

SHER DeepAI Research

Munich, Germany | August 2026

Authors:

Dr. Silvia Moyses-Scheingruber

Rajkumar Verma

Categories:

Trustworthy AI | AI Governance | AI Engineering | AI Lifecycle | Explainable AI (XAI)

Engineering Trustworthy AI Across the AI Lifecycle: Why Good Engineering Is the Foundation of Trustworthy AI

Summary

As AI adoption accelerates, companies face increasing expectations to develop AI systems that are transparent, explainable, accountable, secure, and compliant with evolving governance and regulatory requirements. While AI regulations define what Trustworthy AI should achieve, they provide far less guidance on how these objectives can be implemented in practice.

This article argues that Trustworthy AI is fundamentally an engineering challenge. Trust must be engineered throughout the AI lifecycle - from problem definition and data preparation to model development, deployment, monitoring, and continuous improvement. The article also introduces SHER DeepAI's Engineering Capability Framework for Trustworthy AI, highlighting the core capability domains of **Governance**, **Transparency**, **Observability**, and **Security**. The authors also explain how **Explainable AI (XAI)** supports multiple stages of the AI lifecycle and serves as a foundational capability for building trustworthy AI. Finally, the article outlines how **DeepTAI**, SHER DeepAI's **modular platform for Trustworthy AI** helps to operationalize these engineering capabilities to build, deploy, and operate AI systems that are trustworthy by design.

Table of Content

Summary	1
Table of Content	2
1. Why Trustworthy AI Matters	3
2. The AI Lifecycle: Where Trustworthy AI Is Built	4
2.1. Problem Definition:	5
2.2. Data Collection and Data Preparation:	6
2.3. Feature Engineering:	6
2.4. Model Training	7
2.5. Model Evaluation	7
2.6. Deployment	8
2.7. Monitoring and Continuous Improvement	8
3. Core Engineering Capabilities for Trustworthy AI	9
4. The Role of Explainable AI (XAI) Across the AI Lifecycle	11
5. Operationalizing Core AI Engineering Capabilities with DeepTAI	12
About the Authors	13
About SHER DeepAI	13

1. Why Trustworthy AI Matters

Companies are under increasing pressure to develop, deploy, and operate AI systems that stakeholders can understand, trust, and confidently use. At the same time, they must meet rapidly evolving governance and regulatory requirements for transparency, explainability, accountability, and compliance. Yet while regulations define **what** Trustworthy AI should achieve, they provide far less guidance on **how** these objectives can be translated into practical engineering capabilities across the AI lifecycle.

Artificial intelligence is transforming the way companies innovate, make decisions, and create value. From healthcare, legal services, consulting, manufacturing, and industrial AI to finance, critical infrastructure, and the public sector, AI is becoming deeply embedded in products, services, and business operations. As AI becomes an integral part of business operations and decision-making, companies are increasingly expected to build AI systems that stakeholders can understand, trust, and confidently use.

Growing regulatory and governance requirements - including the EU AI Act, the EU Product Liability Directive, and other emerging international standards - are reinforcing the need for transparency, explainability, accountability, and compliance. Although these regulations and governance frameworks differ across jurisdictions, they are increasingly converging around a common objective: enabling the development and deployment of Trustworthy AI.

As AI adoption accelerates and regulatory expectations continue to evolve, building Trustworthy AI has become one of the defining engineering challenges of the AI era. Achieving Trustworthy AI requires more than governance and compliance—it requires a disciplined engineering approach.

Trust cannot be added at the end of the development process through documentation, risk assessments, or regulatory reviews. Instead, it must be engineered into AI systems from the very beginning. Every engineering decision - from defining the problem and preparing data to developing models, deploying systems, monitoring their performance, and continuously improving them - contributes to whether an AI system is transparent, explainable, robust, accountable, and ultimately trustworthy.

Global AI Regulations Directory

SHER DeepAI's AI Regulations Directory provides an overview of AI regulations, policies, standards, and certification frameworks from around the world.

Scan the QR code to explore the directory.



Just as a building cannot be given a strong foundation after construction has been completed, an AI system cannot become trustworthy once development is finished. Trust must be engineered into every stage of the AI lifecycle.

This article explores why good AI engineering is the foundation of Trustworthy AI. It examines how engineering decisions made throughout the AI lifecycle determine whether AI systems can satisfy both business expectations and the growing demands of governance and regulation. It also introduces the core engineering capabilities that enable companies to operationalize Trustworthy AI in practice.

2. The AI Lifecycle: Where Trustworthy AI Is Built

The development of an AI system usually follows a structured engineering process that begins with problem definition and continues through data collection, data preparation, feature engineering, model training, model evaluation, model deployment, monitoring, and continuous improvement (see Figure 1).

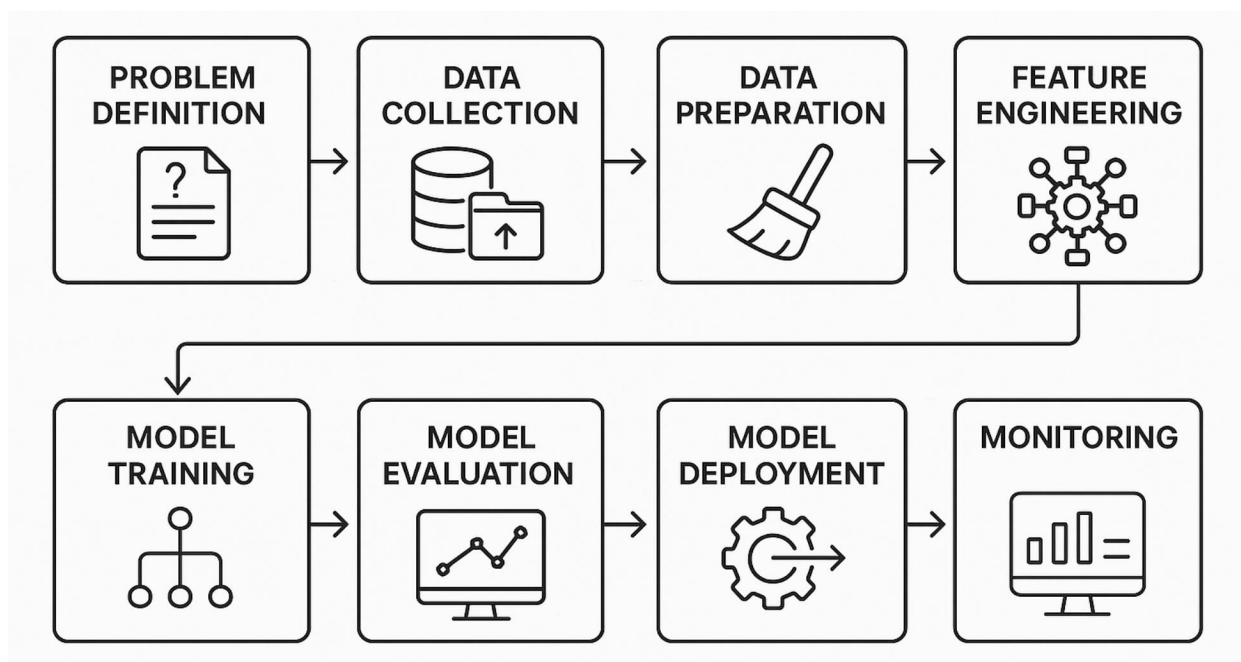


Figure 1. The AI Lifecycle

Many companies today view AI governance primarily as a compliance activity that takes place towards the end of development. Documentation is prepared, risk assessments are completed, and regulatory requirements are reviewed shortly before an AI system is deployed. While these activities are essential, they cannot by themselves create Trustworthy AI.

The trustworthiness of an AI system is determined by the **engineering decisions made throughout its lifecycle**. Decisions about data, feature engineering, model development, evaluation, deployment, and monitoring collectively determine whether an AI system is transparent, explainable, accountable, robust, secure, and capable of meeting both business objectives and regulatory requirements.

The AI lifecycle is therefore much more than a sequence of technical activities. It is the engineering process through which Trustworthy AI is built. Each phase of the lifecycle makes a distinct contribution to the trustworthiness of an AI system, and together they provide the foundation for responsible AI development and operation.

The following sections examine each phase of the AI lifecycle and explain how sound engineering practices enable companies to build AI systems that are trustworthy by design. Throughout the article, we also illustrate how these engineering practices can be operationalized using **DeepTAI**, SHER DeepAI's **modular platform for Trustworthy AI** (see Section 5). **DeepTAI** translates governance and regulatory requirements into practical engineering capabilities that support the design, development, deployment, and continuous operation of trustworthy AI systems across the entire AI lifecycle.

2.1. Problem Definition:

Every AI project begins by defining the problem it is intended to solve. While this is often viewed as a business or technical planning activity, it is also the first governance decision. Clearly defining the purpose, intended users, and expected outcomes of an AI system establishes the foundation for responsible AI development and helps determine whether AI is the appropriate solution. This stage also provides an opportunity to identify key stakeholders, assess potential risks, understand applicable regulatory obligations, and determine the level of human oversight that may be required.

Key Takeaway: Trustworthy AI begins with the first engineering decision. Clearly defining the right problem, identifying potential risks, and establishing governance early create the foundation for AI systems that are trustworthy by design.

DeepTAI in Practice: At this stage, stakeholders involved in developing AI systems can use SHER DeepAI's **Virtual TAI Assistant** to analyse their AI use case, identify the applicable governance and regulatory requirements, and determine the engineering capabilities needed to build Trustworthy AI (see Section 5).

2.2. Data Collection and Data Preparation:

The quality and reliability of an AI system depend on the data used to develop it. Collecting, preparing, and managing data are therefore critical engineering activities that directly influence model performance, fairness, robustness, and ultimately the trustworthiness of the system. Companies should ensure that data is relevant, representative, accurate, and collected in accordance with applicable legal and ethical requirements. Documenting data sources, data transformations, and data quality assessments enhances transparency, traceability, and accountability throughout the AI lifecycle.

Key Takeaway: Well-governed, representative, and high-quality data provides the foundation for reliable, explainable, and trustworthy AI systems.

DeepTAI in Practice: At this stage, stakeholders involved in developing AI systems can use **DeepTAI's Engineering module** to collect, prepare, validate, transform, and govern data throughout the AI lifecycle. The module supports data quality assessment, data annotation, privacy protection, bias analysis, dataset documentation, and continuous data monitoring. Together, these capabilities enable companies to build Trustworthy AI on a foundation of reliable, representative, and well-governed data.

2.3. Feature Engineering:

Feature engineering transforms raw data into meaningful features that enable AI models to learn effectively. It is one of the most influential engineering activities in the AI lifecycle, directly affecting model performance, robustness, fairness, and explainability.

Selecting, creating, and transforming features require both domain expertise and sound engineering judgment to ensure that models learn from relevant, representative, and meaningful information. Thoughtfully designed features not only improve model performance but also make AI systems easier to understand, explain, and validate. Documenting feature engineering decisions further strengthens traceability, supports reproducibility, and enables more effective governance throughout the AI lifecycle.

Key Takeaway: Feature engineering is the foundation of trustworthy AI. Well-designed features enable AI models that are more accurate, robust, explainable, reliable and reproducible.

DeepTAI in Practice: At this stage, stakeholders involved in developing AI systems can use **DeepTAI's Engineering module** to support feature engineering throughout the AI lifecycle. The module helps engineers create, transform, validate, and document and

identify important features. By integrating feature extraction, documentation, versioning, and Explainable AI (XAI) capabilities, DeepTAI helps companies to develop AI models that are more transparent, explainable, and trustworthy by design.

2.4. Model Training

Model training is the process of learning patterns from engineered features to create an AI model capable of making predictions or decisions. Selecting appropriate algorithms, tuning model parameters, and managing training experiments are critical engineering activities that directly influence model quality, reliability, and robustness. Maintaining version control of datasets, features, models, and training experiments ensures reproducibility and provides the traceability needed for governance, auditing, and regulatory compliance.

Key Takeaway: Model training is more than optimizing performance - it is about creating AI models that are reproducible, traceable, and accountable by design.

DeepTAI in Practice: At this stage, stakeholders involved in developing AI systems can use **DeepTAI's Engineering module** to manage model training throughout the AI lifecycle. The module supports experiment management, model and dataset versioning, parameter tracking, and comprehensive documentation of training activities. By ensuring reproducibility, traceability, and governance across the model development process, DeepTAI helps companies build AI models that are robust, transparent, and trustworthy by design.

2.5. Model Evaluation

Model evaluation determines whether an AI model is fit for its intended purpose and provides confidence that it can operate reliably in real-world conditions. While predictive performance is an important measure of model quality, companies should also evaluate robustness, fairness, bias, and explainability. Understanding how a model reaches its decisions helps validate its behaviour, identify potential risks, and build confidence among AI engineers, business users, governance teams, and regulators.

Key Takeaway: High-performing models alone do not create Trustworthy AI. AI models must also be transparent, explainable, fair, robust, and reliable.

DeepTAI in Practice: At this stage, stakeholders involved in developing AI systems can use **DeepTAI's Explainable AI (XAI) and Engineering modules** to evaluate model behaviour beyond traditional performance metrics. DeepTAI supports explainability

analysis, bias assessment, robustness evaluation, and model validation while documenting evaluation results to ensure transparency, traceability, and regulatory readiness. These capabilities help companies build confidence that AI systems are not only accurate, but also understandable, reliable, and trustworthy by design.

2.6. Deployment

Deployment is the stage where an AI model moves from development into real-world operation. At this point, the focus shifts from building the model to ensuring it can be used reliably, securely, and responsibly in production. Companies should establish clear deployment procedures, maintain version control of models and supporting artifacts, and implement logging and monitoring capabilities to ensure traceability and accountability. Deployment is also the point where AI systems begin interacting with users and business processes. Appropriate governance controls, human oversight, and documentation help ensure that AI systems operate as intended and that decisions can be reviewed when necessary.

Key Takeaway: Deployment is not the end of AI development - it marks the beginning of responsible AI operation, where governance, monitoring, and human oversight become essential.

DeepTAI in Practice: At this stage, stakeholders involved in deploying AI systems can use **DeepTAI's Engineering, Governance, and Observability modules** to support the transition from development to production. DeepTAI enables controlled model deployment, version management, logging, observability, human oversight, and comprehensive documentation while maintaining full traceability of deployed models and supporting artefacts. These capabilities help companies to operate AI systems responsibly, transparently, and in accordance with governance and regulatory requirements.

2.7. Monitoring and Continuous Improvement

AI systems require continuous monitoring after deployment to ensure they remain accurate, reliable, and aligned with business objectives. Over time, changes in data, user behavior, or operating environments can affect model performance and introduce new risks. Companies should therefore monitor performance, detect model and data drift, assess ongoing fairness and explainability, and verify that AI systems continue to meet regulatory and governance requirements.

Continuous improvement is equally important. Insights gained from monitoring should be used to refine models, update features, improve data quality, and strengthen governance processes. This creates a feedback loop that enables AI systems to evolve while maintaining trust, transparency, and accountability throughout their lifecycle.

Key Takeaway: Trustworthy AI is not a one-time achievement. It requires continuous monitoring, evaluation, and improvement to ensure AI systems remain reliable, explainable, transparent, accountable, and aligned with evolving business and

DeepTAI in Practice: DeepTAI supports continuous monitoring of AI systems through observability, drift detection, explainability analysis, fairness assessment, and compliance monitoring. These capabilities help companies identify emerging risks early and continuously improve AI systems throughout their lifecycle.

3. Core Engineering Capabilities for Trustworthy AI

The AI lifecycle (see Section 2) defines the engineering process for developing AI systems. Successfully building trustworthy AI, however, requires more than following the lifecycle stages alone. Companies also need **core AI engineering capabilities** that enable them to develop, deploy, operate, and continuously improve AI systems that are transparent, explainable, accountable, secure, and reliable by design. These capabilities provide the traceability, auditability, and operational visibility needed to validate model behaviour, reproduce AI-driven decisions, and demonstrate compliance with governance and regulatory requirements.

While emerging AI regulations, standards, and governance frameworks define **what** companies must achieve, they provide only limited guidance on **how** these objectives can be implemented in practice. The engineering capability domains presented below provide a practical framework for operationalizing Trustworthy AI across the entire AI lifecycle.

The core AI engineering capability¹ domains for trustworthy AI include:

- **Governance²:**
 - Data governance
 - Risk management
 - Role-based access control
- **Transparency³:**
 - Explainable AI (XAI)
 - Logging and Traceability
- **Observability⁴**
 - Model and operational monitoring
 - human oversight
- **Security⁵:**
 - Cybersecurity

Together, these capability domains provide the engineering foundation needed to build, deploy, and operate Trustworthy AI systems at scale.

While all of these capability domains are essential, **Explainable AI (XAI)** plays a unique role because it supports multiple stages of the AI lifecycle - from feature engineering and model evaluation to deployment, monitoring, and continuous improvement. XAI helps developers, business users, regulators, and other stakeholders understand, validate, and trust AI-driven decisions. **Explainable AI (XAI)** is SHER DeepAI's primary area of expertise and a key capability of **DeepTAI**, our modular platform for Trustworthy AI (details see Section 5).

The following section explores the role of Explainable AI (XAI) across the AI lifecycle and demonstrates why it is a foundational engineering capability for developing Trustworthy AI.

¹ The capability domains reflect the core engineering capabilities addressed by SHER DeepAI. Future articles in the *AI Engineering Guide Series* will explore each capability and demonstrate how it is implemented in **DeepTAI**, SHER DeepAI's **modular platform for Trustworthy AI**.

² Governance brings together data governance, risk management, and role-based access control to ensure that AI systems are developed and operated with trusted data, appropriate safeguards, and clearly assigned responsibilities.

³ Transparency encompasses explainable AI (XAI), logging, and traceability to ensure that AI decisions can be understood, system behavior can be reconstructed, and engineering activities can be audited.

⁴ Observability enables continuous assessment of AI system performance and behavior through model and operational monitoring and human oversight, supporting reliable and responsible AI operation.

⁵ Security focuses on protecting AI systems, data, models, and supporting infrastructure from cyber threats and unauthorized access to ensure secure and reliable operation.

4. The Role of Explainable AI (XAI) Across the AI Lifecycle

Explainable AI (XAI) is a core engineering capability that supports the design, development, deployment, and operation of AI systems. Rooted in advanced machine learning research, XAI provides methods for interpreting, validating, and communicating the behaviour of increasingly complex AI models. As AI systems become more sophisticated, explainability plays a critical role in helping stakeholders understand how AI models reach their decisions.

As AI regulations, standards, and governance frameworks place increasing emphasis on transparency, human oversight, risk management, and trust, explainability is evolving from a desirable feature into a foundational capability for engineering Trustworthy AI.

As illustrated in Figure 2, XAI supports multiple stages of the AI lifecycle. During feature engineering, it helps identify meaningful and interpretable features. During model evaluation, XAI enables AI engineers to validate model behaviour, detect potential bias, and build confidence in AI-driven decisions. Once AI systems are deployed, XAI helps developers, business users, auditors, regulators, and decision-makers understand how AI systems reach their conclusions while supporting continuous monitoring as models evolve over time.

This cross-lifecycle role makes Explainable AI a foundational engineering capability for developing and operating Trustworthy AI systems. **Explainable AI** is a core focus of SHER DeepAI's research and engineering work and a central capability of the **DeepTAI platform** (see Section 5).



Figure 2: The AI Lifecycle and the Role of Explainable AI (XAI)

5. Operationalizing Core AI Engineering Capabilities with DeepTAI

DeepTAI is SHER DeepAI's **modular platform for Trustworthy AI**, designed to help companies to operationalize the engineering capabilities required to build, deploy, and operate AI systems that are trustworthy by design. Supporting the entire AI lifecycle, **DeepTAI** enables AI engineers, data scientists, governance teams, and business stakeholders to translate governance and regulatory requirements into practical engineering capabilities.

Built around the four engineering capability domains - **Governance, Transparency, Observability**, and **Security** - DeepTAI provides an integrated engineering environment for developing, deploying, monitoring, and continuously improving AI systems. By embedding these capabilities throughout the AI lifecycle, the platform makes Trustworthy AI easily accessible and guides to build AI systems that are transparent, explainable, accountable, secure, and compliant by design.

Get Started with Our Virtual TAI Assistant

Building trustworthy AI starts with understanding your specific use case and the regulatory and governance requirements that apply to it.

The **TAI Assistant** is SHER DeepAI's AI-powered management advisor for Trustworthy AI. It analyzes your AI use case, identifies applicable governance and regulatory requirements, assess potential risks, and determine the engineering capabilities needed to develop AI systems that are trustworthy by design. Rather than navigating complex AI regulations and governance frameworks independently, companies receive practical guidance tailored to their specific use case and aligned with the AI lifecycle.

Designed for both **business and technical stakeholders** - including business leaders, product managers, AI engineers, data scientists, legal and compliance professionals, and AI governance specialists - the Virtual **TAI Assistant** provides tailored recommendations based on a company's AI use case. It guides stakeholders to the appropriate DeepTAI modules and engineering capabilities, helping them accelerate the development, deployment, and adoption of AI systems that are transparent, explainable, accountable, secure, and compliant by design.

Scan the QR code to analyze your AI use case and discover the engineering capabilities required to build Trustworthy AI.



About the Authors

Dr. Silvia Moyses-Scheingruber is Co-Founder and Commercial Lead of SHER DeepAI. She is committed to advancing the responsible and trustworthy adoption of artificial intelligence. As a management consultant specializing in corporate strategy, business transformation, and innovation, she brings a strategic business perspective to the development and adoption of Trustworthy AI. She is also a lecturer in Strategy at the Chair of Engineering at Munich University of Applied Sciences.

Rajkumar Verma is Co-Founder and Technical Lead of SHER DeepAI. He is an AI engineer specializing in AI engineering, machine learning, Explainable AI (XAI), and the development of engineering capabilities for Trustworthy AI. He leads the design and development of **DeepTAI**, SHER DeepAI's modular platform for Trustworthy AI. His work focuses on making trustworthy AI engineering practical and accessible, enabling developers to build AI systems that are transparent, explainable, secure, and trustworthy by design.

About SHER DeepAI

SHER DeepAI is a research-driven AI lab based in Munich (Germany) dedicated to advancing the responsible and trustworthy adoption of artificial intelligence. Combining expertise in Explainable AI (XAI), AI engineering, AI optimization and AI governance, we develop practical engineering solutions that enable companies to build, deploy, and operate AI systems that are transparent, explainable, efficient, secure, and trustworthy by design.

Interested in learning more about Trustworthy AI?

Whether you are exploring AI governance, Explainable AI (XAI), AI engineering, or the DeepTAI platform, our team would be delighted to continue the conversation.

Visit: www.sherdeepai.com



Connect: SHER DeepAI on LinkedIn

Email: s.moyeses@sherdeepai.com or rkverma@sherdeepai.com