

Trustworthy AI by Design

The CC-FEAI executive blueprint for privacy-preserving, explainable, secure and compliance-ready AI in regulated domains.

Train anywhere. Share no raw data. Explain every consequential decision. Govern continuously.

Dr. Jami Kiran Kumar

Trusted Global Technology Advisor

AI · Data · Cybersecurity · Cloud · Quantum

Derived from the author's doctoral thesis: *Federated & Explainable AI: A Compliance-Centric Framework for Trustworthy AI in Regulated Domains*. Authored and authorised by Jami Kiran Kumar.

Regulated AI needs a system of trust—not another isolated model

Accuracy is necessary. It is not sufficient when decisions affect rights, safety, money, health or essential services.

The problem

- Centralized training increases data concentration and sovereignty risk.
- Opaque predictions weaken oversight and contestability.
- Compliance is often assessed after architecture decisions are fixed.
- Distributed models create new security, drift and accountability risks.

The CC-FEAI response

- Federated learning keeps raw data at origin.
- SHAP and LIME support local and global explanation.
- Governance controls operate through the lifecycle.
- Privacy, robustness and compliance become optimization objectives.

Jami Kiran's central principle

Trust must be engineered into data flow, training, explanation, security, documentation, human authority and monitoring. It cannot be created by a policy document placed around an opaque system after deployment.

Thesis contribution	Executive meaning	Decision consequence
Unified FL + XAI + compliance	Privacy and transparency are designed together.	Fund the complete operating architecture.
Multi-objective trust function	Accuracy competes with privacy, explanation and robustness.	Approve explicit trade-offs, not hidden defaults.
Lifecycle governance	Compliance evidence must persist after launch.	Build monitoring, logs and change control before production.

Why conventional AI architectures fail in high-stakes environments

Privacy

Moving raw data into one training environment increases exposure, legal and sovereignty risk.

Opacity

Complex models may be accurate while remaining difficult to challenge or audit.

Compliance

Static checklists do not control evolving models, data and use contexts.

Security

Poisoned clients, malicious updates and inference attacks target distributed learning.

Drift

Changing populations and operations weaken performance and explanation stability.

Accountability

Unclear roles allow providers, deployers and business owners to displace responsibility.

The regulated-domain test

Domain	Consequential decision	Required evidence
Healthcare	Diagnosis, triage, prioritization	Clinical validity, privacy, subgroup performance, oversight
Financial services	Credit, fraud, pricing, suitability	Reason codes, fairness, traceability, data governance
Public sector	Eligibility, inspection, resource allocation	Lawful basis, contestability, impact assessment, audit
Critical infrastructure	Anomaly response and operational control	Resilience, human authority, security, recovery

Design question: Can the organization demonstrate not only what the model predicted, but which data and version produced it, why it was used, who had authority, what control was applied and how an affected person can obtain review?

The five-layer CC-FEAI architecture

1 · Data & Edge	Local collection, minimization, quality controls, lawful processing and sensitive-data boundaries.
2 · Federated Learning	Local training, client selection, secure updates, FedAvg/FedAvgM aggregation and global distribution.
3 · Explainability	SHAP and LIME explanations, feature stability, decision rationale and stakeholder-specific presentation.
4 · Security	Identity, least privilege, encryption, differential privacy, Byzantine resilience and anomaly detection.
5 · Compliance & Lifecycle	Risk classification, policy mapping, audit evidence, human oversight, drift, incident and change management.

End-to-end control flow

Classify Use, risk, rights and actors	Constrain Data, authority and purpose	Train Local models and secure updates	Explain Prediction and model behaviour	Validate Performance, risk and evidence	Operate Monitor, intervene and improve
---	---	---	--	---	--

Important boundary: Federated learning reduces raw-data movement; it does not by itself guarantee privacy. Model updates can leak information unless protected through secure aggregation, differential privacy, encryption and attack monitoring.

Collaborate on learning without centralizing raw data

Each client trains on local data and contributes a model update. An aggregator combines permitted updates into a global model and returns the revised model to participants. The thesis evaluates FedAvg and FedAvgM in IID and non-IID settings.

$$w(t+1) = \sum_k (n_k / n) \cdot w_k(t+1)$$

Control objective	Mechanism	Evidence
Data localization	Raw records remain at controlled nodes	Data-flow verification and egress controls
Update confidentiality	Secure aggregation and encryption	Protocol design, key management and tests
Client integrity	Identity, attestation and admission policy	Trusted-node inventory and access logs
Poisoning resistance	Robust aggregation and anomaly detection	Attack simulations and rejection metrics
Non-IID performance	Momentum, personalization or weighting	Per-node and subgroup evaluation
Communication efficiency	Client sampling, compression and round control	Bytes, rounds, latency and energy per convergence

What remains local—and what still travels

Raw data stays local, but gradients, weights, metrics, metadata and explanations may move. Every artefact requires classification, minimization and retention rules.

Privacy-preserving architecture is a reduction of exposure—not a declaration of zero risk.

Explain decisions for the person who must act, challenge or audit

SHAP

Attributes a prediction to features using a Shapley-value framework. Useful for local explanations and aggregated model patterns.

LIME

Approximates model behaviour locally with an interpretable surrogate. Flexible, but explanation stability must be tested.

Audience	Need	Explanation design
Decision subject	Understand and contest an outcome	Clear reason codes, material factors, review path
Professional user	Judge whether to rely on output	Drivers, confidence, uncertainty, limitations
Model owner	Diagnose performance and drift	Global importance, subgroup variation, stability
Auditor/regulator	Reconstruct design and operation	Versioned explanation, data lineage, controls and logs

Explanation quality scorecard

Fidelity

Stability

Comprehensibility

Actionability

Completeness

Uncertainty

Audience fit

Guardrail: An explanation can be plausible without being faithful. CC-FEAI therefore treats explanations as evaluable model outputs, not decorative narratives.

Translate obligations into architecture, evidence and accountability

The CC-FEAI compliance layer maps legal and policy duties to lifecycle controls. It supports alignment; it cannot automatically declare a system lawful. Applicability depends on role, jurisdiction, use, risk classification and sector law.

Source	Relevant objective	CC-FEAI evidence
EU AI Act	Risk management, data governance, documentation, logs, transparency, oversight, accuracy, robustness and cybersecurity for applicable high-risk systems	Risk file, data controls, model card, logs, oversight design, test results
GDPR	Lawfulness, purpose limitation, minimization, security, rights and accountability for personal data	Processing record, DPIA where required, data map, access and retention controls
NIS2	Cybersecurity risk-management measures and incident handling for entities in scope	Security architecture, suppliers, continuity, incident and vulnerability records
NIST AI RMF	Govern, Map, Measure and Manage AI risk	Governance roles, context map, TEVV results and risk treatment

Evidence chain

Requirement Applicable obligation	Control Technical or procedural response	Owner Accountable role	Test Verification method	Record Versioned evidence
---	--	----------------------------------	------------------------------------	-------------------------------------

“Compliance-ready” means evidence can support a qualified legal and regulatory assessment. It does not mean the architecture replaces that assessment.

Optimize what the enterprise actually values

The thesis extends the objective beyond predictive loss by representing privacy, explainability, compliance and robustness as co-equal design considerations.

$$\text{Min } J(w) = \alpha \cdot \text{Loss} + \beta \cdot \text{PrivacyRisk} + \gamma \cdot \text{Opacity} + \delta \cdot \text{ComplianceGap} + \epsilon \cdot \text{RobustnessRisk} + \zeta \cdot \text{OperationalCost}$$

Weights express organizational priorities and constraints; they must be governed, justified and stress-tested.

Objective	Candidate measure	Typical trade-off
Accuracy	F1, AUC, calibration, subgroup error	Complexity and explainability
Privacy	ϵ/δ budget, leakage tests, exposure surface	Utility and convergence
Explainability	Fidelity, stability and human comprehension	Compute and response latency
Compliance readiness	Control coverage and evidence quality	Delivery speed and documentation effort
Robustness	Attack success and worst-case degradation	Model utility and operational complexity
Efficiency	Rounds, bandwidth, latency and cost	Participation and model quality

Governance decision: The weights are not merely technical hyperparameters. They encode risk appetite and the distribution of protection among stakeholders.

What the doctoral evaluation demonstrates—and what remains open

95%

Reported FedAvgM accuracy in the controlled experimental setup.

94 / 92 / 93

Reported precision, recall and F1 percentages.

IID + non-IID

Evaluation considered homogeneous and heterogeneous distributions.

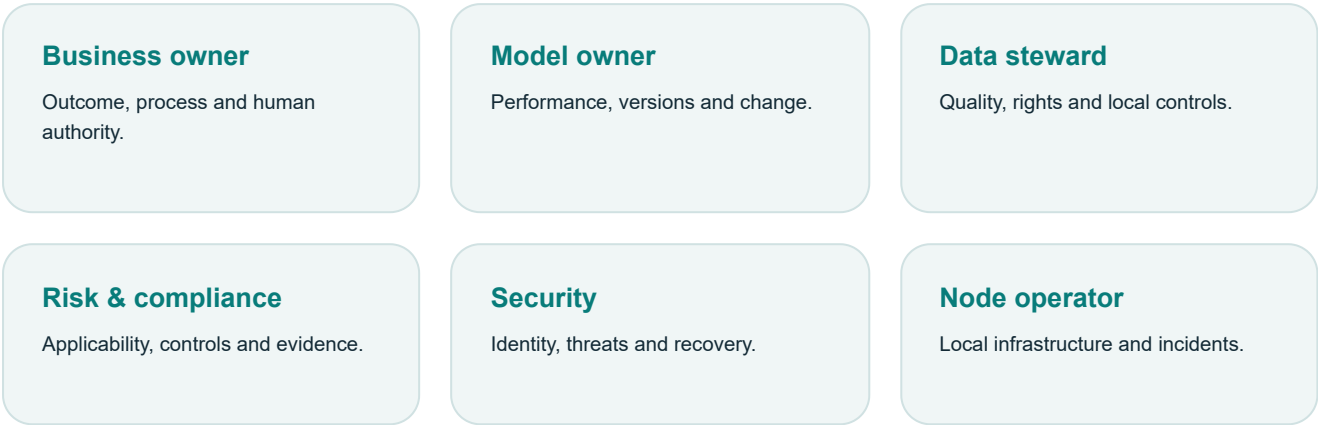
The thesis used a multi-node distributed environment, TensorFlow/PyTorch, FedAvg/FedAvgM, SHAP/LIME and synthetic or enterprise-like structured, time-series and sensitive-regulatory datasets.

Responsible interpretation

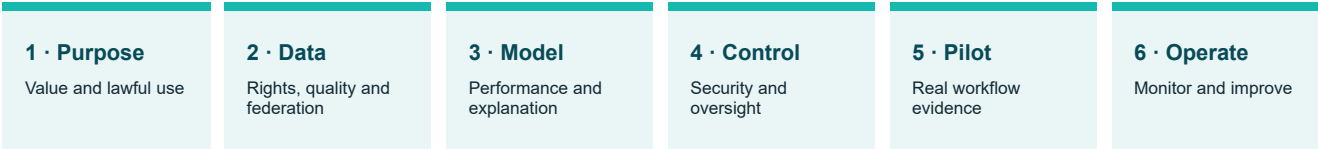
Supported	Not yet established
Technical feasibility of integrating FL, XAI and governance components	Performance across real healthcare, finance or public-sector production systems
Strong controlled classification metrics for the selected setup	Universal 15–20% improvement or superiority across models and datasets
Operational compliance mapping and evidence design	Automatic or legally conclusive compliance
Identified explanation outputs and stability concepts	Human comprehension, contestability and decision-quality effects at scale

Publication integrity: The controlled results justify further pilots and comparative research—not unrestricted production deployment.

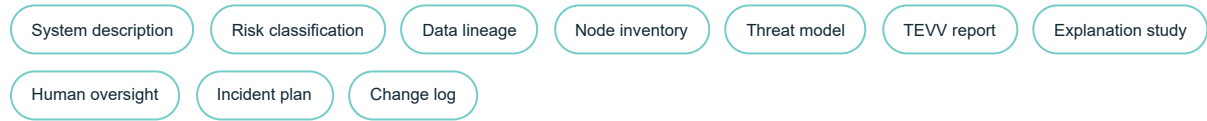
Turn CC-FEAI into an accountable enterprise capability



Lifecycle gates



Minimum production dossier



A 120-day controlled adoption programme

Days 0–20 Classify use, actors, data and harms	Days 21–45 Design nodes, controls and evidence	Days 46–75 Build benchmark and attack tests	Days 76–100 Pilot with users and reviewers	Days 101–120 Gate deployment or redesign
--	--	---	--	--

Twelve questions before build

1. What consequential decision is being supported?
2. Which laws and sector rules apply?
3. Why must data remain distributed?
4. What information can updates leak?
5. Who is authorized to participate?
6. How will poisoned clients be detected?
7. Who needs which explanation?
8. How are fidelity and stability tested?
9. Where can a human intervene?
10. What constitutes compliance evidence?
11. What triggers rollback or retraining?
12. Which result must exist before scale?

Scale rule: Advance only when business value, model performance, privacy protection, explanation quality, security resilience, human oversight and evidence readiness all meet predefined thresholds.

Evidence base and authorship

1. Kumar, J. K. (2026). *Federated & Explainable AI: A Compliance-Centric Framework for Trustworthy AI in Regulated Domains*. Doctoral thesis / author manuscript.
2. McMahan, B., et al. (2017). "Communication-Efficient Learning of Deep Networks from Decentralized Data." *AISTATS / PMLR* 54.
3. Lundberg, S. M., & Lee, S.-I. (2017). "A Unified Approach to Interpreting Model Predictions." *NeurIPS* 30.
4. Ribeiro, M. T., Singh, S., & Guestrin, C. (2016). "Why Should I Trust You?" Explaining the Predictions of Any Classifier." *KDD '16*.
5. European Union (2024). Regulation (EU) 2024/1689, Artificial Intelligence Act.
6. European Union (2016). Regulation (EU) 2016/679, General Data Protection Regulation.
7. European Union (2022). Directive (EU) 2022/2555, NIS2 Directive.
8. Tabassi, E. (2023). *NIST AI Risk Management Framework 1.0*, NIST AI 100-1.

Authorship and authorisation

Author: Dr. Jami Kiran Kumar.

Original contribution: CC-FEAI architecture, compliance-centric lifecycle, multi-objective trust model, controlled evaluation synthesis and executive implementation framework.

Evidence note: Legal mappings support structured assessment but do not constitute legal conclusions. Real-world performance requires domain-specific validation.

Disclaimer: Strategic and educational material only; not legal, clinical, financial, cybersecurity-certification or regulatory advice.