

Responsible Artificial Intelligence Program



At RIMAC, we have established a Policy for the Safe, Responsible, and Ethical Use of Artificial Intelligence Systems

RIMAC	POLICY				Code: POL-6098
	SAFE, RESPONSIBLE AND ETHICAL USE OF ARTIFICIAL INTELLIGENCE (AI) SYSTEMS				Version: 008
	Macroprocess:	Data and Analytics Governance	Process:	AI Growth	Pages 11 of 11

1. OBJECTIVE

This policy establishes the principles, responsibilities, and guidelines for the safe, ethical, and sustainable use of Artificial Intelligence (AI) within the organization. It seeks to ensure that the development, implementation, and use of AI solutions are carried out in strict compliance with applicable regulations - including Law No. 31814 and its Regulations approved by Supreme Decree No. 115-2025-PCM - and with the principles of legality, transparency, accountability, non-discrimination, among others.

2. SCOPE

This policy applies to all areas and levels of the organization and is mandatory for employees, directors, partners, strategic allies, suppliers, developers, technical teams, and third parties that design, develop, use, or manage artificial intelligence systems on behalf of the company, as well as users and persons affected by such systems, whose rights to transparency, explainability, appeal, and redress are guaranteed.

This policy applies to all AI systems used, developed, or acquired by the organization throughout all stages of their life cycle and includes internal developments under any modality (full-code, low-code/no-code, or configuration on domain platforms), third-party systems, material modifications, and non-formalized uses (Shadow AI).

Likewise, all AI-related contracts must include obligations aligned with this policy. Any exception shall require the express authorization of senior management, following a risk assessment.

3. DEFINITIONS

For terms of general corporate use - such as risk appetite, personal data, sensitive data, among others - the definitions in force under the organization's internal policies shall apply, including the Personal Data Protection Policy, the corporate data classification guidelines, and the corporate risk management framework.

The definitions contained in this section are specific to the field of artificial intelligence or have a particular scope within this policy. For the operational definitions of AI assistants, agents, and multi-agent systems, see section 5.9 of this policy.

Click [here](#) to view our policy



Limiting Access to Sensitive AI Capabilities

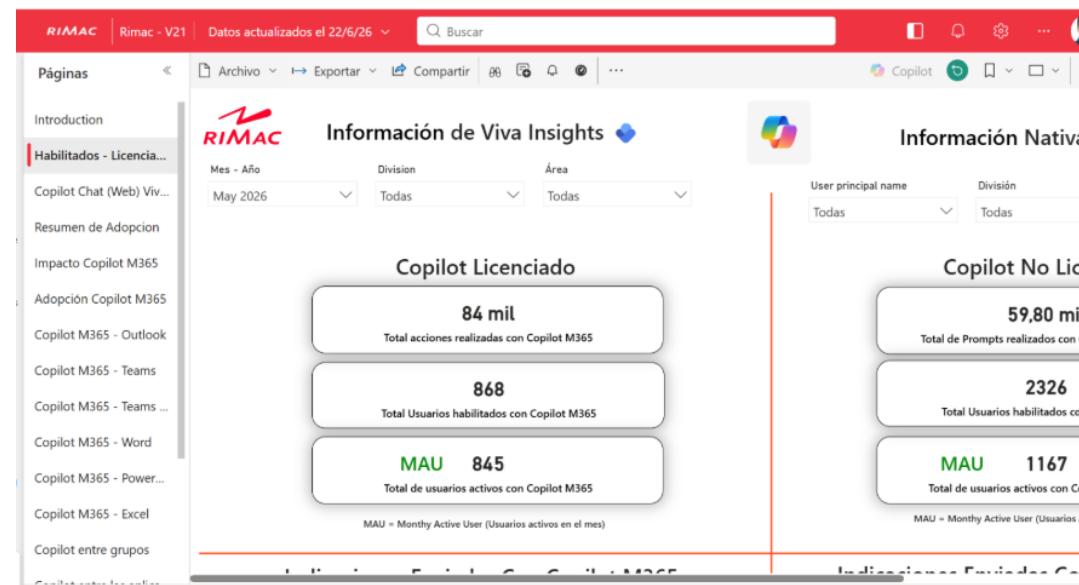
As part of our Responsible Artificial Intelligence Program, we manage access to AI-enabled through differentiated access levels and authorization criteria. Access rights and responsibilities are defined through the AI Operating Model and the corresponding RACI Matrix. In these cases, the following elements are considered:

- Use case and business objective.
- Roles and responsibilities of the participating areas.
- Risks associated with the solution.
- Nature of the data used.
- Dependencies on other corporate systems or capabilities.
- Approvals and deliverables defined for each stage of the lifecycle.

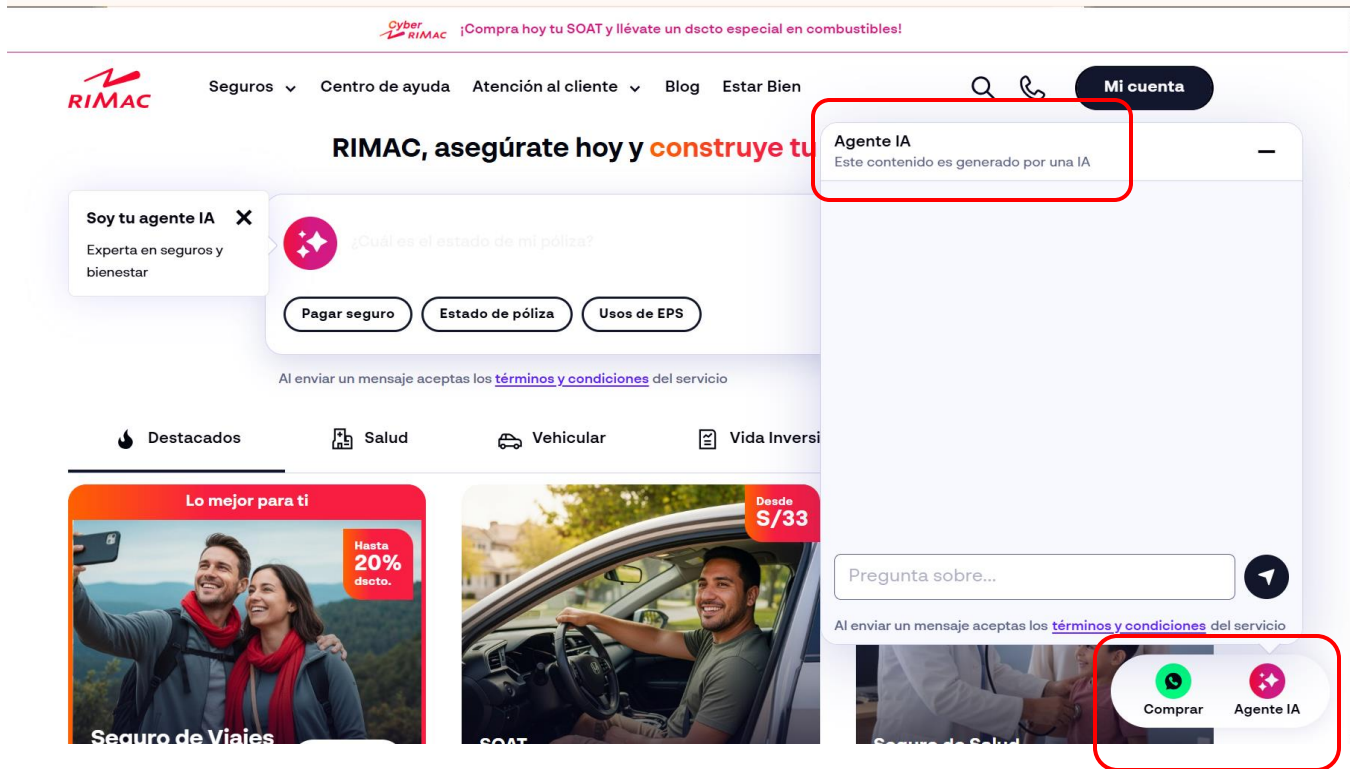
To strengthen oversight, we use a dashboard that allows us to monitor AI tool access and usage across the organization. This includes visibility over assigned M365 and Copilot licenses, active users and usage levels.

The screenshot shows a dashboard with a table titled 'MetricDate'. The table has five columns: 'División', 'Total Usuarios habilitados con Copilot M365', 'Total acciones realizadas con Copilot M365', 'Acciones promedio diario por colaborador', and 'Acciones promedio mensual por colaborador'. The data is organized into a hierarchy starting with 'DIVISION SEGUROS PERSONAS', followed by 'Empleado', 'Jefe', and individual email addresses, ending with a 'Total' row.

División	Total Usuarios habilitados con Copilot M365	Total acciones realizadas con Copilot M365	Acciones promedio diario por colaborador	Acciones promedio mensual por colaborador
DIVISION SEGUROS PERSONAS	211	19.278	4,21	92,68
Empleado	150	13.151	4,04	88,86
Jefe	34	4.022	5,38	118,29
aastengo@rimac.com.pe	1	159	7,23	159,00
amarcelov@rimac.com.pe	1	85	3,86	85,00
angie.milientantalean@rimac.com.pe	1	365	16,59	365,00
bryan.pineda@rimac.com.	1	580	26,36	580,00
Total	868	83.779	4,51	99,15



Distinct labeling of AI-generated content and outcomes of AI-driven decisions



We have AI agents that interact with users and support service-related processes. These agents are clearly labeled as artificial intelligence tools, allowing users to identify when they are interacting with an AI-based solution rather than a human representative.

This labeling strengthens transparency, promotes informed interaction, and supports the responsible use of AI across our platforms.

We have implemented a dashboard to monitor the status and performance of our AI models. This tool enables the identification of potential deviations, performance deterioration, or model drift over time, allowing the responsible teams to review model behavior and define corrective actions when needed. This strengthens the reliability, traceability, and continuous improvement of AI-based solutions.

Trabajar con mecanismos para detectar y corregir desviaciones o deterioro de los modelos de IA en el tiempo.	Enviar captura del dashboard y consultar a William por otras evidencias.
Realizar evaluaciones periódicas de modelos de IA desplegados respecto a equidad o sesgos	Consultar a William por evidencias.

- **Dashboard del COE AI:**
Monitorea los modelos en uso.
- **Campo ESTADO:**
Indica el estado del modelo, en este caso todos están sin observación

Modelos Activos	Modelos por Revisar
20	0

Periodo	Data Scientist	Modelos Activos	Modelos por Revisar
		20	0

Modelos de Clasificación										
Modelo	DS	Estado	AUC	Gini	Efectividad	Muy Alto	Alto	Medio	Bajo	Muy Bajo
1. Simestralidad SOAT	MA	Modelo OK	63.9%	0.28	0.81	1.36	0.80	0.53	0.42	0.25
2. Simestralidad SOAT Empresas	MA	Modelo OK	89.0%	0.78	2.29	-	12.86	1.40	0.28	-
3. Comportamiento de Pago	MA	Modelo OK	71.0%	0.47	72.87	78.6%	77.77%	58.1%	27.19	-
4. Venta Vehicular	MA	Modelo OK	68.3%	0.32	2.45	5.21	3.96	2.53	1.86	0.67
5. Venta Vehicular Digital	MA	Modelo OK	78.0%	0.56	1.07	4.34	1.73	0.94	0.36	0.16
6. Vida Multiproducto	MA	Modelo OK	71.72	0.43	1.60	4.87	2.18	1.33	0.83	0.45
7. Vida Telemarketing	MA	Modelo OK	67.07	0.34	1.72	-	3.14	1.28	0.78	-
8. Excedente de Ahorro	GM	Modelo OK	78.06	0.58	0.01	0.03	0.01	0.00	0.00	0.00
9. Migración de Frecuencia de Pago	GM	Modelo OK	78.69	0.59	0.04	0.19	0.07	0.04	0.02	0.00
10. Proyección Financieras Provincia (IG)	GM	Modelo OK	-	-	0.00	-	0.00	0.00	0.00	-
11. Proyección Financieras Provincia (RG)	GM	Modelo OK	-	-	0.00	-	0.00	0.00	0.00	-
12. Solicitud de Cancelación APC	GM	Modelo OK	80.41	0.61	1.69	5.87	3.04	1.32	0.68	0.22

1-16 | 16 <

Modelos de Recomendación										
Modelo	DS	Medición	AI día?	Ultima	Estado	Efectividad	Producto 1	Producto 2	Producto 3	Tabla
1. Next Best Product No Digital	GM	MENSUAL	NO	2026-04-01	Modelo OK	80.13	40.92	27.68	11.54	GM_NBP_No_Digital_Monitoring
2. Next Best Product Digital	GM	MENSUAL	NO	2026-04-01	Modelo OK	83.13	42.92	25.79	14.42	GM_NBP_Digital_Monitoring
3. First Best Product	GM	MENSUAL	NO	2026-04-01	Modelo OK	24.96	15.8	5.47	3.69	GM_FBP_Monitoring

1-3 | 3 <

Modelos de Regresión										
Modelo	DS	Ultima	Estado	MAPE	sMAPE	Precisión	p1	p2	p3	p4
1. Estimación de Ingresos	JP	2025-09-01	Modelo OK	57.72	30.77	67.35	80.93	82.89	82.47	78.56

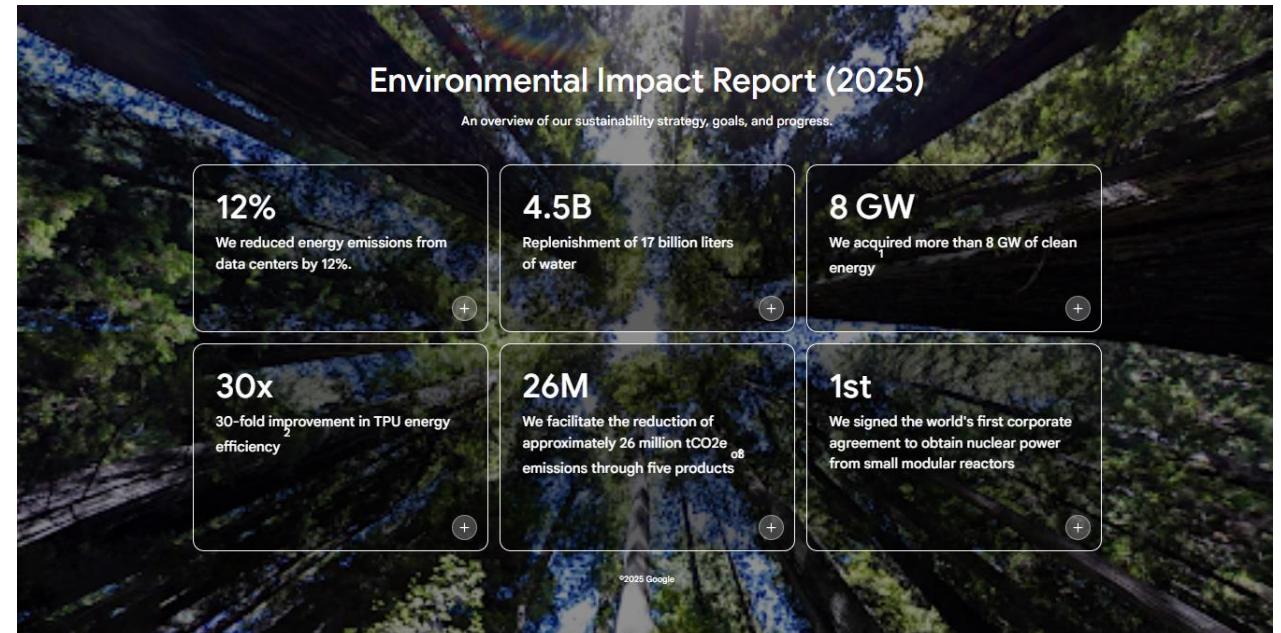
52.64

1-1 | 1

Initiatives (own/with suppliers) to lower the ecological footprint of AI data centers/models

We consider environmental criteria when selecting and working with technology providers that support our AI-enabled tools, prioritizing those that publicly disclose their sustainability commitments, environmental performance, and initiatives to improve the efficiency of their digital infrastructure.

In this regard, Google's 2025 Environmental Impact Report provides evidence of actions aimed at reducing the environmental impact of its operations, including lower data center energy emissions, water replenishment, clean energy procurement, and improvements in the energy efficiency of its AI TPUs. These practices reflect the type of environmental performance we seek to consider when adopting AI-related technologies and digital solutions.



<https://sustainability.google/google-2025-environmental-report/>

Appeals process for users/affected third parties to contest an AI decision or outcome

As part of our **Responsible Artificial Intelligence Program**, we use the **Integrity Channel** as a formal and confidential mechanism to manage reports, concerns, and alerts related to the use of AI. Through this channel, employees, internal users, customers, and third parties may report:

- AI-related incidents
- Risks
- Misuse
- Shadow AI practices
- Or outcomes they consider discriminatory, biased, anomalous, or ethically questionable.

These reports are managed under our corporate Ethics and Compliance model, with the involvement of the relevant technical and AI governance bodies when required.

The screenshot shows the RIMAC Integrity Channel report form. At the top, the RIMAC logo is displayed with the tagline 'Juntos todo es posible'. Below the logo, a note states: 'Las siguientes preguntas son opcionales. Recuerde que la mayor cantidad de información, ayudará a realizar un mejor análisis de su reporte.' The main section is titled 'Datos principales del Reporte' and contains a dropdown menu labeled 'Seleccione el tipo de reporte que desea registrar:'. Below this, step 1 is titled 'Identifique a las personas involucradas en el reporte a registrar. Utilice el siguiente botón para agregar los datos de un nuevo involucrado.' It includes a 'Persona a reportar número 1:' section with input fields for 'Nombres', 'Apellidos', 'Relación con la compañía', and 'Especificar (empresa/cargo/otro)'. A 'Limpiar' button is also present. Step 2 is titled 'Indique en qué lugar sucedió el incidente:' and includes input fields for 'País', 'Estado/Provincia', 'Ciudad', and 'Sede'.

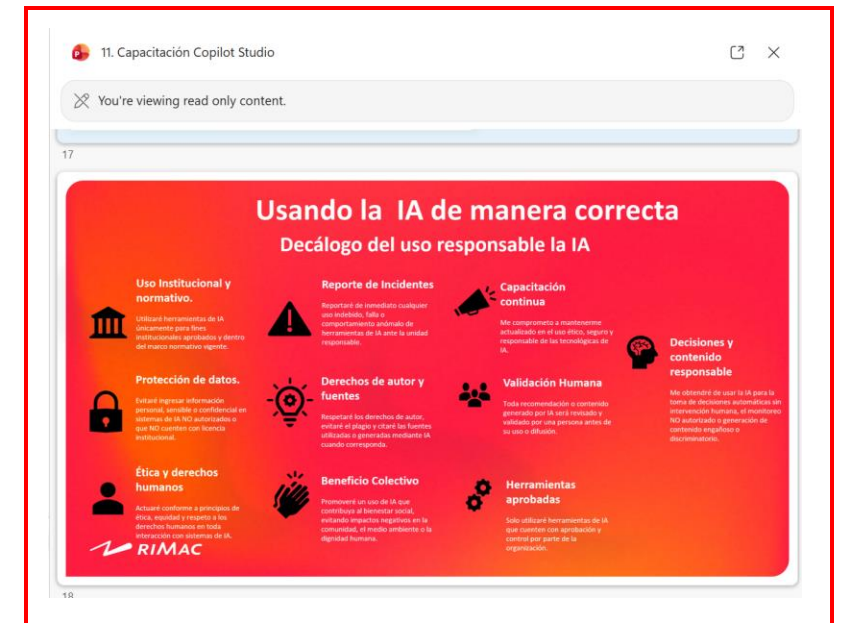
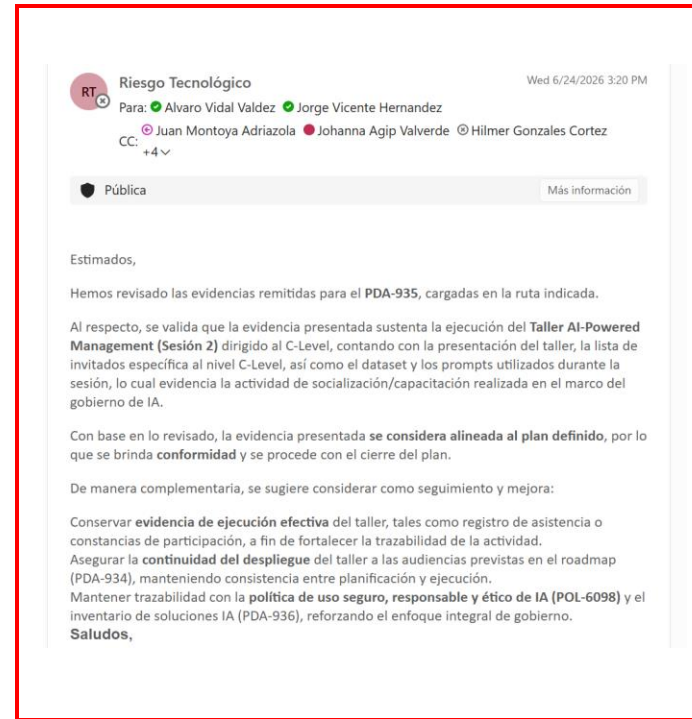
To view our ethics channel, [click here](#)

Training Employees on Responsible and Safe AI Use

We develop **training workshops for employees on the proper, ethical, and safe use of AI tools**. These sessions are aimed at strengthening our internal capabilities for the responsible adoption of technologies such as Copilot and other AI-enabled solutions, addressing key topics such as:

- Information confidentiality
- Data protection
- Cybersecurity risks
- Human oversight
- Prevention of bias
- Appropriate use of AI-generated outputs

Through these workshops, we seek to ensure **that our employees understand both the opportunities and risks associated with AI**, promoting its use in a transparent, secure, and aligned manner with our principles of responsible innovation and sustainability.



Resumen por Copilot

 Uso Interno

Hola Equipo,

Adjunto, un pdf con respuesta al correo. Los enlaces que se mencionan se encuentran en la carpeta compartida:

☐ [05. Gobierno y gestión de IA Generativa \(Chatbots, asistentes\)](#)

Saludos!
Kike

I. Punto: Políticas y procedimientos de desarrollo, autenticación y seguridad

- Diagrama del marco de implementación
- Marco de trabajo de autenticación y seguridad

El detalle de la arquitectura y lineamiento están cubiertos en los enlaces:

- Arquitectura de Gobierno de Agentes de IA
- Lineamiento de Arquitectura – Enfoque Gen AI

Auditoría Interna

Nombre	Modificado por	Modificado	Tamaño de archivo
Arquitectura Agente AIDA Sales v2.pdf	Juan Pablo Quiñones	24/07/2025	321 KB
Arquitectura Agente AIDA Service v1.pdf	Juan Pablo Quiñones	24/07/2025	357 KB
Arquitectura de Gobierno de Agentes de L...	Jorge Vicente Hernan	15/08/2025	105 bytes
Arquitectura TI- Continuidad y Recuperaci...	Jorge Vicente Hernan	21/08/2025	126 bytes
Control documental final - Dashboard (2)...	Jorge Vicente Hernan	27/08/2025	137 KB
Documento Funcional AIDA CDC 2025.pdf	Jorge Vicente Hernan	27/08/2025	96.4 KB
ENLACE A DASHBOARD url	Jorge Vicente Hernan	27/08/2025	115 bytes
Informe_Auditoria_PWC_IA_Generativa_Ag...	Jorge Vicente Hernan	15/08/2025	154 KB

TeamMate+

Seguimiento

Quitar

Guardar

Eliminar

Vista preliminar

Plan de acción

Adjunto

Dimensión

Rimac

Perspectiva

Revisión de incidencia

VISTAS

EXPORTAR

BUSCAR

EDITAR

EXPORTAR

BUSCAR

Mover columnas...

100%

Mover columnas...

Título	Buscar	X	Lanzados	Area	Estado	Auditor Responsable	Revisor	Fecha de implementación de OES	Contacto de negocios	Fecha de implementación reprogramada de OES	Fecha real	Fecha de creación
Rimac												
Rimac Seguro												
SEDAI-017-2025-010 - Asesoría de evidencia de m...			Lanzado	Gobierno de Tecnologías Emergentes	Cerrado - Verificado	ALVARO CANDELA		7/31/2025				1/28/2025

*Título de Observe... SEDA-017-2025-010 - Asesoría de evidencia de monitoreo continuo sobre la gestión de usuarios con aprobaciones en la plataforma web con IA

Natfago	Nombre del Proyecto	Tema SBS	¿Reprogramado? Ob...	Rango	Propiedades	Flujo de trabajo	Asignaciones	Papel de trabajo	Referencias
Fecha de implementación de OES	Fecha de implementación reprogramada de OES	Fecha de cierre OES							
07/31/2025		02/28/2025							

Estado	Comentarios	Fecha	Usuario
Cerrado - Verificado	Se procede a dar por cerrado la presente observación.	2/3/2025	ALVARO CANDELA
En proceso		2/3/2025	ALVARO CANDELA
En revisión		2/18/2025	JAVIER VENDE
Pendiente		1/28/2025	ALVARO CANDELA

SEDAI-017-2025-011 - Aplicaciones AIDA Services y AIDA Sales

LanzadoGobierno de Tecnologías EmergentesCerrado - VerificadoALVARO CANDELA7/31/20251/28/2025

*Título de Observe... SEDA-017-2025-011 - Aplicaciones AIDA Services y AIDA Sales declaradas como internas, permiten el acceso a terceros desde internet

Natfago	Nombre del Proyecto	Tema SBS	¿Reprogramado? Ob...	Rango	Propiedades	Flujo de trabajo	Asignaciones	Papel de trabajo	Referencias
Se evidenció que las aplicaciones AIDA Services y AIDA Sales (inteligencias artificiales), permiten el acceso desde internet, sin necesidad de estar dentro del entorno de RIMAC. Esta validación se realizó desde una red externa.									

SEDAI-017-2025-012 - Asesoría de planes periód...	Lanzado	Gobierno de Tecnologías Emergentes	Cerrado - Verificado	ALVARO CANDELA	3/31/2025		1/28/2025
SEDAI-017-2025-013 - Capacitaciones sobre la gest...	Lanzado	Gobierno de Tecnologías Emergentes	Cerrado - Verificado	ALVARO CANDELA	3/31/2025		1/28/2025

RIMAC strengthens the accountability of its Responsible Artificial Intelligence Program through internal verification processes of its AI management system. These reviews help validate the governance mechanisms, controls, risk management practices, and responsible use procedures applied to AI-related tools and processes.