

A person with dark hair, seen from the side, is sitting at a desk in a control room. They are wearing a black long-sleeved shirt with a Mexican flag patch on the sleeve and a bright yellow high-visibility safety vest. They are looking at several computer monitors. The top row of monitors shows various data visualizations, including a color-coded map on the left and several smaller maps or satellite images on the right. The bottom row of monitors shows a large window with text and data, possibly a report or a software interface. The person's right hand is on a black keyboard.

RESPONSIBLE USE OF AI: SUSTAINABILITY ACTIONS 2025

Responsible Use of Artificial Intelligence¹

Responsible Artificial Intelligence Programs

At Grupo México, we have a permanent commitment to responsible operations, the creation of sustainable value, and the protection of the interests of our employees, shareholders, customers, suppliers, and communities. This commitment requires us to stay at the forefront of new technology adoption, ensuring that its use is at all times aligned with our values and with the highest standards of ethics, security, and integrity.

The following are the programs and content of our extended Responsible Use of Artificial Intelligence Policy which serves as an internal standard or toolkit for our employees.

AI Governance Model

AI governance at Grupo México is organized into five internal levels of responsibility. External stakeholders (including institutional investors, regulators (CNBV, SEC), legal advisors, communities, and suppliers) influence the design and updating of our Responsible Artificial Intelligence Policy through their expectations regarding transparency, ethics, and responsible AI management, which are considered by the Audit and Company Practices Committee when reviewing the governance program.

Level	Responsibility
Level 1. Audit and Corporate Practices Committee of the Board	Approves the Responsible Use of AI Policy and its substantial amendments. Receives quarterly reports on the governance program, or sooner if necessary; the person responsible for producing the reports is the corporate CISO, supported by the work of each Division's CISO, as well as, where applicable, the Data Protection Officer (DPO) Each division's Legal department will designate the person who will serve as the division's DPO.
Level 2. Information Technology (IT) Departments of each division	This level of responsibility is aimed at the executives responsible for information security in each of the divisions (Information Technology Departments). Approve High and Very High-risk systems, define the catalog of authorized AI tools, and approve exceptions to our Responsible Use of AI Policy where applicable and duly justified. Operate the Inventory of AI Models and Tools. Manage the evaluation and approval process. Coordinate bias assessments and data loss prevention (DLP) controls. Produce the monthly executive report. Channel and respond to inquiries and incidents. In addition, they monitor the performance of AI tools in use.
Level 3. Implementation teams	These are department/area leaders, advanced users, and tool owners. They are responsible for safe day-to-day use and for reporting anomalies.
Level 4. Champions	These are individuals designated by each Division who, based on their experience and knowledge of the sector and technology, promote the responsible use of our AI Policy, resolve first-level questions, and serve as a liaison with the Information Technology teams of each division.
Level 5. Users	All employees are responsible for how they use AI in their daily work.

¹ The actions described in this document are fully applicable to our subsidiary Southern Copper Corporation.

Use of Classified Information in AI Systems

The classification of Grupo México's information is established in the Data Governance and Handling Policy, which defines the four applicable levels (Public, Internal, Confidential, and Restricted). This section does not redefine that taxonomy; it only specifies the permitted use of each classification level in AI systems and tools. Before feeding any Grupo México data into an AI system or tool, its classification must be verified in accordance with that policy.

Table. Use of Classified Information in AI Systems

Classification	Description	Examples	Permitted use in AI
Public	No distribution restrictions	Published annual reports, industry statistics, commodity prices, public information on Grupo México's and its subsidiaries' official websites	Any approved AI tool
Internal	For internal use; not intended for the public	Operational KPIs, internal procedures, organizational charts.	Only approved "Enterprise" tools. Prohibited on public AI tools.
Confidential	Its exposure can cause significant harm	Contracts, unpublished financial data, corporate strategy, personnel selection, personal data in general, and intellectual property.	Secure environments with a data processing agreement. PROHIBITED on public tools or through personal devices.
Restricted	Maximum sensitivity; serious regulatory risk	Sensitive personal data, including biometric data, occupational health data, and data related to legal background, information under legal investigation, trade secrets, know-how, regulatory reports pending approval, and formal publication.	PROHIBITED on any AI system without a DPIA + authorization from each division's IT Department + documented legal basis.

Transparency and Communication

AI must be visible both within and outside the organization:

- All content generated wholly or partly with Generative AI that is distributed (internally or externally) must be identified as "Generated with AI assistance."
- High-impact decisions supported by AI models must document the name and version of the model used, and the action taken by the person responsible.

Any content in whose creation an AI system had substantial involvement whether text, code, image, design, or analysis must be labeled with the notice "Generated with AI assistance" or an equivalent approved by the Communications Department and the Legal Department. Substantial involvement will be understood to exist when the AI system produced the structure, wording, or main design of the content, even if a human subsequently reviewed or edited it. Labeling will not be required when AI is used exclusively as an auxiliary search, spell-checking, or formatting tool. It is strictly prohibited to submit to authorities (INDAUTOR, the Mexican Institute of Industrial Property (IMPI), or others that may apply) content generated predominantly by AI as a work of exclusively human authorship, given the possible legal implications of inaccurate statements in registration proceedings. Failure to comply with this provision will be considered a serious offense under the Internal Labor Regulations and may lead to corresponding disciplinary and legal action.

Examples of Risks and Controls in AI Use

AI risk management requires a comprehensive approach based on governance, technological controls, and organizational culture, enabling responsible innovation without compromising compliance, reputation, or decision-making.

The following are examples of risks and controls as mitigation measures.

	Reference/Type	Risks	Controls / Mitigation Measures
1.	Misuse of personal/confidential data	Exposure of personal, financial, or strategic data in AI tools	Mandatory data classification and masking (PII/PCI). Restricting use to approved tools (corporate whitelist)
2.	Data leakage	Information uploaded to public AI that may be retained or reused	"No input" policies for critical information. Implementation of private/enterprise AI environments.
3.	Incorrect results or "hallucinations"	Decisions based on inaccurate outputs	Mandatory human review ("human-in-the-loop") for critical cases. Use of verified sources / <i>grounding</i> of information
4.	Algorithmic bias and discrimination	Unfair or unethical decisions (e.g., HR evaluations).	Periodic <i>bias</i> and <i>fairness</i> assessments. Inclusion of diverse datasets and independent validation.
5.	Regulatory and legal non-compliance	Violation of regulations (data protection, intellectual property).	Prior legal review for AI use cases. Model inventory and decision traceability.
6.	Lack of explainability (black box)	Inability to justify automated decisions.	Use of explainable AI (XAI) models in critical processes. Documentation of logic, assumptions, and key variables.
7.	Excessive reliance on AI	Substitution of professional judgment and loss of control.	Defining usage limits (which decisions cannot be delegated). Training in responsible use and critical thinking.
8.	Cybersecurity risks (<i>prompt</i> injection, adversarial attacks)	Model manipulation or information extraction.	Input/output filters and monitoring <i>prompts</i> . Security testing specific to AI models.
9.	Model drift (model/data drift)	Loss of accuracy due to changes in data or environment.	Continuous monitoring with "sentinels" (performance alerts). Periodic retraining and validation of models.
10.	Unauthorized use or shadow AI	Employees using unapproved tools without oversight.	Official catalog of permitted tools. Usage monitoring and clear disciplinary policies.

Security, Data Quality, and Privacy

Personal data protection and information security must be built into every AI system by design.

Vulnerability assessment prior to model release.

Every AI system (whether Traditional or Generative) must undergo a formal vulnerability analysis process before being put into production or released to internal users. This assessment must confirm, prior to production, that the model or system was classified by risk level, that technical, privacy, security, bias, and fundamental-rights vulnerabilities were evaluated, that residual risks are within the approved tolerance, and that auditable evidence exists to support the decision to release or block deployment.

This analysis must include, at a minimum: (i) *adversarial robustness* testing to identify weaknesses against malicious or unexpected inputs; (ii) risk analysis to identify potential threats of sensitive or confidential information leakage, as well as risks that could involve security breaches of personal data embedded during training; (iii) validation of the model's *resistance to data poisoning techniques* and prompt manipulation (prompt injection); and (iv) a review of biases that could lead to inequitable or discriminatory decisions. The results of the analysis must be formally documented in the model's technical data sheet and approved by each division's IT Department, as applicable.

Data Analysis in Traditional AI

For analytics and predictive AI, consideration must be given to validating the quality and integrity of training data; documenting the origin and transformations of the data (lineage); enforcing least-privilege access control to models and pipelines; and protecting AI models against data poisoning.

Data Privacy in Generative AI and Agents

Privacy Impact Assessment (DPIA): Mandatory for any AI system that: (i) processes personal data at scale; (ii) uses sensitive personal data; (iii) has the ability to evaluate, analyze, or predict personal characteristics for the purpose of making or recommending decisions that could significantly affect individuals' interests, rights, and/or legal situation. Users who use personal data in AI tools must formally document the DPIA and obtain approval from the IT, Legal, and Compliance Departments, as well as the opinion of the DPO of each division or corporate DPO, as applicable.

Ethical Use of AI Systems

Grupo México adopts and promotes the ethical use of artificial intelligence in accordance with the international principles established by the United Nations for the ethical use of AI systems. These principles guide the decision-making and expected behavior of all employees when interacting with AI tools.

In addition, all employees must observe the following ethical conduct guidelines when using AI systems:

- **Do no harm:** AI systems must not be used in a way that causes or aggravates harm to people, groups, communities, or the environment. All AI use must be aligned with Grupo México's commitments regarding human rights, social responsibility, and environmental sustainability.

- **Legitimate purpose and proportionality:** the use of AI systems must have a clear, legitimate, and documented business purpose. The AI capabilities used must be appropriate, necessary, and proportionate to the intended objective, avoiding excessive or unjustified use of the technology.
- **Sustainability:** The environmental, economic, and social impact of the AI solutions implemented must be considered. Employees should favor the use of tools that promote energy efficiency and minimize the carbon footprint, in line with Grupo México's sustainability commitments.
- **Inclusion and participation:** When designing, implementing, or using AI systems that affect groups of employees, communities, or third parties, a participatory and inclusive approach should be promoted, seeking to consult potentially affected groups and incorporating diverse perspectives.

Communication and Escalation Channels

The following are the communication channels enabled to report any related incident for our stakeholders.

- Inquiries channel: Corporate email ai-governance@gmexico.mx.
- Escalation process: Employee -> Direct leader / AI Champion -> Division IT Director / Division CISO / Corporate CISO.
- AI security incident reporting: Critical incidents (data leaks, activation of the emergency shutoff) must be reported within a maximum of 4 hours to each division's IT Department, each division's CISO, the Corporate CISO, and Internal Audit.

For security incidents that involve (or may involve) personal data, the escalation process must include each division's Data Protection Officer (DPO) and, where applicable, the corporate DPO.

Ongoing Training in AI Tools

Grupo México recognizes that the effectiveness of our Responsible Use of AI Policy depends not only on its formal existence but on the level of knowledge and competence of the employees who interact with AI systems. Accordingly, each division's Human Resources and IT Departments, in coordination with the CISO, will maintain an ongoing training program to ensure that every employee authorized to use AI tools has training appropriate to the risk level and complexity of the tools they operate.

This program must include, at a minimum: (i) mandatory onboarding training for every employee accessing a corporate AI tool for the first time, covering our Responsible Use of AI Policy's guiding principles, authorized tools, and permitted and prohibited use cases, and incident reporting channels; (ii) periodic updates (at least annually) when new tools are added to the corporate inventory, when our AI Policy is updated, or when relevant risks or vulnerabilities are identified in operations; (iii) specialized training for those responsible for AI models and for technical staff who develop, configure, or administer AI systems, focused on security controls, risk management, and regulatory compliance; and (iv) a documented record of training completed by each employee as evidence of compliance.

Compliance with the training plan is a prerequisite for being granted access to corporate AI tools.

International Frameworks and Reference Standards

Our Responsible Use of AI Policy and its corresponding implementation have been designed in strict alignment with Mexican regulations and the most relevant international standards on the subject, including the United Nations Principles for the Ethical Use of Artificial Intelligence, ISO/IEC 42001, NIST AI RMF, COSO ERM, and the European Union's Artificial Intelligence Act (EU AI Act), among others, and takes into account the specific obligations arising from our regulatory exposure as a publicly traded company.