

Inside the Black Box:

Explainable AI

Authors:

Finn Murray

Samuel Adebayo, PhD

ISx4 • AI research - 2025



Executive Summary

Artificial Intelligence (AI) systems now influence decisions across finance, healthcare, manufacturing and logistics. However, as these systems grow in complexity, they also become harder to understand and trust. Explainable AI (XAI) addresses this challenge by making AI's inner workings transparent and understandable.

This white paper reframes explainability as a business function – integral to compliance, user trust, and operational efficiency – rather than a mere research afterthought. We introduce a practical five-stage Explainability Framework that can be embedded throughout the AI lifecycle, from data preparation to model deployment. We also explore use cases in finance and legal operations where explainability delivers tangible benefits: satisfying regulatory requirements (e.g. EU AI Act, GDPR), improving human-AI team performance, and enabling confident adoption of AI solutions. By operationalizing XAI, enterprises can strengthen oversight, build user trust, and turn responsible AI practices into a competitive advantage.

Table of Contents

The Need for Explainability	4
Five Stage Explainability Framework	5
Applications across Industries.....	6
Proven Business Impact	7

The Need for Explainability

As organizations adopt AI for high-stakes tasks, they face a widening “trust gap” between AI’s decisions and the people affected by them. Modern AI often involves layers of third-party models (e.g. pre-trained foundation models fine-tuned by others), making it hard for any single team to fully understand why a given prediction was made. This opacity amplifies risk: business leaders, regulators, and customers may lose confidence in AI outcomes that cannot be justified in human terms¹. Without robust explainability, organisations face:

- **Compliance Risk:** Emerging regulations demand transparency. For example, the EU Artificial Intelligence Act (2024) requires that high-risk AI systems be “sufficiently transparent to enable deployers to interpret the system’s output and use it appropriately”². Providers must supply clear documentation on model functionality to satisfy regulators^{3/4}. Likewise, under GDPR, individuals have the right to receive “meaningful information about the logic involved” in automated decisions⁵. Without XAI, organizations struggle to meet these legal obligations, exposing themselves to fines and legal challenges.
- **Operational Blind Spots:** Opaque AI can fail silently or behave unexpectedly. If a model’s decisions can’t be audited, engineers and risk officers cannot diagnose errors, bias, or drift in model behavior. This makes it “challenging to explain decisions” and improve models over time⁶. In regulated industries (finance, healthcare), the inability to trace how an outcome was determined also means AI systems cannot be certified for use.
- **Erosion of Trust:** Users and decision-makers hesitate to adopt AI they don’t trust. A MIT Sloan study notes that stakeholders must believe an AI system is accurate and fair before they will rely on it, and explainability is key to building that trust¹. If an AI loan system or legal prediction tool cannot justify its recommendation, employees may ignore it and customers may feel alienated or treated unfairly. Over time, such AI failures can damage an enterprise’s reputation.

THE NEED FOR EXPLAINABILITY



COMPLIANCE RISK

Relates to regulatory fines and legal challenge (e.g. EU AI Act, GDPR)



OPERATIONAL BLIND SPOTS

Inability to diagnose errors, identify bias, or detect model drift



EROSION OF TRUST

Hesitation from users, stakeholders, and the public to adopt or rely on the system

¹ <https://mitsloan.mit.edu/ideas-made-to-matter/why-companies-need-artificial-intelligence-explainability#:~:text=Why%20it%20Matters>

² <https://artificialintelligenceact.eu/article/13/#:~:text=1.%20High,set%20out%20in%20Section%203>

³ <https://www.euaiact.com/key-issue/5#:~:text=1,features%20of%20HRAIS%20to%20deployers>

⁴ <https://www.euaiact.com/key-issue/5#:~:text=transparent%20in%20their%20functioning%20so,SA1>

⁵ <https://www.fieldfisher.com/en/insights/what-constitutes-meaningful-information-about-automated-decision-making#:~:text=Where%20automated%20decision,informatio%20to%20the%20data%20subject>

⁶ <https://researchprofiles.tudublin.ie/en/publications/a-comparative-analysis-ofshap-lime-anchors-anddice-forinterpretin/#:~:text=Financial%20institutions%20heavily%20rely%20on,techniques%20in%20fraud%20prevention%2C%20but>

Five Stage Explainability Framework

Businesses implementing AI solutions should apply this Five-Stage XAI Framework to operationalise transparency across the AI lifecycle ⁷. Each stage connects technical explainability with business accountability:



1. **Data & Model Documentation:** Transparency begins with careful documentation of the training data and models. The enterprise should record data provenance (sources, timeframes, demographics) and note any biases or limitations before deployment. Models are catalogued with details about their architecture, objectives, and performance. For example, a bank developing a credit risk model logs which features it uses (income, payment history, etc.) and how it was validated. These records ensure traceability during later audits and help identify bias early.
2. **Explanation Generation:** Once a model is in place, XAI techniques generate interpretability artifacts – i.e. raw explanations of the model’s behaviour. Common methods include feature attribution (e.g. SHAP, LIME) to quantify which inputs most influenced a given prediction, and visualizations like Grad-CAM for highlighting important regions in an image. Rule-based approaches or counterfactual examples can also reveal “what-if” scenarios (e.g. “If income had been higher, the loan would be approved”). These technical outputs form the basis for human-understandable explanations.
3. **Interpretation:** Technical explanations must be translated into forms meaningful to the end-user or decision-maker. This stage focuses on context and communication. For instance, an AI system might output a feature weight vector, but a loan officer sees a simple narrative: “Loan denied due to high credit utilization and recent late payments.” Tailoring the explanation to the audience (executive, engineer, customer, etc.) ensures it is actionable, not just descriptive. Effective interpretation often requires collaboration between data scientists and business domain experts to frame insights in the right language.
4. **Evaluation:** Organisations should evaluate the quality of explanations just as rigorously as model accuracy. Key criteria include fidelity (does the explanation correctly reflect the model’s actual reasoning?), comprehensibility (will the target user understand it?), and usefulness (does it improve decision quality or user trust?). This can involve user testing and metrics. For example, does providing an explanation reduce an auditor’s time to validate a transaction anomaly? High fidelity may be measured by how closely an explanation method’s output predicts the model’s behaviour. By measuring these factors, companies ensure the XAI system truly adds value.

⁷ https://www.researchgate.net/publication/389783123_Explainable_AI_XAI_for_trustworthy_and_transparent_decision-making_A_theoretical_framework_for_AI_interpretability

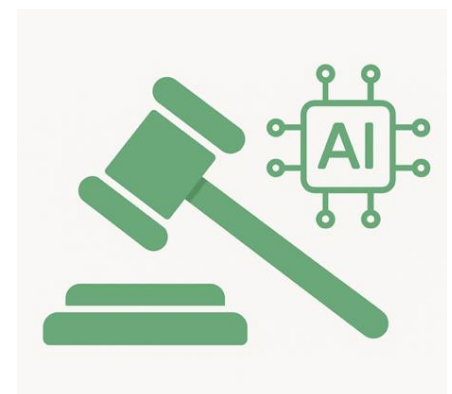
5. **Deployment & Feedback:** Finally, explainability is integrated into production workflows and continuously refined. Explanations should appear wherever AI outputs are consumed – e.g. in a compliance dashboard, a customer-facing app, or an internal report – so that users can see why a decision was made. Importantly, there should be a feedback loop: if users flag an AI decision as questionable, that feedback triggers model review or retraining, and the documentation is updated. This stage treats explainability as a living process, evolving with the model and data. Over time, such feedback helps maintain transparency as the AI system and business context change.

Together, these stages transform explainability from a one-off technical effort into a lifecycle. The AI model's entire journey – from data acquisition to daily operation – is accompanied by appropriate explainability measures, ensuring that the system remains transparent, auditable, and aligned with business goals.

Applications across Industries

Financial Services: Explainable AI enables transparent, data-driven decision making in banking, insurance, and investment sectors. In credit scoring, XAI surfaces the factors that most influence loan approvals, ensuring fairness and auditability. Fraud detection models leverage explainability to show why specific transactions were flagged, improving case resolution rates and regulator confidence⁸.

Legal & Compliance: Legal teams are beginning to deploy AI for tasks like contract review, case law analysis, and compliance auditing – and explainability is critical in these contexts. AI assistants in a law firm might scan hundreds of contracts to flag high-risk clauses. Explainable AI would highlight which clause and which criteria caused each flag (e.g., “Indemnification clause deviates from standard; high liability exposure”). This allows attorneys to quickly validate the alert rather than blindly trust the AI. In case law research or litigation strategy, an AI model might predict the likely outcome of a case based on past precedents. With XAI, it can show the top precedents or evidence that influenced its prediction. For example, a legal team using an AI case evaluator could see: “Predicted 70% chance of plaintiff win – key factors: similar past case Smith v. ACME (2019) and presence of clause X in the contract.” This gives lawyers a starting point to dig deeper and also a level of assurance about the AI's reasoning.



Other High-Risk Operations: Many enterprise applications outside of finance and law also benefit from explainability. In human resources and operations, companies use AI for screening resumes, forecasting supply chain disruptions, or monitoring quality control – areas where errors or bias can have serious consequences. XAI helps ensure these AI-driven decisions align with ethics and policy. For instance, a recruitment AI that explains its rankings of candidates (highlighting experience and skills rather than demographic traits) can prove that it's free of bias, supporting Diversity & Inclusion goals. In manufacturing and engineering, AI systems diagnose equipment failures or product defects; explainability pinpoints the

⁸ <https://ieeexplore.ieee.org/abstract/document/10509682>

sensor readings or image features that led to a failure prediction, so engineers can verify and act on it. Lastly, in cutting-edge autonomous and generative AI systems (such as algorithmic trading bots or AI content generators), explainability provides much-needed control. An algorithmic trading system operating at high frequency can be required to produce an audit trail of which signals or market data drove its buy/sell decisions – a practice increasingly demanded by institutional investors and regulators⁹. Similarly, an enterprise using a generative AI (e.g. for drafting reports or assisting customer service) will embed explainability to trace which sources the AI relied on for its output, guarding against errors or inappropriate content. Across these diverse use cases, the pattern is the same: transparent AI is more robust and trustworthy. It enables domain experts to validate AI decisions, catch problems early, and confidently harness AI in sensitive workflows that would otherwise be off-limits due to risk.

Proven Business Impact

Explainable AI is yielding concrete business benefits wherever it's implemented. By bridging the gap between complex models and human understanding, XAI directly contributes to key enterprise objectives:

- Regulatory Compliance & Risk Mitigation:** Explainability frameworks help organizations prove they are managing AI responsibly under new regulations. For example, Article 13 of the EU AI Act mandates providing deployers with information to interpret high-risk AI systems' outputs^{4/10}. Similarly, GDPR's transparency rules (Articles 13–15) effectively create a “right to an explanation” for individuals impacted by automated decisions⁵. Implementing XAI enables companies to meet these requirements – offering clear documentation and user-facing explanations – thereby avoiding legal penalties and building trust with regulators. In practice, firms with strong XAI can more easily pass audits and demonstrate algorithmic accountability as part of their governance (a major factor as regulatory scrutiny of AI intensifies in finance, healthcare, and beyond).
- Improved Human–AI Performance:** Studies confirm that providing explanations with AI recommendations leads to better human decision-making. In one controlled experiment, domain experts performing real-world tasks (manufacturing inspection and medical diagnosis) were **7.7%** and **4.7%** more accurate when assisted by an explainable AI (which showed visual heatmap explanations) compared to those using a black-box AI¹¹. Explanations helped the experts understand when the AI was correct or mistaken, so they could act accordingly. Likewise, a 2025 study on anomaly detection found that “participants performance and trust were higher when the AI agent provided explanations for its recommendations”¹². In other words, people trust and follow AI advice more when they know why – resulting in faster decisions, fewer errors, and increased willingness to use AI tools. These gains in productivity and accuracy translate to significant business value in settings like fraud investigation, medical triage, and customer service (where AI aids human agents).

⁹ <https://www.n-ix.com/what-is-explainable-ai-xai/#:~:text=internal%20stakeholders%20and%20external%20auditors>

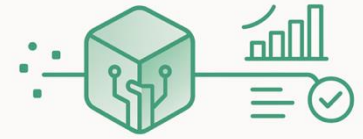
¹⁰ <https://artificialintelligenceact.eu/article/13/#:~:text=,system%20and%20use%20it%20appropriately>

¹¹ <https://www.nature.com/articles/s41598-024-82501-9#:~:text=lesions,significant%20improvement%20in%20task%20performance>

¹²

https://www.researchgate.net/publication/391747863_Effects_of_Explanations_and_Accuracy_on_Human_Performance_and_Trust_in_AI-Assisted_Anomaly_Diagnosis_Tasks#:~:text=participant%20completed%20two%20sessions%20using,not%20on%20trust%2C%20or%20SA

- Operational Transparency & Efficiency:** XAI makes AI systems more manageable and efficient by illuminating their inner logic. For data science and engineering teams, explainability offers a powerful debugging aid – it's easier to pinpoint why a model is making odd predictions if you can see which factors it's relying on. In a comparative study of fraud detection models, methods like SHAP and ANCHORS produced consistent explanations (high stability and similarity metrics) across cases⁶. This consistency means analysts can quickly trust the pattern the model is signalling (e.g. a particular transaction pattern) and focus their investigation, rather than wrestling with inconsistent or opaque outputs. Overall, providing an audit trail for each AI decision reduces the time spent interpreting results or chasing model errors. It also facilitates knowledge transfer – new team members or auditors can understand a model's behaviour without needing to dive into code, simply by reviewing its documented explanations. Companies report that integrating XAI into monitoring tools accelerates issue diagnosis (for instance, catching a data drift problem because the explanation for predictions started to change). Thus, explainability contributes to more efficient operations and continuous improvement of AI systems.



- User Trust and Adoption:** From a change-management perspective, explainability drives broader adoption of AI solutions within an organization. Employees and customers are more inclined to accept an AI-generated insight when it comes with a clear rationale. This boosts user satisfaction and confidence in AI-enabled services. For example, a fintech company deploying an AI investment advisor found that providing users with plain-language reasons for portfolio recommendations significantly increased uptake of the advice feature. Trust generated through XAI also has a feedback effect: as users trust the system more, they engage with it

more deeply, providing more feedback and data, which can further improve the AI. In essence, explainability humanizes the AI interaction – turning scepticism into engagement. Businesses that prioritize XAI thus see higher ROI on their AI investments, as solutions are actually utilized to their full potential rather than abandoned due to mistrust.

Conclusion

Explainable AI is no longer a research luxury—it has become the foundation of trustworthy AI deployment. Enterprises that operationalise explainability (for example, via the lifecycle framework outlined) are positioning themselves to succeed in an era of AI regulations and heightened accountability. By systematically building transparency and user-centric design into AI systems, organizations can strengthen compliance, reduce risks, and foster an environment of trust around AI. The payoff is not only avoiding pitfalls, but positively turning responsible AI into a sustainable competitive advantage. Companies that can confidently explain their AI's actions will lead in innovation—because they can safely deploy AI in domains where others can't, winning trust from customers and regulators alike. In summary, investing in XAI capabilities now is an investment in the long-term viability and value of enterprise AI initiatives.