

	RIESGOS Y CUMPLIMIENTO	Versión 1.0
	POLÍTICA CORPORATIVA SOBRE INTELIGENCIA ARTIFICIAL	2026-07-01
		Página 1 de 12

I. PROPÓSITO Y ALCANCE

Esta Política establece los lineamientos de conducta para el fomento y uso responsable de los sistemas de Inteligencia Artificial (IA), orientados a identificar, analizar y mitigar los potenciales riesgos que puedan derivarse de su uso y desarrollo. Asimismo, se busca salvaguardar los derechos y libertades de las personas, garantizando que la utilización de IA se realice de acuerdo con nuestros valores declarados en el Código de Ética y Conducta (CODEC), los principios éticos, la transparencia y la normativa vigente.

La presente Política se aplica a todos los trabajadores, sin importar la modalidad de su vínculo laboral, y directores del Grupo UNACEM (en adelante “Trabajadores”), lo que incluye tanto a los trabajadores y directores de UNACEM CORP S.A.A., como a los de sus subsidiarias (en adelante “Unidades de Negocio”).

La presente Política es aplicable a la totalidad de iniciativas o proyectos que incorporen tecnología de IA, ya sean desarrollados de manera interna o resulten de la adquisición, implementación o integración de productos, servicios o soluciones de IA de fuentes externas, o de naturaleza mixta.

En caso la legislación local donde opere el Grupo UNACEM considere mayores restricciones que las disposiciones de la presente política, se debe aplicar las disposiciones legales en esa jurisdicción.

II. POLÍTICA

Nuestro CODEC, en la sección 4, Nuestras responsabilidades con el Grupo UNACEM y sus accionistas, letra f, señala:

“Utilizamos la inteligencia artificial para potenciar nuestras capacidades y apoyar la toma de decisiones, siempre bajo criterio y responsabilidad humana. Para aprovechar sus beneficios, las decisiones relevantes cuentan con supervisión humana que, valida la calidad, exactitud y pertinencia de sus resultados, así como la adecuada gestión de riesgos. Usamos estas tecnologías de manera diligente, ética y responsable, mitigando sesgos, promoviendo decisiones informadas, inclusivas y respetuosas de los derechos de terceros.”

III. LINEAMIENTOS

A continuación, se presentan los lineamientos que deben ser considerados en el uso de la IA:

1. Alineamiento estratégico y propósito legítimo

- Los sistemas de IA en el Grupo UNACEM buscan aumentar, complementar y potenciar las aptitudes cognitivas y sociales de las personas, por lo que supervisamos los procesos de trabajo ejecutados por dichos sistemas;
- Actuamos para que los sistemas de IA protejan la dignidad humana y la integridad física y mental de las personas, evitando cualquier daño a los seres humanos;
- Utilizamos la IA de forma responsable, de acuerdo con la estrategia del negocio, con el objetivo de mejorar el bienestar de nuestros grupos de interés y respetando las estructuras de autoridad establecidas;
- Promovemos que los sistemas y entornos de IA sean seguros y robustos desde un punto de vista técnico, que en ningún caso sean destinados a usos malintencionados;

	RIESGOS Y CUMPLIMIENTO	Versión 1.0
	POLÍTICA CORPORATIVA SOBRE INTELIGENCIA ARTIFICIAL	2026-07-01
		Página 2 de 12

e) Promovemos proveedores con centros de datos con baja huella ecológica.

2. Equidad, inclusión y no discriminación

- a) Impulsamos el desarrollo y uso de sistemas de IA de manera segura, responsable, ética, transparente, y respetuosa de los derechos humanos de acuerdo con nuestros lineamientos establecidos en nuestra Política Corporativa de Derechos Humanos.
- b) Promovemos el desarrollo y uso de soluciones con IA que favorezcan decisiones justas, inclusivas y objetivas. Para ello identificamos y mitigamos riesgos de sesgos que puedan afectar a las personas y/o grupos de interés; promoviendo su uso responsable;

3. Transparencia y explicación

- a) Promovemos el uso de la IA de manera transparente y confiable, informando oportunamente a los usuarios internos y externos relevantes sobre sus capacidades, finalidades y funcionamiento, para facilitar su comprensión y adopción responsable;
- b) Impulsamos que las decisiones provistas por la IA sean comprensibles por el nivel adecuado de usuario;
- c) Desarrollamos registros claros que faciliten la trazabilidad y el uso confiable de la IA de acuerdo con la factibilidad técnica, así como las decisiones y responsabilidades durante su ciclo de vida.

4. Las personas y su responsabilidad

- a) Cada solución aplicada a un sistema basado en IA cuenta con un responsable (o dueño) claro que impulse su uso adecuado;
- b) Aprovechamos la IA con criterio, validando sus resultados dentro de nuestro rol y responsabilidad. El uso de soluciones con IA no reemplaza la responsabilidad de los trabajadores frente a sus decisiones en la organización y los marcos normativos aplicables;
- c) Las decisiones críticas siempre dispondrán de un nivel apropiado de supervisión humana.

5. Gobierno en el desarrollo de la IA

La Función Corporativa responsable de la IA (ver Anexo 2) habilita el uso responsable, seguro y escalable de la IA mediante estructuras de responsabilidad a nivel Corporativo y de Unidad de Negocio, así como procesos de acompañamiento y seguimiento corporativo para:

- a) Habilitar y supervisar que el desarrollo de soluciones basadas en IA esté alineado con la estrategia y dispongan de los recursos necesarios.
- b) Fomentar el uso y facilitar la supervisión de las iniciativas de IA en función de tres niveles:
 - i. Básico
 - ii. Operacional
 - iii. Ultra
- c) Facilitar el uso confiable de agentes de IA mediante supervisión y aprobación humana significativa en los puntos de control clave;
- d) Establecer KPI y reportar los logros, beneficios y capacidades desarrolladas a través de la implementación de soluciones con IA;

	RIESGOS Y CUMPLIMIENTO	Versión 1.0
	POLÍTICA CORPORATIVA SOBRE INTELIGENCIA ARTIFICIAL	2026-07-01
		Página 3 de 12

- e) Impulsar Casos de Negocio, según la categoría correspondiente (ver letra b y Anexo 2), con el aporte de las funciones de negocio, tecnología, ciberseguridad, cumplimiento, legal y riesgos, para facilitar la aprobación y uso responsable de soluciones de IA;
- f) Supervisar que los proveedores de soluciones de IA sean evaluados considerando aspectos de seguridad, privacidad, cumplimiento y propiedad intelectual, y de acuerdo con el Código de Conducta de Proveedores;
- g) Gestionar responsablemente el uso y los riesgos asociados a la IA mediante el modelo interdependiente de Tres Líneas de Defensa, alineado con los estándares del Instituto de Auditores Internos (IIA) y las mejores prácticas internacionales, en donde:
 - i. La Primera Línea de Defensa está conformada por quienes utilizan, desarrollan o brindan soporte al uso de la IA, y tienen la responsabilidad de gestionar sus riesgos conforme a las políticas corporativas;
 - ii. La Segunda Línea de Defensa la comprende dos funciones:
 - o La Función Corporativa responsable de la IA, encargada de definir las reglas, supervisar, evaluar, aprobar, capacitar, monitorear, supervisar y reportar los impactos y aspectos vinculados a la promoción, desarrollo y aplicación de la IA (Ver Anexo 2);
 - o La Dirección Corporativa de Riesgos y Cumplimiento, encargada de definir, monitorear y supervisar y reportar la aplicación de la política corporativa de gestión integral de riesgos;
 - iii. La Tercera Línea de Defensa la comprende la Función de Auditoría Interna Corporativa y está encargada de realizar evaluaciones independientes sobre la gobernanza, ambiente de control interno, cumplimiento y efectividad de las políticas sobre el uso de la IA.

6. Usos de la IA

Fomentamos que los trabajadores experimenten con IA y dispongan del soporte corporativo. Para ello, el uso de sistemas de IA podrá realizarse en función del nivel de riesgo asociado:

- a) Uso permitido: Aquellos que no involucren datos sensibles ni generen impactos significativos en usuarios o decisiones de negocio. Contamos con un listado de herramientas de IA autorizadas para su uso corporativo (ver Anexo 3),
- b) Uso restringido: Aquellos que involucren datos personales, procesos críticos o decisiones automatizadas, los cuales requerirán controles adicionales y aprobación previa por las áreas de Tecnología de la Información, Ciberseguridad y Privacidad de Datos¹;
- c) Uso prohibido (cualquiera de las siguientes):
 - i. Aquellos que impliquen el uso de datos confidenciales, sensibles o estratégicos sin autorización (ver numeral 7 y 8);
 - ii. El uso en plataformas de IA no autorizadas;
 - iii. El uso de soluciones IA que incluyan la toma de decisiones automatizadas sin supervisión humana;

¹ El responsable de Privacidad de Datos incluirá otras áreas especializadas cuando corresponda.

	RIESGOS Y CUMPLIMIENTO	Versión 1.0
	POLÍTICA CORPORATIVA SOBRE INTELIGENCIA ARTIFICIAL	2026-07-01
		Página 4 de 12

iv. El uso que vulnere derechos fundamentales o normativas vigentes.

7. Escalabilidad y ciberseguridad

- a) Integramos la IA con sistemas y aplicaciones corporativas para impulsar su uso tecnológicamente escalable, resiliente y sostenible ante actualizaciones y mantenimiento;
- b) Fortalecemos la ciberseguridad en la IA para habilitar su uso confiable, prevenir vulnerabilidades, proteger la información y asegurar la continuidad operativa;
- c) En ese sentido, solo podrán utilizarse las herramientas y aplicaciones de IA que cumplan con ambas:
 - i. Cuenten con la autorización del *Chief Information Officer* Corporativo (CIO) y del *Chief Information Security Officer* Corporativo (CISO) o a quien ellos deleguen (Ver Anexo 3);
 - ii. Y que hayan sido instaladas por el equipo de TI;

8. Protección de datos

- a) Facilitamos el acceso a los datos necesarios para usar IA de forma responsable, bajo criterios establecidos por el *Chief Data Protection Officer* Corporativo (CDPO), de acuerdo con la necesidad, mínimo privilegio y normativa aplicable;
- b) Impulsamos el uso responsable de datos del Grupo UNACEM en soluciones de IA, con controles de permisos, trazabilidad factible y revisión periódica que nos apoya en prevenir accesos o usos de datos no autorizados.

9. Gestión del riesgo

- a) Aplicamos a la utilización de la IA el apetito de riesgo definido en la Política Corporativa de Gestión Integral de Riesgo y la Función Corporativa responsable de la IA coordinará con la Dirección Corporativa de Riesgos y Cumplimiento consultas o revisiones sobre este marco.
- b) Nuestras iniciativas de IA, según corresponda (ver niveles en Anexo 2), son sustentadas en Business Cases adecuados, que permiten evaluar sus riesgos, definir medidas de mitigación. Aquellas de alto riesgo presentan una Evaluación de Impacto en IA bajo un enfoque de doble materialidad, considerando posibles sesgos e impactos éticos, sociales o sobre derechos fundamentales, antes de su implementación;
- c) La metodología para la identificación, evaluación y la gestión integral de riesgos se hará en base a lo descrito en el Manual Corporativo de Gestión Integral de Riesgos;
- d) La taxonomía de riesgos incluye, en adición a los riesgos sobre el derecho de personas, sesgo, datos, seguridad y cumplimiento, aquellos riesgos tradicionales de la IA, los específicos y amplificados de la IA generativa, así como los riesgos de inyección y los riesgos de datos sintéticos y distorsiones;
- e) Las Unidades de Negocio establecerán Planes de Respuesta a Incidentes o Crisis de IA. Esto incluirá protocolos de escalamiento, clasificación (por ejemplo, ataques o brecha de datos), determinación de causa raíz y resolución de fallas, así como establecimiento de escenarios de crisis y simulaciones;
- f) El CIO, CISO y CDPO reportarán periódicamente la gestión de la tecnología, ciberseguridad y protección de datos bajo el ámbito de su aplicación en la IA;

	RIESGOS Y CUMPLIMIENTO	Versión 1.0
	POLÍTICA CORPORATIVA SOBRE INTELIGENCIA ARTIFICIAL	2026-07-01
		Página 5 de 12

10. Capacitación

Promovemos y difundimos el uso de la IA mediante capacitación práctica y periódica para los trabajadores involucrados, fortaleciendo sus capacidades en ética, mitigación de sesgos, protección de datos, seguridad, uso responsable de IA generativa y *prompting* seguro, conforme a lo definido por el dueño de la solución y el Plan de Capacitación del Grupo UNACEM.

11. Auditorías y monitoreo

Realizamos auditorías anuales en sistemas o proyectos de IA de alto riesgo para fortalecer su uso confiable, verificando la gestión de datos, sesgos, beneficios del caso de negocio y cumplimiento normativo aplicable.

Asimismo, los responsables de las soluciones de IA operacional y ultra evaluarán oportunidades de actualización cada dos años y cada año, respectivamente, para fortalecer su desempeño, confiabilidad y alineamiento frente a cambios tecnológicos, regulatorios u operativos.

IV. RESPONSABLE DE LA POLÍTICA Y SU REVISIÓN

El Gerente General Corporativo es responsable de esta política, la cual es parte del Sistema Integral de Riesgos y Cumplimiento. El Comité de Riesgos y Cumplimiento supervisa la ejecución de su contenido a través del monitoreo, evaluaciones y reportes realizados por el Director Corporativo de Riesgos y Cumplimiento.

La Dirección Corporativa de Riesgos y Cumplimiento revisará periódicamente esta Política para garantizar la correcta actualización a la normativa vigente que pueda ser aplicable en el ámbito de la IA, así como cuando ocurra algún cambio importante en el entorno del Grupo UNACEM o por lo menos cada dos años.

La Función Corporativa responsable de la IA y los Gerentes Generales de las Unidades de Negocio deben adoptar todas las medidas necesarias para asegurar que se adopten y cumplan con los lineamientos estipulados en la presente política.

Los miembros del Grupo UNACEM tienen la responsabilidad individual de cumplir con las reglas y lineamientos aquí establecidos, así como de buscar orientación ante cualquier duda².

V. INFRACCIONES A LA POLÍTICA

El incumplimiento de los lineamientos contenidos en esta Política será considerado una conducta inapropiada, y las Unidades de Negocio del Grupo UNACEM, así como el Corporativo, se reservan el derecho de tomar acciones disciplinarias laborales, civiles y/o penales según corresponda en base a dicho incumplimiento.

VI. DOCUMENTOS DE REFERENCIA

- a) Código de Ética y Conducta
- b) Política Corporativa de Gestión Integral de Riesgos

² Acudir al CIO, CISO, CDPO, responsable Legal o al Encargado de Cumplimiento y de Protección de Datos Personales, Gerente Legal Corporativo, al Gerente Corporativo de Cumplimiento o al Director Corporativo de Riesgos y Cumplimiento.

	RIESGOS Y CUMPLIMIENTO	Versión 1.0
	POLÍTICA CORPORATIVA SOBRE INTELIGENCIA ARTIFICIAL	2026-07-01
		Página 6 de 12

- c) Política Corporativa para el Tratamiento de Datos Personales
- d) Código de Conducta de Proveedores
- e) Política Corporativa de Cadena de Abastecimientos (*Supply Chain*)
- f) Política Corporativa de Clasificación y tratamiento de la Información
- g) Política Corporativa de Evaluación y Gestión de Riesgos de Seguridad de Información y Ciberseguridad
- h) Política Corporativa de Organización de la Seguridad de la Información y Ciberseguridad
- i) Política Corporativa de Derechos Humanos (DDHH)
- j) Manual Corporativo de Gestión Integral de Riesgos
- k) Reglamento Interno de Trabajo de la Unidad de Negocio
- l) Legislación de IA aplicable a cada Unidad de Negocio

VII. REVISIONES Y APROBACIÓN

	POLITICA CORPORATIVA DE USO DE LA INTELIGENCIA ARTIFICIAL			VERSIÓN
Área	DIRECCION CORPORATIVA DE RIESGOS Y CUMPLIMIENTO			
Elaborado por	Fernando Dyer Director Corporativo de Riesgos y Cumplimiento Silvana Pérez Yalán Gerente Corporativo de Cumplimiento y Oficial Corporativo de Protección de Datos	Fecha de Elaboración	29.05.2026	1.0
Revisado por	Equipo Proyecto Salto	Fecha de Revisión	10.06.2026	1.0
Revisado por	Luis Budge Gerente Corporativo de Ciberseguridad y Corporate Chief Data Officer	Fecha de Revisión	10.06.2026	1.0
Revisado por	Juan José Dorich Gerente Corporativo de Gestión Integral de Riesgos	Fecha de Revisión	10.06.2026	1.0
Revisado por	José Luis Perry Gerente Corporativo Legal	Fecha de Revisión	10.06.2026	1.0
Revisado por	Álvaro Morales VP Corporativo de Finanzas	Fecha de Revisión	10.06.2026	1.0
Revisado por	Marlene Negreiros VP Corporativo de Talento y Cultura	Fecha de Revisión	10.06.2026	1.0
Revisado por	Eduardo Sánchez VP Corporativo Industrial	Fecha de Revisión	12.06.2026	1.0
Revisado por	Pedro Lerner Gerente General Corporativo	Fecha de Revisión	12.06.2026	1.0
Aprobado por	Comité de Riesgos y Cumplimiento	Fecha de Aprobación	16.06.2026	1.0
Aprobado por	Directorio	Fecha de Aprobación	01.07.2026	1.0

	RIESGOS Y CUMPLIMIENTO	Versión 1.0
	POLÍTICA CORPORATIVA SOBRE INTELIGENCIA ARTIFICIAL	2026-07-01
		Página 7 de 12

ANEXO 1 - DEFINICIONES

1. **Inteligencia Artificial (IA):** Es la disciplina orientada a la investigación y desarrollo de mecanismos y aplicaciones de sistemas y programas capaces de simular capacidades cognitivas humanas;
2. **Sistema basado en IA:** Es un sistema basado en una máquina que, de forma explícita o implícita, infiere, a partir de la entrada que recibe, cómo generar salidas tales como predicciones, contenido, recomendaciones o decisiones que puedan influir entornos físicos o virtuales;
3. **Datos Personales:** Toda información sobre una persona natural que la identifica o la hace identificable a través de medios que pueden ser razonablemente utilizados;
4. **Tratamiento de Datos Personales:** Cualquier operación o procedimiento técnico, automatizado o no, que permite la recopilación, registro, organización, almacenamiento, conservación, elaboración, modificación, extracción, consulta, utilización, bloqueo, supresión, comunicación por transferencia o por difusión o cualquier otra forma de procesamiento que facilite el acceso, correlación o interconexión de los Datos Personales;
5. **Titular de Datos Personales:** Persona natural a quien corresponden los Datos Personales;
6. **Normativa de IA:** Conjunto de normativas vigentes y requisitos obligatorios en el ámbito de la IA;
7. **Función Corporativa responsable de la IA:** Órgano colegiado encargado de dirigir y supervisar las acciones del Grupo UNACEM en el ámbito de la IA. En el momento de emitir esta política corporativa, este rol lo asume el equipo ejecutivo del Proyecto Salto llamado “Equipo Salto” (referencia al Sr. Jaime Bedoya – Gerente de Innovación y a la Sra. Elena Mujica - Gerente Corporativo de Transformación Cultural, Efectividad y Cambio);
8. **CISO (Chief Information Security Officer) Corporativo:** responsable corporativo de la estrategia de seguridad de la información y la supervisión de su implementación en el Grupo UNACEM. Al momento de emitir esta política corporativa, este rol lo tiene el Sr. Luis Budge – Corporate Chief Information Security Officer;
9. **CDPO (Chief Data Protection Officer –** responsable Corporativo de Protección de Datos Personales de la estrategia de protección del tratamiento de datos personales en el Grupo UNACEM, de acuerdo con las funciones establecidas en la legislación sobre Protección de Datos Personales. Al momento de emitir esta política corporativa, este rol lo tiene la Sra. Silvana Pérez – Gerente Corporativo de Cumplimiento y Corporate Chief Data Protection Officer;
10. **CIO (Chief Information Officer) Corporativo:** responsable corporativo encargado de la estrategia de tecnología de la información y la supervisión de su implementación a en el Grupo UNACEM. Al momento de emitir esta política corporativa, este rol lo tiene el Sr. Javier Chang – Corporate Chief Information Officer.
11. **Alucinación (Hallucination):** Generación por parte de un sistema de IA de información falsa, inexacta o fabricada presentada como factual, sin base en los datos de entrenamiento o en la realidad.
12. **Inyección de Prompts (Prompt Injection):** Técnica de manipulación de las entradas de un sistema de IA generativa con el fin de evadir sus controles de seguridad o alterar su comportamiento previsto.

	RIESGOS Y CUMPLIMIENTO	Versión 1.0
	POLÍTICA CORPORATIVA SOBRE INTELIGENCIA ARTIFICIAL	2026-07-01
		Página 8 de 12

13. **Deepfake:** Contenido audiovisual (imagen, video o audio) generado o alterado mediante IA para simular de manera realista la apariencia o voz de personas reales.
14. **Drift de Modelo:** Degradación progresiva del desempeño de un modelo de IA debido a cambios en los datos, el entorno operativo o el comportamiento de los usuarios a lo largo del tiempo.
15. **Red-Teaming:** Metodología de prueba adversarial en la que un equipo especializado intenta deliberadamente provocar fallos, vulnerabilidades o comportamientos no deseados en un sistema de IA, con el objetivo de identificar riesgos antes de su despliegue.
16. **Guardrails:** Controles técnicos automatizados implementados en sistemas de IA para restringir sus outputs dentro de límites predefinidos de seguridad, ética y calidad.
17. **Apetito de Riesgo de IA:** Nivel y tipo de riesgo relacionado con la inteligencia artificial que el Grupo UNACEM está dispuesto a aceptar en la búsqueda de sus objetivos estratégicos de innovación y eficiencia operativa.
18. **Datos Sintéticos:** Datos generados artificialmente mediante algoritmos que replican las características estadísticas de datos reales, utilizados para entrenamiento o prueba de modelos de IA.

	RIESGOS Y CUMPLIMIENTO	Versión 1.0
	POLÍTICA CORPORATIVA SOBRE INTELIGENCIA ARTIFICIAL	2026-07-01
		Página 9 de 12

ANEXO 2 – FUNCIONES DE LA FUNCIÓN CORPORATIVA RESPONSABLE DEL GOBIERNO DE LA IA

- Garantizar que las iniciativas, soluciones o proyectos que incorporen tecnología de IA, sean desarrollados internamente o adquiridos, cumplan con los principios y reglas establecidas en la presente Política.
- Supervisar el desarrollo de todas las soluciones de IA del Grupo UNACEM, y realizar un seguimiento de su ciclo de vida completo, con el fin de velar por: (i) el cumplimiento de la normativa que resulte de aplicación y de los requerimientos establecidos a estos efectos por parte de las áreas corporativas del Grupo UNACEM; (ii) el respeto por parte de dichas soluciones de los valores éticos y morales; y, (iii) la robustez de las soluciones de IA tanto desde el punto de vista técnico como social.
- La supervisión que ejerce de esta Función Corporativa tiene tres niveles:

a) NIVEL Básico - Exploración y uso libre

Comprende prototipos internos, pruebas de concepto, herramientas de productividad personal o de equipo, y cualquier uso de IA que no procese datos sensibles ni genere decisiones con efectos sobre terceros.

Régimen aplicable:

- Se debe informar al champion del área/equipo, pero no se requiere autorización previa. El owner de la solución IA opera con plena autonomía.
- Dentro de los treinta (30) días naturales desde el inicio, el owner completa el Formulario de Registro del repositorio corporativo de iniciativas de IA.
- El Equipo SALTO puede revisar en cualquier momento las iniciativas notificadas y recategorizar motivadamente una iniciativa a un nivel superior si detecta factores de riesgo no considerados.

b) NIVEL Operacional - Implementación supervisada

Comprende soluciones de IA que procesan datos de clientes, empleados o terceros; que apoyan decisiones de negocio con impacto real en personas; o que se integran en procesos operativos de la organización.

Régimen aplicable:

- El owner de la solución IA registra la iniciativa en el repositorio corporativo antes de su desarrollo o adquisición, mediante el Formulario de Registro.
- La revisión técnica y ética es realizada por el AI Champion del área correspondiente antes de iniciar, figura designada y certificada por la Función. No requiere intervención directa del Equipo SALTO salvo solicitud expresa.
- El equipo SALTO mantiene acceso en lectura al registro y puede solicitar revisiones motivadas en cualquier momento durante el ciclo de vida de la solución.

c) NIVEL Ultra - Aprobación centralizada

Comprende soluciones de IA que producen decisiones automatizadas con efectos jurídicos o impactos significativos sobre derechos fundamentales; sistemas de IA incorporados en productos o servicios externos; y cualquier uso de IA de alto impacto reputacional, regulatorio o social para el Grupo.

	RIESGOS Y CUMPLIMIENTO	Versión 1.0
	POLÍTICA CORPORATIVA SOBRE INTELIGENCIA ARTIFICIAL	2026-07-01
		Página 10 de 12

Régimen aplicable:

- i. El desarrollo, adquisición o implementación de toda iniciativa de nivel ultra requiere aprobación previa y expresa del equipo SALTO.
 - ii. El equipo SALTO monitorea activamente el ciclo de vida completo de la solución, incluyendo su desarrollo, implementación, operación y eventual retirada.
 - iii. Se realizarán revisiones de impacto periódicas con una frecuencia mínima anual, o cada vez que se produzca un cambio material en la solución o en el contexto de su uso.
 - iv. Las iniciativas de nivel ultra están sujetas a auditoría interna o externa según se determine, y sus resultados se reportan a la dirección del Grupo.
 - v. Las decisiones automatizadas producidas por sistemas de nivel ultra que tengan efectos jurídicos o impactos significativos sobre personas no podrán basarse exclusivamente en el sistema de IA. Deberá existir siempre un mecanismo de supervisión humana efectiva.
4. Monitorear el registro de las iniciativas de IA que se llevan a cabo.
 5. Asesorar a las distintas áreas/negocios en todas las cuestiones legales y éticas que afectan a proyectos de IA.
 6. Evaluar durante el desarrollo, implementación y uso de los sistemas de IA, promoviendo el cumplimiento de lo siguiente:
 - a) Privacidad desde el diseño y por defecto: respeto de la privacidad, mediante la aplicación de criterios de privacidad desde el diseño y por defecto en el desarrollo, adquisición, implementación y uso de sistemas de IA, en cumplimiento de la normativa aplicable.
 - b) Solidez Técnica y Seguridad: las medidas de seguridad aplicadas en los aspectos técnicos y de seguridad asociados a la iniciativa de IA, cumplen con los estándares de seguridad requeridos, minimizando riesgos en el uso de IA.
 - c) Transparencia: información clara y precisa para los usuarios cuando sus datos sean tratados mediante sistemas de IA, facilitando la comprensión integral y la trazabilidad de los procesos, decisiones y algoritmos asociados a IA. Asimismo, se promueve la sensibilización y la capacitación digital y en materia de IA.
 - d) Supervisión Humana y Derechos Fundamentales: respeto de los derechos fundamentales en el uso de la IA.
 - e) Diversidad, no discriminación y equidad: inclusión, diversidad e igualdad de trato entre las personas durante todo su ciclo de vida.
 - f) Bienestar social y ambiental: integración de los valores de bienestar social, eficiencia y desarrollo sostenible.
 - g) Responsabilidad: controles y mecanismos adecuados para garantizar la responsabilidad y rendición de cuentas en el desarrollo y uso de la IA.
 - h) Gobierno adecuado del dato: cumplimiento de los principios de minimización de datos, calidad y exactitud.

	RIESGOS Y CUMPLIMIENTO	Versión 1.0
	POLÍTICA CORPORATIVA SOBRE INTELIGENCIA ARTIFICIAL	2026-07-01
		Página 11 de 12

- i) Derecho de autor y derechos conexos: respeto de los derechos de autor, normas sobre propiedad intelectual y normativas vigentes.
 - j) Decisiones automatizadas y perfilamiento: las decisiones que produzcan efectos jurídicos o impactos significativos sobre los usuarios no estarán basadas exclusivamente en sistemas de IA.
- 7. Comunicar información a las áreas implicadas, de un modo claro y oportuno, sobre las funcionalidades, finalidad y limitaciones de los sistemas de IA. Ser transparentes acerca del hecho de que se está trabajando con un sistema de IA.
- 8. Facilitar la trazabilidad y la auditabilidad de los sistemas de IA, especialmente en contextos o situaciones críticas.
- 9. Invitar a que participen a las áreas correspondientes en todo el ciclo de vida de los sistemas de IA.
- 10. Promover la formación y la educación, de manera que todas las áreas implicadas sean conocedoras de la IA Responsable y reciban formación en la materia.
- 11. Para los sistemas de IA de nivel Ultra, la Función Corporativa responsable de la IA (o a quien delegue), previo a su despliegue en producción, supervisará la realización de las siguientes actividades de validación:
 - a) Red-teaming: pruebas adversariales conducidas por un equipo independiente para identificar vulnerabilidades, modos de fallo y comportamientos no deseados del sistema de IA.
 - b) Pruebas de seguridad (safety testing): evaluación sistemática de comportamientos riesgosos, incluyendo generación de contenido dañino, fugas de datos y evasión de controles.
 - c) Implementación de guardrails técnicos: controles automatizados que restrinjan los outputs del sistema de IA dentro de los límites de seguridad, ética y calidad predefinidos.
 - d) Para todos los sistemas de IA en producción de niveles Operacional y Ultra, se implementarán adicionalmente:
 - e) Monitoreo continuo de drift: mecanismos de detección de la degradación progresiva del desempeño del modelo debido a cambios en los datos, el entorno operativo o el comportamiento de los usuarios.
 - f) Definición de límites operativos (boundary definition): establecimiento formal y documentado de los parámetros dentro de los cuales cada sistema de IA debe operar, incluyendo tipos de datos admisibles, rangos de confianza aceptables y restricciones de uso.

Los resultados de todas las actividades de validación se documentarán y formarán parte del expediente técnico de cada sistema de IA.

	RIESGOS Y CUMPLIMIENTO	Versión 1.0
	POLÍTICA CORPORATIVA SOBRE INTELIGENCIA ARTIFICIAL	2026-07-01
		Página 12 de 12

ANEXO 3 – LISTA DE HERRAMIENTAS Y APLICACIONES DE IA AUTORIZADAS

Solamente las herramientas y aplicaciones de IA instaladas por el equipo e TI del Grupo UNACEM están autorizadas para su uso.

La siguiente lista corresponde a aquellas que podrán ser solicitadas para su uso. Para ello, deberá enviarse una solicitud de instalación a la mesa de ayuda de TI.

I. Herramientas de IA de Propósito General:

La siguiente lista corresponde a aquellas que podrán ser solicitadas para su uso. Para ello, deberá enviarse una solicitud de instalación a la mesa de ayuda de TI.

1. (Microsoft) Copilot
2. (Anthropic) Claude

II. Funcionalidades de IA embebidas en sistemas corporativos

La siguiente lista corresponde a aquellas que podrán ser solicitadas para su uso que deberá ser sustentado en el Business Case de la solución basada en IA:

3. SAP Joule u otras capacidades de IA habilitadas en SAP S/4HANA
4. Salesforce Einstein u otras capacidades de IA de Salesforce
5. Monday sidekick u otras capacidades de IA en Monday.com

III. Plataformas de acceso a modelos fundacionales

La siguiente lista corresponde a aquellas que podrán ser solicitadas para su uso que deberá ser sustentado en el Business Case de la solución basada en IA:

1. AWS bedrock

Nota: Esta lista será actualizada periódicamente incorporando otras herramientas, funcionalidades y/o plataformas disponibles a ser solicitadas para su uso.