

# 15th Webinar on Data Science and Analytics

Time: 3:00 PM to 5:00 PM

Date: 23-Mar-24

## Enhancing Model Transparency: The Role of Explainable AI in Actuarial Risk Assessment

**Vamsidhar Ambatipudi**, B.Tech, MBA(IIM), FIAI, CERA, FRM, PRM, CSQA



# Agenda

- Introduction to Actuarial Risk Assessment
- Fundamentals of AI and Machine Learning
- Why Explainability Matters in AI
- Introduction to Explainable AI (XAI)
- XAI Techniques and Tools
- Applying XAI in Actuarial Models
- Challenges and Limitations of XAI



# Introduction



- Actuarial risk assessment
  - Analyzing and quantifying the risks associated with future contingent events, primarily in the insurance and finance sectors.
  - Goal → ensure that financial entities like insurance companies can meet their future obligations to policyholders by setting appropriate premiums and reserves
- Key Components
  - Probability Estimation: Calculating the likelihood of various events, such as accidents, natural disasters, or health issues.
  - Financial Modeling: Assessing the financial impact of these events on an organization.
  - Premium Calculation: Determining the price of insurance policies based on risk assessment to ensure profitability and solvency.

# Introduction



- Traditional methods used in actuarial risk assessment
  - Actuarial tables: Historical data on mortality, morbidity, and other relevant factors used to calculate risk.
  - Life tables: Summarize mortality rates and life expectancies for different demographic groups.
  - Premium calculation formulas: Determine insurance premiums based on risk factors and expected claims.
  - **Limitations:** While effective for standard risks with ample historical data, these methods struggle with complex, nonlinear risks and patterns.
- **Introduction of AI/ML Models**
  - **Rise of Computational Power:** Advancements in computing technology enabled the use of more complex mathematical models to predict risk.
  - **AI/ML Adoption:** Artificial Intelligence (AI) and Machine Learning (ML) models offer sophisticated techniques that can analyze vast datasets, recognize patterns, and make predictions or decisions with minimal human intervention.

# Introduction



- **Benefits of AI/ML in Actuarial Science**
  - **Enhanced Accuracy:** Process and learn from larger datasets than traditional methods, leading to more accurate risk assessments.
  - **Efficiency:** Automates the analysis process, significantly reducing the time required to assess risks.
  - **Flexibility:** Capable of identifying subtle patterns and nonlinear relationships in the data that traditional models might miss.
  - **Innovation in Products and Services:** Enables the development of new insurance products and personalized pricing models based on individual risk profiles.

# Introduction



- **Challenges and Considerations**
  - **Data Privacy and Security:** The use of extensive personal data raises concerns about privacy and the need for secure data handling practices.
  - **Regulatory Compliance:** Ensuring AI/ML models comply with existing and emerging regulations regarding transparency, fairness, and accountability.
  - **Model Explainability:** The "black box" nature of some AI/ML models can make it difficult to understand and explain how decisions are made, posing challenges for accountability and trust.
- **Current Trends**
  - Increasing adoption of AI/ML in actuarial practices, driven by the need for more precise risk assessment and the potential for creating more customized insurance products.
  - Growing focus on developing explainable AI (XAI) models to address concerns about transparency and interpretability.

# Introduction



- **Challenges and Considerations**
  - **Data Privacy and Security:** The use of extensive personal data raises concerns about privacy and the need for secure data handling practices.
  - **Regulatory Compliance:** Ensuring AI/ML models comply with existing and emerging regulations regarding transparency, fairness, and accountability.
  - **Model Explainability:** The "black box" nature of some AI/ML models can make it difficult to understand and explain how decisions are made, posing challenges for accountability and trust.
- **Current Trends**
  - Increasing adoption of AI/ML in actuarial practices, driven by the need for more precise risk assessment and the potential for creating more customized insurance products.
  - Growing focus on developing explainable AI (XAI) models to address concerns about transparency and interpretability.

# Fundamentals of AI and Machine Learning



- **Artificial Intelligence (AI)**
  - AI refers to the simulation of human intelligence in machines that are programmed to think and learn like humans.
  - **Goal** is to enable machines to perform tasks that typically require human intelligence, such as visual perception, speech recognition, decision-making, and language translation.
- **Machine Learning (ML)**
  - A subset of AI, that involves the development of algorithms that allow computers to learn and improve from experience without being explicitly programmed for each task.
  - ML models learn patterns from data and make predictions or decisions based on new, unseen data.

# Fundamentals of AI and Machine Learning



- **Common Machine Learning Algorithms**
  - **Linear Regression:** Used for predicting a continuous value (e.g., house prices).
  - **Logistic Regression:** Used for binary classification tasks (e.g., spam or not spam).
  - **Decision Trees:** A model that uses a tree-like graph of decisions and their possible consequences.
  - **Random Forests:** An ensemble method using multiple decision trees to improve prediction accuracy.
  - **Neural Networks:** Inspired by the human brain, these algorithms are capable of capturing complex patterns through layers of processing units.
  - **Clustering Algorithms (e.g., K-means):** Used in unsupervised learning to group data into clusters based on similarity.

# Poll Question



- "Which AI/ML model do you believe has the most potential to revolutionize actuarial risk assessment?"
- A) Linear Regression
- B) Decision Trees and Random Forests
- C) Neural Networks
- D) Clustering Algorithms (e.g., K-means)

# Fundamentals of AI and Machine Learning



- **"Black Box" Models**
  - A model whose internal workings are not visible or understandable to users, even though it can produce accurate predictions or decisions.
  - These models, especially complex neural networks and deep learning models, are highly opaque, making it difficult to trace how inputs are transformed into outputs.
- **Challenges Posed by "Black Box" Models**
  - Lack of Transparency
  - Trust Issues
  - Regulatory and Ethical Concerns
  - Difficulty in Error Diagnosis

# Introduction to Explainable AI (XAI)



- AI & ML techniques that make the actions, decisions, or predictions of AI models understandable to human users
- **Objectives of XAI**
  - **Transparency:** Make the decision-making process of AI models as transparent as possible to users.
  - **Trust:** Build trust among users and stakeholders by making AI systems' operations understandable and predictable.
  - **Compliance:** Ensure AI systems meet growing regulatory and legal requirements for accountability and transparency.
  - **Debugging and Improvement:** Facilitate the debugging and continuous improvement of AI models by understanding their decision-making processes.

# Poll Question



- "Which feature of Explainable AI (XAI) do you value the most for actuarial applications?"
- A) Transparency of decision-making process
- B) Ability to audit and validate model predictions
- C) Potential to uncover bias and ensure fairness
- D) Insight into model limitations and areas for improvement

# Interpretability vs. Explainability



- **Interpretability**

- Degree to which a human can understand the cause of a decision made by an AI model
- Often relates to the model's internal mechanisms being comprehensible to some extent.
- Crucial in models where understanding the decision process is as important as the decision itself.

- **Explainability**

- Additionally, Conveys the decision-making process in human terms, often through post-hoc analysis.
- Provides clear, understandable explanations to users about how AI models arrive at their decisions or predictions.
- Often achieved through additional tools or techniques that elucidate the model's output, even if the model itself is a "black box."

# Interpretability vs. Explainability



- **Key Differences**

- **Scope:** Interpretability is about the inherent transparency of the model, while explainability can include external methods to make opaque models more understandable.
- **Focus:** Interpretability focuses on the model's internal workings, whereas explainability is concerned with making the model's outputs understandable to users, regardless of the model's complexity.
- **Application:** Interpretability is often a feature of simpler models (e.g., linear regression, decision trees), whereas explainability techniques are used for complex models (e.g., deep neural networks) that are not inherently interpretable.

# XAI Techniques and Tools



- **Model-Agnostic Techniques**

- Techniques that can be applied to any machine learning model to provide explanations, regardless of the model's internal workings.
- Versatility in use across different types of models and flexibility in application.
- **Examples:**
  - **LIME (Local Interpretable Model-agnostic Explanations):** Generates explanations for individual predictions by approximating the local decision boundary around the prediction.
  - **SHAP (SHapley Additive exPlanations):** Utilizes game theory to assign an importance value to each feature for a particular prediction, indicating how much each feature contributes to the output.

# XAI Techniques and Tools



- **Model-Specific Techniques**

- Techniques designed to work with specific types of models, leveraging knowledge of the model's internal architecture for explanations.
- Can provide deeper insights into the model's decision-making process due to access to its internal structure.
- **Examples:**
  - **Decision Trees:** Inherently interpretable, as they provide a clear path of decisions from features to the final prediction.
  - **Neural Network Attention Mechanisms:** Highlight parts of the input that are important for the model's prediction, specific to neural networks.

# LIME (Local Interpretable Model-agnostic Explanations)



- Technique that explains the predictions of any ML classifier in an interpretable manner, by approximating the model locally with an interpretable model
- **How LIME Works**
  - **Local Perturbation:** LIME generates a new dataset consisting of perturbed samples around the observation of interest.
  - **Prediction:** It then uses the original complex model to predict outcomes for these new samples.
  - **Weight Assignment:** Each perturbed sample is assigned a weight based on its proximity to the original observation—the closer the sample, the higher the weight.
  - **Local Model Training:** LIME trains a simple, interpretable model (e.g., linear regression or decision tree) on this new dataset, using the weights as part of the training process.
  - **Interpretation:** The coefficients of the simple model serve as explanations for the prediction at the original observation.

# LIME (Local Interpretable Model-agnostic Explanations)



- **Example of LIME**
  - Imagine a complex model that predicts whether a loan application will be approved. For a particular application that is predicted to be rejected, LIME can be used to understand why.
  - LIME might create numerous simulated loan applications by tweaking the original application's features (e.g., income, credit score).
  - After obtaining predictions for these perturbed applications from the complex model, LIME trains a simple model to approximate the complex model's behavior around the original application.
  - If the simple model indicates that the most significant factor for rejection is a low credit score, LIME provides this insight as an explanation.

# SHAP (SHapley Additive exPlanations)



- A game theory-based approach to explain the output of any machine learning model.
- It assigns each feature an importance value for a particular prediction.
- **How SHAP Works**
  - **Shapley Values Concept:** Shapley values aim to fairly distribute the "payout" (prediction) among the "players" (features).
  - **Feature Contribution:** SHAP computes the contribution of each feature to the difference between the actual prediction and the average prediction over the dataset.
  - **Combination of Contributions:** It considers all possible combinations of features to determine the marginal contribution of each feature.
  - **Aggregation:** The Shapley values are aggregated to provide a global understanding of feature importance across the model.

# SHAP (SHapley Additive exPlanations)



- **Example of SHAP**

- Consider a model predicting house prices with features like location, size, and age.
- For a specific house predicted to be worth \$500,000, which is higher than the average prediction of \$300,000, SHAP analyses how much each feature contributes to this \$200,000 difference.
- It might reveal that location contributes \$150,000 due to the house's placement in a high-demand area, size contributes \$40,000 because of its spaciousness, and age contributes \$10,000 due to it being relatively new.
- These contributions explain the prediction in a manner that is easy to understand and quantifies the impact of each feature.

# Case studies



- **Case Study 1: Life Insurance Underwriting**
  - **Application:** An AI model was developed to predict life expectancy based on lifestyle factors, medical history, and real-time health data.
  - **XAI Application:** LIME was used to explain individual risk assessments to applicants, detailing which factors most influenced their risk score.
  - **Benefits:** Improved customer trust and satisfaction by transparently showing how decisions were made, leading to higher acceptance rates of insurance offers.
- **Case Study 2: Property and Casualty Insurance Claim Predictions**
  - **Application:** A machine learning model predicted the likelihood and severity of claims for property and casualty insurance.
  - **XAI Application:** SHAP values were used to highlight the specific conditions or characteristics that led to the prediction of high-risk claims.
  - **Benefits:** Enabled insurers to tailor their communication and risk mitigation strategies for high-risk policyholders, improving risk management and reducing fraudulent claims.

# Poll Question



- "What do you think is the biggest challenge in adopting Explainable AI (XAI) within actuarial practices?"
- A) Complexity vs. Explainability Trade-off
- B) Data Privacy and Security Concerns
- C) Ensuring Regulatory Compliance
- D) The "Black Box" Nature of Advanced AI/ML Models

# Challenges and Limitations of XAI



- **Complexity vs. Explainability Trade-off**
  - High model complexity often leads to lower explainability.
  - Advanced models like deep neural networks offer superior performance but are harder to interpret.
  - Simpler models (e.g., linear regression, decision trees) are more interpretable but might not capture complex patterns as effectively as more complex models.
- **Balancing Act**
  - Achieving the right balance between model complexity and explainability is crucial, especially in fields requiring both high accuracy and transparency, such as healthcare and finance.
  - Decision-makers often face the challenge of choosing between a slightly less accurate model that is fully interpretable and a more accurate but opaque model.

# Future Directions of XAI



- **Hybrid Models**
  - Combining interpretable models with complex algorithms to leverage the strengths of both: accuracy from complex models and transparency from simpler ones.
- **Automated Explainability**
  - Development of tools and frameworks that automatically generate explanations alongside predictions, making XAI more accessible to non-experts.
- **Interactive Explanations**
  - Advancements in creating dynamic, user-interactive explanation interfaces that allow users to query and understand AI decisions in real-time.
- **Personalized Explanations**
  - Tailoring explanations to the user's level of expertise or role, ensuring that they are meaningful and useful across different stakeholders.