

Explainable Ai Frameworks For Decision Support In Data Governance To Enhance User Trust And Compliance

¹Ododoade Idowu Adewuyi

²Abdullah Oladoyin Akinde

³Grace Oluwaseun Ikudehinbu

¹Project Management, Northeastern University, Portland Maine

akindeabdullaholadoyin@gmail.com

²Department of Computer Science, Austin Peay State University; Clarksville, TN, United States

³Department of Accounting, Southern Illinois University Edwardsville, Edwardsville, IL United States

¹<https://orcid.org/0009-0005-8190-4486>

DOI: <https://doi.org/10.63084/cognexus.v1i04.248>

Abstract

This study sets out to understand the contribution of Explainable AI (XAI) techniques to data governance decision support with respect to organisational trust, accountability, transparency and compliance outcomes. To address this objective, a systematic literature review was performed on scientific publications and relevant policy documents. The retrieved articles were analysed with relevance coding to XAI methods, data governance applications and trust-related implications. A thematic analysis was then applied to group results into categories that allow for contrasting and comparing the different XAI approaches and their impact on data governance. The study revealed that four major XAI approaches exist (feature attribution, rule extraction, counterfactual explanations, and interpretable-by-design models), and they all contribute differently to governance workflows. The feature attribution methods, which include SHAP (SHapley Additive exPlanations) or LIME (Local Interpretable Model-agnostic Explanations), contribute to the transparency and accuracy of trust through the recognition of decision pathways, especially in relation to privacy-related risks. Rule extraction improves auditability and accountability by generating human-readable logic in finance and health governance to reinforce compliance. The counterfactual method allows fairness audits to be improved through bias mitigation, but there is still uncertainty concerning the conformation to discrimination law. On the other hand, the Interpretable models directly integrated transparency into the design of system architecture for public administration. User studies, trust scales, usability indicators, and audit outcomes were all evaluated, and they all showed improved stakeholder confidence and compliance readiness. IBM's AI Fairness 360, and Microsoft's Responsible AI dashboard which shows improved audit breaches and better organisational decision-making, are real-life cases, that back these findings. In general, trust, accountability and compliance are attributes strengthened by XAI, which has been proven to improve data governance. Companies must take the double approach: standardise formats and train stakeholders. Policy makers should, while making a regulatory framework, include XAI.

Keywords: Explainable AI (XAI), Data Governance, Trust and Transparency, Accountability and Compliance, Fairness in Decision Support

Introduction

Background and Rationale

Explainable Artificial Intelligence (XAI) today, is at the forefront of AI research. It focuses on the challenge of overcoming machine learning complexity, and it is a model that makes AI interpretable by humans (Longo et al., 2024). As noted by Doshi-Velez and Kim (2017), XAI is transparent, traceable, and accountable for the system. In high-stakes areas such as health care, finance, and governance, AI interpretation will become more important, as it will be more important where opaque decision-making can undermine trust and compliance. Recently, research has shown that explanation mechanisms such as feature attribution or counterfactuals can not only improve user trust but also help regulators in their jurisdictions comply with legislatures (Adadi & Berrada, 2018; Teixeira et al., 2025).

Following the trend and popularity of XAI, data governance has come to be a paramount part of organisational and regulatory concerns. Data governance is referred to as the “policies, standards, and processes used to help ensure that data is handled in an ethical, secure, and compliant manner” (Thirunagalingam, 2022). Specifically in relation to AI, governance is required to trace data provenance and maintain privacy and accountability through automated decisions. The use of transparency and auditability of AI processes is mandated by the General Data Protection Regulation (GDPR) in the European Union and similar tools globally (Voigt and Vondem Bussche, 2017).

There is a substantial research gap between XAI and data governance that needs to be bridged. Despite much of the XAI techniques research having been done in many technical contexts, XAI has not been applied for governance decision support. For example, while feature attribution methods could provide insight into “the precise features that contributed to the model’s prediction of compliance risk”, organizations may still not find it easy to translate these technical explanations into regulatory requirements (Samek et al., 2019; Samek, 2022). Interpretable-by-design models have semantic transparency, but they may not scale well.

As a result, there is a disconnect between the technical explainability of these tools and their organisational needs, creating challenges in operationalising trust, accountability, and compliance outcomes.

XAI and governance have only been studied together in parallel, rather than as integrated frameworks. There remain few studies that systematically assess the impact of different explanation designs on governance-related consequences, such as audit reinforcement, calibration of trust, and bias reduction (Chamola et al., 2023; Utomo et al., 2026; Olateju et al., 2024; Potluri, 2025; Pahune et al., 2025; Zhang & Wang, 2025; Samek et al., 2019; Samek, 2022). The absence of a common set of standards differentiating between various types of XAI discards the possibility of using XAI as a unifying tool

to support governance. The gap above needs to be filled with technical and regulatory perspectives that, while mathematically correct, would also be meaningful for compliance officers, auditors, and users.

This paper contributes to the field by integrating the fragmented technical literature on post-hoc and intrinsic explainability to existing regulatory guidelines. This paper identifies many different governance applications in isolated industry segments and does not simply reiterate the need for transparency, but presents a unified perspective on governance that maps each of the mathematical explanation types to the specific requirements of the law (e.g., GDPR, EU AI Act) and offers an actionable, multi-layered conceptual framework to provide developers with increased oversight and audit readiness.

The remainder of this manuscript is organised as follows: Section 2 provides the theoretical background, Section 3 the systematic selection process, Section 4 the synthesised results, and Section 5 provides a summarised framework and theoretical, managerial, and policy implications.

Research Questions and Objectives

This study aims to map the explainable AI(XAI) methods used in data-governance decision support with a particular focus on the effects of explanation design on trust, accountability, and compliance.

Research Questions

Which XAI approaches (feature attribution, rule extraction, counterfactual explanations, example-based explanations, interpretable-by-design models) are used for data-governance decision support?

How do these explanations improve trust-related outcomes (trust calibration, understanding, perceived transparency), and what evaluation methods are used (user studies, usability metrics, trust scales, audit outcomes)?

What compliance and governance requirements are most frequently linked to XAI (privacy, accountability, auditability, bias/fairness), and what gaps exist between technical explainability and regulatory/organisational needs?

THEORETICAL BACKGROUND

Overview of Explainable AI (XAI)

XAI refers to the ways and frameworks in which AI decision-making is understood by humans. Compared to older “black-box” models, XAI provides greater visibility into the effects of inputs on the final model outputs (Chinnaraju, 2025). Chamola et al. (2023) show that this interpretability is important in highly ethical or socially sensitive settings such as healthcare, finance and governance.

XAI is based on several key techniques. “Feature attribution methods” such as SHAP and LIME, are used to highlight the variables most important in a model’s output and thus allow for more precise identification of decision pathways (Utomo et al., 2026). “Rule extraction approaches” extract human-

readable rules from complex models and provide auditors and regulators with a consistent means to assess compliance (Pahune et al., 2025). The “counterfactual explanations” provide alternative scenarios in which small changes in input may impact outcome, which increases the fairness and accountability (Olateju et al., 2024). Decision trees or linear models are “interpretable-by-design” models, and enable transparency in their architecture rather than retrofitting explanations (Gupta, 2022).

XAI is not only important in technical terms. It acts as a governance agent, enabling regulators to check for fairness, auditors to follow responsibility, and end-users to calibrate trust in AI systems (Potluri, 2025). As Shri (2025) argues, transparency-driven models increasingly require interpretability as a compliance requirement, indicating the shift from efficiency-based AI to trust-centred governance.

Data Governance and Decision Support Systems

Data governance can be understood as the policies, standards, and processes that allow for the ethical, secure, and compliant use of data. Traditional governance is concerned with the categorisation, auditing, and with legal adherence being the main concern (Gupta, 2022). However, these understanding conceptualize data as something more or less fixed and do not take into consideration the dynamic and adaptive character of AI systems (Zhang & Wang, 2025).

Among them are access control to data, or the trade-off between openness and privacy; data quality management, or the maintenance of accuracy and consistency among systems; privacy, particularly with requirements like those in the GDPR; and auditability, or a transparent documentation of data flows and decision-making processes (Gudepu & Eichler, 2024). In reality, this has translated to unethical governance in the form of prejudiced hiring tools and biased healthcare diagnostic tools (Gupta, 2022).

These AI-based decision support systems represent an opportunity to strengthen governance. The use of AI will simplify lineage tracing, anomaly tracing and compliance checks, allowing for better oversight and less human error (Pahune et al., 2025). But these systems create opacity without explanation. As noted by Chamola et al. (2023) technical accuracy is required for the purposes of governance as well as auditability and ethics of AI decisions.

Thus, AI can augment efficiency and scale up monitoring, but can also create different risks of opacity. So it becomes necessary for these systems to have explanation as part of them to provide trustworthiness and social legitimacy to forms of governance (Shri, 2025).

XAI and Data Governance Intersection

According to Utomo et al. (2026), the combination of explainability and data governance models can enhance trust, accountability, and transparency, and fill in the gaps of a traditional compliance-driven approach.

XAI methods (SHAP and LIME), are able to track the flow of information in a way that enables the auditor to follow the decision pathways and assign culpability for errors (Hermosilla et al., 2025). The method of rule extraction of rules is a form of accountability because it enables governing actors to interrogate the logical processes of the algorithm in the context of regulation (O'Brien et al., 2025; Pahune et al., 2025). These methods, because factual accounts enhance fairness and bias perceptions, allow organizations to show alignment with moral expectations (De Schutter & De Cremer, 2024; Goethals et al., 2023; Olateju et al., 2024).

The importance of interpreting AI's decisions in the governance process from a socio-economic perspective is their social effect. Reversibility implies that healthcare professionals can validate AI-powered medical diagnoses, finance is guaranteed as following regulations when implementing machine learning for automated lending, and citizens can contest the decisions by algorithms in public administration (Chamola et al., 2023; Maleki Varnosfaderani & Forouzanfar, 2024; Mennella et al., 2024). Potluri (2025) argues that "explainability must be a compliance requirement of policy-aware systems" while Shri (2025) advocates for the utility of models of adaptive governance by employing the principle of use in an explainable manner.

Though there has indeed been progress, issues remain. The concern for transparency instead of efficiency remains lagged in a lot of governance structures, creating a trust deficit in the long run (Zhang & Wang, 2025). Few have conducted empirical analyses of the ways in which XAI might influence governance, and most of the rest are theoretical. This is a space in which hybrid models need to be deployed that marry technical interpretability with institutional accountability to ensure that AI systems are efficient computationally and legitimate socially.

Baseline Taxonomy of XAI Techniques in Governance

To anchor the systematic analysis that follows, Table 1 provides a structural taxonomy of the primary XAI approaches evaluated in current literature, outlining their basic mechanism and primary data governance objective.

Table 1: Taxonomy and Primary Governance Objectives of Evaluated XAI Approaches

XAI Approach	Core Technical Mechanism	Primary Governance Objective
Feature Attribution	Computes local feature weights (e.g., SHAP, LIME).	Algorithmic transparency and privacy validation.
Rule Extraction	Translates black-box logic into human-readable IF-THEN rules.	Procedural auditability and compliance verification.
Counterfactuals	Generates minimal input changes to	Bias detection, fairness auditing,

	show alternative outcomes.	and discrimination checks.
Interpretable-by-Design	Employs structurally transparent architectures (e.g., Decision Trees).	Inherent accountability and post-hoc-error-free governance

METHODOLOGY

Systematic Literature Review Approach

The present work uses a systematic literature review (SLR) approach to provide a comprehensive account of studies on Explainable AI (XAI) and data governance. The advantages of an SLR include the ability to use a standardised, replicable, and open method to find, assess, and integrate pertinent research (Carrera-Rivera et al., 2022; Smela et al., 2023).

Academic databases, including IEEE Xplore, ACM Digital Library, SpringerLink and ScienceDirect, and policy-related repositories like OECD and EU publications, were used in gathering data. Keywords included “Explainable AI,” “XAI,” “data governance,” “decision support,” “trust”, “compliance,” and “accountability.” Boolean operators were used to refine queries and catch terminology variations (Azarian et al., 2023). The search period was bounded between January 2018 and May 2026 to capture the most recent advancements in post-hoc explainability and evolving global data regulations (such as the operationalisation of GDPR).

To ensure reproducibility, a formalised multi-stage selection process was applied:

Formulation and Duplicate Removal: Initial database search yielded 412 records. Automated and manual title deduplication removed 43 duplicate entries.

Screening: The remaining 369 titles and abstracts were screened against initial relevance boundaries; 146 articles were excluded for focusing strictly on unrelated computer-vision or non-governance technical tasks.

Eligibility: 223 full-text articles were rigorously assessed.

Quality Assessment Criteria: Studies were excluded if they failed to meet three baseline quality indicators: (a) clearly defined XAI methodologies, (b) explicit articulation of the governance, regulatory, or organisational application context, and (c) peer-reviewed institutional framework or empirical backing. Non-peer-reviewed white papers lacking methodological transparency were excluded. Following this quality filter, 174 papers were excluded, leaving a final synthesis cohort of 49 studies. This selection flow is mapped precisely in the PRISMA flowchart (Figure 1).

Data Extraction and Synthesis Process

The individual papers chosen were each marked or coded based on relevance to various XAI

techniques including, Feature Attribution, Rule Extraction, Counterfactuals, Interpretable-by-design, Governance Applications including privacy, auditability and accountability, and Effects on trust including transparency, user confidence and fairness. Extraction templates included information on study methods, outcome measures, and reported outcomes.

The synthesis employed thematic analysis and grouped studies according to types of XAI methods and governance domains. A comparison matrix was created to help understand the impacts of explanation design in governance outcomes and to identify patterning, strengths, and gaps in the literature.

From the databases, 412 articles were retrieved, and after review of titles and abstracts for irrelevant articles, the authors excluded 189 articles as being either not applicable or not adequate. All remaining 223 articles were reviewed in aggregate. In this process, 174 methodological articles which were either non-relevant or lacked complete data were excluded from analysis. A total of 49 studies were included in the study.

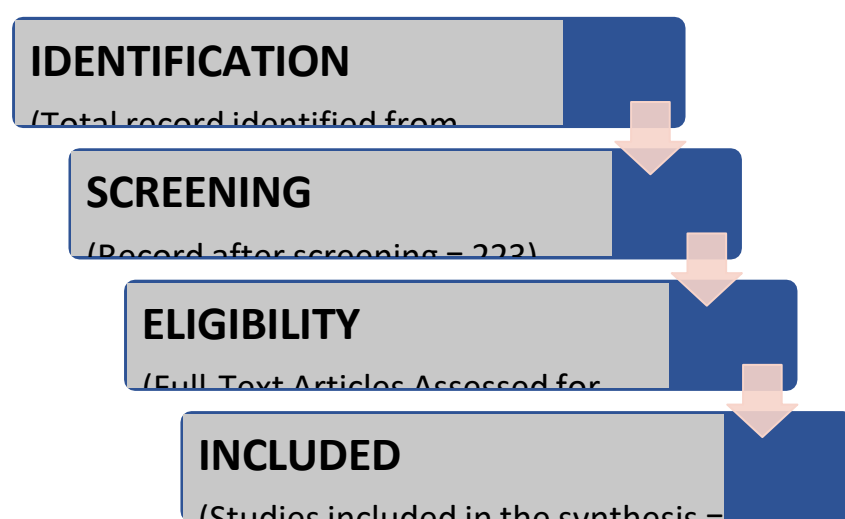


Figure 1: PRISMA Flowchart

Limitations

The following limitations were acknowledged in this study. Firstly, relying on published academic literature alone may exclude industry practices or proprietary frameworks that are not publicly documented. Second, the diversity of assessment approaches within studies makes comparisons difficult because trust outcomes are often measured on several scales and metrics. Third, XAI and governance systems are rapidly changing, and findings can be inaccurate as new approaches and rules emerge.

RESULTS

XAI Approaches in Data Governance

There have been a number of Explainable AI (XAI) approaches that have been applied to data governance workflows. These methods differ in their technical design and governance relevance, but they together show that explanation can improve compliance, accountability and trust in AI-driven decision support systems.

Feature Attribution Methods

One of the most widespread forms of feature attribution was SHAP (SHapley Additive Explanations) and LIME (Local Interpretable Model-agnostic Explanations) for governance. (Mathew et al., 2025; Utomo et al., 2026; Chamola et al., 2023; Linardatos et al., 2020; Saarela & Podgorelec, 2024; Teixeira et al., 2025). These strategies measure the influence each input feature has on model predictions, enabling auditors and compliance officers to track decision-making.

SHAP has been used to determine which factors have the greatest impact on risk scores in privacy governance (Antonini et al., 2024; Lin & Wang, 2025). SHAP explanations for healthcare compliance, for example, revealed the disproportionate influence of age and genetics on risk assessments that led

regulators to reconsider implications for fairness (Gupta, 2022). Just almost the same as financial governance, LIME explanations explain why some loan applications were identified as high risk, which in turns enables transparency in credit decision-making (Shri, 2025).

IBM's AI Fairness 360 (AIF360) toolkit contains a full set of fairness metrics and bias mitigation algorithms for fairness audits to help organisations show compliance with anti-discrimination laws. SHAP and LIME, to a large extent, are used as interpretability techniques and are also often used in conjunction with AIF360 in practice so as to increase transparency and trust in fairness assessments (Varshney, 2024; Choudhary, 2025). KPMG's 2024 global survey indicated that 72% of companies are already selectively implementing AI in financial reporting, predicting the rate to 99% within three years, and Gartner (2024) said that 41% of internal audit teams are considering generative AI.

Rule Extraction Approaches

Rule extraction technologies extract human-readable rules from complex models and provide clear explanations that correspond to governance requirements (Pahune et al., 2025). These approaches are particularly useful for data quality governance where compliance officers are asked to verify the integrity of algorithms.

For instance, in the case of financial governance, rules have been extracted from credit scoring models which allow for specific decisions that auditors can analyse (Bahlool et al., 2026). This has made audit readiness to be more conducive, since regulators can now see whether decision rules are consistent with legal requirements (Shri, 2025). As reported by Chamola et al. (2023), the rule extraction method

has been used in healthcare governance to detect anomalies across patient data and to ensure that diagnostic models follow privacy regulations.

In 2025, European banks began to include rules-based explanations in fraud detection systems to allow auditors to trace suspicious transactions back to specific decision rules. Scholars and industry publications indicate that such approaches increase regulator confidence and reduce compliance risks by making decisions based on AI more transparent and accountable (Bellamy et al., 2019; Kokina & Davenport, 2017).

Counterfactual Explanations

Counterfactual explanations offer alternative scenarios in which small changes in input could change outcomes (Olateju et al., 2024). These explanations are relevant in contexts of governance where fairness and bias mitigation are paramount.

Counterfactuals for different user attributes affect access permissions are provided in access control governance. For example, in hiring governance, counterfactuals showed that altering gender or ethnicity variables affected hiring outcomes and enabled organizations to identify and mitigate bias (Potluri, 2025). Counterfactuals have been used in privacy governance to show the effects of removing sensitive attributes (e.g. ethnicity) on risk scores in order to support fairness audits.

Google's What-If Tool (WIT) provides counterfactual explanations for model predictions to help organizations test fairness across demographic groups. In practice, healthcare providers and other regulated industries have sought counterfactual interpretations of predictive AI systems to explain consistent outcomes across populations and thus comply with regulatory requirements for non-discrimination (Pandey, 2025; Wexler et al., 2019).

Interpretable-by-Design Models

Interpretable-by-design models emphasize transparency in architecture, meaning explanations are not retrofitted but inherent (Chinnaraju, 2025). Decision trees, linear regression and rule-based classifiers are some of the examples. These models may suffer from a lower predictive accuracy than deep learning, but are especially useful in transparency driven governance workflows.

For privacy governance, access rules are based on interpretable models that can be audited by itself, so there is less need for post-hoc explanations (Zhang & Wang, 2025). Credit scoring has been scored using decision trees in financial governance, providing explicit decision pathways that regulators can easily audit.

In public administration, more agencies have adopted decision tree models for citizen eligibility assessments. These models have a clear structure that auditors can trace all decisions, and thus improve transparency and reduce citizen concerns about opaque outcomes (Molnar, 2025; Kokina & Davenport, 2017).

In the year 2025, Li as an author, claimed that " by allowing auditors to essentially trace the path of decisions made in the analysis of citizen eligibility, decision tree models and other forms of interpretable AI applications can contribute to increased transparency in public administration". On the other hand, these solutions reduce the transparency of government processes, while providing enhanced accountability and thus lessen the friction amongst citizens regarding the fairness of AI-driven governance.

Use Cases Across Governance Workflows

Several governance workflows where XAI methods have been applied were identified in this study. They are;

Access Decisions: Feature attributions and counterfactuals make it clear to users why they are being provided access or not and so add to the accountability of identity management systems (Ribeiro et al., 2016; Wexler et al., 2019).

Privacy Risk Assessment: SHAP and LIME show which variables are sensitive, allowing the compliance officer to assess the privacy risks in the context of GDPR (Lundberg & Lee, 2017; Ribeiro et al., 2016).

Data Quality Governance: Methods of rule extraction identify anomalies in the lineage of data, allowing governance frameworks to remain sound (Molnar, 2025).

Audit and Compliance: They offer open decision logics, which helps these models be more audit ready and regulation compliant (Kokina & Davenport 2017; Li, 2025).

Microsoft's Responsible AI Dashboard offers governance audits of feature- attribution, counterfactuals, and fairness measures (Microsoft, 2025). It was reported that the dashboard was being used by multinational companies to legitimize HR recruitment AIs, to find biases in candidate selection, and to find ways to fix that (Naoum et al., 2026; Soleimani et al., 2025). These findings also imply, that feature attribution methods provide very fine-grained insight into the rationale behind a specific decision, set up the explanation of rules extraction, provide counterfactuals for fairness auditing, and help in making the model interpretable as a part of the system's design.

But there are also challenges as most of the XAI techniques are not accessible to non-technical stakeholders that require technical know-how (Chamola et al., 2023). Also, standard disclosures are needed by regulatory frameworks but the formats and levels of clarity of XAI techniques vary

immensely (Shri, 2025). There is a need to fill in this void in a blend of technical innovative and regulatory comprehension.

Trust Outcomes and Evaluation Methods

Using Explainable AI (XAI) to enhance data governance workflows creates trust. Many have argued that an explanation increases trustworthiness, perceived transparency and accountability and thus legitimacy of compliance and governance.

XAI Contributions to Trust Outcomes

Trust Calibration

Trust calibration is a process that helps to align trust between the user and the actual reliability of AI systems. Without explanation, users can over-reliance on or dismiss AI outputs, which results in poor governance. In addition, features attribution methods such as SHAP and LIME enable better calibration by clarifying what variables influence decisions (Utomo et al., 2026). In governance environments, calibrated trusts ensure compliance officers are able to use AI recommendations in the appropriate way. For example, in healthcare governance, SHAP explanations helped clinicians to distinguish diagnostic outputs, thereby preventing blindness and over-scepticism (Gupta, 2022).

Understanding and Transparency

Rule extraction and interpretable models make user understanding more intuitive by providing clear decision criteria. Rule-based explanations for credit risk were made available in financial governance and then able to permit auditors to check for compliance with regulatory requirements (Shri, 2025). Transparency is particularly important in public administration, where citizens need to be able to challenge algorithmic decisions. Case studies have shown that the use of interpretable models in assessments for citizen eligibility improved transparency and reduced complaints about opaque outcomes (Chamola et al., 2023).

Perceived Fairness and Accountability

Counterfactual explanations also enhance the illusion of fairness by demonstrating how different inputs can lead to different outcomes. Counterfactuals suggest that it altered hiring decisions, which became a proxy for gauging employee performance, in trying to identify and penalize bias (Potluri, 2025).

Accountability is more readily achieved when there is a visible line of accountability for the outcomes of the algorithms, a necessity that is beginning to be seen as important within governance structures (Pahune et al., 2025).

Deloitte's 2024 report found that those companies that had begun to implement healthy AI were more aware of stakeholder trust in the governance process. As with IBM's AI Fairness 360 toolkit, combining metrics of fairness with methods of interpretability such as feature attributions or counterfactuals has been shown to reinforce fairness assessments and increase accountability of organizations (Bellamy et al., 2018).

Evaluation Methods in Studies

Several evaluation methods used to measure trust outcomes within data governance are identified include;

User Studies

Controlled experiments are the most common evaluation method. Participants interact with XAI systems and report trust levels, often through surveys or interviews (Utomo et al., 2026). For example, clinical trials conducted in healthcare governance showed that clinicians felt more confident when diagnostic systems were presented with SHAP explanations.

Trust Scales

User confidence is often quantified via standard measures, such as the Trust in Automation Scale (Olateju et al., 2024). They are used to describe factors such as reliability, predictability and transparency. Trust scales, for example, in financial governance, showed that auditors assessed rule-based explanations as more trustworthy than black-box outputs.

Usability Metrics

Usability measures explanation clarity, relevance, and cognitive load. Studies indicate that explanations need to balance detail and simplicity; too technical outputs overwhelm non-expert stakeholders (Chamola et al., 2023). Usability measures are important in governance contexts to ensure that

explanations are accessible to all stakeholders, including regulators, auditor and citizens.

Audit Outcomes

Audit readiness is another form of assessment. Companies evaluate XAI justification in terms of compliance with norms of transparency and accountability mandated by regulations (Zhang & Wang, 2025). GDPR, for instance, requires companies to provide "meaningful information" about automated decision making. These requirements have been contrasted to XAI techniques such as rule extraction and interpretable models for better auditability.

Challenges in Measuring Trust

Despite these advances, there are still issues associated with evaluating trust outcomes. These include;

The multidimensional nature of trust, which explains that trust is influenced by cultural, organisational, and regulatory factors, making it difficult to measure consistently.

The stakeholder diversity, which requires the tailoring of explanations to stakeholder needs. This is because technical explanations may suffice for auditors but overwhelm end-users (Chamola et al., 2023).

There is limited empirical evidence as many evaluations rely on small-scale user studies, limiting generalizability. In addition, large-scale and longitudinal studies are needed to validate findings across diverse governance contexts.

The issue of regulatory alignment as regulators often struggle to interpret technical outputs, though, XAI methods improve transparency. Given as an example, SHAP explanations may be mathematically rigorous but difficult for non-technical stakeholders to understand (Shri, 2025).

Compliance and Governance Requirements

Compliance Needs Related to XAI

Privacy and Risk Management

Privacy is a central compliance requirement in AI governance. Regulations such as the EU General

Data Protection Regulation (GDPR) mandate organizations to provide “meaningful information” about automated decisions (European Union, 2024; Voigt & Von dem Bussche, 2017). Explainable AI (XAI) methods enhance privacy protection by clarifying how variables influence outcomes. Tools like SHAP identify which attributes drive risk scores, enabling regulators to verify appropriate use of personal data (Utomo et al., 2026). Counterfactual explanations further support compliance by showing how removing sensitive attributes—such as ethnicity or health status—affects outcomes, aiding fairness audits (Olateju et al., 2024).

Privacy and risk tracking are a set of compliance-requiring behaviours in the AI Act in the EU and the NIST AI Risk Management Framework (AI RMF) in the NIST (EU, 2024; NIST, 2023). In these frameworks, specific documentation and mitigation requirements for high-risk AI are required (EU, 2024; NIST, 2023). XAI techniques, such as feature attribution metrics (SHAP/LIME), can support these frameworks by measuring the contribution of attribute properties to the risk score, allowing compliance officers to provide evidence that they have adhere to data minimization principles.

Accountability and Global Principles

Accountability requires that organisations can trace the responsibility for algorithmic outcomes. International mandates, including the OECD AI Principles (OECD, 2024) and relevant ISO/IEC AI governance standards (such as ISO/IEC 38505-1), dictate that organisations must maintain clear lines of accountability for algorithmic workflows. Rule-based explanations have been used in financial governance to show accountability in credit scoring, providing explicit decision criteria for auditors to ask questions of model logic (Pahune et al., 2025) and ensuring that decision pathways are aligned with regulatory standards (Shri, 2025). Including interpretable-by-design models as part of the system’s architecture also has the advantage of representing systems more accurately than post-hoc explanations are able to do (Chinnaraju, 2025).

Auditability

Auditability is an aspect that has been very important in compliance frameworks like GDPR and other regulations. Decision trees and interpretable models enhance auditability through transparent decision pathways (Zhang & Wang, 2025) while features attribution also aids auditability in that auditors can duplicate decisions pathways and determine who is to be held accountable for the outcomes (Chamola et al., 2023). The practical impact shows that organizations using XAI methods are more prepared for

audits.

Fairness

Fairness is also becoming a compliance requirement as well, especially regarding recruitment, lending, and healthcare governance. Counterfactual explanations and fairness sensitive models help organizations in recognizing and addressing bias (Potluri, 2025). Fairness audits employ XAI to show that there is no discrimination, and that equal opportunity regulations are being obeyed in the hiring regime. IBM's Artificial Intelligence Fairness 360 toolkit offers objective measures of fairness as well as counterfactual explanations, allowing businesses to show adherence to anti-discrimination legislation.

Identified Gaps Between Technical Explainability and Regulatory Needs

Important gaps still exist between technical explainability and regulatory or organizational requirements.

The output of many XAI techniques is mathematically correct, yet technically complex for non-technical stakeholders to interpret. SHAP values, for instance, offer a very detailed view, but may be too technical for a regulator without a strong data science background.

Normative explanations are often required to be uniform, while explanations produced by XAI methods are heterogeneous in form and clarity. The absence of standardization makes it difficult to conduct compliance audits

Counterfactual explanations can be fair but not necessarily discriminate by legal definition. Regulators need justifications that correlate with legal definitions rather than with technical parameters.

Several governance use cases of XAI are still hypothetical and lack empirical evidence. The utility of XAI in contexts of compliance has yet to be proven in large-scale, longitudinal studies.

Even when technical, XAI approaches are valid, they are typically not implementable in organizational governance workflows.

SYNTHESIS OF FINDINGS

This systematic review established that Explainable AI (XAI) methods are important in enhancing governance outcomes in terms of trust, accountability, transparency and compliance. XAI methods were defined in the literature as feature attribution, rule extraction, counterfactual explanations, and interpretable-by-design models. Each method is useful to governance workflows and does so

independently, but in different ways depending on the context and stakeholder requirements.

Feature attribution methods such as SHAP and LIME were found to be highly useful for clarifying decision paths, particularly in regard to privacy risk assessment and audit readiness. These methods provide more detailed information about the effect of variables on outcomes, which is helpful to calibrate

trust and to ensure transparency. However, their technical complexity can often prevent access for non-expert stakeholders.

The extraction of rule steps proved to be very consistent with accountability and auditability requirements. These methods provide a human-readable way for auditors and regulators to confront algorithmic logic directly. Rule extraction in financial governance improved audit readiness by clarifying credit scoring criteria, while rule extraction in healthcare governance improved detection of anomaly in patient data lineage.

Counterfactual explanations are especially useful in fairness audits, as they explain the impact of changing variables on outcomes. Counterfactuals found potential bias in hiring decisions in recruitment governance, which facilitated corrective action for the organizations. Counterfactuals promote perceived fairness, but are prone to fall outside of legal definitions of discrimination, which highlight the gap between technical explainability and regulatory needs.

“Interpretable-by-design” models incorporated transparency as part of the system itself, rather than having to find explanations after the fact. These models were also seen as particularly effective within public administration, where accountability to citizens and transparency are vital. But they are not as accurate as deep learning models, and they cannot scale in complex governance environments.

Comparatively, the effectiveness of these approaches are summarised as follows:

XAI Method	Governance Use Cases	Strengths	Limitations
Feature Attribution (SHAP, LIME)	Privacy risk assessment, audit readiness	Granular insights, transparency, trust calibration	Technical complexity, limited accessibility
Rule Extraction	Data quality governance, financial audits	Human-readable rules,	May oversimplify

		accountability, auditability	complex models
Counterfactual Explanations	Recruitment fairness, access control	Bias detection, fairness, perceived accountability	Misalignment with legal definitions
Interpretable Models	Public administration, citizen eligibility	Inherent transparency, auditability, trust	Lower accuracy, limited scalability

Comparing the two methods reveals that they have different trade-offs for their governance utility. Though both SHAP and LIME provide post-hoc local feature attribution, SHAP is preferred in governance contexts as its game-theoretic approach yields efficiency, symmetry, and monotonicity. It is therefore consistent mathematically through compliance audits, whereas LIME depends on local perturbations, which can produce non-reproducible explanations for the same data input.

In contrast, rule extraction techniques are superior if absolute regulatory auditability over model optimization is a priority, such as in areas with high compliance risk such as fraud detection and data lineage verification. In such situations, the value of converting complex neural pathways into readable logic outweighs the modest predictive performance.

But, counterfactual explanations can be problematic in fairness audits because it can indicate how changing a small amount of data impacts algorithmic outcomes (e.g., changing a gender flag), but the algorithmic parameters may not match current case-law definitions of systemic discrimination or protected characteristics.

Interpretable-by-design models are therefore more suited to basic civic administrative contexts, such as assessment of public citizen eligibility, where absolute structural transparency and elimination of post-hoc explanation errors are legal requirements to maintain institutional legitimacy, even if it limits the general predictability of the model.

Overall, the findings suggest that while XAI approaches may provide some value to governance outcomes, they do not necessarily perform as well in practice given the requirements of stakeholder and regulatory contexts. The gap between technical outputs and legal or organizational requirements remains a crucial challenge.

The XAI-Data Governance Decision Support Framework (XAI-DGF)

In order to bridge the identified gaps between technical explainability outputs and regulatory needs, a new framework is proposed. It is the Explainable AI-Data Governance Framework (XAI-DGF). As

shown in Figure 2, this conceptual framework in operation will show how distinct XAI methods translate through the operational governance mechanisms to achieve calibrated trust and compliance outcomes.

The framework operates across four interconnected layers:

Technical XAI Layer: Consists of the core explainability pipelines—Feature Attribution (SHAP/LIME), Rule Extraction, Counterfactuals, and Interpretable-by-Design models.

Operational Governance Layer: Maps these technical outputs directly onto localised workflows, specifically Access Decisions, Privacy Risk Assessments, & Data Quality Management.

Sociotechnical Evaluation Layer: Evaluates the clarity and utility of explanations through mixed-method assessments, including standardised Trust Scales (e.g., Trust in Automation Scale), Usability Metrics (cognitive load), and User Studies.

Strategic Outcome Layer: Realises the final institutional goals of calibrated user trust, strict regulatory compliance readiness (e.g., GDPR), verifiable accountability, and minimised legal/algorithmic bias risk.

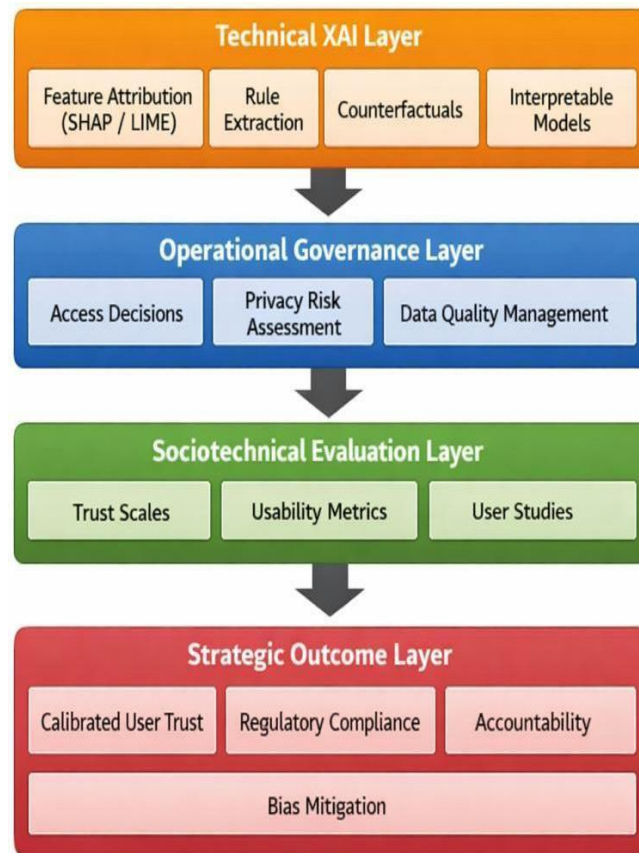


Figure 2: The Conceptual Architecture of the XAI-Data Governance Framework (XAI-DGF)

Discussion of Multi-Layered Implications

Theoretical Implications

From a theoretical perspective, this study recontextualizes explainability as an active, structural mechanism of sociotechnical governance rather than a static computer-science logging tool. By demonstrating how distinct post-hoc styles correspond to varying organizational frameworks, we provide a unified taxonomy that bridges the conceptual gap between mathematical transparency and human trust dynamics.

Managerial and Enterprise Implications

For enterprise managers and data controllers, these findings dictate the adoption of a hybrid XAI infrastructure. Organizations should not rely on a singular technical tool; instead, data management

leaders must deploy multi-tiered explanation delivery pipelines, such as utilizing granular SHAP attribution for technical data engineering audits while leveraging human-readable rule extractions to ensure operational compliance clarity for business division leads.

Policy and Regulatory Implications

For policymakers, this study highlights that future AI statutory mandates must move beyond generic demands for "meaningful information." Regulatory bodies should actively standardize compliant reporting documentation frameworks, leveraging established paradigms like Microsoft's Responsible AI dashboard or formalized Model Cards, to establish uniform criteria that reconcile technical model transparency with clear, enforceable legal definitions of non-discrimination and auditable corporate stewardship.

Agenda for Future Research

Based on the synthesis of the current literature, five critical avenues are proposed to advance the integration of XAI within data governance architectures:

Longitudinal Governance Frameworks: Future research must transition from cross-sectional evaluations to long-term studies that track how XAI integration affects compliance decay and organizational trust calibration over multiple fiscal or regulatory cycles.

Empirical Validation of Mixed Frameworks: There is a critical need to empirically test and quantify the operational efficacy of the proposed XAI-DGF architecture within live enterprise environments to observe real-world performance trade-offs.

Sector-Specific XAI Architectures: Researchers should develop tailored, domain-specific explainability implementations designed explicitly around the nuanced legal demands of distinct operational environments, such as medical diagnostics or automated financial lending.

Standardised Governance Metrics: The community must establish unified, mathematical evaluation metrics that can calculate an explanation's regulatory compliance readiness, moving past subjective usability scales.

Cross-Jurisdictional Regulatory Harmonisation: Comparative policy analyses are needed to investigate how technical XAI outputs can be structurally designed to simultaneously satisfy diverging legal standards across global frameworks, such as navigating the intersections of the EU AI Act and the NIST AI RMF.

CONCLUSION

In conclusion, this systematic review addresses a critical gap in literature by integrating fragmented post-hoc explainability methodologies with formal compliance mechanisms, offering a unified architectural perspective on trustworthy decision support.

For XAI to be successfully incorporated into governance process there needs to be an organizational commitment and policy support. In practice, organizations could look at hybrid forms, which is, the combination of various forms of XAI and tailoring the explanations to distinct user profiles. Protocols for explaining, well-trained compliance officers, and transparency from the very beginning of the development of any system are critical to effectiveness.

Policymakers can also play an important role in fostering the enactment of XAI. Regulations have to follow in order to make explainability a point of compliance but also help shape what kind of explanations and what criteria for evaluation. For instance, Google's Model Cards and Microsoft's Responsible AI dashboard are examples of how regulation might look by having technical documentation that correlates technical output with governance expectations. Policymakers can be a, if not the, facilitator to ensure considerably easier cooperation amongst AI researchers, regulators, and industry to ensure that XAI serves to improve the legitimacy of governance, to minimize the problem of non-compliant audits, and to build up public trust in AI-driven decisions.

REFERENCES

- Adadi, A., & Berrada, M. (2018). Peeking inside the black-box: a survey on explainable artificial intelligence (XAI). *IEEE Access*, 6, pp. 52138-52160, 2018, doi:10.1109/ACCESS.2018.2870052
- Antonini, A. S., Tanzola, J., Asiain, L., Ferracutti, G. R., Castro, S. M., Bjerg, E. A., & Ganuza, M. L. (2024). Machine Learning model interpretability using SHAP values: Application to Igneous Rock Classification task. *Applied Computing and Geosciences*, 23, 100178. <https://doi.org/10.1016/j.acags.2024.100178>

- Azarian, M., Yu, H., Shiferaw, A. T., & Stevik, T. K. (2023). Do We Perform Systematic Literature Review Right? A Scientific Mapping and Methodological Assessment. *Logistics*, 7(4), 89. <https://doi.org/10.3390/logistics7040089>
- Bahlool, R., Hewahi, N., & Elmedany, W. (2026). Performance, Fairness, and Explainability in AI-Based Credit Scoring: A Systematic Literature review. *Journal of Risk and Financial Management*, 19(2), 104. <https://doi.org/10.3390/jrfm19020104>
- Bellamy, R. K. E., Dey, K., Hind, M., Hoffman, S. C., Houde, S., Kannan, K., Lohia, P., Martino, J., Mehta, S., Mojsilovic, A., Nagar, S., Ramamurthy, K. N., Richards, J., Saha, D., Sattigeri, P., Singh, M., Varshney, K. R., & Zhang, Y. (2018, October 3). *AI Fairness 360: an extensible toolkit for detecting, understanding, and mitigating unwanted algorithmic bias*. arXiv.org. <https://arxiv.org/abs/1810.01943>
- Carrera-Rivera, A., Ochoa, W., Larrinaga, F., & Lasa, G. (2022). How-to conduct a systematic literature review: A quick guide for computer science research. *MethodsX*, 9, 101895. <https://doi.org/10.1016/j.mex.2022.101895>
- Chamola, V., Hassija, V., Sulthana, A. R., Ghosh, D., Dhingra, D., & Sikdar, B. (2023). A review of trustworthy and explainable artificial intelligence (XAI). *IEEE Access*, 11, 78994–79015. <https://doi.org/10.1109/ACCESS.2023.3300138>
- Chinnaraju, A. (2025). Explainable AI (XAI) for trustworthy and transparent decision-making: A theoretical framework for AI interpretability. *World Journal of Advanced Engineering Technology and Sciences*, 14(3), 170–207.
- Choudhary, A. (2025, December 20). *Machine learning explainability: Implementing model interpretability with LIME, SHAP, and AI Fairness 360*. Medium. <https://medium.com>
- De Schutter, L., & De Cremer, D. (2024). How Counterfactual Fairness Modelling in Algorithms Can Promote Ethical Decision-Making. *International Journal of Human-Computer Interaction*, 40(1), 33–44. <https://doi.org/10.1080/10447318.2023.2247624>
- Deloitte. (2024, January 15). *New Deloitte survey finds expectations for Gen AI remain high, but many*

are feeling pressure to quickly realize value while managing risks. Deloitte Global.
<https://www.deloitte.com/global/en/about/press-room/gen-ai-survey.html>

Doshi-Velez, F., & Kim, B. (2017, August). *Towards a rigorous science of interpretable machine learning*. In *Proceedings of the 2017 ICML Workshop on Human Interpretability in Machine Learning (WHI 2017)*. arXiv. <https://arxiv.org/abs/1702.08608>

European Union. (2024). Regulation (EU) 2024/1689 of the European Parliament and of the Council of 13 June 2024 laying down harmonised rules on artificial intelligence (Artificial Intelligence Act).

Gartner (2024). Gartner survey shows 41 percent of internal audit teams use or plan to use generative AI this year. Gartner. Available online at: <https://www.gartner.com/en/newsroom/press-releases/2024-03-11-gartner-survey-shows-41-percent-of-internal-audit-teams-use-or-plan-to-use-generative-ai-this-year>

Goethals, S., Martens, D., & Calders, T. (2023). *PreCoF: counterfactual explanations for fairness*. *Machine learning*, 1–32. Advance online publication. <https://doi.org/10.1007/s10994-023-06319-8>

Gudepu, B. K., & Eichler, R. (2024). The role of AI in enhancing data governance strategies. *International Journal of Acta Informatica*, 3(1), 169–187. <https://www.yuktabpublisher.com/index.php/IJAI/article/view/196>

Gupta, N. N. (2022). How inadequate data governance frameworks lead to unethical outcomes in artificial intelligence systems. *International Journal of Science and Research Archive*, 7(1), 580–590. <https://doi.org/10.30574/ijrsra.2022.7.1.0274>

Hermosilla, P., Berríos, S., & Allende-Cid, H. (2025). Explainable AI for Forensic Analysis: A Comparative Study of SHAP and LIME in Intrusion Detection Models. *Applied Sciences*, 15(13), 7329. <https://doi.org/10.3390/app15137329>

Kokina, J., & Davenport, T. H. (2017). The emergence of artificial intelligence: How automation is changing auditing. *Journal of Emerging Technologies in Accounting*, 14(1), 115–122. <https://doi.org/10.2308/jeta-51730>

- KPMG (2024). AI transforming financial reporting globally with near universal adoption expected in the next three years. KPMG. Available online at:
<https://kpmg.com/xx/en/media/press-releases/2024/05/ai-transforming-financial-reporting-globally-with-near-universal-adoption-expected-in-the-next-three-years.html>
- Li, C. (2025). AI-driven governance: Enhancing transparency and accountability in public administration. *Digital Society & Virtual Governance*, 1(1), 1–16.
<https://doi.org/10.6914/dsvg.010101>
- Lin, L., & Wang, Y. (2025). SHAP Stability in Credit Risk Management: A Case Study in Credit Card Default Model. *Risks*, 13(12), 238. <https://doi.org/10.3390/risks13120238>
- Linardatos, P., Papastefanopoulos, V., & Kotsiantis, S. (2020). Explainable AI: A Review of Machine Learning Interpretability Methods. *Entropy (Basel, Switzerland)*, 23(1), 18.
<https://doi.org/10.3390/e23010018>
- Longo, L., Brcic, M., Cabitza, F., Choi, J., Confalonieri, R., Del Ser, J., Guidotti, R., Hayashi, Y., Herrera, F., Holzinger, A., Jiang, R., Khosravi, H., Lecue, F., Malgieri, G., Páez, A., Samek, W., Schneider, J., Speith, T., & Stumpf, S. (2024). Explainable Artificial Intelligence (XAI) 2.0: A manifesto of open challenges and interdisciplinary research directions. *Information Fusion*, 106, 102301. <https://doi.org/10.1016/j.inffus.2024.102301>
- Lundberg, S., & Lee, S. (2017, May 22). *A unified approach to interpreting model predictions*. arXiv.org. <https://arxiv.org/abs/1705.07874>
- Maleki Varnosfaderani, S., & Forouzanfar, M. (2024). The Role of AI in Hospitals and Clinics: Transforming Healthcare in the 21st Century. *Bioengineering*, 11(4), 337.
<https://doi.org/10.3390/bioengineering11040337>
- Mathew, D. E., Ebem, D. U., Ikegwu, A. C., Ukeoma, P. E., & Dibiaezue, N. F. (2025). Recent emerging techniques in explainable artificial intelligence to enhance the interpretable and understanding of AI models for human. *Neural Processing Letters*, 57(1).
<https://doi.org/10.1007/s11063-025-11732-2>

- Mennella, C., Maniscalco, U., De Pietro, G., & Esposito, M. (2024). Ethical and regulatory challenges of AI technologies in healthcare: A narrative review. *Heliyon*, 10(4), e26297. <https://doi.org/10.1016/j.heliyon.2024.e26297>
- Microsoft (2025). Assess AI Systems and Make Data-Driven Decisions with Azure Machine Learning Responsible AI Dashboard - Azure Machine Learning. Microsoft Learn. <https://learn.microsoft.com/en-us/azure/machine-learning/concept-responsible-ai-dashboard?view=azureml-api-2>
- Molnar, C. (2025). *Interpretable machine learning: A guide for making black box models explainable* (3rd ed.). Lulu.com. . <https://christophm.github.io/interpretable-ml-book>
- Naoum, R. F., Szakadáti, T., & Balogh, G. (2026). Artificial Intelligence (AI) in human resource management (HRM): a systematic review of its dual impact on diversity, equity, and inclusion (DEI). *Management Review Quarterly*. <https://doi.org/10.1007/s11301-025-00580-y>
- National Institute of Standards and Technology (NIST). (2023). AI Risk Management Framework (AI RMF 1.0). NIST Risk 100-1.
- O'Brien, T., Green, J., Saunders, G., Yanagihara, H., & Leon, D. (2025). *Algorithmic accountability frameworks: Ensuring transparency and fairness in AI-driven public sector decision-making*.
- OECD. (2024). Recommendation of the Council on Artificial Intelligence. OECD/LEGAL/0449.
- Olateju, O., Okon, S. U., Olaniyi, O. O., Samuel-Okon, A. D., & Asonze, C. U. (2024). Exploring the concept of explainable AI and developing information governance standards for enhancing trust and transparency in handling customer data. *SSRN Electronic Journal*. Available at SSRN: <https://ssrn.com/abstract=5047713>
- Pahune, S., Akhtar, Z., Mandapati, V., & Siddique, K. (2025). The importance of AI data governance in large language models. *Big Data and Cognitive Computing*, 9(6), 147.

<https://doi.org/10.3390/bdcc9060147>

Pandey, P. (2025, January 18). *Using What-If Tool to investigate Machine Learning models*. Towards Data Science. <https://towardsdatascience.com/using-what-if-tool-to-investigate-machine-learning-models-913c7d4118f/>

Potluri, R. (2025). [Recruitment governance and bias detection]. *AI & Society*.

Ribeiro, M.T., Singh, S. and Guestrin, C. (2016) “Why Should I Trust You?”: Explaining the Predictions of Any Classifier. Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, San Francisco, 13-17 August 2016, 1135-1144. <https://doi.org/10.1145/2939672.2939778>

Saarela, M., & Podgorelec, V. (2024). Recent Applications of Explainable AI (XAI): A Systematic Literature Review. *Applied Sciences*, 14(19), 8884. <https://doi.org/10.3390/app14198884>

Samek, W. (2022). Explainable deep learning: Concepts, methods, and new developments. In W. Samek, G. Montavon, A. Vedaldi, L. K. Hansen, & K.-R. Müller (Eds.), *Explainable deep learning AI* (pp. 7–33). Amsterdam: Elsevier. <https://doi.org/10.1016/B978-0-323-91181-9.00002-2>

Samek, W., Montavon, G., Vedaldi, A., Hansen, L. K., & Müller, K.-R. (2019). *Explainable AI: Interpreting, explaining and visualizing deep learning*. Cham: Springer.

Shri, G. (2025). Adaptive data governance models using explainable AI. *International Journal of Emerging Trends in Computer Science and Information Technology*

Smela, B., Toumi, M., Świerk, K., Francois, C., Biernikiewicz, M., Clay, E., & Boyer, L. (2023). Rapid literature review: definition and methodology. *Journal of market access & health policy*, 11(1), 2241234. <https://doi.org/10.1080/20016689.2023.2241234>

Soleimani, M., Intezari, A., Arrowsmith, J., Pauleen, D. J., & Taskin, N. (2025). Reducing AI bias in recruitment and selection: an integrative grounded approach. *The International Journal of Human Resource Management*, 36(14), 2480–2515.

<https://doi.org/10.1080/09585192.2025.2480617>

Teixeira, B., Carvalhais, L., Pinto, T., & Vale, Z. (2025). Explainable AI framework for reliable and transparent automated energy management in buildings. *Energy and Buildings*, 347, 116246. <https://doi.org/10.1016/j.enbuild.2025.116246>

Thirunagalingam, A. (2022). Enhancing data governance through explainable AI: Bridging transparency and automation. *International Journal of Sustainable Development Through AI, ML and IoT*, 1(2). <https://ijsdai.com/index.php/IJSDAI/index>

Utomo, S., Pratap, A., Karthikeyan, P., Ayeelyan, J., Hsu, H. C., & Hsiung, P. A. (2026). When explainable artificial intelligence meets data governance: Enhancing trustworthiness in multimodal gas classification. *Information Fusion*, 125, 103440. <https://doi.org/10.1016/j.inffus.2025.103440>

Varshney, K. (2024, November 19). *Introducing AI Fairness 360*. IBM Research. <https://research.ibm.com/blog/ai-fairness-360>

Voigt, P., & Von dem Bussche, A. (2017). *The EU General Data Protection Regulation (GDPR): A practical guide*. Cham: Springer International Publishing. <https://doi.org/10.1007/978-3-319-57959-7>

Wexler, J., Pushkarna, M., Bolukbasi, T., Wattenberg, M., Viégas, F., & Wilson, J. (2019). The What-If Tool: Interactive probing of machine learning models. *arXiv preprint arXiv:1907.04135*. <https://arxiv.org/abs/1907.04135>

Zhang, L., & Wang, W. (2025). Data governance and digital trust in smart markets. *Journal of Electronic Commerce Management*, 1(1), 85–104. https://www.joecm.ir/article_222689_729c1681423218ae78413c432cdc76e4.pdf