

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	Page 1-11

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.

- **AI First:** An organizational strategy that places artificial intelligence at the core of the design, development, and operation of products, services, and processes. Under this approach, AI is not considered a complementary tool, but rather the main driver of digital transformation, operational efficiency, and competitive differentiation. AI First organizations prioritize the use of AI solutions and autonomous systems to automate decisions, scale capabilities, and improve the user experience, ensuring that each initiative is aligned with institutional strategic objectives. This strategy requires robust governance that ensures traceability, human oversight, explainability, security, and regulatory compliance in accordance with applicable regulations, while also promoting an organizational culture of responsible and ethical AI adoption.
- **AI Stewards (business owners):** Roles designated within each business unit, at the head, management, or equivalent level, responsible for ensuring effective artificial intelligence governance from the business perspective. They are responsible for direct human oversight (human-in-the-loop) and for ensuring the integrity, security, and regulatory compliance of the AI systems within their scope of responsibility. The designated AI Steward is solely responsible for defining, periodically evaluating, and documenting compliance with the clear and measurable objectives of each AI system promoted by their business unit, as well as

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	Page 1-11

demonstrating its impact on organizational processes and results. The AI Steward acts as the business owner of the AI system and bears ultimate responsibility for its use, purpose, value generated, risks assumed, and compliance, regardless of any technical execution delegated to other areas. This responsibility may not be delegated and must be exercised in accordance with institutional standards and goals.

- **Digital literacy:** The process through which the company's employees and users acquire and develop essential digital skills to interact with digital systems and services, including the safe and responsible use of artificial intelligence in the workplace and business environment.
- **AI Growth Area:** The area responsible for the organizational governance of artificial intelligence within the company, as established in Section IV (Responsibilities) of this policy.
- **Integrity Channel:** A mechanism for reporting conduct that may be illegal, unethical, or in violation of professional standards; that is, conduct inconsistent with RIMAC's Code of Conduct, as well as risks of bias, misuse, or unintended impacts arising from the use of AI systems.
- **AI use case:** A specific application or function for which an AI system is used (for example, a customer service chatbot, AI assistance for decision-making in benefit determination, predictive maintenance analytics, etc.).
- **AI life cycle:** All stages in the existence of an AI system, including design, data collection, development, testing, implementation, operation, monitoring, maintenance, continuous improvement, and decommissioning. It also includes reuse, significant modification, and the management of legacy systems or non-formalized developments (Shadow AI), across all project stages, from internal developments, pilots, and proofs of concept to the operation and continuous improvement of systems deployed in production.
- **AI Risk or Regulatory Classification:** The designation of an AI system's risk level pursuant to the Emerging Technologies and GenAI Risk Management Methodology (MAN-6038), considering the potential impact its use may have on fundamental rights, safety, access to critical services, fairness, privacy, or people's well-being. It is based on criteria established by Law No. 31814, its regulations, and international standards such as the European Union Artificial Intelligence Act. Responsibility for this classification lies with the Technology Risk Area.
- **AI CoE:** The Center of Excellence responsible for the governance and technical execution of artificial intelligence systems, as established in Section IV (Responsibilities) of this policy.
- **AI Committee:** A formal governance body responsible for overseeing and ensuring institutional accountability in the use of artificial intelligence systems, as established in Section IV (Responsibilities) of this policy.
- **AI Impact Assessment:** A systematic process for identifying, analyzing, and mitigating the ethical, technical, legal, and social risks associated with the implementation of an AI system, particularly when it is classified as high-risk or involves the processing of personal data.
- **AI Governance:** The set of policies, procedures, roles, and controls intended to ensure that AI is developed, used, and supervised responsibly, in accordance with laws, ethical principles, and institutional objectives.

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	Page 1-11

- **Generative AI:** Artificial intelligence technology that uses advanced models - such as deep neural networks and language models - to generate original content (text, images, audio, video, code, etc.) from training data. Generative AI may operate autonomously or with assistance and poses specific challenges relating to ethics, transparency, explainability, security, governance, and intellectual property. Its implementation requires human oversight mechanisms, impact assessments, data traceability, and regulatory compliance.
- **Artificial Intelligence (AI):** A set of computer systems that, through techniques such as machine learning, logical reasoning, natural language processing, computer vision, and other related disciplines, are capable of performing tasks that traditionally require human intelligence. These tasks include data analysis, decision-making, interaction with people, and content generation. AI may operate autonomously or with assistance, and its design must consider ethical, security, transparency, and accountability principles.
- **AI Inventory:** An up-to-date list of all AI systems in use or under development within the organization, including key details such as purpose, input data, provider or developer, implementation status, risk classification, and assigned responsible parties.
- **Proactive Accountability Principle:** The institutional obligation to adopt technical, organizational, and ethical measures before developing or using an AI system in order to prevent adverse impacts on security, business continuity, privacy, or people's rights.
- **AI Risk:** The possibility that the design, development, deployment, or operation of an artificial intelligence system may generate unintended negative impacts on people, processes, assets, the organization's reputation, or regulatory compliance as a result of automated or assisted decisions, levels of autonomy, data use, system learning, or failures in control and human oversight mechanisms.

These risks are managed in accordance with the corporate Risk Management and Ethics and Compliance frameworks, applying controls proportionate to the identified level of impact and risk.

- **Shadow AI:** The use, development, acquisition, or implementation of Artificial Intelligence systems, models, tools, or functionalities outside the formal frameworks defined by the organization, without having been identified, assessed, or authorized in accordance with this Policy and the applicable institutional processes.

Shadow AI constitutes an ethical, operational, and compliance risk regardless of its scale, purpose, or degree of autonomy and must be identified, reported, and managed in accordance with the mechanisms established by the organization.

- **High-Risk System:** An AI system whose operation may have a significant impact on fundamental rights, personal safety, access to essential services, or the stability of critical processes. It requires enhanced governance, human oversight, explainability, impact assessment, and redress mechanisms.
- **Artificial Intelligence System:** A set of applications or technological solutions designed to operate with varying degrees of autonomy and using artificial intelligence techniques - such as machine learning, natural language processing, computer vision, logical reasoning, among others - to analyze and learn from data and generate outputs, recommendations, or decisions. These systems may directly influence organizational, operational, or strategic processes and must be subject to principles of transparency, human oversight, explainability, security, ethics, and governance.

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	Page 1-11

- **Human Oversight (human-in-the-loop, HITL):** Human intervention or the ability of individuals to intervene in automated AI processes, ensuring that relevant decisions remain under human control and accountability.
- **Algorithmic Transparency:** The ability of an AI system to explain, in an understandable manner, how its models operate, the criteria used for decision-making, and the results generated, enabling internal and external audits.

Clarifying note:

Whenever this policy refers to the Ethical Reporting Channel, such reports shall be made through the organization's current Integrity Channel, in accordance with established institutional procedures.

4. ROLES AND RESPONSIBILITIES

For the effective implementation of this policy, clear responsibilities are defined among the organization's various stakeholders:

4.1. ROLE: Business Owner (AI Steward)

- Ensure that the AI systems within their division are aligned with strategic objectives, risk appetite, and applicable regulatory frameworks.
- Allocate the resources and technical personnel required for the operation, oversight, and continuous improvement of their AI systems.
- Oversee the performance and results of AI systems, intervening in the event of errors, biases, or deviations.
- Coordinate with the AI Committee, AI Growth, and other governance areas to ensure compliance with this policy.
- Be accountable for the impact, traceability, and responsible adoption of the AI systems under their responsibility.

4.2. ROLE: AI Supervisor (AI Steward delegate)

- Lead the life cycle of the designated AI systems without replacing the AI Steward's ultimate responsibility.
- Coordinate the implementation of AI initiatives, ensuring compliance with schedules, budgets, and deliverables.
- Manage feedback, appeal, and transparent communication mechanisms with customers and end users.
- Document lessons learned and promote the continuous improvement of the systems under their responsibility.
- Promptly escalate to the AI Steward any risks, incidents, and decisions that exceed their scope.

4.3. AI Committee

- Approve high-impact or high-risk AI use cases.
- Oversee the management of incidents, appeals, and redress mechanisms.
- Validate the alignment of AI systems with the organization's ethical, legal, and strategic principles.
- Act as the institutional point of contact for auditors, regulators, and stakeholders.
- Make decisions regarding cross-cutting impacts on data, processes, technical capabilities, and organizational change associated with AI adoption.

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	Page 1-11

4.4. Area: AI Growth

- Define the principles, policies, standards, and acceptance criteria for AI initiatives.
- Oversee the life cycle of AI systems, including Shadow AI and legacy systems.
- Maintain the corporate inventory and functional catalog of AI solutions.
- Orchestrate the use-case portfolio, assessing value versus risk and escalating high-impact cases to the AI Committee.
- Designate and coordinate with AI Stewards and require performance, trustworthiness, and risk indicators.
- Promote the responsible adoption of AI under the AI First approach and organizational literacy.
- Coordinate independent reviews of compliance with this policy with Internal Audit.

4.5. Area: AI Center of Excellence (AI CoE)

- Maintain the mandatory technical framework for the design, validation, and operation of AI systems.
- Exercise technical governance over the AI life cycle and provide technical evidence for audits and regulatory reviews.
- Implement security, resilience, observability, and technical incident response controls.
- Coordinate the interoperability and traceability of AI systems with IT, Architecture, Data, DevOps, Security, and Risk.
- Technically remediate identified Shadow AI developments.

4.6. Technical Areas, AI CoE, and Business Tech-Savvy Personnel:

- Apply this policy and the technical standards throughout the life cycle of the AI systems under their management.
- Conduct the applicable technical, ethical, legal, and operational risk assessments.
- Ensure mandatory alignment with the enterprise architecture, avoiding duplicate capabilities, unapproved technologies, and technical debt.
- Detect and report Shadow AI cases for formalization and assessment.
- Coordinate with the AI CoE and AI Stewards to integrate their initiatives into the established governance frameworks.

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	Page 1-11

4.7. Area: Data

- Ensure the availability, quality, traceability, and protection of the data used by AI systems.
- Classify data according to its sensitivity and criticality in accordance with internal policies and applicable regulations.
- Document the origin, use, and life cycle of data used in AI initiatives.
- Detect and report Shadow AI cases involving the use of data outside formal controls.
- Participate in audits and risk assessments associated with the use of data in AI.

4.8. Area: Processes

- Identify opportunities for automation and process improvement through AI solutions.
- Validate the alignment of AI systems with critical processes and the customer experience.
- Participate in the definition and validation of AI use cases from a process perspective.
- Document the processes involved in AI initiatives to ensure traceability.
- Monitor performance and risk indicators in AI-automated processes.

4.9. Area: Talent Management

- Design and implement mandatory AI training and competency development programs.
- Promote an AI First culture and training on ethics, bias, and good practices.
- Incorporate AI responsibilities into job profiles, recruitment processes, and performance evaluations.
- Lead the agentic organization strategy that coordinates collaboration between people and AI solutions, as well as performance and risk indicators in AI-automated processes.

4.10. Area: Central Architecture

- Define and oversee the architecture patterns applicable to AI systems.
- Ensure the interoperability, scalability, and alignment of AI systems with the enterprise architecture.
- Validate the integration of AI solutions into organizational processes and platforms.
- Facilitate technical audits of AI system architecture.

4.11. Information Security CoE

- Define, implement, and monitor security controls applicable to AI data and systems, in accordance with RIMAC policies, applicable regulations, and international standards.
- Identify, assess, and address information security risks associated with the use of AI, jointly with the Business Owner.
- Incorporate AI considerations into vulnerability testing, incident management, and responses to emerging threats.
- Ensure the continuous improvement of the security and resilience of AI systems.

4.12. Area: Ethics and Compliance

- Act as the second line of defense, providing independent advice on ethical, integrity, and conduct risks in AI initiatives.

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	Page 1-11

- Assess the impact of AI initiatives against applicable regulatory programs (Ethics, Crime Prevention, Personal Data Protection, among others).
- Conduct independent reviews of AI initiatives involving relevant ethical risks and escalate them to the AI Committee or other governance bodies.
- Manage ethical alerts related to the use of AI, including Shadow AI, through the Integrity Channel.
- Support responses to authorities and audits related to the ethical use of AI.

4.13. Legal Area

- Interpret the regulations applicable to the use of AI systems and issue the corresponding legal opinions.
- Advise on the assessment of legal risks and regulatory obligations associated with AI initiatives.
- Manage relations with regulatory authorities in coordination with the responsible areas.

4.14. Business Units and AI Initiative Owners

- Act as the first line of defense for AI initiatives under their management.
- Define the purpose and assume operational, ethical, and legal responsibility for the AI systems under their management.
- Identify, manage, and report the risks associated with their AI initiatives.
- Promptly request advice from Ethics and Compliance and the Legal Area, as applicable.
- Ensure compliance with this policy throughout the AI system's life cycle.

4.15. Operational Risk Area

- Define the methodology for assessing non-financial risks applicable to AI systems.
- Report to the Comprehensive Risk Management Committee (GIR) on AI systems classified as "high risk" under Law No. 31814 when they are considered a significant or prioritized change.

4.16. Technology Risk Area

- Define the technology and emerging risk assessment methodology applicable to AI systems, including its activation criteria.
- Oversee and support the business areas and AI Growth in identifying, assessing, and managing these risks.
- Ensure the implementation of appropriate mitigation controls.

4.17. Employees and Internal Users

- Comply with this policy in the day-to-day use of AI solutions.
- Participate in mandatory AI training.
- Report irregularities, incidents, or ethical or security concerns through the Integrity Channel.
- Acknowledge that lack of knowledge of this policy does not exempt anyone from compliance.

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	Page 1-11

4.18. Third Parties and AI Providers

- Contractually align their deliverables and conduct with this policy and the principles of responsible AI use.
- Undergo the due diligence process before contracting.
- Comply with contractual obligations regarding data protection, auditing, incident reporting, and redress mechanisms for affected end users.

4.19. End Users and Affected Citizens

- Provide transparency regarding the use of AI in services that affect them.
- Provide understandable explanations of automated decisions.
- Enable appeal, correction, and redress mechanisms for adverse impacts.
- Protect their personal data and fundamental rights.

Each person or team involved must understand that their role is crucial to ensuring that AI delivers value to the organization without compromising ethics, security, or regulatory compliance. The responsibilities defined herein shall be incorporated into job descriptions, performance evaluations, or other applicable internal documents to ensure clarity of expectations and organizational traceability.

5. POLICY PROVISIONS

5.1 Guiding Principles

The use and development of Artificial Intelligence within the organization shall be governed by the following guiding principles, in accordance with Law No. 31814 and Supreme Decree No. 115-2025-PCM:

- Legality and respect for fundamental rights:** All AI-related activities must comply with applicable laws and regulations (principle of legality), while respecting human rights, dignity, privacy, and fundamental freedoms. AI must not infringe rights or freedoms; its use shall be framed by the values and rights protected by the Constitution and applicable regulations. The use of AI for unlawful purposes is prohibited, including mass surveillance without a legal basis, subliminal manipulation, social scoring, discrimination, or any practice prohibited by national or international regulations. Regulatory compliance prior to the deployment of AI systems shall be validated in accordance with the applicable institutional and legal frameworks.
- Ethical development and human accountability:** Ethics is the fundamental basis for the responsible use of artificial intelligence. The design, development, and use of AI systems must ensure that ultimate accountability always rests with people, avoiding the exclusive delegation of decisions to automated systems. No AI system exempts the organization or its members from acting with integrity or from assuming responsibility for the results obtained. The use of AI must align with the institution's risk appetite and corporate values, ensuring effective human oversight and appropriate governance mechanisms to prevent, detect, and manage unintended impacts. The organization shall also maintain an ethical reporting mechanism integrated into the current Integrity Channel to channel incidents, risks, misuse, or ethical concerns related to AI systems, in accordance with applicable institutional models and guidelines. In complex or high-impact cases, ethical deliberation forums shall be promoted.

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	Page 1-11

- c) **Transparency and explainability:** The organization shall promote the use of artificial intelligence in a transparent and explainable manner, proportionate to each system's level of risk and impact. Users and stakeholders must have clear and accessible information regarding the purpose, scope, limitations, and relevant risks associated with the use of AI. Decisions, outputs, or recommendations generated by AI systems must be understandable and explainable in light of the context and risk, ensuring the traceability of automated decisions and the availability of information for audits, requests for explanation, or appeal mechanisms, in accordance with applicable institutional frameworks and guidelines. Where warranted by the level of impact, additional transparency measures shall be adopted, including applicable communication and labeling practices.
- d) **Accountability and human oversight:** The organization shall ensure accountability for all decisions arising from the use of AI by assigning clear roles and responsibilities to identifiable individuals and governance bodies. The management of incidents, exceptions, and reviews shall be documented in a traceable manner. All AI systems, particularly high-impact systems, must have effective human oversight that enables automated decisions to be monitored, interrupted, or overridden and critical cases to be escalated. Auditable records shall be maintained, and internal and external audits shall be facilitated. Under no circumstances may professional, ethical, or legal responsibility be delegated exclusively to AI. The organization must ensure the capacity and authority to intervene at any point in the system's life cycle.
- e) **Non-discrimination and inclusion:** The organization shall ensure the prevention, detection, and mitigation of bias and discriminatory outcomes at every stage of the AI system life cycle, actively promoting fairness, inclusion, and accessibility for the full diversity of users and stakeholders. AI systems must be designed and operated so that they neither cause discrimination nor perpetuate inequalities, ensuring equitable access to their benefits.

When a material bias is identified that cannot reasonably be mitigated, the initiative must be adjusted, escalated to the appropriate governance bodies, or discontinued in accordance with the principles of responsible AI use, applicable regulations, and relevant international ethical standards.

- f) **Privacy and personal data protection:** The organization shall ensure that the processing of personal data in AI systems complies with applicable regulations, including Law No. 29733 and relevant international standards. AI systems must be designed and operated under a privacy-by-design and privacy-by-default approach, ensuring that privacy is considered throughout every stage of the life cycle. The organization shall maintain traceability of data processing operations and adopt risk-proportionate measures necessary to safeguard data subjects' rights in accordance with applicable institutional guidelines.
- g) **Security and reliability:** AI systems must be robust, secure, and reliable throughout their life cycle and operate in accordance with their intended purpose even under adverse conditions or foreseeable misuse. The organization shall adopt a comprehensive information security and risk management approach for AI systems, aligned with applicable regulatory frameworks and international standards. Risk-

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	Page 1-11

proportionate measures shall be adopted to ensure system resilience, integrity, and reliability in accordance with current institutional guidelines.

- h) **Sustainability:** The organization shall systematically assess the social, environmental, and economic effects of using AI, promoting the contribution of these technologies to sustainable development and social well-being. Alignment with the Sustainable Development Goals (SDGs) shall be prioritized by optimizing resource consumption, minimizing the environmental footprint, and preventing adverse impacts on vulnerable communities. Responsible innovation, intergenerational equity, and transparency in the measurement and reporting of environmental and social impacts shall be promoted by integrating these criteria into decision-making and continuous improvement.
- i) **Innovation and continuous improvement:** The organization shall foster responsible AI innovation by promoting controlled experimentation and the adoption of new technologies under ethical and security principles. Regular cycles for reviewing and updating policies, models, and controls shall be established to ensure that AI evolves safely, efficiently, and in alignment with the organization's strategic objectives and values.
- j) **AI First approach:** The organization adopts an AI First strategy in which artificial intelligence is regarded as a central pillar for the design, development, and operation of products, services, and processes. This approach involves prioritizing intelligent agents and autonomous systems to scale capabilities, automate decisions, and generate strategic value, while always ensuring alignment with institutional objectives, ethical principles, and regulatory compliance. Implementing this strategy requires robust governance, traceability, effective human oversight, and an organizational culture focused on the responsible use of AI.
- k) **Organizational governance and redress mechanisms:** The organization shall maintain a clear and coherent AI governance structure, with defined roles and responsibilities for the oversight, review, and decision-making of AI systems, ensuring alignment with ethical principles, risk management, and regulatory compliance. The organization shall also ensure the availability of accessible and effective redress mechanisms through which users and stakeholders may submit complaints, request explanations, appeal automated decisions, or request corrections where appropriate, in coordination with the current Integrity Channel and in accordance with applicable institutional models.

5.2 AI Governance and Risk Management

AI governance is structured through the Capabilities function - AI Growth, Data, AI (AI CoE), and Processes - under the Capabilities Vice Presidency: AI Growth leads organizational and strategic governance; the AI CoE leads technical governance and execution; and Data and Processes ensure, respectively, responsible data management and alignment with processes. The corporate Technology, Information Security, Risk, Legal, and Ethics and Compliance areas participate within their respective remits as validation and control functions, without replacing the governance model defined by Capabilities.

The organization shall establish robust and auditable governance for the use of Artificial Intelligence, ensuring clearly defined roles, structures, and responsibilities to oversee compliance with this policy, not only from a non-financial risk perspective but also from an

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	Page 1-11

ethical, human accountability, and people and rights protection perspective. Senior management shall support and set the example for the ethical, safe, and lawful use of AI by assigning authorities responsible for directing AI system initiatives and managing their risks.

In line with the NIST AI Risk Management Framework (AI RMF) and the principles of ISO/IEC 27005, a structured, continuous, and adaptive approach shall be adopted to identify, assess, mitigate, and monitor risks associated with AI systems, including ethical, social, and people-impact risks that, by their nature, require analysis complementary to traditional non-financial risk assessment.

AI life-cycle operating model: The organization operates the end-to-end life cycle of AI systems through a formal operating model embodied in two complementary and synchronized documents: SGAI-MO-001 (AIMS Operating Model), a management document that defines the life-cycle phases, roles by phase, and the relationship with AIMS artifacts in accordance with NTP-ISO/IEC 42001:2025, Clause 8.1 and Control A.6; and PRO-7323 (Capabilities Operating Model), a corporate procedure within RIMAC's process system that details the specific activities, responsible parties, and deliverables for implementing AI solutions. Both documents must be read together and kept aligned.

a) **AI system mapping**

An up-to-date inventory shall be maintained of all AI-based systems used, developed, or acquired by the organization, including legacy systems and non-formalized developments (Shadow AI). For each system, the AI system functional document must include:

- Purpose, scope, and improved process
- Data used and update periods, where required
- Automated or assisted decisions
- Interaction with end users
- Providers or technical owners
- Risk classification and required level of human oversight
- Monitoring performed (traceability and observability).
- Plan to ensure business continuity in the event the solution is decommissioned

This mapping provides an understanding of the context in which AI is used and serves as the basis for effective risk assessment.

b) **Risk assessment and measurement**

A multidimensional risk assessment shall be conducted for each identified AI system, considering:

- Technical risks: algorithmic bias, overfitting, model adjustments, cybersecurity vulnerabilities, connection to systems with technical debt, access to uncontrolled information, adversarial attacks, and jailbreaks.
- Ethical and social risks: discrimination, manipulation, lack of explainability, and impacts on vulnerable users.
- Legal and regulatory risks: compliance with Law No. 31814, personal data protection, and users' rights.

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	Page 1-11

- Reputational and operational risks.

Performance and trustworthiness metrics (accuracy, error rate, fairness, robustness, and traceability) shall be defined, and the system shall be measured against these criteria. The risk level classification of AI systems shall be performed in accordance with the AI Risk Management Methodology (MAN-6038) of the Technology Risk Area, under the corporate non-financial risk management framework. AI Growth is responsible for identifying systems that require formal assessment and referring them to Technology Risk. For AI systems classified as high-risk or that directly affect customers, employees, or third parties, the assessment must explicitly consider ethical and fundamental rights impacts, supplementing the non-financial risk assessment in accordance with the current methodology.

c) Criteria for triggering a technology risk assessment of AI-based systems

AI systems that meet one or more of the following conditions shall be subject to a mandatory technology risk assessment by the Technology Risk Area in accordance with the AI Risk Management Methodology (MAN-6038):

- They are used in the company's critical processes, in accordance with the prioritized BPE list, and their results directly affect customer rights, benefits, or other relevant conditions, or business operations.
- They involve the processing of personal, sensitive, or critical data in accordance with the company's DP001 Data Privacy Guideline and Law No. 29733 - Personal Data Protection Law.
- They are identified as high-risk use cases in accordance with the regulatory criteria established in Law No. 31814 and its Regulations approved by Supreme Decree No. 115-2025-PCM.
- They are made available to or operated by customers, suppliers, business partners, or allies within the framework of RIMAC's processes or services, when such use may affect people's rights, personal data processing, regulatory compliance, or the organization's reputation.
- They could affect RIMAC's ethical principles and corporate values, including uses that undermine integrity, fairness, or stakeholder trust, or that create conflicts with the organization's Code of Ethics and Conduct.

When an AI-based system meets one or more of these conditions, the Technology Risk Area, in coordination with AI Growth and the areas involved, shall lead the corresponding risk assessment in accordance with the current methodology. The outcome of this assessment shall determine controls proportionate to the identified risk level and within the company's risk appetite.

d) Risk management and mitigation

Based on the foregoing assessment, controls proportionate to the risk level shall be implemented, including:

- Selection and/or adjustment of models and data to reduce bias

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	Page 1-11

- Strengthening technical security and data protection
- Mandatory human oversight (HITL) for critical decisions
- Operating limits, additional validations, and ethical review
- Legal, ethical, and privacy impact assessments (PIA)
- Appeal and redress mechanisms for affected users

Risk management shall follow a continuous cycle of identification, analysis, treatment, and monitoring, incorporating these requirements into internal procedures and third-party contracts.

e) Participation of Data and Processes in risk management

The Data Area must ensure the availability, integrity, representativeness, and protection of data used and generated by AI systems, including both enterprise solutions and AI tools embedded in work environments (for example, desktop assistants and copilots). It must also support the identification and mitigation of risks related to data quality, traceability, and privacy, considering the distributed use of AI by employees. The Processes Area must validate that the use of AI - including agents, automations, and AI tools embedded in the work environment - is aligned with the organization's operational and strategic processes, participating in the identification of operational risks, impacts on users' daily work, and the definition of appropriate controls.

f) Ongoing governance

An AI committee shall be established as a formal governance body, or this function shall be assigned to an existing committee (for example, innovation, risk, or ethics). This body shall coordinate with the Comprehensive Risk Management Committee (GIR) and operate within the risk appetites, methodologies, and decision-making authorities defined by it, without replacing its functions or redefining acceptable risk levels. This body shall also be responsible for:

- Periodically reviewing the performance of AI systems
- Verifying adherence to this policy and its guiding principles
- Resolving ethical dilemmas and validating new AI projects
- Overseeing incidents, appeals, and redress mechanisms
- Coordinating with AI Stewards, Ethics and Compliance, and other relevant areas as needed (Data, Processes, Technology)

The committee shall act as the point of contact for auditors, regulators, and internal and external stakeholders and shall be orchestrated by AI Growth.

Senior management must ensure that AI risk management processes are integrated into the organization's overall risk management system and aligned with international standards such as ISO/IEC 42001 and the NIST AI RMF. This ensures a unified, scalable, and effective approach to addressing both traditional risks and risks specific to artificial intelligence.

5.3 Information Security and Privacy in AI

Information and data used by AI systems must be protected in accordance with internal information security policies, applicable regulations, and standards such as SBS 504, and aligned with the requirements of ISO/IEC 27001 and the NIST Cybersecurity Framework.

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	Page 1-11

Technical and organizational controls shall be implemented to prevent unauthorized access, information leakage, model manipulation, and other threats that may compromise the confidentiality, integrity, availability, and traceability of AI data and systems. In particular:

- a) Algorithm training and input data shall be managed under strict access controls. Only authorized and trained personnel may handle sensitive data. Techniques such as masking, pseudonymization, or encryption in transit and at rest shall be applied, as appropriate, to safeguard sensitive data.
- b) The integrity of data used by AI must be ensured, as incorrect or maliciously altered data may lead to erroneous or risky decisions.
- c) AI systems shall undergo security testing and periodic assessments (for example, penetration testing for applications that interact online) to identify and remediate vulnerabilities. Any incorporated AI software component (libraries, frameworks, APIs) must undergo a multidisciplinary assessment before being used in production environments, with the Security Area responsible for reviewing whether the component creates any technical vulnerability or risk.
- d) If an AI system identifies or manages security incidents (for example, AI systems used in cybersecurity), such incidents must be recorded and managed in accordance with the organization's general incident management procedure, including traceability, root-cause analysis, and remediation.
- e) Controls shall be implemented to prevent attacks such as data poisoning, model manipulation, prompt injection, and jailbreaking. These risks shall be assessed during the design, testing, and operation stages.
- f) In coordination with the GenAI CoE, relevant AI system use cases shall be incorporated into Security Incident Response management.

With respect to privacy, the organization shall adopt a Privacy by Design approach for all AI projects involving personal data. This means:

- g) **Data minimization:** Only personal data strictly necessary for the legitimate purposes of the AI system shall be collected and processed. Unnecessary attributes (for example, sensitive data) shall not be included where the objective can be achieved using less intrusive information.
- h) **Privacy impact assessments:** Whenever an AI system that processes personal data on a large scale or processes sensitive data is developed or implemented, a Personal Data Protection Impact Assessment (also known as a PIA) shall be conducted. This shall help identify privacy risks to data subjects and define mitigation measures before implementation.
- i) **Consent and legal bases:** If an AI system requires personal data from customers, employees, or other individuals, the existence of a valid legal basis for such processing shall be verified (the data subject's express consent or another legal basis permitted by law). Where consent is required, it must be obtained in advance and be freely given, informed, and express, clearly indicating the AI-related uses that will be made of the data.

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	Page 1-11

The processing of this data must be performed in accordance with the framework established by Macroprocess M013 - Data and Analytics Governance, ensuring that the information used is reliable, cleansed, and properly managed by the designated responsible parties. In this context, the Data Area shall be responsible not only for the proper management of data but also for monitoring its use, ensuring compliance with data governance policies and the traceability of information involving AI.

- j) **Data subject rights:** ARCO rights (Access, Rectification, Cancellation, and Objection), or their equivalents, held by personal data subjects shall be respected at all times. If an AI system makes automated decisions that significantly affect a person (for example, denying a service), the affected person must be informed and provided with mechanisms to request human review of the decision in accordance with personal data protection regulations.
- k) **Retention and deletion:** Personal data used by AI systems shall not be retained beyond the period necessary. Specific retention policies shall be defined, and once the applicable period has expired or the purpose has been fulfilled, the data shall be securely deleted or anonymized.

Compliance with information security and privacy controls shall be regularly audited by the corresponding area in coordination with the AI Committee and AI Growth. The organization also encourages the adoption of recognized certifications or standards in this field (for example, ISO/IEC 27701 compliance for privacy information management) to strengthen trust in our AI practices.

For further details on data management, quality, and remediation, refer to the following policies:

- POL-4687 Third-Party Data Collection
- POL-6021 Data Governance and Quality
- POL-4688 Methodology for Addressing Mass Remediation Incidents
- Macroprocess M013 - Data and Analytics Governance

5.4 Transparency, Explainability, and Communication

Transparency is essential to building trust in the use of AI. The organization is committed to transparency both internally (with its employees, committees, and AI Stewards) and externally (with customers, users, regulators, and the general public) regarding its Artificial Intelligence systems. In practice, this means:

- a) **Disclosure of AI system use:** Clear and accessible notice shall be provided whenever a company product, service, or process is assisted by Artificial Intelligence. For example, if a customer interacts with a virtual assistant or automated chatbot, the customer must be informed that the interaction is with AI rather than a person. This disclosure shall be mandatory in high-impact interfaces, understood as points of interaction where the use of AI may have a material impact on the user's rights, decisions, outcomes, or expectations, and shall include, at a minimum:
 - Identification of the system as AI

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	Page 1-11

- Purpose of the system
 - User rights in relation to automated decisions
- b) **Information on purpose and scope:** Every AI system in production shall have an associated use policy or information notice describing its purpose, the tasks it performs, the scope of its capabilities, and its known limitations. For AI systems classified as high-risk under methodology MAN-6038, or where there is a direct impact on people's rights, this information shall be more detailed and shall include the types of data used, decision criteria, and human oversight mechanisms. Example: If AI is used to support credit decisions, users shall be informed that an algorithm is used to analyze information and what types of data it considers.
- c) **Explainability of decisions:** To the extent possible, meaningful explanations shall be provided as to how an AI system reached a conclusion or recommendation, particularly where the output directly affects a person (for example, rejection of an application or classification as a certain type of customer). Developers must design models by balancing performance and explainability. For complex models (black boxes), complementary tools or methods shall be implemented to extract explanations (for example, LIME or SHAP techniques, or executive summaries of the logic used).
- d) **Documentation and traceability:** Each AI solution shall have technical and functional documentation recording its operating logic, data sources, variables used, and known limitations. For complex systems, the versions of the algorithms, the training data used, and the results of performance and bias testing shall also be documented. Keeping this documentation up to date shall enable clearer explanations and facilitate audits and technical reviews.
- e) **Communication in non-technical language:** Information about AI intended for users or the general public shall be presented in plain language and avoid technical jargon. Frequently asked questions or other educational materials shall be published to help different audiences understand, in general terms, how AI operates in our services and what measures are taken to ensure its responsible use.
- f) **Respect for intellectual property and originality of content:** If AI is used to generate content (text, images, designs, code, etc.), intellectual property rules must be respected. AI models shall not be used to infringe third-party copyrights; rather, it must be verified that training data or generated material does not violate licenses. Any original creation generated by AI in the workplace shall belong to the organization to the extent permitted by law, and precautions shall be taken not to publish AI-generated content that could be mistaken for human-created content without disclosing its origin, in line with the principle of transparency.

Transparency and explainability strengthen stakeholder trust in the use of AI. The organization shall always favor open and communicative approaches, responding to concerns that may arise regarding how and why AI made certain decisions. Formal channels shall be enabled for:

- Requests for explanation
- Appeals of automated decisions
- Reporting errors, bias, or adverse impacts

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	Page 1-11

These channels shall be managed by the AI Committee and AI Growth, or the designated area, and shall be available to employees, customers, and end users.

5.5 Fairness, Non-Discrimination, and Ethics in AI

The organization reaffirms its commitment to non-discrimination and fairness in all its operations, expressly extending this commitment to Artificial Intelligence. To ensure that our AI systems operate fairly and ethically:

- a) **Bias-free design and data:** During the design and development stages of AI models, the datasets used for training and testing shall be carefully reviewed to identify and mitigate potential biases. Efforts shall be made to ensure that data is representative and diverse and to avoid training algorithms with information that reflects historical prejudice or unfair disparities. If bias is identified in the data (for example, underrepresentation of a population group), correction techniques shall be applied or data sources shall be supplemented. This review shall be documented and reviewed by the AI Committee and AI Growth.
- b) **Fairness validations:** In addition to traditional performance metrics, developers shall assess fairness metrics. For example, they shall measure whether a model produces significant differences in its predictions for different subgroups (gender, age ranges, geographic areas, etc.). If unjustified disparities are found, the model must be adjusted or even discarded. Any automated decision affecting people must be justified by data and criteria relevant to the case and never by protected or discriminatory attributes.
- c) **Human review of sensitive outcomes:** In sensitive areas (for example, AI-assisted recruitment processes, credit assessments, or automated medical diagnoses), a double-check shall be implemented: AI recommendations shall be reviewed by a person before the final decision is made, at least on a random basis or when warranted by the case. This introduces human oversight (HITL) capable of detecting potential bias or algorithmic errors that could undermine fairness in the specific case.
- d) **Training on ethics and bias:** The organization shall train its technical teams and other relevant teams on artificial intelligence ethics, including the identification of cognitive and data biases and the potential social and organizational impacts arising from the use of AI systems. This training shall be developed with the participation of Ethics and Compliance to incorporate integrity, conduct, and responsible AI-use perspectives consistent with corporate values and applicable regulatory frameworks. Internal users of AI systems (for example, analysts who use AI outputs to support decisions) shall also be made aware of the importance of questioning anomalous results, exercising professional judgment, and reporting potential bias, misuse, or unexpected AI-tool behavior through established institutional mechanisms.
- e) **Integrity Channel:** The organization's Integrity Channel shall be used as the formal and confidential mechanism through which employees, customers, or third parties may report anomalous conduct or outcomes in AI systems that they consider discriminatory, biased, or ethically questionable. Reports shall be managed in accordance with the corporate Ethics and Compliance model, with the participation of the appropriate technical and AI-governance bodies, and the necessary corrective actions shall be adopted where the issue is confirmed.

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	Page 1-11

- f) **Ethical culture:** Beyond formal rules, the organization shall foster a culture of ethics and integrity in the adoption of artificial intelligence, consistent with the current corporate Code of Ethics. Everyone involved is expected to act prudently and responsibly, placing corporate values and people's rights above commercial or efficiency objectives. If any AI use, although technically possible and lawful, conflicts with the organization's ethical principles, it must be reassessed or discontinued.
- g) **Stakeholder participation:** The organization shall promote consultation with and participation by diverse groups in the design, validation, and monitoring of AI systems, particularly those that may directly affect end users. This includes:
- Review of use cases by multidisciplinary committees
 - Inclusion of intergenerational equity, gender, accessibility, and cultural diversity perspectives
 - Assessment of the systems' social and reputational impact

By ensuring fairness and preventing discrimination in AI, we not only comply with the law but also strengthen our reputation as a fair, people-centered organization, increasing the acceptance and effectiveness of our artificial intelligence solutions.

5.6 Training and Awareness

The introduction of Artificial Intelligence systems entails significant changes in processes and decision-making. It is therefore essential that everyone in the organization be informed and trained regarding the appropriate use of these technologies. Accordingly, the following actions shall be implemented:

- a) **Training programs:** Periodic training shall be provided to employees at all levels on AI fundamentals, benefits, risks, and associated responsibilities. These programs shall cover topics such as basic AI and machine learning concepts, interpretation of AI outputs, ethical considerations, data protection in AI projects, and related internal procedures. The content and depth of training shall be tailored to the audience - for example, technical teams shall receive more specialized training, while business areas shall receive practical, user-focused information.
- b) **Awareness of safe and ethical use:** Internal communication campaigns (talks, infographics, newsletters) shall reinforce key messages regarding the safe, responsible, and ethical use of AI. The guiding principles of this policy shall be highlighted so that employees internalize them (for example, reminders to 'consider transparency' when designing an AI solution or 'do not overlook human intervention' when using an AI recommendation in a decision).
- c) **Guides and manuals:** Practical manuals or guides shall be developed as quick-reference resources for employees. For example, a 'Good Practices for AI Projects' guide listing steps to follow and mistakes to avoid, or a Frequently Asked Questions document explaining what to do in certain situations (what to do if an algorithm appears biased; how to proceed if a potential security incident is found in an AI system; etc.).

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	Page 1-11

These guides shall be available on the corporate intranet or through other accessible channels.

- d) **Specialized training for developers and data scientists:** Technical personnel involved in AI shall receive ongoing training on advanced standards and tools. This includes training on the NIST AI Risk Management Framework, fairness and explainability testing methodologies, software security standards (for example, Secure AI Development), and applicable local and international AI regulations. Relevant certifications shall be encouraged (such as certifications in AI ethics, responsible data science, etc.).
- e) **Assessment of effective training:** The organization shall periodically measure the effectiveness of these training initiatives through knowledge surveys, online assessments, or simulations. The purpose is to identify gaps in understanding and continuously improve training programs.
- f) **Organizational culture and scaled learning:** A culture of continuous and collaborative learning shall be fostered, in which every person is an active participant in the responsible use of AI. Scaled organizational learning shall be promoted by leveraging experience, incidents, and good practices for collective improvement. AI Stewards and AI Growth shall facilitate reflection, knowledge-sharing, and continuous-improvement forums, integrating acquired knowledge into processes and strategic decisions.

Through these training and awareness actions, the organization shall ensure that the internal culture supports and advances the objectives of this policy, making each person an active participant in the responsible use of AI in accordance with our values.

The training and awareness activities described in this section are coordinated and led by AI Growth, in conjunction with Talent Management, AI Stewards, and the corresponding technical areas, in accordance with the roles and responsibilities defined in Section IV of this policy.

5.7 Impact Assessment and High-Risk Systems

Based on the precautionary principle and in compliance with the Regulations of Law No. 31814, the organization shall pay particular attention to AI systems classified as 'high-risk.' The classification of AI systems as 'high-risk' shall be performed in accordance with applicable regulations (Law No. 31814 and its Regulations) and aligned with RIMAC's corporate risk management methodology, integrating it into the Operational Risk and Technology Risk frameworks in accordance with the criteria and validations established by the Risk areas.

The criteria determining when an AI system requires a technology risk assessment are defined in section 5.2 of this policy, in accordance with the Technology Risk Area's emerging technologies and AI risk methodology. Classification of a system as high-risk may arise from the outcome of that assessment or from the direct application of the regulatory criteria established in Law No. 31814 and its Regulations.

The following additional control measures shall apply to such 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	Page 1-11

- a) **Prior AI Impact Assessments (AIAs):** Before developing or implementing a high-risk AI tool, an AI Impact Assessment shall be conducted. This comprehensive assessment shall analyze the system's potential effects on privacy, non-discrimination, personal safety, and other rights. Its purpose is to identify and mitigate potential risks before the system enters production, ensuring that the technological solution does not infringe fundamental rights or perpetuate inequality or bias.
- b) **Voluntary versus mandatory nature:** Under Supreme Decree No. 115-2025-PCM, impact assessments for high-risk AI are voluntary in the private sector. Nevertheless, our organization adopts these assessments as a mandatory internal requirement due to their value for responsible management. Any exceptional decision not to conduct an AIIA must be documented and justified, approved by the AI Committee and Compliance, and recorded for audit purposes.
- c) **Documentation and retention of results:** All findings, analyses, and corrective measures arising from an AI impact assessment shall be documented in a formal report. The report shall include a description of the system assessed, identified risks, implemented recommendations, and the signatures of those responsible for its preparation. These reports shall be retained for at least three (3) years so they are available for traceability and accountability if required by internal audits or competent authorities.
- d) **Periodic reassessment:** High-risk AI systems shall be subject to periodic impact reassessments (for example, annually or whenever material changes are introduced to the system or its use context). This ensures that new risks or previously unforeseen effects can be detected and managed in a timely manner.
- e) **Enhanced oversight and human controls:** In keeping with the human oversight principle described above, additional controls shall be implemented for high-risk systems. Specially trained personnel shall be assigned to closely monitor these systems. The points in the process at which a person must intervene to validate or approve an AI decision before it is implemented shall be clearly defined, ensuring that the final decision on sensitive matters rests with an accountable person. For example, if AI were used to assist in deciding on surgery or approval of a high-value loan, the final decision must be verified and assumed by a qualified human professional.
- f) **Incident recording and notification:** Any incident or adverse outcome arising from a high-risk AI system (for example, an incorrect prediction that caused harm or a security failure in a critical smart device) must be reported immediately to senior management and the AI Committee. The root cause shall be investigated and relevant external bodies shall be notified as required by regulations (for example, data protection authorities if personal data is involved, or other sector regulators). Incidents affecting information security must be reported immediately to Information Security; the mandatory contacts are the CISO (Information Security Manager) and the CSIRT Lead (Blue Team Chapter Lead).

By implementing these measures, the organization seeks not only to comply with current Peruvian regulations governing high-risk AI, but also to adopt exemplary practices aligned with international frameworks (such as the European Union's proposed AI regulation), anticipating future obligations and strengthening trust in our operations.

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	Page 1-11

5.8 Monitoring, Auditing, and Continuous Improvement

Effective compliance with this policy requires sustained monitoring and continuous improvement. AI and its regulation evolve rapidly; therefore, the organization shall implement mechanisms to ensure that our practices remain current and effective:

- a) **Continuous system monitoring:** The responsible areas (IT, Data, Processes, and the business owner) shall continuously monitor AI systems in operation. This includes monitoring key performance indicators (accuracy, response times, error rates) and risk indicators (for example, model drift and the emergence of bias over time). Where material deviations from approved parameters are detected, measures must be taken (model retraining, parameter adjustment, investigation of recent data, etc.). Monitoring shall be managed by AI Stewards and periodically reported to the AI Committee.

The Data Area is responsible for systematically and periodically monitoring the quality, traceability, and security of data used and generated by AI systems, reporting findings, and proposing corrective actions.

The Processes Area is responsible for defining and monitoring the performance of processes automated or assisted by AI, ensuring continuous improvement and alignment with the business's strategic objectives.

Both teams collaborate in internal and external audits, ensuring traceability, documentation, and regulatory compliance of AI agents.

- b) **Periodic internal audits:** The organization shall include specific reviews of compliance with this policy in its internal audit plan. Internal auditors (or external auditors, where independent expertise is required) shall assess samples of AI system projects to verify whether the corresponding risk assessments were conducted, security and privacy controls were applied, ethical principles are being observed in day-to-day operations, and documentation is current. Audits may include technical model testing (for example, fairness assessments), code reviews, interviews with responsible parties, etc. Audit findings shall be documented and communicated to senior management, together with action plans to remediate any identified noncompliance.
- c) **Policy review and update:** This policy is not static and must be reviewed at least annually or whenever material changes occur (for example, new applicable laws or regulations, the incorporation of radically different AI technologies, or serious incidents that reveal opportunities to improve the guidelines). The review shall be conducted by the AI Committee together with Compliance and AI Growth and must be approved by senior management. Any update shall be communicated to all personnel and shall take effect immediately unless a transition period is specified.
- d) **Compliance indicators:** Indicators shall be defined to measure the degree of compliance with this policy within the organization. Examples of possible metrics include the percentage of AI projects with a completed risk assessment, the number of employees trained in AI compared with the plan, the number of AI-related incidents or near misses reported, and the average response time to correct identified bias. These

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	Page 1-11

indicators shall be reported to senior management within the corporate risk management system or relevant committee so that executive-level visibility is maintained.

- e) **Feedback and participation:** The organization shall encourage employees to provide suggestions or raise concerns regarding this policy and its implementation. Open channels shall be maintained to receive feedback, which may identify areas of the policy requiring greater clarity or additional management support (for example, where teams find a guideline difficult to apply, the need for additional training or adjustment of the guideline shall be assessed). Effective use of AI in the company shall result from collective learning; therefore, the voices of those who work with these tools every day are essential.

In summary, by applying a Plan-Do-Check-Act (PDCA) cycle to our AI management, we ensure that high control standards are maintained. Continuous improvement shall ensure that our institutional AI policy remains relevant and effective in a dynamic technological and legal environment.

5.9 Governance of Assistants, Intelligent Agents, and Autonomous Systems

The organization recognizes the rapid growth of solutions based on intelligent agents and autonomous multi-agent systems, which present unique challenges in governance, oversight, traceability, and alignment with institutional values. Accordingly, the following specific guidelines are established for their design, implementation, operation, and control:

- a) **Definition of AI assistants/agents:** An AI agent shall mean any system that acts autonomously and has the capacity for perception, reasoning, decision-making, and execution of actions without constant human intervention. This includes individual agents, generative agents, autonomous assistants, software agents embedded in processes, and multi-agent systems that interact with one another.

Each agent must have a clear and measurable objective aligned with the business unit's goals and demonstrate its impact on organizational processes or results. Periodic assessment of compliance with these objectives shall be the responsibility of the designated AI Steward.

- b) **Specialized governance:** Intelligent agents shall be subject to enhanced governance controls led and supervised by AI Growth, including:
- Objective-alignment assessments using tools such as the Goal Alignment Canvas to ensure that the agent's objectives are aligned with institutional values, policies, and goals.
 - Mandatory human oversight for critical decisions, ensuring human intervention or validation at high-impact points.
 - Trust & Capability Indicators to assess the agent's autonomy, reliability, and operating limits.
 - Monitoring of emergent behavior and goal drift, with mechanisms to detect deviations in the agent's behavior from its original purpose.

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	Page 1-11

c) **Cross-functional validation and formal approval:** Every intelligent agent must be validated and approved by the areas involved, including:

- The technical area responsible for development (Security and Risk, Data, AI, Processes, QA, etc.) and AI architecture.
- Ethics and Compliance, in its role of identifying ethical and compliance risks and advising on their mitigation.
- The functional or business area that will use the agent.
- The AI Committee, in the case of high-impact agents.

Validation shall include a review of risks, ethical alignment, traceability, human oversight, and regulatory compliance. No agent, at any stage of development or operationalization (proof of concept or MVP), may be deployed without documented formal approval.

d) **Validation and approval by Data and Processes: Every AI system must be validated and approved by the Data and Processes Areas, in addition to the technical and legal areas.**

- Data verifies the quality, protection, and traceability of the data used by the agent.
- Processes validates the agent's alignment with business objectives and processes, ensuring that its implementation generates value and complies with ethical and regulatory standards.

e) **Mandatory technical framework: Agents must be designed and built in accordance with a technical framework defined by the specialized AI CoE team for the respective implementation technologies, including:**

- Secure and scalable architecture standards.
- Training, testing, and validation protocols.
- Security, privacy, and explainability controls.
- Audit mechanisms and technical documentation.
- Governance tools such as the Goal Alignment Canvas, Agentic Trust Scan, and ARAT.

This framework shall be periodically updated and shall be mandatory for all internal developments or external procurement of intelligent agents.

f) **Management of multi-agent systems: For solutions composed of multiple agents that interact with one another, mechanisms shall be established to:**

- Coordinate objectives and prevent conflicts among agents with different goals or functions.
- Detect unforeseen behaviors, emergent effects, or systemic risks arising from interactions among agents.
- Audit distributed decision-making, ensuring traceability, accountability, and transparency in collaborative or autonomous environments.

g) **Registration and traceability: Every agent must be recorded in the institutional AI Inventory, including:**

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	Page 1-11

- Functional purpose and scope
- Level of autonomy and type of decisions it may make
- Assigned risk classification
- Designated human owner and person responsible for the operation of the assistant/agent (AI Steward)
- Control, monitoring, and oversight mechanisms applied

h) **Generative AI System:** A set of applications that leverage generative artificial intelligence models to create content, automate processes, and enable intelligent experiences.

- **AI Assistant:** An application designed for conversational interaction (text or voice) that interprets user intent, queries authorized sources, and returns responses or recommendations.

Key characteristics:

- Does not take direct actions on systems or execute transactions.
- Its function is to answer queries and assist with specific tasks.
- May include optional human oversight (HITL) and basic traceability.

- **AI Agent:** An application capable of perceiving its environment, planning, deciding, and acting autonomously to achieve defined objectives.

Key characteristics:

- Executes end-to-end actions and transactions under controls and governance.
- Invokes tools, APIs, and other agents and/or assistants.
- Operates autonomously, incorporating HITL where required by risk.

- **LLM Model (NLP, CV, Voice):** A set of foundational and specialized models that enable the following capabilities:

- NLP: Language understanding and generation, information extraction, and reasoning.
- Voice (ASR/TTS): Speech recognition and synthesis.
- CV: Image and video interpretation.

These models are combined to provide multimodal experiences and automate complex tasks (e.g., reading documents, analyzing images, and responding by voice).

- **Multi-Agent Systems:** A structure in which multiple specialized agents and assistants collaborate, compete, or coordinate under centralized or distributed orchestration models to solve complex tasks.

Benefits:

- Scales processing capacity.
- Combines diverse capabilities (NLP, CV, transactions).

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	Page 1-11

- Manages risks through roles, policies, and interruption mechanisms.
- **Workflow (AI Workflow):** An orchestrated sequence of steps that integrates:
 - Model calls (LLM/CV/Voice).
 - Tools and APIs.
 - Agents/assistants.
 - Human validations (HITL).
- Defines:
 - Measurable inputs and outputs (KPIs, SLAs).
 - Decisions and actions.
 - Controls (auditing, traceability, intervention thresholds).

The AI Steward or business expert shall be directly responsible for agents developed by the business unit, ensuring their proper operation, alignment with institutional objectives, and regulatory compliance. This role must oversee the performance, impact, and evolution of the agents under its management.

Regulatory and ethical compliance: Intelligent agents must comply with the guiding principles established in this policy and with applicable national regulations (Law No. 31814 and its Regulations) and international regulations. Their ethical, legal, and social impact shall be assessed separately, considering their capacity to act without direct oversight, their level of autonomy, and their potential to make decisions with significant consequences. 5.10 Regulatory Framework and Reference Standards

This policy is aligned with the requirements of the following national and international regulatory instruments, which establish obligations and good practices for the use of Artificial Intelligence, information security, and personal data protection:

- a) **Law No. 31814** - Law Promoting the Use of Artificial Intelligence for the Economic and Social Development of the Country (Peru, 2023): Establishes the national legal framework for the safe, transparent, responsible, and human-rights-respecting use of AI. Its principles and provisions (e.g., risk-based approach, centrality of human dignity, and promotion of digital talent) have been incorporated into this policy.
- b) **Supreme Decree No. 115-2025-PCM - Regulations of Law No. 31814 (Peru, 2025):** Sets out in detail the obligations and principles for the safe and ethical use of AI within Peru. It includes guiding principles (non-discrimination, personal data privacy, protection of fundamental rights, respect for intellectual property, security and reliability, sustainability, human oversight, transparency, accountability, and ethical development) and classifies AI systems by risk level, establishing special obligations for high-risk systems (e.g., algorithmic transparency, documented records, and effective human oversight). All these elements are reflected in our policy.
- c) **Personal Data Protection Law No. 29733 and its Regulations (Peru):** Protects individuals' rights to the proper processing of their personal data. Our policy incorporates the principles of consent, purpose limitation, proportionality, security, and confidentiality required by this law, as well as mechanisms for exercising ARCO Rights,

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	Page 1-11

ensuring that any AI system handling personal data does so in strict compliance with privacy regulations.

- d) **International Information Security Standards (ISO/IEC 27001:2022 and ISO/IEC 27002:2022):** Provide a framework for implementing and maintaining effective information security management systems. We align our AI security controls with these standards, including access management, change control, encryption, event monitoring, business continuity, and other measures recommended by ISO 27001.
- e) **International Security Risk Management Standard (ISO/IEC 27005:2022):** Guides the performance of security risk assessments. We have adopted its guidelines to identify threats and vulnerabilities in our AI systems, assess impact and likelihood, and treat risks through documented mitigation plans, ensuring a systematic and traceable risk management process.
- f) **International Privacy Information Management Standard (ISO/IEC 27701:2019):** An extension of ISO 27001 focused on privacy. Our Privacy by Design approach and AI data protection measures follow the practices recommended by this standard, ensuring that personal data processing involving AI addresses both legal requirements and organizational privacy management controls.
- g) **NIST Artificial Intelligence Risk Management Framework (United States, 2023):** A global reference for managing AI risks. We have integrated its four core functions - Govern, Map, Measure, and Manage - into our AI governance process, enabling us to address the unique challenges posed by these systems in a structured manner and improve the trustworthiness of AI solutions by incorporating trust considerations into their design, development, use, and evaluation.
- h) **Other International Regulations and Standards:** We remain attentive to emerging regulations in other jurisdictions that may apply to our operations (e.g., the EU General Data Protection Regulation - GDPR - where data of European citizens is processed, or future AI regulations of the European Union or other countries where we operate). We also consider international principles such as the OECD AI Principles and the UNESCO Recommendation on the Ethics of Artificial Intelligence (2021) as guidance for global good practices, ensuring that our policy is consistent with international efforts toward the ethical development of AI.
- i) **Policy breaches and disciplinary measures:** Noncompliance with this policy is considered a breach of internal rules and may result in disciplinary action determined according to the seriousness of the case and in accordance with internal regulations and applicable labor law, including removal of the employee from their duties, without prejudice to any applicable civil or criminal action.

The processing and management of data used by artificial intelligence systems must be performed in accordance with the current regulatory framework, including Law No. 31814, ISO/IEC 42001, Macroprocess M013 - Data and Analytics Governance, and internal policies POL-4687, POL-6021, and POL-4688. (See the Related Documents section for further details.)

5.11 Acquisition and Use of Third-Party AI Solutions

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	Page 1-11

- Any acquisition of technologies that use AI capabilities requires prior review by the Capabilities team.
- The use of free AI capabilities or AI capabilities licensed for personal use is not permitted.
- Where third-party AI processes confidential or regulated data, it must be reviewed by the corresponding control functions (Legal, Privacy, Technology Risk, and Information Security).

Any procurement, licensing, subscription, or activation of artificial intelligence functionality provided by third parties - including SaaS platforms whose primary functionality is AI, foundational-model APIs, and AI modules embedded in or activatable within existing corporate software - requires prior approval from the Capabilities team before the corporate purchasing, allocation, or activation process begins. This approval is based on the Provider Catalog (SGAI-CAT-PROV-001), which defines which providers are authorized and under what conditions of use according to the applicable data classification and personal data category (Guideline DP001, Law No. 29733). The assessment and approval procedure is documented in SGAI-PRO-001.

Purchasing, contracting, license allocation, and access management are carried out in accordance with RIMAC's standard corporate software acquisition processes, following the same guidelines, controls, and workflows applicable to any other software. Approval by the Capabilities team does not replace or duplicate these processes; it precedes them as an AIMS compliance gate.

The use of AI solutions through individual subscriptions, personal accounts, or unauthorized arrangements for work purposes constitutes Shadow AI and is expressly prohibited. Employees who require access to AI tools must channel their request through the official catalog or, if the tool is not listed, request its assessment under SGAI-PRO-001. By default, each employee shall have access to one general-purpose AI tool. Any exception must be documented and justified by the AI Steward of the relevant unit and approved by the AI Committee, regardless of whether the requested tool appears in the Provider Catalog (SGAI-CAT-PROV-001).

6. RELATED DOCUMENTS

Code	Name
MAN-4792	RIMAC Code of Conduct
POL-3672	Personal Data Protection Policy
M013	DATA AND ANALYTICS GOVERNANCE MACROPROCESS

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	Page 1-11

POL-4687	THIRD-PARTY DATA COLLECTION
POL-6021	DATA GOVERNANCE AND QUALITY
POL-4688	METHODOLOGY FOR ADDRESSING MASS REMEDIATION INCIDENTS
Guideline DP001	PERSONAL AND SENSITIVE DATA CATEGORIES
MAN-6038	AI Risk Management Methodology

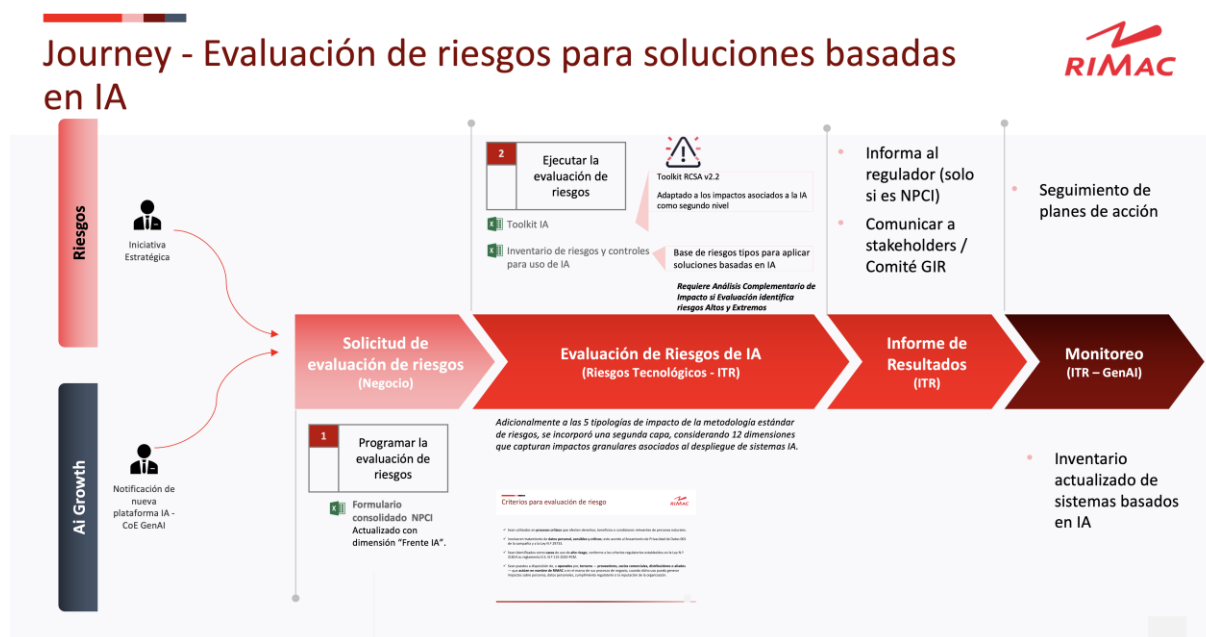
RIMAC	POLICY				Code: POL-6098
	SAFE, RESPONSIBLE AND ETHICAL USE OF ARTIFICIAL INTELLIGENCE (AI) SYSTEMS				Version: 008
	Macroprocess:	Data and Analytics Governance	Process:	AI Growth	Page 1-11

7. APPENDICES

APPENDIX I: Risk Classification of Artificial Intelligence Systems

Artificial intelligence systems are risk-classified in accordance with the AI Risk Management Methodology (MAN-6038) defined by the Technology Risk Area.

Diagram: Risk assessment process for AI systems



UNACCEPTABLE-RISK AI

Certain artificial intelligence applications are prohibited because they are inherently incompatible with legal, ethical, or constitutional principles. These include:

- AI uses that are expressly unlawful under Peruvian law.
- AI that violates constitutional protections, human rights, or established public policies.
- AI used for social scoring of citizens that results in punitive or discriminatory measures.
- AI that enables mass surveillance without legal authority.
- AI that claims to predict criminal behavior based on race, ethnicity, or other protected characteristics, among others.

If an employee or area identifies or receives a request to implement an AI system of this type, the activity must be stopped immediately and the case reported to AI Growth so that it may be escalated to the AI Committee, where the corresponding investigation and measures shall be undertaken.

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	Page 1-11

For each AI system assessed, AI Growth and Technology Risk shall maintain written justification for the assigned risk classification, including the criteria met under Section 5.2(c), the outcome of the assessment conducted under MAN-6038, and the traceability of the decision. This documentation forms part of the AI Inventory (SGAI-INV-001) and shall be available for internal and external audits in accordance with the retention periods defined in this Policy. Entities must maintain written justification for the risk classification assigned to each AI system, including the criteria met and the reasoning supporting the determination.

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	Page 1-11

APPENDIX II: Checklist for Contractual Clauses with AI System Providers

This appendix serves as a guide for the organization's contract managers and legal advisors when drafting or reviewing contracts (purchase orders, service agreements, and MOUs) with providers of artificial intelligence systems or services. Its purpose is to protect the company's interests, ensure AI reliability, and comply with Law No. 31814 and its Regulations.

1. Scope and Purpose of the Service

- Clearly define the AI system or service, its functionality, scope, and purpose within the organization.
- Establish limitations to prevent unauthorized or unintended uses.

2. Data Ownership and Rights

- The organization retains full ownership of the data provided and the outputs generated by the AI system.
- Prohibit secondary use of data without written authorization.
- Restrict the transfer of data to third parties or outside the country without express authorization, in accordance with the Personal Data Protection Law (Law No. 29733).

3. Data Protection and Privacy

- Mandatory compliance with Law No. 29733 and supplementary regulations.
- Minimum requirements:
 - Encryption of data in transit and at rest.
 - Secure storage in approved jurisdictions.
 - Prohibition on mining, selling, or reusing personal or sensitive data.
- For sensitive data (health, biometric data), require specific agreements such as confidentiality agreements or equivalent instruments.

4. Confidentiality

- Include a clause prohibiting disclosure of corporate data or sensitive AI-system information.
- This obligation must survive termination of the contract.

5. Compliance with Security Controls

- The provider must comply with standards such as ISO/IEC 27001, ISO/IEC 42001, and the NIST CSF.
- Cloud solutions must hold recognized certifications (e.g., SOC 2, ISO 27017).
- Periodic security maintenance and updates are required.
- The company may conduct independent audits.

6. Audit Rights

- The entity has the right to audit the provider's performance, including security, data governance, and bias testing.
- This must include on-site inspections and remote document reviews.

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	Page 1-11

7. Testing and Validation

- Provider cooperation in initial and ongoing testing.
- Immediate correction of bias, accuracy failures, or vulnerabilities.
- Bias-mitigation clause supported by technical model documentation.

8. Service Levels and Support

- Define service-level agreements (SLAs) for availability, performance, and response times.
- Define maintenance and update procedures, with prior notice.

9. Change Management

- Require advance notice and institutional approval for model or software changes that affect outcomes.
- Prohibit deployment of major versions without prior testing and acceptance.

10. Deliverables and Documentation

- User and administrator manuals.
- Model documentation (model cards, data sheets).
- Integration guides and data schemas.
- Continuous documentation updates.

11. Intellectual Property Rights

- The provider retains rights to pre-existing intellectual property but grants the company the necessary licenses.
- For custom developments, the company must have rights to use, modify, and maintain the solution.

12. Indemnification

- The provider shall indemnify the company for:
 - Intellectual property infringement claims.
 - Damage caused by the AI system, negligence, security breaches, or noncompliance.
 - Unlawful use of data or improper acquisition of training data.

13. Limitations of Liability

- Avoid nominal or unreasonably low liability caps.
- Exclude serious breaches (confidentiality, intellectual property, and security) from liability caps.

14. Termination Rights

- The entity may terminate the contract for breach or for convenience.
- The provider must ensure continued access to data and deliverables after termination.

15. Transition Assistance

- The provider is required to facilitate an orderly transition:

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	Page 1-11

- Return of data in a usable format.
- Export of models, where applicable.
- Support for new providers without disruption.

16. Applicable Laws and Regulations

- The contract is governed by Peruvian law.
- The provider must comply with:
 - Law No. 31814 and its Regulations.
 - Personal Data Protection Law No. 29733.
 - Applicable ISO/IEC standards.
 - UNESCO and OECD ethical principles.

17. Penalties

- Include penalties for SLA breaches or security failures.

18. Governance and Reporting

- Periodic reports on performance, issues, and updates.
- For multi-year contracts, require annual certification of compliance.

19. Subcontractors

- The provider is responsible for its subcontractors.
- Institutional approval is required for subcontractors that handle sensitive data.

20. Dispute Resolution

- Avoid clauses requiring disputes to be resolved outside the company or in unfavorable forums.

	POLICY				Code: POL-6098
	SAFE, RESPONSIBLE AND ETHICAL USE OF ARTIFICIAL INTELLIGENCE (AI) SYSTEMS				Version: 008
	Macroprocess:	Data and Analytics Governance	Process:	AI Growth	Page 1-11

APPENDIX III: Annual AI Compliance Report - Template for Private Companies

[Company Name] - Annual AI Compliance Report - [Reporting Year]

1. Executive Summary

Brief description of AI-related activities during the year. Include:

- Institutional commitment to responsible AI governance.
- Key achievements (e.g., implementation of new systems, certifications, audits).
- Material changes to internal policies or AI use.

2. AI System Inventory

Table listing all AI systems in use, under development, or decommissioned. Include prototypes and pilots.

System Name	Description / Purpose	Status (Pilot / Production / Decommissioned)	Risk Level (Low / Medium / High)	Assessment Date	Provider / Internal
AIDA CDC	Customer service chatbot	In production since November 2024	Medium	Assessed in January 2025	Internal
SHERLOCK	Claims fraud detection	Development (RFP)	High	Preliminary assessment	Provider: XYZ AI

Attach a list of decommissioned systems, including the date and reason for decommissioning.

3. Risk Assessment and Mitigation Measures

For each medium- or high-risk system, include:

- Date of the latest assessment.
- Key risks identified (e.g., bias, privacy, misuse).
- Mitigation measures implemented.
- Changes in risk classification and justification.
- Confirmation of continuous monitoring.

Example:

- **Fraud - Sherlock**
 - Assessment: July 2025
 - Risks: Bias in decisions; unauthorized access to sensitive data.
 - Mitigation: Human review of critical cases; bias testing before deployment.
 - Monitoring: Scheduled for November 2025.

4. Regulatory Compliance Checklist

Self-assessment of compliance with the internal AI policy and Peruvian legal framework.

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	Page 1-11

Policy Area	Compliance (Yes / No / In Progress)	Comments
Institutional AI Committee	Yes	Committee active since March 2025
Internal AI policy	In progress	Approval expected in December 2025
Training of key personnel	Yes (85%)	Refresher training scheduled for Q1 2026
Data classification and protection	Yes	In accordance with Law No. 29733
Incident reporting	Yes	No serious incidents recorded
High-risk system registry	Yes	Updated quarterly

5. AI-Related Incidents

Record relevant events (failures, bias, security breaches).

Example:

- **August 2025 - Error in customer risk classification**
 - Nature: False fraud alert.
 - Impact: Manual review of 120 cases.
 - Corrective action: Adjustment of decision threshold and model retraining.

If there were no incidents:

'No AI-related incidents were recorded during the reporting period.'

6. Improvements in AI Governance

Measures adopted to strengthen governance:

- Implementation of internal AI review forms.
- Participation in sector working groups (e.g., APESEG, industry associations).
- Workshops on ethics and algorithmic bias.
- Integration of AI asset tracking into IT management systems.

7. Future Plans

Activities planned for the next year:

- New AI projects under assessment (e.g., AI for claims prediction).
- Updates to or decommissioning of existing systems.
- Certification under ISO/IEC 42001 or the AIGN Trust Label.
- Request for external advice on AI governance.

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	Page 1-11

8. Additional Comments (Optional)

Success stories, challenges encountered, and lessons learned. Example:

'The implementation of AIDA reduced average customer service time by 35%, improving satisfaction by 18 percentage points.'

Signature and Validation

- **Prepared by:** [Name, Position] - AI Governance Lead / CIO / CISO.
- **Approved by:** [Name, Position] - General Manager / Chief Technology Officer.
- **Date:** [dd/mm/yyyy]

Attachments

- Risk assessment reports.
- Compliance matrices.
- Incident reports.
- Training evidence.
- AI system registry.

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	Page 1-11

1. APPROVAL WORKFLOW RESPONSIBLE PARTIES

STAGE	AREA	POSITION	NAME
Owner	AI Growth	AI Growth Master	Jorge Vicente
Editor	AI Growth	AI Growth Specialist	Alan Fermin
Editor	AI Growth	AI Growth Specialist	Joel Espinoza
Reviewer	AI CoE	AI COE LEAD	William Cortez
Reviewer	Legal Advisory and Contracts	Legal Manager	Rodolfo Grados
Reviewer	Regulatory Compliance	ETHICS AND COMPLIANCE MANAGER	Sara Sologuren Montenegro
Reviewer	Information Security	INFORMATION SECURITY MANAGER	Karla Parra
Reviewer	Technology Risk	SECURITY AND DATA TECHNOLOGY RISK MANAGER	Jose Fernandez Salas
Reviewer	IT Architecture	DIGITAL CHANNELS DOMAIN ARCHITECT	Jordi Tanta Diaz
Reviewer	IT Architecture	IT ARCHITECTURE VICE PRESIDENT	Gustavo Zuazo
Reviewer	Business Process Engineering	PROCESS COE LEAD	Karin Arenaza
Reviewer	Structures and Productivity	Structures and Productivity Lead	Viviana Porto
Reviewer	Data	Data Manager	Lourdes Rojas
Approver 1	AI Growth	AI Growth Master	Jorge Vicente
Approver 2	Capabilities	Capabilities Vice Presidency	Miguel Portugal

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	Page 1-11

Approval 3	Transformation Division	Executive Vice President	Giselle Larco
------------	-------------------------	--------------------------	---------------

2. CHANGE CONTROL

DOCUMENT CREATION			
DATE PREPARED	DESCRIPTION	V	PREPARED BY
26/12/2025	First version	01	Jorge Vicente