

Theoretical and Technical Foundations of IGS

Contemporary AI Governance Frameworks as Benchmarks

The IGS is designed to be compatible with, and operationalizable under, the major AI governance frameworks that are emerging as global reference industry standards. While these frameworks differ in legal force, scope, and jurisdiction, they share several structural themes, such as:

- Risk-based Oversight
- Lifecycle Oriented Controls
- Management-system Thinking
- Documentation and Traceability Expectations
- Explicit Human Accountability (for AI aided decision making processes and outcomes)

IGS does not replace these frameworks. It provides an internal governance architecture that organizations, especially SMEs, can use to implement and document compliance in accordance with their standards.

1. EU AI Act

The EU AI Act is the most comprehensive risk-based regulatory framework for AI that is currently enforceable by law. It classifies AI systems into risk tiers (unacceptable, high-risk, limited-risk, and minimal-risk); obligations scale in proportion to the potential impact to health, safety, and fundamental human rights. For example, for high-risk systems, the Act requires:

- A documented risk management system as a continuous process across the lifecycle of AI use
- Systemic identification, analysis, and mitigation of risks
- Robust data governance and technical documentation
- Logging, transparency, and human oversight mechanisms
- Post-market monitoring and incident reporting

Recent policy discussions and the “Digital Omnibus” initiative have focused on adjusting timelines and easing aspects of implementation for certain high-risk areas, but the core architecture and obligations remain static.

2. NIST AI Risk Management Framework (RMF)

The NIST AI RMF is a voluntary, sector-agnostic framework intended to help organizations manage AI risks and build trustworthy systems for everyday use. It centers on four core functions:

- Govern
- Map
- Measure
- Manage

Each function is broken into categories and sub-categories that describe outcomes for risk-aware AI development and deployment.

In 2024, NIST released a Generative AI Profile (NIST AI-600-1) to tailor the AI RMF to generative and large-model use cases, further emphasizing context specific risk management practices.

3. ISO/IEC 42001: AI Management System

ISO/IEC 42001:2023 is the first international standard for an artificial intelligence management system (AIMS). It specifies requirements for establishing, implementing, maintaining, and continually improving an AI management system within an organization providing or using AI-based products or services. Key concepts include:

- Integrating AI into broader management-system structures (such as ISO 9001)
- Defining policies, roles, and responsibilities for AI governance
- Addressing ethical, transparency, and accountability decisions
- Embedding continuous improvement practices into all operations using AI assistance

As with other ISO/IEC standards, ISO 42001 serves as an organizational reference standard for deployment of a traceable, auditable, and repeatable processes integrated with continuous-improvement practices for any business procedure using AI systems.

4. TRAIGA (Texas Responsible Artificial intelligence Governance Act)

The TRAIGA is one of the most significant state-level AI laws in the United States. Passed via House Bill 149 in 2025, it establishes governance, disclosure, and use-case restrictions primarily focused on governmental and other high-impact AI uses within Texas' regulatory scope. Core features include:

- Governance obligations for state level entities deploying AI
- Disclosure and transparency requirements when AI is used in public-facing or decision-making contexts
- Prohibitions on specific practices (such as social scoring by governmental entities, certain biometric uses for policing citizens)
- Penalties and enforcement clauses for misuse

While TRAIGA is not as comprehensive as the AI Act, it signals growing societal expectations that organizations, especially those interacting with the public, must document AI use, implement transparency and disclosure mechanisms, and maintain governance structures proportional to risk-based assessments.

5. Additional Reference Frameworks

In addition to the binding and voluntarily binding frameworks above, two additional benchmarks influence national and corporate AI policies globally.

- Organization for Economic Co-operation and Development (OECD)'s AI Principles: The first intergovernmental standard on AI, adopted in 2019 and updated again in 2024. The Principles promote innovate, trustworthy AI that respects human rights and democratic values, and they also provide guidance on responsible stewardship, transparency, robustness, and accountability.
- Council of Europe Framework Convention on Artificial Intelligence: The first legally binding international treaty on AI, opened for signature in September 2024. It requires that AI systems be consistent with human rights, democracy, and the rule of law. It adopts a risk-based approach that complements the EU AI Act and other regional efforts.

These instruments reinforce the normative direction that:

- AI systems must be aligned with fundamental rights and democratic values.
- Accountability, transparency, and meaningful human oversight are not optional.
- Risk-based, lifecycle governance is the emerging global baseline.

Together, these frameworks define what responsible AI governance must achieve; however, they do not prescribe how organizations, particularly SMEs, should structure their internal systems to operationalize those requirements.