el usuario;
- Monitoreo de límites de velocidad y umbrales de uso para funciones clave;
- Puntuación de riesgo basada en datos históricos;
- Prevenir la entrega de phishing, suplantación de identidad y malware por correo electrónico al rechazarlos o ponerlos en cuarentena. mensajes entrantes que coinciden con los controles de amenazas automatizados;
- Prevenir el acceso a destinos web maliciosos bloqueando la resolución de DNS y/o el tráfico web para dominios y categorías asociados con malware, phishing, DGAs y comando y control; y
- Detectar y bloquear ataques basados en la red (por ejemplo, intentos de fuerza bruta, bloqueos basados en reputación, escaneo de vulnerabilidades y patrones de explotación) utilizando firmas IPS y políticas de firewall.

Uso de autenticación de correo electrónico SPF/DKIM/DMARC para reducir la suplantación de dominios de Awaze y mejorar la detección y el rechazo de correos electrónicos de suplantación de identidad y phishing.

Indicadores de precisión y posible tasa de error

La confianza en nuestras herramientas es sólida, con una baja tasa de falsos positivos. Sin embargo, los clientes pueden apelar a nuestro equipo de soporte si creen que se ha producido un falso positivo en nuestros procesos de moderación de contenido o en nuestras herramientas antiabuso.

Correo electrónico: Las detecciones se basan en inteligencia de amenazas automatizada, comprobaciones de autenticación, análisis de contenido y análisis de malware. Pueden producirse falsos positivos (por ejemplo, remitentes masivos legítimos o dominios recién registrados) que se mitigan mediante la inclusión en listas blancas, el ajuste de políticas y la revisión de mensajes en cuarentena o rechazados.

Web/DNS: Los bloqueos de DNS y firewall se basan en fuentes de amenazas, categorización, indicadores de comportamiento (por ejemplo, detección de DGA) y firmas IPS. Pueden producirse falsos positivos (por ejemplo, dominios mal categorizados o infraestructura compartida) y se gestionan mediante excepciones/listas de permitidos y ajustes periódicos de reglas.

Awaze revisa las tendencias y ajusta las políticas cuando es necesario para reducir los falsos positivos y mantener la protección.

- Los espacios de trabajo deben pasar una evaluación antiabuso antes de poder aprovechar muchas funciones, especialmente aquellos que permiten la comunicación saliente.
- La limitación de velocidad, el escaneo de contenido y la detección de patrones de spam están implementados en muchos canales de nuestro producto
- Si el contenido no infringe nuestros términos, pero un cliente desea administrar su espacio de trabajo según sus propios términos, puede eliminar contenido o bloquear usuarios como lo considere conveniente.
- Las anulaciones manuales por parte del servicio de atención al cliente (CS) están disponibles para bloqueos automatizados.
- Los empleados de Awaze (por ejemplo, Atención al cliente) tienen herramientas para aprobar o denegar la reincorporación de los empleados bloqueados. intentos de correo electrónico.
- Aplicamos controles en capas (puerta de enlace de correo electrónico + filtrado DNS + firewall + IPS) para que no haya un solo control únicamente confiado.
- Utilizamos listas blancas/excepciones (cuando está justificado) y ajustamos las políticas en función de las condiciones operativas. Impacto y riesgo de seguridad.
- Mantenemos registros de auditoría e informes de eventos de seguridad para respaldar la investigación y la respuesta a incidentes.
- Las políticas de seguridad y las capacidades de detección se mantienen mediante actualizaciones de los proveedores y recursos internos. revisar

Aplicamos autenticación de correo electrónico a nivel de dominio (SPF, DKIM y DMARC) como protección adicional para reducir la suplantación de identidad y mejorar la confiabilidad de las decisiones de filtrado de correo electrónico automatizado.

6. Quejas recibidas

Número de quejas que recibimos a través de nuestros sistemas internos de gestión de quejas

Mecanismo de quejas internas

Número de quejas presentadas	0
Base de la queja	N / A
Decisiones tomadas a raíz de una queja	N / A
Tiempo medio para atender una queja	N / A