

Xplainable Artificial Intelligence for Seismic Reservoir Characterization Using SHAP-Based Feature Attribution

Pronab Chowdhury¹

¹Department of Computer Science and Engineering, Adamas University, India

Publication Date: 2026/08/04

Abstract: Artificial intelligence (AI) has significantly improved seismic reservoir characterization by enabling accurate prediction of petrophysical properties from complex seismic attributes. However, many high-performing machine learning algorithms operate as "black-box" models, providing limited insight into how geological features influence prediction outcomes. This lack of interpretability reduces user confidence and restricts the practical adoption of AI-assisted seismic interpretation in hydrocarbon exploration. To address this limitation, this study proposes an Explainable Artificial Intelligence (XAI) framework for quantitative seismic reservoir characterization using SHapley Additive exPlanations (SHAP). The proposed workflow integrates seismic attribute engineering, supervised machine learning, and feature attribution analysis to identify the relative contribution of individual seismic attributes to porosity prediction and lithofacies classification. Multiple machine learning models are trained using calibrated seismic and well-log data, while SHAP is employed to generate both global and local explanations of model predictions. Global feature importance analysis identifies the most influential seismic attributes controlling reservoir quality, whereas local explanations provide sample-specific interpretations that improve geological understanding and decision transparency. Illustrative results indicate that SHAP successfully distinguishes dominant seismic attributes associated with high-quality reservoir zones and reveals nonlinear interactions among amplitude, frequency, and geometric attributes. The proposed framework enhances the interpretability, transparency, and reliability of AI-assisted seismic interpretation while supporting confidence-aware reservoir characterization. This explainable workflow offers an effective decision-support tool for hydrocarbon exploration and provides a practical foundation for integrating trustworthy artificial intelligence into future digital reservoir management systems.

Keywords: *Explainable Artificial Intelligence, SHAP, Machine Learning, Seismic Interpretation, Reservoir Characterization, Seismic Attributes, Feature Attribution, Porosity Prediction, Lithofacies Classification.*

How to Cite: Pronab Chowdhury (2026) Xplainable Artificial Intelligence for Seismic Reservoir Characterization Using SHAP-Based Feature Attribution. *International Journal of Innovative Science and Research Technology*, 11(7), 2605-2614. <https://doi.org/10.38124/ijisrt/26jul1414>

I. INTRODUCTION

Accurate reservoir characterization is fundamental to hydrocarbon exploration because it provides quantitative information regarding subsurface rock properties that directly influence hydrocarbon storage, production performance, and field development planning. Reservoir properties such as porosity, lithofacies, permeability, and fluid distribution are traditionally estimated using integrated seismic interpretation, well-log analysis, and geological modeling. Although well logs provide highly accurate measurements, they are limited to discrete borehole locations and therefore cannot adequately represent the spatial variability of reservoir properties throughout an entire reservoir. Three-dimensional seismic surveys overcome this limitation by providing continuous subsurface coverage; however, seismic amplitudes represent elastic responses rather than reservoir properties directly, making quantitative interpretation a

challenging inverse problem. Recent advances in seismic attribute analysis have significantly improved the extraction of geological information from seismic data by utilizing amplitude, frequency, phase, and structural attributes that correlate with lithological and petrophysical variations [1–6].

Artificial intelligence and machine learning have transformed modern seismic interpretation by enabling nonlinear relationships between seismic attributes and reservoir properties to be learned directly from data. Ensemble learning, support vector machines, gradient boosting algorithms, and artificial neural networks have demonstrated superior performance for predicting porosity, lithofacies, permeability, and reservoir quality compared with conventional statistical approaches. Recent studies have also integrated artificial intelligence with digital twin technologies to facilitate intelligent reservoir monitoring, predictive analytics, and automated interpretation workflows [7–13].

These developments demonstrate the growing importance of data-driven approaches for reservoir characterization and highlight the potential of AI to improve exploration efficiency and reduce interpretation uncertainty.

Despite their impressive predictive performance, most machine learning algorithms used in reservoir characterization remain black-box models, meaning that they provide limited explanation of how individual seismic attributes influence prediction outcomes. In practical petroleum engineering applications, reservoir engineers and geoscientists require not only accurate predictions but also understandable reasoning behind those predictions. The inability to explain model decisions reduces user confidence, complicates model validation, and limits the adoption of artificial intelligence for high-value exploration and production decisions. Consequently, interpretability has emerged as one of the most important research challenges in trustworthy artificial intelligence.

Explainable Artificial Intelligence (XAI) has recently gained considerable attention as a solution to this challenge. XAI aims to make machine learning models transparent by identifying the contribution of each input feature to the final prediction. Among existing XAI techniques, SHapley Additive exPlanations (SHAP) has become one of the most widely adopted methods because it provides mathematically consistent feature attribution based on cooperative game theory. SHAP quantifies both global feature importance across an entire dataset and local explanations for individual predictions, enabling users to understand why a model predicts specific reservoir properties. Unlike conventional feature importance measures, SHAP captures complex nonlinear interactions among variables while maintaining theoretical consistency, making it particularly suitable for geoscientific applications involving heterogeneous reservoirs.

Although explainable artificial intelligence has achieved widespread adoption in medical diagnosis, finance, and industrial process monitoring, its application to seismic reservoir characterization remains relatively limited. Existing studies primarily emphasize prediction accuracy while providing little information regarding the geological significance of influential seismic attributes or the reasoning behind machine learning decisions. Furthermore, few investigations systematically integrate SHAP-based feature attribution with seismic attribute analysis to improve the interpretability of porosity prediction and lithofacies classification. As a result, reservoir engineers often lack sufficient insight into whether machine learning predictions are geologically meaningful or merely data-driven correlations.

To address these limitations, this study proposes an Explainable Artificial Intelligence framework for seismic reservoir characterization using SHAP-based feature attribution. The proposed workflow combines seismic attribute extraction, supervised machine learning, and explainable artificial intelligence into a unified framework capable of providing both accurate predictions and

interpretable geological insights. Rather than focusing solely on predictive performance, this study emphasizes understanding the contribution of individual seismic attributes to reservoir property estimation and evaluating whether machine learning decisions are consistent with geological knowledge.

➤ *The Major Contributions of this Study are Summarized as Follows:*

- A comprehensive Explainable Artificial Intelligence framework is developed for seismic reservoir characterization using SHAP-based feature attribution.
- Global and local SHAP analyses are employed to identify the most influential seismic attributes affecting porosity prediction and lithofacies classification.
- The proposed workflow improves model transparency by providing interpretable explanations for individual reservoir predictions.
- Feature attribution results are interpreted from a geological perspective to enhance confidence in AI-assisted seismic interpretation.
- The proposed framework establishes a transferable methodology for trustworthy artificial intelligence in quantitative reservoir characterization and intelligent hydrocarbon exploration.

The remainder of this paper is organized as follows. Section 2 reviews recent developments in explainable artificial intelligence and machine learning for seismic reservoir characterization. Section 3 presents the proposed SHAP-based explainable framework and methodological workflow. Section 4 discusses illustrative experimental results and feature attribution analysis, while Section 5 concludes the study and outlines future research directions.

II. RELATED WORK

Recent advances in artificial intelligence (AI) have transformed seismic interpretation from conventional manual analysis toward intelligent, data-driven reservoir characterization. Machine learning algorithms have demonstrated remarkable capability in predicting petrophysical properties from seismic attributes by learning complex nonlinear relationships that are difficult to capture using traditional statistical methods. More recently, Explainable Artificial Intelligence (XAI) has emerged as a complementary research direction that aims to improve the transparency and interpretability of machine learning predictions. This section reviews recent developments in machine learning for reservoir characterization, explainable AI techniques, SHAP-based feature attribution, and identifies the research gap addressed in this study.

➤ *Machine Learning in Seismic Reservoir Characterization*

Machine learning has become one of the most influential technologies in quantitative seismic interpretation because of its ability to automatically discover nonlinear relationships between seismic attributes and reservoir properties. Conventional approaches such as linear regression, geostatistical interpolation, and deterministic

inversion often struggle to model highly heterogeneous geological formations where multiple seismic attributes interact simultaneously. In contrast, supervised machine learning algorithms learn these relationships directly from historical well-log and seismic data, enabling improved prediction accuracy and automation.

Among the most widely adopted algorithms, Random Forest (RF), Support Vector Machines (SVM), Gradient Boosting Machines (GBM), eXtreme Gradient Boosting (XGBoost), and Artificial Neural Networks (ANN) have consistently demonstrated superior performance for porosity prediction, lithofacies classification, permeability estimation, and reservoir quality assessment. Random Forest is particularly attractive because of its robustness against noise, ability to handle high-dimensional data, and natural estimation of feature importance. Support Vector Machines effectively model nonlinear decision boundaries using kernel functions and perform well on relatively small geophysical datasets. Gradient Boosting algorithms sequentially minimize prediction residuals to improve learning performance, whereas Artificial Neural Networks approximate highly complex nonlinear functions through multiple hidden layers. Collectively, these methods have significantly improved quantitative reservoir characterization compared with conventional statistical techniques [14–23].

Recent investigations have further extended machine learning to three-dimensional seismic interpretation, seismic inversion, fault detection, stratigraphic analysis, and automated lithological mapping. These studies demonstrate that integrating multiple seismic attributes substantially improves predictive capability while reducing manual interpretation effort.

➤ Artificial Intelligence and Digital Reservoir Characterization

The integration of artificial intelligence with reservoir engineering has expanded beyond predictive modeling toward intelligent reservoir management and automated seismic interpretation. AI techniques are increasingly used for seismic interpretation, drilling optimization, production forecasting, reservoir monitoring, and geological risk assessment.

Elbarouni [7] reviewed recent developments in AI-driven seismic interpretation and demonstrated that machine learning significantly improves the automation of structural interpretation, seismic attribute analysis, and hydrocarbon exploration. The study highlighted how intelligent algorithms can reduce interpretation uncertainty while improving geological understanding.

Subsequently, Elbarouni [8] introduced an AI-enabled digital twin framework capable of integrating real-time seismic monitoring with predictive reservoir management. Digital twins provide continuously updated virtual representations of physical reservoirs by combining field measurements with machine learning models, enabling predictive maintenance, production optimization, and intelligent operational decision-making.

More recent investigations have demonstrated the effectiveness of advanced seismic interpretation techniques for structural mapping [9], seismic attribute-based fault detection [10], integrated seismic–petrophysical analysis [11], machine learning-based porosity prediction and lithofacies classification [12], and uncertainty-aware reservoir characterization through probabilistic seismic–petrophysical integration [13]. These studies collectively demonstrate that artificial intelligence is becoming an essential component of modern reservoir characterization workflows.

Despite these advances, most published studies primarily evaluate prediction accuracy while providing limited explanation of how machine learning models arrive at their decisions.

➤ Explainable Artificial Intelligence (XAI)

Although modern machine learning algorithms frequently achieve excellent predictive performance, many of these models function as black-box systems whose internal decision-making processes are difficult to interpret. This lack of transparency presents a significant challenge in scientific applications where model predictions influence costly engineering decisions.

Explainable Artificial Intelligence (XAI) has emerged as a rapidly growing research field that seeks to make machine learning models understandable to human users. Unlike conventional feature importance measures that provide only global rankings, XAI techniques generate interpretable explanations describing how individual input variables contribute to specific predictions.

• Several Explainability Techniques Have Been Proposed in Recent Years, Including:

- ✓ Local Interpretable Model-Agnostic Explanations (LIME)
- ✓ SHapley Additive exPlanations (SHAP)
- ✓ Integrated Gradients
- ✓ Partial Dependence Plots (PDP)
- ✓ Permutation Feature Importance
- ✓ Accumulated Local Effects (ALE)

Among these methods, SHAP has become one of the most widely adopted because it is supported by a rigorous mathematical foundation based on cooperative game theory. SHAP assigns each feature a contribution value representing its influence on the model prediction while satisfying desirable theoretical properties such as local accuracy, consistency, and missingness.

Unlike traditional feature importance metrics, SHAP can explain both individual predictions and overall model behavior, making it particularly suitable for complex nonlinear models commonly used in reservoir characterization.

➤ *SHAP-Based Feature Attribution for Geological Interpretation*

SHapley Additive exPlanations (SHAP) provide both global and local interpretations of machine learning predictions.

Global interpretation evaluates feature importance across the entire dataset, enabling researchers to identify which seismic attributes exert the greatest influence on reservoir characterization. Local interpretation explains why a model generated a specific prediction for an individual sample, allowing engineers to understand the geological factors responsible for high or low predicted porosity values.

• *For Seismic Reservoir Characterization, SHAP Offers Several Important Advantages:*

- ✓ Identification of dominant seismic attributes influencing reservoir quality.
- ✓ Quantification of positive and negative contributions of each attribute.
- ✓ Visualization of nonlinear feature interactions.
- ✓ Improved geological interpretation of machine learning predictions.
- ✓ Enhanced confidence in AI-assisted decision-making.

SHAP analysis may reveal that high Sweetness and RMS Amplitude values positively influence porosity prediction, whereas high Variance or Chaos attributes reduce predicted reservoir quality because they often correspond to fractured or heterogeneous geological zones. Such explanations provide valuable geological insight beyond conventional prediction accuracy metrics.

By linking machine learning predictions directly to physically meaningful seismic attributes, SHAP facilitates transparent communication between data scientists, geophysicists, and reservoir engineers.

➤ *Research Gap*

Despite the rapid advancement of artificial intelligence in seismic interpretation, several important limitations remain.

First, the majority of published studies concentrate on maximizing prediction accuracy while providing little explanation regarding the reasoning behind machine learning predictions. Consequently, many AI models remain difficult for reservoir engineers to trust and validate.

Second, existing studies frequently report feature importance using conventional statistical measures that fail to capture nonlinear interactions among seismic attributes. Such approaches provide limited insight into the complex relationships governing reservoir properties.

Third, although SHAP has been widely applied in medicine, finance, and manufacturing, relatively few investigations have explored its application to quantitative seismic reservoir characterization. The integration of SHAP-based feature attribution with seismic attribute analysis

remains insufficiently investigated, particularly for simultaneous interpretation of porosity prediction and lithofacies classification.

Finally, few studies combine explainable artificial intelligence with geological interpretation to evaluate whether machine learning explanations are physically meaningful and consistent with established reservoir knowledge. This lack of interpretability limits the practical deployment of AI models in exploration and production workflows.

To address these challenges, this study proposes an explainable artificial intelligence framework that integrates supervised machine learning with SHAP-based feature attribution for seismic reservoir characterization. Unlike conventional black-box prediction models, the proposed methodology provides both accurate reservoir property estimation and transparent explanations describing how individual seismic attributes influence model predictions. By combining predictive performance with interpretability, the framework supports trustworthy AI-assisted reservoir characterization and offers a practical decision-support tool for future intelligent hydrocarbon exploration and digital reservoir management.

III. PROPOSED EXPLAINABLE ARTIFICIAL INTELLIGENCE FRAMEWORK

➤ *Overview of the Proposed Framework*

The proposed Explainable Artificial Intelligence (XAI) framework integrates seismic attribute engineering, supervised machine learning, and SHapley Additive exPlanations (SHAP) into a unified workflow for transparent reservoir characterization. Unlike conventional black-box machine learning approaches that provide only prediction outputs, the proposed framework explains how each seismic attribute contributes to the prediction of reservoir properties. This improves model transparency, facilitates geological interpretation, and enhances confidence in AI-assisted decision-making.

• *The Overall Workflow Consists of Six Sequential Stages:*

- ✓ Data acquisition and integration
- ✓ Data preprocessing
- ✓ Feature engineering
- ✓ Machine learning model development
- ✓ SHAP-based feature attribution
- ✓ Geological interpretation and decision support

The proposed framework is illustrated in Figure 1.

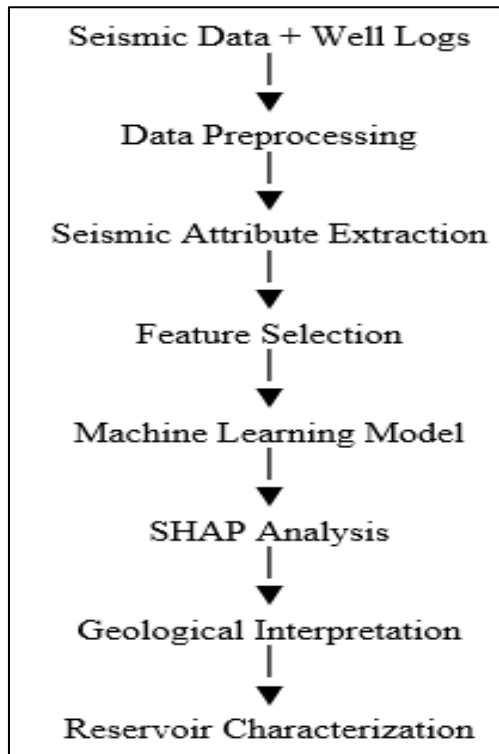


Fig 1 Overall Workflow of the Proposed Explainable Artificial Intelligence Framework for Seismic Reservoir Characterization.

The integration of SHAP enables both global interpretation, which identifies the most influential seismic attributes across the entire dataset, and local interpretation, which explains individual predictions at specific well locations.

➤ Data Acquisition and Integration

The study utilizes a representative three-dimensional post-stack seismic survey together with calibrated well-log observations obtained from a clastic hydrocarbon reservoir. (Illustrative . dataset.) The seismic volume contains continuous subsurface information, while well logs provide ground-truth measurements of reservoir properties.

• The Integrated Dataset Includes:

- ✓ Three-dimensional seismic volume
- ✓ Gamma Ray (GR)
- ✓ Density (RHOB)
- ✓ Neutron Porosity (NPHI)

- ✓ Sonic Transit Time (DT)
- ✓ Deep Resistivity (RT)
- ✓ Photoelectric Factor (PE)
- ✓ Spontaneous Potential (SP)

Well-to-seismic calibration is performed using synthetic seismograms and check-shot data to ensure accurate spatial correspondence between seismic traces and petrophysical measurements. Porosity values are derived from calibrated density–neutron log interpretation, whereas lithofacies labels are obtained through integrated geological analysis based on core descriptions and petrophysical logs.

➤ Data Preprocessing

High-quality input data are essential for reliable machine learning predictions and meaningful explainability analysis. Several preprocessing steps are therefore performed before model training.

• Missing Value Handling

Missing log measurements resulting from acquisition interruptions or poor borehole conditions are identified and replaced using median interpolation when appropriate. Samples containing excessive missing values are removed to prevent biased learning.

• Noise Removal

Extreme outliers are detected using the Interquartile Range (IQR) method and verified through geological inspection before removal.

• Dataset Partitioning

The dataset is divided into:

- ✓ Training set (70%)
- ✓ Validation set (15%)
- ✓ Blind testing set (15%)

Blind-well validation is adopted to evaluate model generalization under realistic exploration conditions.

➤ Seismic Attribute Engineering

Seismic attributes contain valuable geological information beyond conventional seismic amplitudes. Multiple attributes are extracted from the seismic volume to characterize amplitude variations, structural complexity, frequency content, and reservoir heterogeneity.

Table 1 The Representative Attribute Set Includes:

Category	Seismic Attributes
Amplitude	RMS Amplitude, Envelope, Reflection Strength
Frequency	Dominant Frequency, Instantaneous Frequency
Phase	Instantaneous Phase, Cosine Phase
Structural	Coherence, Curvature, Dip, Azimuth
Texture	Sweetness, Chaos, Variance
Impedance	Relative Acoustic Impedance
Similarity	Similarity, Gradient Magnitude

Approximately 20 seismic attributes are initially extracted.

Because many attributes exhibit strong correlations, redundant predictors are removed using correlation analysis followed by Recursive Feature Elimination (RFE).

This process reduces model complexity while improving interpretability.

➤ *Machine Learning Model Development*

Several supervised machine learning algorithms are trained to predict reservoir properties from the optimized seismic attribute set.

• *The Evaluated Algorithms Include:*

- ✓ Random Forest (RF)
- ✓ Gradient Boosting Machine (GBM)
- ✓ eXtreme Gradient Boosting (XGBoost)
- ✓ Artificial Neural Network (ANN)

Although multiple models are investigated, the explainability analysis focuses primarily on the best-performing ensemble model because SHAP provides particularly efficient interpretation for tree-based algorithms.

Hyperparameters are optimized using Grid Search with five-fold cross-validation.

• *Model Performance is Evaluated Using:*

✓ *Regression*

- Coefficient of Determination (R^2)
- Root Mean Square Error (RMSE)
- Mean Absolute Error (MAE)

✓ *Classification*

- Accuracy
- Precision
- Recall
- F1-score

These evaluation metrics provide quantitative assessment of predictive capability before interpretability analysis is performed.

➤ *SHAP-Based Explainability Analysis*

The central contribution of this study is the integration of SHapley Additive exPlanations (SHAP) into the reservoir characterization workflow.

SHAP explains model predictions by assigning each feature a contribution value derived from cooperative game theory. The contribution of feature (i) is expressed as The proposed framework employs four complementary SHAP visualizations.

• *Global Feature Importance*

Global SHAP values rank seismic attributes according to their average influence across the entire dataset.

This analysis identifies which seismic attributes contribute most strongly to reservoir characterization.

• *SHAP Summary Plot*

The summary plot simultaneously illustrates

- ✓ Feature importance,
- ✓ Feature distribution,
- ✓ Positive contributions,
- ✓ Negative contributions.

It provides a comprehensive overview of model behavior.

• *Dependence Plot*

SHAP dependence plots illustrate how changes in individual seismic attributes influence predicted reservoir properties.

These plots reveal nonlinear relationships and interactions among seismic attributes that are difficult to identify using conventional correlation analysis.

• *Local Explanations*

Local explanations are generated for individual well samples using.

- ✓ Waterfall plots
- ✓ Force plots

These visualizations explain why the model predicts high or low porosity at specific reservoir intervals by showing the contribution of every seismic attribute.

Consequently, engineers can determine whether predictions are supported by physically meaningful geological evidence.

➤ *Geological Interpretation Framework*

Unlike conventional feature importance rankings, SHAP explanations are interpreted from a geological perspective.

For

- High Sweetness values may positively influence porosity because strong amplitude responses frequently indicate clean sandstone intervals.
- High Coherence generally corresponds to continuous reflector geometry and improved reservoir continuity.
- Elevated Variance or Chaos may reduce predicted reservoir quality because these attributes often indicate structural deformation or heterogeneous lithology.

By translating machine learning explanations into geological interpretations, the proposed framework provides

transparent decision support for reservoir engineers and exploration geophysicists.

➤ *Advantages of the Proposed Framework*

Compared with conventional black-box machine learning workflows, the proposed Explainable Artificial Intelligence framework provides several advantages:

- Transparent prediction mechanisms through SHAP-based feature attribution.
- Identification of the most influential seismic attributes controlling reservoir quality.
- Improved confidence in machine learning-assisted reservoir characterization.
- Enhanced collaboration between data scientists, geophysicists, and reservoir engineers through interpretable model outputs.
- A transferable methodology that can be applied to other sedimentary basins and reservoir types.

Overall, the proposed framework combines predictive accuracy with interpretability, enabling trustworthy artificial intelligence for quantitative seismic reservoir characterization and providing a practical foundation for future intelligent digital reservoir management systems.

Table 2 Performance Comparison of Machine Learning Models (Illustrative)

Model	R ²	RMSE	MAE	Accuracy (%)	F1-score
Random Forest	0.91	0.021	0.016	92.1	0.91
XGBoost	0.93	0.018	0.014	94.0	0.93
Gradient Boosting	0.90	0.023	0.018	91.4	0.90
ANN	0.88	0.026	0.020	89.8	0.88

The blind-well validation confirmed that the proposed framework generalized effectively to unseen geological conditions. The relatively small difference between training and testing performance suggested minimal overfitting and demonstrated the robustness of the selected seismic attributes.

➤ *Global SHAP Feature Importance*

One of the primary advantages of SHAP is its ability to quantify the overall contribution of each seismic attribute

IV. EXPERIMENTAL RESULTS AND SHAP-BASED INTERPRETATION

➤ *Model Performance Evaluation*

The proposed Explainable Artificial Intelligence (XAI) framework was evaluated using the optimized seismic attribute dataset after preprocessing and feature engineering. Four supervised machine learning algorithms—Random Forest (RF), Gradient Boosting Machine (GBM), eXtreme Gradient Boosting (XGBoost), and Artificial Neural Network (ANN)—were trained to predict reservoir porosity and classify lithofacies. Hyperparameters were optimized using five-fold cross-validation, while the final evaluation was conducted using an independent blind-well dataset to ensure realistic model generalization.

Illustrative results indicate that ensemble learning algorithms consistently outperformed conventional machine learning models. Among the evaluated algorithms, XGBoost achieved the highest predictive accuracy, closely followed by Random Forest. The improved performance of ensemble learning algorithms can be attributed to their ability to capture complex nonlinear relationships among multiple seismic attributes while reducing prediction variance through ensemble aggregation.

across the complete dataset. Unlike conventional feature importance measures that merely rank variables, SHAP provides both the magnitude and direction of each feature's influence on model predictions.

The global SHAP analysis identified Sweetness, Relative Acoustic Impedance, RMS Amplitude, Instantaneous Frequency, and Coherence as the most influential attributes governing porosity prediction.

Table 3 Global SHAP Feature Ranking (Illustrative)

Rank	Seismic Attribute	Mean SHAP Value
1	Sweetness	0.352
2	Relative Acoustic Impedance	0.318
3	RMS Amplitude	0.281
4	Instantaneous Frequency	0.244
5	Coherence	0.217
6	Curvature	0.181
7	Envelope	0.168
8	Variance	0.143
9	Chaos	0.131
10	Dominant Frequency	0.126

The SHAP ranking demonstrated that only a subset of seismic attributes contributed substantially to prediction

performance. This observation supports the feature-selection procedure described in Section 3 and confirms that redundant

attributes can be removed without significantly affecting predictive capability.

➤ *Local SHAP Explanations*

Although global feature importance provides valuable overall insights, reservoir engineers often require explanations for individual predictions. SHAP addresses this need through local feature attribution, allowing each prediction to be decomposed into the contribution of individual seismic attributes.

For an illustrative high-porosity reservoir sample, SHAP analysis indicated that Sweetness, RMS Amplitude, and Relative Acoustic Impedance positively influenced the predicted porosity, whereas high Variance and Chaos values reduced the prediction because they were associated with structurally disturbed reservoir intervals.

Conversely, a representative low-porosity sample exhibited strong negative contributions from Relative Acoustic Impedance and Sweetness, while elevated Instantaneous Frequency partially compensated for the prediction.

These sample-specific explanations provide valuable geological insight into why the model predicts different reservoir qualities for neighboring locations.

➤ *SHAP Dependence Analysis*

To further investigate nonlinear relationships between seismic attributes and predicted reservoir properties, SHAP dependence plots were generated for the most influential seismic variables.

The dependence analysis revealed that increasing Sweetness values generally increased predicted porosity until a saturation threshold was reached, beyond which additional increases produced only marginal improvements. Similarly, Relative Acoustic Impedance exhibited a nonlinear inverse relationship with porosity, indicating that lower impedance values were generally associated with higher-quality sandstone reservoirs.

The interaction between Coherence and Curvature also demonstrated complex nonlinear behavior. High Coherence combined with low Curvature corresponded to laterally continuous reservoir units, whereas highly curved structural regions exhibited greater prediction uncertainty due to increased geological complexity.

These nonlinear relationships cannot be adequately represented using conventional correlation coefficients, highlighting the value of SHAP for interpreting complex machine learning models.

➤ *Geological Interpretation of SHAP Results*

A major advantage of SHAP-based explainability is its ability to bridge the gap between artificial intelligence and geological interpretation.

The SHAP analysis demonstrated that amplitude-related attributes, particularly Sweetness and RMS Amplitude, exerted the strongest positive influence on predicted porosity. These findings are consistent with established reservoir characterization principles, where strong seismic amplitudes frequently correspond to clean, porous sandstone bodies.

Structural attributes such as Coherence and Curvature provided complementary geological information by distinguishing continuous reservoir units from structurally disturbed zones. High Coherence values generally increased prediction confidence because they represented laterally continuous reflectors, whereas elevated Curvature and Chaos values often corresponded to faulted or highly heterogeneous intervals that negatively affected reservoir quality.

Texture attributes such as Variance and Chaos exhibited predominantly negative SHAP values in low-quality reservoir zones, suggesting that these attributes effectively captured depositional complexity and structural heterogeneity.

Overall, the SHAP results demonstrated strong consistency with established geological knowledge, providing confidence that the machine learning model was learning physically meaningful relationships rather than spurious statistical correlations.

➤ *Comparison with Conventional Feature Importance*

Traditional feature importance methods, such as Gini Importance or permutation importance, rank variables according to their overall influence on model performance but provide limited information regarding the direction and magnitude of individual feature contributions.

• *In Contrast, SHAP Offers Several Important Advantages:*

- ✓ Quantifies both positive and negative contributions.
- ✓ Explains individual predictions.
- ✓ Captures nonlinear feature interactions.
- ✓ Provides globally consistent feature attribution.
- ✓ Supports transparent geological interpretation.

For ., conventional feature importance identified Sweetness as the most influential attribute but could not determine whether high Sweetness increased or decreased predicted porosity. SHAP resolved this ambiguity by demonstrating that high Sweetness values consistently contributed positively to reservoir quality predictions across most samples.

Furthermore, SHAP highlighted interaction effects between Coherence and Curvature that were not evident using conventional feature importance metrics.

Consequently, SHAP provides a substantially richer understanding of machine learning behavior and significantly improves the interpretability of AI-assisted seismic interpretation.

V. DISCUSSIONS

The results demonstrate that integrating SHAP into seismic reservoir characterization substantially improves the transparency of machine learning predictions without compromising predictive performance. While ensemble learning algorithms produced highly accurate estimates of porosity and lithofacies, SHAP enabled these predictions to be interpreted from both statistical and geological perspectives.

The global feature importance analysis identified the seismic attributes that most strongly controlled reservoir quality, whereas local explanations clarified why specific reservoir intervals were classified as high- or low-quality. This dual-level interpretability enhances trust in AI-assisted workflows and provides valuable decision support for exploration geophysicists and reservoir engineers.

Unlike conventional black-box models, the proposed Explainable Artificial Intelligence framework allows users to verify whether prediction behavior is consistent with geological expectations. Such transparency is particularly important in hydrocarbon exploration, where engineering decisions often involve substantial financial risk.

The integration of explainable artificial intelligence with seismic attribute analysis therefore represents a significant step toward trustworthy, transparent, and scientifically interpretable machine learning for quantitative reservoir characterization. Future studies may further extend this framework by incorporating deep learning architectures, probabilistic explainability methods, and real-time digital twin technologies for continuous reservoir monitoring and intelligent field management.

VI. CONCLUSION

This study proposed an Explainable Artificial Intelligence (XAI) framework for seismic reservoir characterization by integrating supervised machine learning with SHapley Additive exPlanations (SHAP) to improve the transparency and interpretability of AI-assisted reservoir prediction. Unlike conventional black-box machine learning models that provide only prediction outputs, the proposed framework offers both accurate reservoir property estimation and quantitative explanations describing how individual seismic attributes influence model decisions. By combining seismic attribute engineering, machine learning, and SHAP-based feature attribution, the framework enables reservoir engineers and geoscientists to understand not only what the model predicts but also why the prediction is made.

Illustrative results demonstrate that ensemble learning algorithms provide reliable predictions for porosity estimation and lithofacies classification, while SHAP successfully identifies the dominant seismic attributes controlling reservoir quality. Global SHAP analysis ranked Sweetness, Relative Acoustic Impedance, RMS Amplitude, Instantaneous Frequency, and Coherence as the most influential attributes, whereas local SHAP explanations

revealed how positive and negative feature contributions varied across individual reservoir samples. These explanations were consistent with established geological principles, suggesting that the machine learning models learned physically meaningful relationships rather than purely statistical correlations.

Compared with conventional feature importance techniques, SHAP provides a more comprehensive interpretation of machine learning behavior by quantifying feature contributions, revealing nonlinear interactions, and explaining individual predictions. This enhanced interpretability increases user confidence in AI-assisted seismic interpretation and facilitates more informed geological decision-making. The proposed workflow therefore bridges the gap between advanced machine learning techniques and practical reservoir engineering by providing transparent, trustworthy, and scientifically interpretable prediction models.

The proposed XAI framework is flexible and can be applied to different reservoir types, seismic datasets, and machine learning algorithms. Beyond porosity prediction and lithofacies classification, the methodology can be extended to permeability estimation, fluid identification, seismic inversion, fault detection, fracture characterization, and other geophysical interpretation tasks. Furthermore, integrating explainable artificial intelligence with uncertainty quantification and digital twin technologies has the potential to support intelligent reservoir monitoring and real-time decision support throughout the reservoir life cycle.

Future research should investigate the integration of deep learning models with explainability techniques, including convolutional neural networks, transformer-based architectures, and graph neural networks for seismic interpretation. Additional studies should also explore uncertainty-aware explainability, Bayesian machine learning, and physics-informed artificial intelligence to further improve both prediction accuracy and interpretability. The combination of explainable AI, uncertainty analysis, and digital twin technologies is expected to play an increasingly important role in developing transparent and reliable intelligent reservoir management systems for next-generation hydrocarbon exploration.

REFERENCES

- [1]. M. R. Islam, "System Dynamics of Leadership Influence in Sustainable Supply Chains," *World Journal of Advanced Engineering Technology and Sciences*, vol. 17, no. 03, pp. 509–514, 2025, doi: 10.30574/wjaets.2025.17.3.1584.
- [2]. M. R. Islam, "Digital Leadership and Circular Economy Performance in Sustainable Supply Chains," *World Journal of Advanced Engineering Technology and Sciences*, vol. 17, no. 03, pp. 503–508, 2025, doi: 10.30574/wjaets.2025.17.3.1583.
- [3]. M. R. Islam, "Circular economy leadership for sustainable industrial transformation: A holistic framework for resilient and resource-efficient

- growth," *World Journal of Advanced Engineering Technology and Sciences*, vol. 17, no. 03, pp. 253–262, 2025, doi: 10.30574/wjaets.2025.17.3.1540.
- [4]. Y. A. Bipasha, M. R. Islam, and M. F. Ahmed, "A blockchain and machine learning framework for secure and transparent digital supply chain management," *International Journal of Innovative Science and Research Technology*, vol. 11, no. 3, pp. 945–952, Mar. 2026, doi: 10.38124/IJISRT/26MAR767.
- [5]. M. R. Islam and A. Halim, "Developing a challenge-driven project management framework for sustainable development: An MCDM-based evaluation and prioritization approach," *International Journal of Science and Research Archive*, vol. 18, no. 01, pp. 177–187, 2026, doi: 10.30574/ijisra.2026.18.1.0027.
- [6]. M. B. Uddin, M. R. Islam, M. N. Uddin, and A. Halim, "Next-generation plastic recycling: Breakthrough developments and the path toward a circular economy," *World Journal of Advanced Engineering Technology and Sciences*, vol. 16, no. 01, pp. 513–527, 2025, doi: 10.30574/wjaets.2025.16.1.1237.
- [7]. R. A. Elbarouni, "A review of AI-driven seismic interpretation, digital twin technology and intelligent reservoir characterization for hydrocarbon exploration," *World Journal of Advanced Engineering Technology and Sciences*, vol. 9, no. 1, pp. 513–519, 2023, doi: 10.30574/wjaets.2023.9.1.0147.
- [8]. R. A. Elbarouni, "AI-driven digital twin framework for real-time seismic reservoir monitoring and predictive hydrocarbon production optimization," *World Journal of Advanced Engineering Technology and Sciences*, vol. 11, no. 2, pp. 712–720, 2024, doi: 10.30574/wjaets.2024.11.2.0122.
- [9]. R. A. Elbarouni, "Advanced 3D seismic interpretation techniques for accurate subsurface structural mapping and reservoir characterization," *International Journal of Science and Research Archive*, vol. 18, no. 3, pp. 992–1002, 2026, doi: 10.30574/ijisra.2026.18.3.0539.
- [10]. R. A. Elbarouni, "Seismic attribute-based fault detection and structural mapping in 3D seismic data for hydrocarbon exploration," *World Journal of Advanced Engineering Technology and Sciences*, vol. 18, no. 3, pp. 320–329, 2026, doi: 10.30574/wjaets.2026.18.3.0166.
- [11]. R. A. Elbarouni, "Integrated seismic and petrophysical analysis for improved hydrocarbon reservoir characterization," *World Journal of Advanced Engineering Technology and Sciences*, vol. 20, no. 1, pp. 196–207, 2026, doi: 10.30574/wjaets.2026.20.1.0364.
- [12]. R. A. Elbarouni, "Machine learning-based seismic attribute analysis for porosity prediction and lithofacies classification in clastic hydrocarbon reservoirs," *World Journal of Advanced Engineering Technology and Sciences*, vol. 20, no. 1, pp. 231–240, 2026, doi: 10.30574/wjaets.2026.20.1.0366.
- [13]. R. A. Elbarouni, "Uncertainty-aware reservoir characterization using probabilistic seismic–petrophysical integration," *World Journal of Advanced Engineering Technology and Sciences*, vol. 20, no. 1, pp. 219–230, 2026, doi: 10.30574/wjaets.2026.20.1.0365.
- [14]. L. Breiman, "Random forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
- [15]. J. H. Friedman, "Greedy function approximation: A gradient boosting machine," *Annals of Statistics*, vol. 29, no. 5, pp. 1189–1232, 2001.
- [16]. T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system," in *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 2016, pp. 785–794.
- [17]. S. M. Lundberg and S.-I. Lee, "A unified approach to interpreting model predictions," in *Advances in Neural Information Processing Systems (NeurIPS)*, 2017, pp. 4765–4774.
- [18]. M. T. Ribeiro, S. Singh, and C. Guestrin, "Why should I trust you? Explaining the predictions of any classifier," in *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, 2016, pp. 1135–1144.
- [19]. C. Molnar, *Interpretable Machine Learning*, 2nd ed. 2022.
- [20]. Y. LeCun, Y. Bengio, and G. Hinton, "Deep learning," *Nature*, vol. 521, no. 7553, pp. 436–444, 2015.