

cortAlx

Artificial Intelligence by THALES

Reconciling AI sovereignty and exportability

with purpose-built models and data strategies

Reconciling AI sovereignty and exportability with purpose-built models and data strategies

Authors:

Boussad ADDAD, Rémi BLANCHETTE, Dave COUTURE, Fabien FLACHER,
Katarzyna KAPUSTA, Juliette MATTIOLI, Gabriel RANGONI

Summary

In highly sovereign domains such as Defense, artificial intelligence has demonstrated its strategic value and is now in high demand. Given the strong dependence of data-driven AI on high-quality datasets, closer collaboration between system providers and armed forces—who possess and control the largest pools of operational data—is essential.

A central challenge lies in defining the conditions under which such collaboration can occur safely and effectively, both at national and international levels, to develop high-performance AI models.

This position paper presents a pragmatic and technical framework for reconciling AI sovereignty with the exportability of advanced AI-enabled systems. It identifies the key technological enablers that make sovereign AI operationally viable. These include secure data-sharing mechanisms (such as transfer learning and collaborative learning), privacy-preserving techniques, and traceability tools like watermarking.

A cornerstone of this approach is the development of a sovereign AIOps toolchain, designed to ensure robust data governance, model integrity, and continuous retraining—even in disconnected or sensitive environments.

The paper also addresses cybersecurity risks specific to AI, proposing countermeasures to prevent data leaks, model theft, and adversarial manipulation. Finally, it emphasizes the importance of infrastructure choices, deployment models, and legal frameworks to support trusted international cooperation without compromising sovereignty.

1. Context

1.1. SOVEREIGNTY

AI sovereignty refers to the sustained ability of a nation—or group of nations—to independently govern the development, deployment, and regulation of artificial intelligence within its borders. This autonomy protects against unsolicited external influence and spans four key dimensions:

- **Technological** sovereignty involves mastering advanced AI capabilities across hardware and software, securing access to critical infrastructure, and cultivating national expertise to reduce reliance on foreign providers.
- **Regulatory** sovereignty requires establishing legal frameworks that promote responsible AI use, safeguard privacy and fundamental rights, stimulate domestic innovation, and mitigate foreign extraterritorial jurisdiction's reach when it cannot be avoided.
- **Economic** sovereignty ensures control over the commercial and industrial exploitation of AI, supports national champions, and enables balanced participation in global value chains.
- **Geopolitical** sovereignty encompasses influence over international standards and strategic positioning in global tech competition. It also includes cultural and linguistic considerations—particularly relevant in the context of large language models.

Achieving AI sovereignty is essential to maintain autonomous control over AI-enabled systems. Technologically, this demands a focus on autonomy, trustworthiness, reliability, confidentiality, and cyber resilience. The more of these elements a nation masters, the better positioned it is to engage with externally dependent parts of the AI supply chain from a position of strength.

1.2. AI SOVEREIGNTY CHALLENGES

Sovereignty is a central concern in the defense sector, where Thales operates extensively. AI-enabled

defense systems inherit all the sovereignty requirements of other advanced defense technologies, while introducing additional, domain-specific challenges. Nations and organizations acquiring such systems typically seek to maximize autonomous control across four foundational pillars: Client Autonomy, Information Confidentiality, AI Trustworthiness, and Reliable Performance.

Information Confidentiality

Ensuring confidentiality in sovereign AI systems requires a comprehensive, lifecycle-wide approach. By adhering to standards such as IEEE-7000 and implementing robust security measures, organizations can protect sensitive information from unauthorized access and foreign intelligence threats. This safeguards not only individual privacy or export control compliance but also strategic advantages and intellectual property. Confidentiality mechanisms must shield data, models, and insights from surveillance and espionage.

Thales' sovereign AI systems are designed to uphold confidentiality through rigorous protection against unauthorized access and foreign interference.

Trustable AI

Trustworthiness in sovereign AI systems is achieved through a holistic strategy that secures the integrity of models, data, processes, operations, institutions, and supply chains. These interconnected safeguards ensure that AI systems remain authentic, uncompromised, and aligned with their intended functions. This foundation is essential for building trust, enabling widespread adoption, and ensuring positive societal impact.

Thales' TrUE AI framework—Transparent, Understandable, and Ethical—embodies this commitment to trustable AI.



Figure 1 – Thales TrUE AI framework

1.3. AI SOVEREIGNTY CALLS FOR MORE COLLABORATION

AI systems are distinguished by their capacity to evolve over time—provided they are trained regularly with rich, diverse, and up-to-date datasets. This dynamic adaptability offers a decisive advantage over other technologies.

In today's fast-paced AI innovation landscape, model design and implementation also present significant opportunities for continuous improvement, particularly as hardware capabilities advance. Two key stakeholders drive this dual lever of progress: the users, who operate systems in real-world conditions and provide feedback, and the developers, who design and maintain the AI models and their hosting environments notably based on that return of experience.

Sustained collaboration between these actors is essential to ensure long-term system performance. Users contribute operational data that enhances model retraining, while developers refine architectures and implementation to optimize reliability and efficiency. This cooperation must be enabled by dedicated tools and processes—even in environments where sovereignty is a critical concern.

This principle applies domestically and is equally relevant internationally. AI models benefit from broader training datasets sourced from diverse operational contexts worldwide. Model quality is strongly correlated with the volume, diversity, and accuracy of training data. Techniques such as frugal learning, neuro-symbolic approaches, and the use of synthetic data can mitigate—but not eliminate—the need for real data. Ultimately, richer datasets expand operational domains and improve system outputs.

From an economic standpoint, wider model distribution through exports helps offset the substantial investments required to maintain technological superiority. However, sharing operational datasets in defense contexts demands extreme caution from both the country of origin and the country of destination, under strict secret protection and trade compliance legal framework.

We believe that with the right technologies and a pragmatic approach, reconciling sovereignty and exportability becomes achievable.

In practice, stakeholders recognize the mutual benefits of collaboration, though their priorities differ:

Countries of origin and destination expect:

- Immediate access to state-of-the-art systems with proven reliability.
- The ability to enhance performance using their own operational data.
- Strong guarantees that their data will be protected from unauthorized use—even within cooperative frameworks.
- Robust system resilience against external interference.

Countries of origin, in the context of export, also expect:

- Clear understanding and potential adaptation of system performance prior to export.
- Direct or indirect benefits from improvements gained through diverse operational use.
- Strict enforcement of agreements with destination countries, including secret protection and trade compliance related commitments if any, supported by supplier-provided AI procedures and tooling.

Suppliers, in turn, expect access to operational data necessary to meet these expectations—both domestically and internationally.

Despite differing priorities, alignment is possible. With appropriate workflows, sovereign AI operations toolchains, and advanced cybersecurity technologies, all parties can collaborate effectively.

Overall collaboration scheme

Diagram 1 below outlines a structured collaboration framework designed to reconcile AI sovereignty with exportability.

1. Initial Model Development

The AI model is initially developed and trained using data owned or accessible by the system provider (e.g., Thales).

2. Sovereign Data Integration by Country of Origin

The country of origin authorizes the supplier to enhance the model using its own operational data, excluding highly sensitive information. This process includes robust access controls and data lineage tracking. Such integration may occur at the outset and continue throughout the model's lifecycle, based on operational feedback.

3. Collaborative Performance Optimization

Following retraining with additional data, the model achieves its highest collaborative performance level.

4. Local-Only Retraining by Country of Origin

In certain cases, the country of origin may request the ability to further train the model using data it does not wish to share. The supplier may provide secure AI training tools to support this local-only retraining.

5. Export Model Preparation

To facilitate international cooperation, the country of origin may authorize the supplier to use a subset of its data to train a version of the model intended for export.

6. Destination Country Collaboration

Similar to domestic use, the destination country may collaborate with the supplier to retrain the model

using its own data, tailoring performance to its specific operational context.

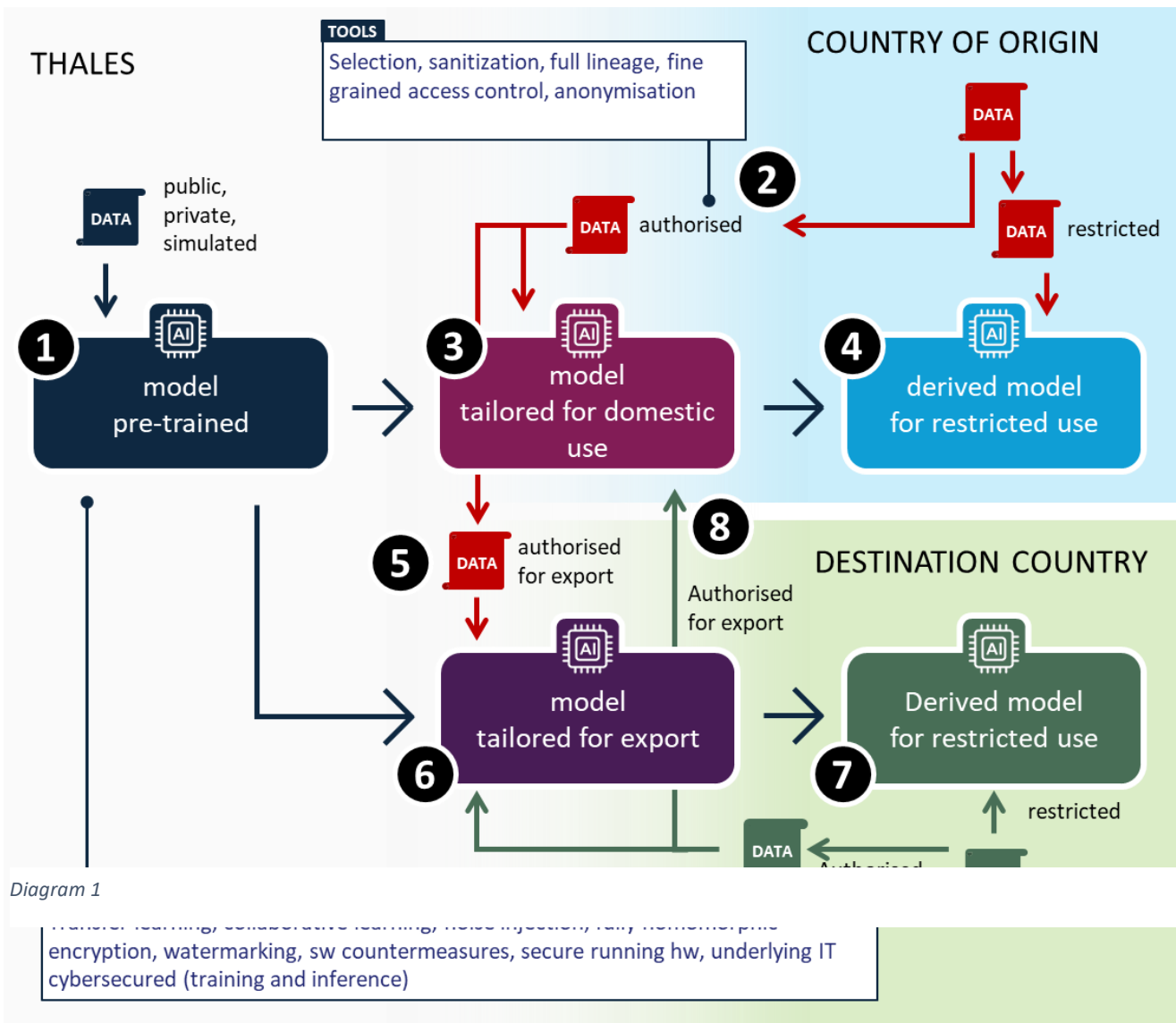
7. Local-Only Retraining by Destination Country

The destination country may also request the ability to further train the model with data it does not intend to share with the supplier.

8. Reciprocal Data Contribution

The destination country may authorize the supplier to integrate part of its data into the domestic model, reinforcing its collaboration with the country of origin. This two-way exchange maximizes system performance.

Enabling these successive training steps and secure data exchanges requires the deployment of specific technologies:



- **Collaborative AI Toolchains:** Platforms that support secure data preparation, sanitization, sharing, and management across stakeholders.
- **Secure AI Models:** Architectures equipped with traceability and protection mechanisms, leveraging advanced cybersecurity and cryptographic techniques.
- **Enforced Secure Processes:** Tooling and workflows that ensure compliance and security throughout the development and training lifecycle, without incurring excessive costs.

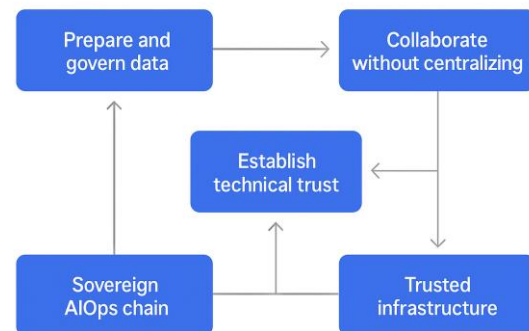
The following chapters provide technical details on how these capabilities can be implemented.

Beyond the technical dimension, this collaborative model significantly reshapes both operational and business relationships between suppliers and customers. It necessitates dedicated framework agreements that define the conditions under which data and models are accessed, used, modified, and shared—ensuring systems remain in fair and optimal operational condition for all parties involved.

This may require bilateral or multilateral General Security Agreements and a set of export control licenses with potential strict associated conditions that should be anticipated and tailored with regards to the need. This may require important preparation and coordination work between the supplier, the country of origin, and, later on, with the country of destination. Any third country's or extraterritorial potentially applicable trade compliance's jurisdictions should be assessed and taken into account prior to and during the cooperation. Should AI be or not be part of the project, this framework preparation phase and its monitoring all along the projects shall not be underestimated in order to secure the fluidity and success of the cooperation.

2. Approach overview

The proposed method follows an industrial and incremental logic, structured around five key technical pillars that together enable sovereign and exportable AI systems:



Framework for data collaboration in sovereign environments

Prepare and Govern the Data

The first step is to mitigate risks at the source. Before any data sharing occurs, a **sanitation pipeline** must be implemented to retain valuable learning information while removing sensitive elements—such as unnecessary metadata, geographic indicators, or critical outliers. This process improves data quality and reduces reliance on heavy confidentiality techniques.

Strong governance accompanies this preparation: versioning and lineage tracking ensure traceability of transformations; access is controlled via role-based (RBAC) and attribute-based (ABAC) mechanisms; secrets are encrypted and managed securely. A versioned and encrypted feature store guarantees consistency between training and production environments. The result is a set of cooperative datasets with documented and auditable risk levels.

Collaborate Without Centralizing

High performance does not require centralizing data. Two preferred approaches are transfer learning and collaborative learning.

Transfer learning involves sharing base models trained on non-sensitive data, which are then locally adapted using sovereign datasets. This method allows pooling of scientific investment without moving critical data.

Collaborative learning enables distributed training (e.g., federated or split learning), where each participant computes local updates and shares only model parameters. These exchanges are secured through cryptographic techniques (e.g., MPC, FHE), integrity checks, and mechanisms for revocation or unlearning to manage evolving partnerships.

Establish Technical Trust

Effective collaboration depends on verifiable guarantees. Confidentiality is ensured through encryption at rest and in transit, secure execution environments (e.g., TEE), and privacy-preserving techniques such as homomorphic encryption and differential privacy.

Intellectual property and trade compliance authorized end use are protected via model watermarking and signed logs that capture the training context. Model robustness is reinforced through adversarial training, trigger detection, and regular testing against attacks such as inversion, extraction, and poisoning. These mechanisms transform trust from a declaration into an auditable technical contract.

Industrialize with a Sovereign AIOps Chain

Sovereignty is operationalized through an AIOps toolchain that automates and traces every step of the AI lifecycle. It includes:

- Supply chain integrity (e.g., signed logs, SBOM/AIBOM)
- On-prem orchestration with codified pipelines and ethical controls
- Offline-friendly deployment
- Continuous monitoring (e.g., drift detection, supervised retraining)

Policies are translated into code to ensure compliance with operational constraints (ODD) and activate safe modes in case of breach. This sovereign AI factory enables the construction, evaluation, and maintenance of cooperative models without loss of control.

Deploy on Trusted Infrastructure

Infrastructure must align with both technological and jurisdictional requirements. A multi-layered architecture—combining secure hardware, virtualization, and containerization—is essential. Trust zones and deployment models (on-prem, hybrid, edge-centric) must be selected to match sovereignty needs.

Model and data exchanges are conducted via encrypted and signed packages, orchestrated by the AIOps chain. This ensures confidentiality, traceability, and resilience against threats.

The proposed method is based on an industrial and gradual logic, structured around five technical pillars: data preparation and governance, choice of collaboration paradigms, trust mechanisms, sovereign AIOps chain, and controlled infrastructure.

3. Enabling framework and technologies

Organizations and nations may hesitate to share training data, fearing it could compromise their autonomy or expose confidential information. While sharing pre-trained models might appear to be a safer alternative, it can also inadvertently reveal insights about the underlying data—unless done securely.

Yet, the promise of enhanced performance through collaboration remains a strong incentive to explore secure methods for sharing AI knowledge and data. This need is driving the development of new technologies that support sovereign-compatible cooperation.

Transfer learning and collaborative learning stand out as the two foundational paradigms for secure data collaboration. Thales identifies these approaches as essential to building sovereign-ready AI platforms. They enable multiple stakeholders to jointly train models without exchanging raw data, thereby preserving sovereignty while maintaining high standards of privacy and security.

These frameworks can be reinforced by a set of trust enablers, including:

- **Data sanitization:** removing or anonymizing sensitive elements before sharing.
- **Cryptographic safeguards:** securing data and model exchanges.
- **Anonymization techniques:** protecting individual data points.
- **Traceability mechanisms:** ensuring transparency and accountability throughout the AI lifecycle.

Together, these technologies form the backbone of a secure and sovereign AI collaboration environment.

3.1. TRAINING DATA SANITIZATION

Regardless of the collaboration framework adopted for AI model training, the first essential step is to assess the extent to which available and relevant data can be securely shared with third parties—whether they are model suppliers or export partners.

This evaluation must maintain consistency between the authorization to share trained models and the authorization to share the underlying training data. In many cases, AI-enabled equipment approved for export delivers more value to the acquirer than its individual components, including the raw data used to train the model. Once trained, the model becomes a functional embodiment of the knowledge it has absorbed. Therefore, the data strictly used to develop export-authorized equipment may often be considered suitable for export as well.

However, some datasets may contain information that does not directly contribute to model performance but could reveal sensitive details—such as characteristics of the data capture chain, sensor specifications, or mission locations and timelines. These elements must be excluded or anonymized during the data preparation phase to prevent unnecessary exposure.

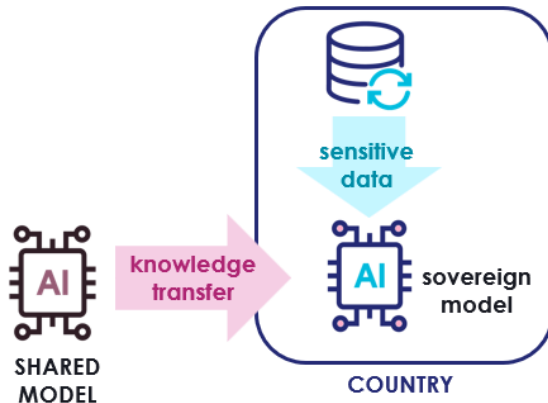
Properly executed data sanitization not only protects confidentiality but also reduces the need for heavy privacy-enhancing techniques during training. It also helps identify and eliminate false or contradictory data that could degrade model performance.

Once this critical phase is complete, access to the sanitized training dataset must remain tightly controlled—both during the training process and during inference. In certain contexts, this protection may extend to limiting access even among trusted partners.

3.2. TRANSFER LEARNING

Transfer learning offers a balanced approach that combines international cooperation with national sovereignty. By sharing foundational models trained on public or anonymized datasets, countries and organizations can collectively enhance their AI capabilities while retaining control over sensitive local

data.



This collaborative method begins with the development of shared base models, built from non-sensitive public data or anonymized private datasets. These models reflect pooled expertise and resources, forming a common technological foundation that would be prohibitively expensive for individual actors to build alone.

Each nation or organization can then fine-tune these shared models using its own sovereign data. This dual approach—leveraging collective knowledge while preserving local control—enables performance optimization without compromising confidentiality or regulatory compliance. It creates a virtuous cycle: international cooperation accelerates foundational model development, and all participants benefit from the resulting improvements.

The format used to share models plays a critical role. Open standards such as ONNX (Open Neural Network Exchange) should be prioritized, as they allow models trained in one framework (e.g., PyTorch or TensorFlow) to be executed in another. This ensures interoperability across tools, platforms, and hardware.

By reducing dependency on specific vendors, ONNX strengthens technological sovereignty and gives organizations greater control over their infrastructure and data flows. Its open and standardized nature supports flexible deployment—whether on-premises or in preferred cloud environments—while meeting data residency and compliance requirements. It also enhances transparency and auditability, reinforcing both security and regulatory efforts.

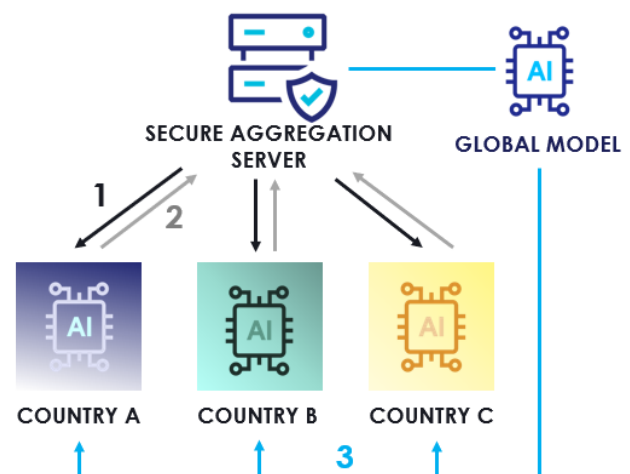
3.3. COLLABORATIVE LEARNING

Collaborative learning offers a powerful framework for enabling cooperation between nations and organizations while preserving data sovereignty. It

allows multiple stakeholders to jointly train AI models on distributed datasets without ever centralizing or explicitly exchanging sensitive, export controlled or classified information.

Training is conducted in a decentralized manner: each participant maintains its own local model and shares only model-related updates (gradients) with an aggregator server. The most widely adopted form of collaborative learning is Federated Learning, where all participants use the same model architecture. This approach comes in various configurations depending on data distribution (horizontal or vertical), the presence of a central coordinating entity, and the number of participating devices (e.g., numerous IoT devices under a single owner, or multiple servers across different organizations).

Another promising technique is Split Learning, which distributes model training responsibilities so that only intermediate activations—not raw data—are exchanged. This is particularly suited to scenarios where data or model components must remain strictly segregated.



Thales enhances these baseline collaborative learning pipelines with advanced capabilities such as cryptographic signing, provenance tracking, and immutable model registries. These features ensure the integrity, authenticity, and traceability of all AI components throughout their lifecycle.

While research continues to refine the security guarantees of Federated Learning, existing techniques already offer substantial protection. For example, privacy-preserving methods can be integrated to reduce the risk of data leakage through AI-specific attacks on the central aggregation server—thus

safeguarding sensitive information even in cross-border collaborations.

To manage ownership rights, ML watermarking can be employed. This technique allows each participant to embed a verifiable contribution into the final model, ensuring transparent attribution and enabling secure collaboration. Combined with unlearning capabilities, it also provides flexibility to revoke contributions if needed.

Altogether, collaborative learning—reinforced by cryptographic safeguards and traceability mechanisms—empowers organizations to train high-performance models with confidence, knowing that every stage of the process is anchored in transparency, auditability, and resilience.

3.4. TRUST ENABLERS

To enable secure and sovereign AI collaboration, trust must be embedded at the technical level. This requires a combination of cryptographic protections, secure hardware, and traceability mechanisms that ensure data integrity, model authenticity, , export controlled and ownership rights.

Cryptography and Secure Hardware

Cryptographic techniques play a central role in securing AI systems. They protect data and models both at rest and in transit. Advanced methods such as Homomorphic Encryption and Multi-Party Computation (MPC) extend this protection to data in use—during training or inference—though they may introduce computational overhead. In some cases, lighter alternatives like obfuscation can be used to balance performance and security.

Secure hardware also contributes to trust. Technologies such as Trusted Execution Environments (TEEs) provide isolated environments for sensitive computations, shielding them even from the host operating system. These hardware-based protections reinforce sovereignty by ensuring that critical operations remain confidential and tamper-resistant.

Traceability and Ownership Protection

Collaboration across organizations and nations raises important questions about the traceability of contributions and the protection of intellectual property. Traditional digital rights management (DRM) mechanisms—such as copyrights or patents—are

often ill-suited to the machine learning context, where models can be copied, fine-tuned, or extracted via AI-specific attacks.

To address this, models can be encrypted or obfuscated to prevent unauthorized replication. Additionally, ML watermarking offers a tailored solution: it embeds persistent, verifiable marks into models that remain detectable even after retraining. These watermarks serve as proof of ownership and can be used to assert rights over a model's origin and integrity.

Together, these trust enablers form the foundation of a secure and auditable AI collaboration environment—one that respects sovereignty, protects sensitive assets, enable export control compliance and ensures accountability across all stages of the AI lifecycle.

4. AI toolchain adapted to critical data and cooperation

AI sovereignty cannot be achieved through policy declarations alone—it requires a robust industrial infrastructure. At the heart of this infrastructure lies a dedicated AIOps toolchain, which acts as an automated build-and-run factory. It secures data, validates supply chains, and ensures operational resilience through continuous monitoring and retraining.

More importantly, this toolchain provides a common platform for allied nations to consolidate fragmented efforts into a unified, sovereign AI capability—strong enough to compete with global technology leaders. In the race for AI leadership, a sovereign AIOps chain is not a luxury; it is a strategic necessity, despite the complexity it entails.

AI differs fundamentally from previous waves of technological innovation. While hardware engineering follows predictable physical laws and software development relies on deterministic logic, AI introduces new layers of complexity. Learning models depend on data distributions, stochastic optimization, and iterative experimentation. Their performance may degrade over time, requiring constant monitoring, retraining, and redeployment.

This fragility makes AI systems highly dependent on their operational backbone—and dangerously exposed if that backbone relies on foreign platforms. In this context, an AIOps toolchain is not merely a technical convenience; it is the foundation of AI sovereignty.

Unlike traditional DevOps pipelines, an AIOps toolchain addresses the unique challenges of AI projects. It manages data versioning, tracks experiments and models, automates compliance with ethical and regulatory standards, and supports robust performance monitoring and retraining.

In defense applications, the toolchain must be tailored to mission-critical and sensitive environments, including air-gapped or disconnected systems. It must support on-premises deployment, enforce strict

access controls, and ensure full traceability of data and model artifacts.

Security policies, operational procedures, and compliance gates are embedded into automated workflows, supporting not only technical requirements but also legal, export control, and sovereignty obligations.

Properly designed, this toolchain delivers three strategic capabilities:

- Security and control of data assets
- Integrity of the AI supply chain
- Continuous sovereign operation

It serves as both a productivity accelerator and a trusted foundation, ensuring that AI systems are industrial-grade, secure, traceable, and aligned with internal and external standards for ethical and trustworthy AI.

4.1. SECURITY AND CONTROL OF DATA ASSETS

In AI systems, data is both the most strategic asset and the most vulnerable. To ensure its protection, a sovereign AIOps toolchain must provide a comprehensive set of capabilities that secure data throughout its lifecycle.

Data Governance and Lineage

A local, versioned data lake must record the origin, approvals, licenses, and full lineage of all datasets. Automated quality checks help prevent license contamination and support the creation of reproducible datasets—essential for audits and export reviews. For example, this enables proof that a defense model was trained exclusively on authorized sources.

Secrets Management and Hardware Security

Sensitive credentials must be securely stored using hardware-backed vaults such as Hardware Security Modules (HSMs) or Trusted Platform Modules (TPMs). Encryption per environment, network segmentation, and runtime scanning reduce the risk of credential leakage and lateral movement within sensitive infrastructures.

Fine-Grained Access Control

Access to AI inputs and outputs must be governed by Role-Based Access Control (RBAC) and Attribute-Based Access Control (ABAC) models. These mechanisms

ensure that users can only access data, run jobs, or test models within their authorized sensitivity levels. The system also evaluates model sensitivity independently of training data, elevating restrictions when emergent capabilities or deployment contexts require it.

Policies-as-code align user attributes (e.g., role, clearance, project, environment) with the sensitivity of operations and artifacts. This prevents unauthorized access, restricts high-impact actions by under-cleared users, and enables automatic revocation and full traceability for audits or incident response.

Feature Management

A local feature store ensures consistency between offline and online environments, supports encryption at rest, and extends fine-grained access controls. It prevents training-serving mismatches, avoids leakage of sensitive signals, and clarifies ownership and update policies—for example, restricting healthcare-related features to specific roles.

4.2. INTEGRITY OF THE AI SUPPLY CHAIN

AI development relies on complex and often opaque supply chains—including open-source frameworks, pre-trained models, and third-party services. Without rigorous oversight, these dependencies can introduce hidden vulnerabilities that compromise sovereignty and security.

A sovereign AIOps toolchain must validate every component, enforce compliance, and prevent malicious or unauthorized insertions into critical systems. This is achieved through a combination of secure artifact management and robust experiment tracking.

Secure Artifact Supply Chain

A local registry must store all code, models, and executables with cryptographic signing and provenance attestations (e.g., SLSA-style). The toolchain should also generate AI Bills of Materials (AI-BOMs) to establish end-to-end trust. These measures prevent supply chain attacks and ensure that deployed models and containers match their verified sources.

Experiment Tracking and Reproducibility

A self-hosted tracker (e.g., MLflow) logs all training and evaluation runs, including code versions, parameters,

dataset snapshots, environment configurations, and resulting artifacts. Deterministic builds are enforced through pinned dependencies, and each run is accompanied by a generated SBOM or AI-BOM. This enables reproducibility, facilitates comparison, and supports efficient exploration and design.

Model Registry with Provenance

Models must be stored alongside their training context, data references, evaluation reports, and digital signatures. This ensures full lifecycle governance and allows only reviewed, traceable models to progress toward deployment. Multi-party approvals and controlled rollback mechanisms reinforce trust and accountability.

Orchestration and Automation

On-prem orchestration tools (e.g., Airflow, Argo) execute codified pipelines with embedded policies-as-code. These pipelines enforce quality gates, approvals, and compliance checks consistently and transparently. Automation eliminates manual errors and ensures that privacy, safety, and regulatory requirements are met at every stage.

4.3. CONTINUOUS SOVEREIGN OPERATIONS

AI models require ongoing attention—they cannot simply be deployed and left unattended. Real-world conditions evolve constantly: operational environments shift, and unexpected events can significantly alter data patterns and system behavior. These changes may lead to performance degradation or prediction drift, which must be detected and addressed promptly.

A robust and sovereign AIOps toolchain is essential to manage these challenges. It must support continuous monitoring of models in production, enabling early detection of performance drops and contextual shifts. When such changes occur, the system should facilitate efficient retraining and redeployment using updated datasets.

Fail-safe mechanisms are also critical to ensure operational safety in case of sudden performance anomalies. Key functionalities include:

Deployment and Runtime Management:

Deployment templates must support diverse targets—embedded systems, edge devices, private clusters—with offline-friendly rollout and rollback strategies. This narrows the gap between lab and field environments and enables secure updates without internet access (e.g., via signed packages delivered through physical media).

Monitoring, Drift Detection, and Feedback Loops:

On-prem telemetry tools (e.g., Prometheus, OpenTelemetry) and drift detection frameworks (e.g., Evidently, Alibi Detect) track data shifts, prediction anomalies, and safety events. These tools provide early warnings, trigger retraining under defined constraints, and accelerate incident response.

ODD Compliance and Fail-Safe Orchestration:

Operational Design Domain (ODD) constraints—such as geofencing, environmental conditions, sensor health, latency, and rules of engagement—must be codified as policies and enforced at runtime. This ensures that models operate within their declared boundaries and activates controlled, auditable fail-safe responses when thresholds are approached or breached.

Together, these capabilities enable continuous sovereign operation of AI systems. They ensure that models remain reliable, secure, and compliant throughout their lifecycle—even in disconnected or sensitive environments. This operational backbone is essential to maintain trust, agility, and sovereignty in mission-critical deployments.

5. Cybersecured AI model

Cybersecurity is a cornerstone of trustworthy AI. Unlike traditional software systems, AI solutions are inherently data-driven and follow a lifecycle that includes data collection, model architecture design, training, deployment, and often retraining or fine-tuning. This expanded lifecycle increases the attack surface, exposing models to threats at multiple stages.

AI systems also present unique vulnerabilities. Beyond conventional risks, they are susceptible to AI-specific attacks such as adversarial examples—inputs crafted to deceive the model—and model inversion techniques that attempt to reconstruct training data from model outputs.

In the context of reconciling sovereignty and exportability, a key concern is the potential exposure of sovereign information through exported AI systems. This risk drives the development of privacy-enhancing technologies, which allow data to be exploited for learning while anonymizing individual samples. These techniques help determine which parts of the information can be securely shared across borders.

Other critical cybersecurity challenges include:

- Protecting ownership rights over data and models.
- Preventing behavioral deviations caused by external manipulations.

5.1. REVISITING THE SECURITY ANALYSIS

The classic Confidentiality–Integrity–Availability (CIA) triad remains a foundational model in information security. In the context of AI, it is often extended with a fourth dimension: Ownership, which addresses the risks of theft or misuse of AI knowledge. While privacy is closely linked to confidentiality, privacy-enhancing techniques go beyond protecting personal data—they aim to prevent any form of information leakage, which is sometimes misunderstood.

A robust AI security strategy begins by identifying the assets to be protected—typically the training data, input data, and the model itself. This is followed by the construction of a threat model, which evaluates

potential attack vectors and their impact. Based on this analysis, appropriate countermeasures are applied. These fall into two categories:

- Model-enhancement techniques, which modify the AI system itself and often require retraining.
- Model-agnostic techniques, such as query filtering, which can be implemented externally without altering the model.

Although AI cybersecurity is still an emerging field, it is evolving rapidly—especially in the area of generative AI. Over the past few years, several authoritative bodies have published guidelines and recommendations, including ETSI TC SAI, NIST, MITRE ATLAS, and ANSSI. These efforts are paving the way for future certification schemes that will formalize AI security validation and harmonize cybersecurity metrics across the industry.

5.2. AVOIDING DATA LEAKS FROM SHARED MODELS

Data leakage remains one of the most critical threats to the confidentiality of AI systems—especially when models are shared across organizations or borders. While traditional risks stem from unsecured data supply chains, more subtle and sophisticated threats arise from the models themselves.

Indeed, machine learning models often retain traces of the data they were trained on. Attackers with access to a deployed model may exploit this through techniques such as:

- **Model inversion**, which attempts to reconstruct original training samples from model outputs.
- **Membership inference**, which determines whether a specific data point was part of the training set.
- **Property inference**, which seeks to extract general characteristics of the training dataset.

These attacks can compromise sensitive information even if the model has been fine-tuned on new data before being shared. In 2023, Thales' Friendly Hackers team demonstrated this risk by successfully reconstructing confidential training data during the

CAID challenge—highlighting the importance of robust protections.¹

To mitigate such risks, several strategies can be employed:

- **Differential Privacy**: This technique introduces controlled noise into inputs, outputs, or training algorithms, making it statistically impossible to determine whether a specific data point was used during training. While effective, it may impact model performance and should be applied after other cybersecurity measures, such as data sanitization.
- **Machine Unlearning**: This emerging approach allows the removal of specific data influences from a trained model without requiring full retraining. It supports the “right to be forgotten” and is particularly useful when data becomes obsolete or when models need to be exported without retaining sensitive information.

Currently, verifying that a model does not leak undesirable information relies primarily on differential privacy guarantees and empirical penetration testing. Because quantifying privacy levels remains a research challenge, it is essential to exclude highly sensitive data—such as outliers—during the preprocessing phase.

5.3. SHARING MODELS WHILE KEEPING THEM SECRET

Some organizations or nations may wish to make AI models available to external partners while keeping their internal architecture, parameters, or input/output data confidential. Conversely, recipients may want to use these models without revealing their own queries or data inputs.

Ensuring the confidentiality of shared models presents a distinct challenge from protecting training data. The goal is to allow models to be executed on untrusted devices without exposing their internal structure or the data they process. This is conceptually similar to protecting software code during execution.

One theoretical solution is Fully Homomorphic Encryption (FHE), which enables computations to be performed directly on encrypted data. While recent

¹ AI Hacking: Thales Wins the DGA Challenge and Showcases Its Enhanced AI Security Solutions for Military and Civilian Applications | Thales Group

advances have improved its applicability to AI, FHE remains computationally intensive and is currently best suited to smaller, less complex models.

As a more practical alternative, obfuscation techniques can offer a moderate level of protection at a lower computational cost, though they provide weaker security guarantees. Another viable option is the use of Trusted Execution Environments (TEEs)—secure hardware enclaves that isolate sensitive computations. Technologies such as Intel SGX, AMD SEV, and ARM TrustZone allow models to run securely, even on potentially untrusted platforms.

While TEEs are more mature than FHE implementations, they come with trade-offs, including dependency on specific hardware vendors and susceptibility to side-channel attacks. Nonetheless, they represent a valuable tool for enforcing sovereignty and confidentiality in cross-border AI deployments.

5.4. FACING AI-SPECIFIC ATTACKS ON INTEGRITY

Ensuring the integrity of AI models is critical, particularly in sovereign and defense contexts. Among the most concerning threats are data poisoning and backdoor attacks, which can compromise model behavior and reliability.

- **Data poisoning** involves the deliberate manipulation of training datasets to cause the model to behave incorrectly or unpredictably.
- **Backdoor attacks** embed hidden triggers during training that, when activated by specific inputs, cause the model to produce malicious or unintended outputs. These backdoors can propagate through transfer learning or collaborative training, making them especially dangerous when models are reused from external sources.

The most effective mitigation strategy is to secure the data supply chain from the outset. Sampling data from controlled environments significantly reduces the risk of manipulation. Additional countermeasures can be applied during inference, including:

- **Data sanitization**, which filters out potentially harmful inputs.
- **Trigger detection**, which identifies and neutralizes backdoor activation patterns.

Together, these techniques help preserve model integrity and reinforce trust—especially when models are shared or adapted across sovereign boundaries.

5.5. ENSURING ROBUSTNESS OF THE DEPLOYED MODELS

Neural networks are inherently vulnerable to adversarial examples—inputs that are subtly modified to mislead the model while appearing unchanged to human observers. This threat is particularly pronounced in computer vision applications, where minute pixel alterations can cause a model to misclassify an image—for instance, interpreting a tank as an ambulance due to imperceptible perturbations.

Such vulnerabilities expose deployed models to remote attacks, especially when adversaries possess insights into the model’s architecture or weaknesses. This poses a serious challenge to sovereignty and trust, particularly in defense and allied contexts.

To mitigate these risks, several countermeasures can be implemented:

- **Adversarial training:** Reinforces model robustness by exposing it to adversarial samples during the training phase.
- **Input preprocessing:** Applies filters (e.g., low-pass filtering) to remove adversarial noise from inputs without requiring model retraining.
- **Ensemble techniques:** Combines multiple models or algorithms to improve predictive reliability and reduce susceptibility to targeted attacks.

A key concern with adversarial examples is their transferability—an attack that succeeds on one model often works on others, even if they differ in architecture. This underscores the importance of robustness verification, especially when integrating third-party models into sovereign systems.

5.6. SECURING COLLABORATIVE LEARNING

Even when raw data is never explicitly exchanged, collaborative learning frameworks—particularly Federated Learning (FL)—must be reinforced with privacy-enhancing techniques to minimize the risk of information leakage through the final model.

In centralized FL setups, a key vulnerability lies in the aggregation server, which orchestrates training and collects updates from participants. These updates may

inadvertently reveal sensitive information, making the server a prime target for attackers.

To mitigate this risk, advanced cryptographic methods such as Fully Homomorphic Encryption (FHE) or Multi-Party Computation (MPC) can be applied during aggregation. These techniques ensure that contributions remain confidential, even if the server is untrusted. While FHE introduces performance overhead, it is well-suited to aggregation tasks, which are computationally lighter than full model training or inference.

However, encryption alone does not eliminate the possibility of privacy leaks from the final model. Therefore, additional safeguards must be implemented during training, such as noise injection to anonymize individual contributions. These privacy-preserving techniques help protect sensitive data—even in cross-border collaborations.

Thales has developed **SaferLearn**, a flexible framework for secure collaborative learning. It supports multiple learning protocols and integrates layered security mechanisms. SaferLearn enables the exploitation of distributed sensitive data without requiring a trusted central entity or a unified model architecture across participants.

The framework also includes tools to:

- Measure the quality of individual contributions.
- Detect poisoning attempts early in the training process.

In sovereignty-sensitive contexts, SaferLearn provides a robust foundation for continuous improvement of exported models—leveraging user data securely and responsibly.

5.7. PROTECTING OWNERSHIP RIGHTS OF MODEL AND DATA PROVIDERS

Protecting the intellectual property of model and data providers is essential in collaborative and export-oriented AI environments. Two main types of model stealing attacks threaten ownership rights:

- **Direct extraction:** Attackers exploit system vulnerabilities to access model parameters, either through side-channel attacks or by retrieving them directly from device memory.
- **Substitute model training:** Attackers query the target model and use the responses as labels to

train a replica, effectively cloning its behavior without access to its internal structure.

To counter these threats, several mitigation strategies can be deployed:

- Securing model hyperparameters and parameters.
- Detecting abnormal query patterns during model operation.
- Obfuscating inference outputs to limit information exposure.
- Implementing robust intellectual property (IP) management protocols.

A particularly effective technique is machine learning watermarking. Inspired by multimedia watermarking, this method embeds a secret signature into the model during training. This signature can later be revealed by the model owner to verify authorship. Watermarks may take the form of subtle parameter modifications or specific behavioral triggers activated by crafted inputs.

Ownership verification relies on demonstrating knowledge of the embedded watermark—either by extracting it using a secret key or by submitting a predefined input that elicits a unique response.

In collaborative learning contexts, watermarking can be adapted to track individual contributions. Each participant may embed a distinct mark into the shared model, providing proof of involvement. If necessary, these contributions can be removed using unlearning techniques, allowing for flexible and granular management of shared intellectual assets.

This approach is particularly valuable when collaboration dynamics evolve over time. Thales develops and promotes watermarking solutions tailored to secure and transparent multi-party AI development—reinforcing trust and accountability across sovereign boundaries.

6. IT infrastructure and operationalization

Building sovereign-compatible AI systems requires strategic infrastructure choices. To enable secure cross-border deployment and collaboration—without compromising national data or autonomy—organizations must adopt an infrastructure-first approach. The underlying architecture, from data centers to networks, determines who truly controls the data and the models.

Where data resides is not a mere technical detail—it is a strategic lever that influences trust, compliance, and eligibility for the public sector and export-sensitive missions. Without infrastructure aligned with sovereignty requirements, nations risk regulatory breaches and may miss opportunities for international cooperation.

Recognizing this, governments and alliances are forming trusted technology partnerships and embedding digital sovereignty safeguards. At Thales, we view investment in sovereign-ready AI infrastructure as foundational to global collaboration that respects national laws and strategic autonomy.

A robust and flexible architectural foundation must support:

- Diverse deployment models.
- Federated and collaborative learning pipelines.
- Secure data and model exchange.
- Effective monitoring and governance.

Strong governance and clear incident response protocols are essential to ensure compliance and accountability. With the right infrastructure, organizations gain the confidence to deploy AI solutions that navigate sovereignty and export control requirements—unlocking innovation without compromise.

6.1. HARDWARE

The **architectural foundation** of a sovereign-compatible AI platform must be multi-layered and meticulously designed to meet both technical performance and jurisdictional control requirements.

It begins with robust hardware components—servers, CPUs, GPUs/TPUs, and networking infrastructure—overlaid with virtualization layers that create secure, isolated computing environments. Containerization and orchestration frameworks (e.g., Kubernetes) manage AI workloads consistently across deployment models. At the top of the stack, an AIOps layer streamlines model training, deployment, and real-time monitoring, ensuring operational excellence.

This modular architecture allows policies to be enforced at every layer—from hypervisors to service meshes—effectively controlling cross-border data flows. Trust zones must be aligned with sovereignty boundaries, enabling the creation of project-specific, enterprise, partner, national, or international environments that reflect the required level of trust and legal oversight. With clearly defined sovereignty boundaries, organizations can ensure that sensitive data and assets remain fully protected.

Strategic selection of hardware components is essential. Choosing CPUs, GPUs, and TPUs that comply with export controls and data residency requirements helps maintain sovereignty. Trusted execution and hardware resilience must be prioritized to guarantee the integrity and security of AI systems.

Leading nations are increasingly adopting Trusted Execution Environments (TEEs)—secure enclaves embedded in CPUs or system architectures—to protect sensitive models and data during runtime. Technologies such as Intel SGX, AMD SEV, and ARM TrustZone create encrypted, isolated memory regions where code can execute securely, even from the host operating system. This allows models to be deployed on foreign or untrusted servers without exposing their internal parameters.

Similar protections are emerging in AI-specific accelerators: some GPUs now support encrypted memory or hypervisor-level isolation, extending confidential computing capabilities to high-performance AI hardware.

In certain contexts, opting for CPU-only operation may be a sovereign choice. While advanced AI training typically requires GPU/TPU acceleration, these components are often subject to export restrictions and supply chain risks. For lower-risk or less resource-intensive applications, CPU-based environments offer a viable alternative—reducing foreign dependency while maintaining operational control.

Ultimately, hardware decisions must be guided by a risk-tiering strategy: high-end GPUs for approved, monitored projects, and CPUs or legacy components for others. This approach balances performance, sovereignty, and compliance across diverse operational scenarios.

6.2. DEPLOYMENT MODEL

Deployment models are at the heart of building AI platforms that align with sovereignty requirements. Each model offers distinct capabilities while presenting its own set of challenges.

On-Premises / Sovereign Cloud

This model ensures that AI workloads remain entirely within national borders, leveraging government-operated data centers or certified sovereign clouds. It guarantees compliance with domestic regulations and data sovereignty mandates. This approach offers maximum control and security—an essential requirement for sensitive domains such as government, defense, and critical infrastructure. However, it often comes with higher operational complexity and limited scalability.

Hybrid Deployment

The hybrid model blends national infrastructure with selected global cloud resources. Sensitive data is processed locally, while less critical or compute-intensive tasks are offloaded to trusted foreign cloud partners. Governments frequently collaborate with major cloud providers to establish internal clouds that combine local control with elastic scalability. Through strict data partitioning, robust technical safeguards, and integration with approved regions, this model ensures that protected data remains secure. It strikes a balance between operational efficiency and sovereignty by selectively harnessing global cloud capabilities.

Edge-Centric Deployment

When centralized or partner cloud solutions are not viable, edge computing becomes a strategic alternative. This model decentralizes AI processing to local endpoints—such as IoT devices, field units, or 5G base stations—bringing computation closer to where data is generated. It supports ultra-low latency and adheres to “data stays local” policies, enabling real-time responsiveness and operational continuity even in disconnected environments. While managing updates across distributed nodes poses challenges,

containerized and orchestrated software stacks with over-the-air updates offer effective mitigation. By transmitting only insights rather than raw data, this model reinforces sovereignty and complements sovereign cloud strategies.

6.3. GOVERNANCE

Effective governance is a cornerstone of secure and sovereign AI collaboration. Leveraging container templates enables the implementation of fine-grained opt-in and opt-out controls, along with robust mechanisms to enforce the “right to be forgotten.” These capabilities empower contributors to withdraw their data or models at any time, ensuring their complete and verifiable removal from collaborative environments.

Advanced features such as cryptographic signing, provenance tracking, and immutable model registries provide transparent attribution of contributions and establish an auditable chain of custody throughout the AI development lifecycle. By adopting a Policy-as-Code approach, organizations can automate the enforcement of participation rules, seamlessly translating legal agreements and partnership frameworks into actionable technical controls.

This comprehensive governance framework ensures that every collaborative AI initiative remains compliant with applicable policies, respects the autonomy and legal obligations of all stakeholders, and facilitates transparent recognition and management of individual contributions. By integrating these practices, organizations can build collaborative AI systems that are secure, compliant, and fully accountable by design.

6.4. DATA AND MODEL EXCHANGE

Secure mechanisms for data and model exchange are fundamental to building sovereign-compatible AI platforms—particularly in federated and collaborative learning environments. Ensuring trust during the exchange of datasets and models requires robust, end-to-end security measures.

Each dataset or model, including its weights, is encapsulated and encrypted using strong symmetric encryption algorithms such as AES-256. To guarantee authenticity and integrity, the package is digitally signed with the sender’s private key (e.g., RSA or ECC). This ensures that only authorized recipients with the appropriate decryption keys can access the contents,

while any tampering during transmission is immediately detectable through signature verification.

The encrypted package includes a cryptographic signature, enabling recipients to confirm the origin and integrity of the data or model. This approach ensures that even when transmitted over untrusted networks or through third-party cloud infrastructures, the confidentiality and provenance of the assets remain intact.

By embedding these security practices into the core of AI collaboration workflows, organizations can foster trust, uphold data sovereignty, and mitigate risks—laying the groundwork for secure, compliant, and resilient AI ecosystems.

6.5. MONITORING AND INCIDENT MANAGEMENT

Robust monitoring, effective incident response, and strong governance are foundational pillars of a sovereign-compatible AI platform—especially when deploying models in production. These capabilities ensure operational integrity, security, and regulatory compliance, enabling trusted AI solutions for partners and clients alike.

Even with strong preventive controls in place, challenges such as model drift, data poisoning, or adversarial attacks can still emerge. This makes real-time monitoring of model behavior and performance a critical line of defense. Continuous observation allows organizations to detect shifts in data distributions, the emergence of bias, or anomalous activities that may signal security threats. Leveraging secure cloud platforms and AI Ops systems—such as those offered by Thales—organizations can access automated dashboards and proactive alerts designed to identify risks like model extraction attempts or potential data leakage.

Rapid detection must be paired with swift and coordinated response. Integrating AI platforms with the broader Security Operations Center (SOC) is essential. AI systems should emit security events directly into the organization's Security Information and Event Management (SIEM) infrastructure, enabling centralized visibility and response. Hardware-based protections, including Hardware Security Modules (HSMs) and Trusted Platform Modules (TPMs), further enhance security by preventing unauthorized model extraction and ensuring data integrity. These components can also generate

hardware-level security signals to alert the SOC in case of violations.

Effective AI operations also require well-defined Incident Response (IR) plans tailored to the complexities of jurisdictional boundaries. These plans must clearly outline notification protocols, permissible information sharing, and escalation paths across different legal frameworks. In federated environments involving multiple stakeholders and jurisdictions, incident handling becomes even more complex. Cross-border data access may conflict with sovereignty laws, making localized forensic analysis and regional mini-SOCs essential. Clear governance over leadership handoffs and joint response coordination ensures compliance and operational continuity.

Ultimately, enabling sovereign-compatible AI collaboration demands a resilient architectural foundation, flexible deployment models, and advanced federated learning capabilities.

7. Legal frameworks and policies to safeguard sovereignty.

The global landscape of AI governance is evolving rapidly, with significant regional variations reflecting different cultural, economic, and political priorities. While all regions aim to balance innovation with responsible deployment, their approaches diverge in structure and emphasis.

- Europe leads with the AI Act, a comprehensive, risk-based regulatory framework set to take effect in August 2026. It establishes clear obligations based on the level of risk posed by AI systems, setting a global benchmark for responsible AI.
- The United States follows a more decentralized and sectoral approach, focusing on safety standards and voluntary guidelines while fostering innovation through public-private collaboration.
- Canada has proposed the AI and Data Act, which adopts a risk-based model supported by voluntary codes of practice and federal investment in AI research and development.
- The United Kingdom maintains a flexible, sector-specific regulatory model, though discussions are ongoing about transitioning toward a more unified governance structure.

Asia presents a diverse regulatory landscape:

- China emphasizes transparency, algorithmic accountability, and data protection.
- Singapore and Japan are advancing structured regulatory frameworks that promote innovation while ensuring ethical AI use.
- South Korea's AI Framework Act, effective January 2026, aligns closely with EU principles, combining regulatory oversight with strong support for industrial innovation.

Despite these differences, common principles are emerging across jurisdictions: transparency, fairness, data privacy, accountability, and risk management. These shared values aim to ensure that AI deployment aligns with societal well-being and national interests.

Global Cooperation and AI Sovereignty: Navigating the Balance

Effective AI governance requires coordinated action among stakeholders to shape the future of AI in ways that support sovereignty, national security, economic competitiveness, and societal well-being. However, unilateral pursuits of AI sovereignty by individual states risk fragmenting global standards and undermining international cooperation. Protectionist policies can stifle beneficial trade, limit research collaboration, and distort fair competition. As such, international cooperation is essential to ensure interoperability of technical standards and alignment of governance frameworks.

Multilateral organizations—including the G7, OECD, UNESCO, NIST, ISO, and the Council of Europe—are actively working to establish common principles around safety, transparency, ethics, and security. The debate over global AI standards has become a key arena of geopolitical competition, where the challenge lies in balancing national sovereignty with the need for global coordination and harmonization.

Digital sovereignty itself encompasses more than just data control—it includes regulatory oversight of digital platforms and the ability to shape domestic technological ecosystems. Leading economies such as the European Union assert digital sovereignty through comprehensive regulatory frameworks like the AI Act, which harmonizes ethical and legal standards across member states.

International perspectives increasingly emphasize transnational cooperation to address shared challenges in AI governance. Notable initiatives include:

- UNESCO's global AI ethics guidelines, which aim to establish universally accepted principles.
- The EU's collaboration with NGOs such as the IEEE to promote responsible AI development.
- The AI Alliance, launched in 2023 by IBM and Meta, now comprising 140 members across 23 countries, promoting responsible innovation grounded in scientific integrity, trust, diversity, safety, and competitiveness.
- Impact AI, a French "think-and-do tank," brings together stakeholders to foster ethical and inclusive AI practices.
- Thales, as a founding member of the European AI Trustworthy Association, plays a key role in supporting industrial players in mastering critical

AI technologies, thereby reinforcing technological autonomy, competitiveness, and sovereignty.

These collective efforts underscore the importance of building a shared foundation for AI governance—one that respects national interests while fostering global trust, innovation, and ethical alignment.

8. Thales and AI Sovereignty: Expertise Rooted in Operational Reality

Thales is directly engaged with the challenges and opportunities outlined in this document through its diverse business activities. For several years, artificial intelligence has been embedded across its product portfolio—particularly in critical systems and defense solutions deployed or intended for deployment in multiple countries worldwide. This operational experience has enabled Thales’ research and engineering teams to develop AI components, tools, and methodologies tailored to both sovereign and international contexts.

Some of these innovations have been referenced in previous sections to help the broader ecosystem better understand the complexities of sovereign AI and the practical solutions that can be implemented.

Today, Thales brings together more than 800 AI experts within its dedicated internal organization, **cortAlx**, established in 2024. These experts are distributed across three core teams:

- Labs, focused on foundational research and model development;
- Factory, responsible for toolchains, cybersecurity, and critical infrastructure;
- Sensors, integrating AI into embedded and edge systems.

Together, these teams ensure that AI capabilities are not only cutting-edge but also operationally viable, secure, and aligned with sovereignty requirements. Thales’ AI workforce is strategically located across France, the United Kingdom, Canada, Singapore, and the United Arab Emirates, reflecting the company’s global footprint and commitment to supporting regional sovereignty through local expertise.

9. About Thales

Thales is directly involved in the challenges and opportunities described in this document through its diverse activities. For several years, artificial intelligence has been integrated throughout its product portfolio—particularly in critical systems and defense solutions deployed or intended to be deployed in many countries around the world. This operational experience has enabled Thales’ research and engineering teams to develop AI components, tools, and methodologies adapted to both sovereign and international contexts.

Some of these innovations were mentioned in previous sections to help the ecosystem better understand the complexities of AI used for electronic warfare and the practical solutions that can be implemented.

Today, Thales brings together over 800 AI experts within its internal dedicated organization, **cortAlx**, created in 2024. These experts are divided into three main teams:

- **Labs:** focused on fundamental research and model development.
- **Factory:** responsible for tool chains, cybersecurity, and critical infrastructure.
- **Sensors:** integrating AI into embedded and edge systems.

Together, these teams ensure that AI capabilities are not only state-of-the-art but also operationally viable, secure, and compliant with sovereignty requirements. Thales’ **cortAlx** teams are strategically located across France, the United Kingdom, Canada, Singapore, Germany and the United Arab Emirates, reflecting the company’s global footprint and its commitment to supporting regional sovereignty through local expertise.





4, rue de la Verrerie 92190
Meudon FRANCE

Tél. + 33(0)1 57 77 80 00

www.thalesgroup.com

