

Improving the accuracy of critical downhole equipment detection using unsupervised machine learning and deep learning methods

Mojtaba Farasat¹, Abobakr Sori², Mojtaba Rahimi^{1,3*}, Amir Farasat⁴, Seyyed Ali Eftekhari⁵

Received: 2026 Jun 04, Revised: 2026 Jul. 03, Accepted: 2026 Jul. 15, Published: 2026 Jul. 30



Journal of Geomine © 2026 by University of Birjand is licensed under CC BY 4.0

ABSTRACT

The subsurface equipment industry is crucial to sustainable oil extraction, ensuring efficiency, safety, and the absence of unplanned shutdowns. Conventional approaches to determining critical equipment are becoming less effective in the face of the current amount and complexity of operational data. This study presents a new unsupervised technique using deep clustering to fill this research gap. Based on actual sensor data from Electric Submersible Pumps (ESPs) over two years, the approach entails extensive preprocessing, such as cleaning, normalization, and handling missing values. The proposed system uses a deep clustering method based on a 1D convolutional autoencoder neural network and validates its performance against the traditional K-Means algorithm. The quality of the clusters was assessed using the objective function (OBJ) and mean squared distance (MSD) measures. The outcome proves the proposed approach to be highly superior, with an MSD of 845 and an OBJ of 6.23, compared to 1386 and 0.96 for K-Means. This shows a highly improved cluster density and data point cohesion. Further analysis revealed specific equipment behavior patterns, where Cluster 4 corresponds to critical and unstable equipment that needs to be addressed immediately, and Cluster 3 corresponds to stable and low-risk equipment. This paper proves that deep clustering with the help of autoencoders is a precise and efficient technique for critical equipment detection in a complex industrial setting, with immense potential for optimizing predictive maintenance approaches in the oil and gas industry.

KEYWORDS

Deep clustering, predictive maintenance, critical equipment detection, unsupervised learning, oil and gas industry

NOMENCLATURE

1D-CNN	One-Dimensional Convolutional Neural Network
ADAM	Adaptive Moment Estimation
AI	Artificial Intelligence
ML	Machine Learning
DTW	Dynamic Time Warping
FCNN	Fully Connected Neural Network
GMM	Gaussian Mixture Model
IoT	Internet of Things
KLD / KL	Kullback-Leibler Divergence
KNN	K-Nearest Neighbors
MSD	Mean Squared Distance
OBJ	Objective Function
PCA	Principal Component Analysis
PdM	Predictive Maintenance
RCAM	Reliability-Centered Asset Management
RCM	Reliability-Centered Maintenance
RNN	Recurrent Neural Network
RUL	Remaining Useful Life
WCSS	Within-Cluster Sum of Squares

ESP	Electric Submersible Pump
BIM	Building Information Modeling
TS	Time Series

I. INTRODUCTION

The oil and gas industry underpins global energy supply, economic growth, and geopolitical security (Kashubsky, 2015; Salimovna, 2024). Extraction operations depend on the reliable functioning of numerous mechanical, electronic, and control components. System failures can trigger production halts, escalating operational costs, worker safety hazards, and irreversible environmental damage (Al-Mekhlafi et al., 2025; Egila et al., 2025; Fang, 2025).

Subsurface equipment—particularly Electric Submersible Pumps (ESPs), downhole control valves, and pressure-temperature-flow sensors—operates under extreme conditions: pressures reaching thousands of psi, temperatures up to 150°C, and

¹ Department of Petroleum Engineering, Kho.C., Islamic Azad University, Khomeinishahr, Iran, ²Transport Phenomena Research Center, Chemical Engineering Faculty, Sahand University of Technology, Tabriz, Iran, ³Stone Research Center, Kho.C., Islamic Azad University, Khomeinishahr, Iran, ⁴Iranian Offshore Oil Company, Tehran, Iran, ⁵Department of Mechanical Engineering, Kho.C., Islamic Azad University, Khomeinishahr, Iran
✉ M. Rahimi: mrahimi@iau.ac.ir

corrosive fluid environments (Ramirez and Martinez, 2017). Repair or replacement costs, often requiring workover operations, combined with direct production impacts, make early fault detection and critical equipment designation a top priority (Bloch, 2012; Kakoulli-Constantinou and Papadima-Sophocleous, 2020; Lehr and Furlan, 2017; Pastorek et al., 2019; Salah et al., 2025). Critical equipment is defined as components whose failure jeopardizes system reliability, hydrocarbon production rates, and maintenance budgets. Reliability-Centered Management (RCM) traditionally begins with critical equipment identification. While conventional RCM and its software extensions (e.g., RCAM-Power Distribution System Analysis (Bertling, 2002)) and maintenance optimization integrated with condition monitoring (Tjernberg, 2023) have contributed valuable concepts, substantial limitations persist. These include heavy reliance on human expertise, oversimplified sensitivity analyses, exclusion of key performance criteria, inability to process large-scale complex operational data, and failure to adapt to the dynamic oil business environment where reservoir conditions constantly shift.

Recent advances in sensing technology, Internet of Things (IoT), and big data analytics have generated abundant, high-resolution equipment data, enabling novel Predictive Maintenance (PdM) strategies incorporating artificial intelligence and machine learning models. Deep learning architectures—Fully Connected Neural Networks, Convolutional Neural Networks, and Recurrent Neural Networks—have been applied to predict Remaining Useful Life (RUL) of rig machinery using Principal Component Analysis (PCA) and signal processing (Abisoye et al., 2025; Latrach, 2023; Visconti et al., 2024). Case studies on AI-driven PdM for water injection pumps emphasize that data preprocessing and feature selection are critical steps (Guidotti et al., 2025; Mohamed Almazrouei et al., 2023; Visconti et al., 2024). Integrating Digital Twins, Building Information Modeling, ML, and IoT shows significant promise for pumping systems, though dynamic data processing remains a bottleneck (Hallaji et al., 2022; Shahin et al., 2023).

At the application level, hybrid PCA-XGBoost methods have predicted ESP failures up to seven days in advance with considerable accuracy (Abdalla et al., 2022). Industry 4.0 lean techniques have substantially improved corrective maintenance for hydraulic systems using multi-level operational models (Torre and Bonamigo, 2024). Unsupervised clustering combined with PCA of vibration signals has enabled useful comparative analyses among algorithms (Amruthnath and Gupta, 2018). ML-based historical failure pattern learning from subsurface equipment has reduced maintenance costs by up to 3.6% (Bani and Nwosu, 2024). Hybrid approaches using Affinity Propagation

and Dynamic Time Warping achieved 100% recall in identifying drilling irregularities (Liu et al., 2025), while transformer-based attention mechanisms surpassed expert judgment for coiled tubing operations (Wu et al., 2025).

Despite these successes, critical research gaps remain. Most studies address surface or general equipment, or employ supervised ML requiring large labeled datasets—rare in real oilfield contexts where data is noisy, sparse, unlabeled, and exhibits highly nonlinear variable relationships. Even subsurface-focused studies (Bani and Nwosu, 2024; Wu et al., 2025) lack unsupervised deep learning insights and remain task-specific. Our contribution is threefold: (1) deep learning-based equipment clustering for automatic, effective classification based on individual equipment value to system reliability; (2) design of an optimal 1D Convolutional Autoencoder for robust data compression and intricate pattern identification in multivariate time series; and (3) end-to-end deep model training enabling seamless learning from raw data to final prioritization.

Unlike classical (Baldi, 2012; Bertling, 2002) and recent semi-deep or constrained approaches (Abdalla et al., 2022; Amruthnath and Gupta, 2018; Bani and Nwosu, 2024; Hallaji et al., 2022; Latrach, 2023; Liu et al., 2025; Mohamed Almazrouei et al., 2023; Torre and Bonamigo, 2024; Wu et al., 2025), our framework overcomes their limitations, effectively operating on real industrial data. Expected benefits include reduced operating and maintenance costs, extended equipment service life, improved operational safety, lower environmental risks, and enhanced productivity—constituting a significant contribution toward fully intelligent RCM procedures.

Our approach diverges from prior work in four key ways. First, we designed a 1D convolutional autoencoder (kernel size 3) specifically for ESP data, capturing patterns across consecutive 5-minute intervals where parameters like pressure, temperature, and flow evolve gradually—subtle trends often signal impending trouble. Second, unlike typical two-stage feature extraction-then-clustering pipelines, our framework jointly optimizes reconstruction and clustering objectives, learning representations suited for both tasks simultaneously. Third, to the best of our knowledge, this is the first application of deep unsupervised clustering to ESP operational data from multiple wells over two years; previous ESP work relied on supervised methods (Abdalla et al., 2022) or simpler clustering techniques, whereas our method works without labels while capturing complex nonlinear relationships. Fourth, the method prioritizes practicality and interpretability: each cluster is characterized by risk level and maintenance priority. For instance, Cluster 4 exhibited critical behavior (high fluctuations, elevated risk), while Cluster 3 remained stable and low-risk, providing engineers actionable insights for RCM planning. Notably, using only

three key parameters (pressure, temperature, flow) yielded better results (MSD = 845, OBJ = 6.23) than using all eight parameters (MSD = 0.67, OBJ = 1.53), confirming that thoughtful feature selection outperforms brute-force data inclusion.

In essence, our work delivers a practical, unsupervised solution for critical equipment identification that bypasses expensive labeled data requirements, handles real-world operational complexity, and produces engineer-friendly, interpretable outputs.

II. MATERIALS AND METHODS

A. Simulation System

To conduct the simulations in this study, a computer system with sufficient hardware capabilities for image processing was required. The system is equipped with an Intel Core i7-6700HQ processor (base frequency of 2.60 GHz, quad-core) and 12 GB of RAM. It runs the Windows 10 operating system. All simulations were performed using MATLAB software.

B. Research Data

1) Data Source and Collection

The dataset for this research originates from 18 Electrical Submersible Pumps (ESPs) operating across five well sites within a single Iranian oil field, monitored from November 2022 to August 2024 as part of a digital oilfield performance initiative. Each ESP incorporates a Zenith E7 downhole sensor suite, which captures pressure (0–5800 psi, $\pm 0.5\%$ FS), temperature (0–300°C, $\pm 1^\circ\text{C}$ post-calibration), vibration (100 mV/g accelerometers), motor current ($\pm 1\%$), and applied voltage ($\pm 0.5\%$). These measurements travel along the power cable to an ABB surface control panel, where a PLC samples the feed every 5 minutes before archiving it in a centralized SCADA system. Over the 24-month window, we collected 1,024,356 valid logs, each containing eight primary parameters (Pi, Pd, Tm, Vt, Vz, V, Ct). The raw CSV exports, however, contained $\sim 3.7\%$ missing entries from sensor glitches and $\sim 2.1\%$ outliers caused by transient operational disturbances. Importantly, the logs capture the full spectrum of ESP behavior: normal production (inlet pressure 800–3200 psi, intake temperature 45–85°C, flow 200–600 m³/day), transient phases like startups and load changes, and critical anomalies such as current spikes exceeding 120%, temperature excursions past 120°C, or pressure drops greater than 500 psi within half an hour. For confidentiality, the company anonymized all proprietary metadata before release. We then applied the preprocessing routine detailed in the Proposed Method (Section 2.3) to handle missing values and outliers. Table 1 provides a complete specification of its structure and statistics.

The current study is designed to address the key gaps identified in prior work by proposing an intelligent,

integrated framework that leverages deep unsupervised learning, with a specific focus on subsurface equipment in the oil and gas industry. This approach advances the identification and prioritization of critical equipment to unprecedented levels of accuracy, speed, and adaptability.

C. Proposed Method

The proposed framework operates through an integrated, joint optimization mechanism, fundamentally distinguishing it from conventional sequential pipelines. The workflow begins with rigorous data preparation, where raw operational data comprising 8 parameters undergoes cleaning, missing value imputation, and normalization as described in Section 2.3.2. This preprocessed data is then fed into the encoder of a 1D Convolutional Autoencoder. The encoder, composed of 1D convolutional layers with a kernel size of 3 and subsequent max-pooling layers, effectively captures local temporal dependencies while progressively reducing the feature dimensionality from 8 to a compact 3-dimensional latent vector z_i per sample. This latent space is specifically engineered to preserve essential operational characteristics while filtering out noise.

The core interaction occurs as these latent representations serve directly as the input to the clustering component—they are not utilized as pseudo-labels for supervised training. Instead, the autoencoder and clustering module are jointly optimized within a single objective function. The clustering head initializes K cluster centers μ_j using the K -Means++ algorithm applied to the initial latent space. It then computes soft assignment probabilities q_{ij} using the Student's t -distribution to measure the similarity between each latent vector and every cluster center. Concurrently, an auxiliary target distribution p_{ij} is derived from q_{ij} to accentuate high-confidence assignments and enhance cluster purity. The joint loss function is formulated as the sum of the Reconstruction Loss (ensuring the decoder can faithfully reconstruct the original input from z_i) and the Clustering Loss, which is the KL divergence between q_{ij} and p_{ij} . These two components are balanced by a weighting factor $\lambda=0.1$.

During backpropagation, gradients simultaneously update the autoencoder's convolutional weights—forcing the feature extractor to produce cluster-friendly representations—and the cluster centers μ_j to refine centroid positions. After the model converges, the final latent space is fixed. The Elbow method is applied to determine the optimal number of clusters, yielding $K=5$. Subsequently, Algorithm 1 selects the most critical clusters based on a priority metric that holistically assesses within-cluster variance, the mean values of operational parameters, and associated risk indicators. The overall process flow of the proposed method is illustrated in Fig. 1.

Table 1. Data description

Feature Name	Unit	Data Type	Description
Date/Time	-	Datetime	Timestamp of data recording, continuous with 5-minute intervals
Pi	Psi	Numeric (Float)	Inlet pressure to the equipment
Pd	Psi	Numeric	Outlet or downstream pressure of the equipment
Ti	°C	Numeric	Downhole temperature
Tm	°C	Numeric	Surface or ambient temperature
Vt	mm/s	Numeric	Pump vibrations
Vz	G (probable)	Numeric	Vertical z-component of fluid flow velocity
V	V (volt)	Numeric	Voltage or combined performance parameter
Ct	mA	Numeric (Integer)	Equipment current consumption in milliamperes

1) Data Preprocessing

Data preprocessing is one of the critical stages in this research, aimed at improving the quality of real-world data collected from oil production parameters in an oil well. Due to their real nature and potential human errors, these data include missing values, noise, and qualitative inconsistencies that can negatively impact the accuracy of the fuzzy clustering model.

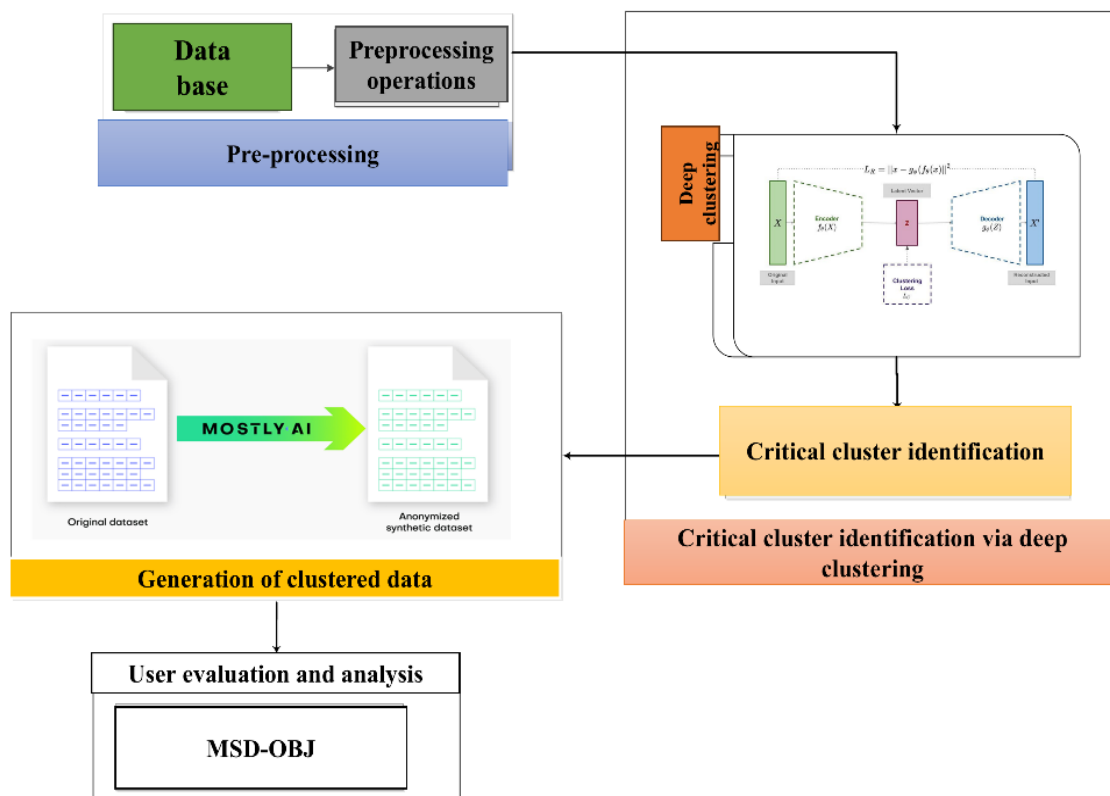
The preprocessing process includes the following steps:

Imputation of Missing Values

Missing values are identified and filled using the K-Nearest Neighbors (KNN) algorithm. This method predicts and substitutes appropriate values based on the similarity of existing records.

Noise Filtering and Outlier Detection

Following the KNN implementation, we refined the data quality to ensure reliable clustering. To dampen high-frequency noise, pressure and temperature typically shift gradually. We applied a moving average filter (window size 10). Statistical outliers were identified using the Interquartile Range method (bounds: $Q1 \pm 1.5 \times IQR$); genuine anomalies, such as sensor glitches, were preserved, whereas physically impossible readings (e.g., negative pressures) were discarded. Petroleum engineering rules further ensured consistency: discharge pressure below inlet pressure (barring pump shutoff) and current draw outside rated limits were manually verified. Finally, a Savitzky–Golay filter (window 11, order 3) smoothed the data while safeguarding meaningful peaks and trends. All steps were run in MATLAB R2022b (Signal Processing Toolbox), with a detailed log maintained for full transparency and reproducibility.

**Fig. 1.** Workflow of the Proposed Method

Dataset Partitioning and Validation Strategy

The preprocessed dataset (1,024,356 records, Nov 2022--Aug 2024) was split chronologically into 70% training, 15% validation, and 15% test sets to preserve temporal dependencies. Training employed early stopping (patience of 20 epochs) with validation checks every 10 epochs. For robustness, 5-fold cross-validation was run on the training set, holding out 20% per fold; final parameters were selected via average fold performance.

Data Normalization: To standardize the scale of features and prevent disproportionate influence on the clustering algorithm, the data are scaled to the [0, 1] range using Min-Max Scaling (whitening normalization). The following equation 1 expresses normalization for a data point:

$$y_i = y_{min} + (y_{max} - y_{min}) * \frac{(x_i - x_{min})}{(x_{max} - x_{min})} \quad (1)$$

Since this transformation is linear, the data distribution remains unchanged.

Final Preparation: The preprocessed data are organized into a feature matrix and prepared for model input. The final number of features was 8 key parameters, although a comparison was also conducted with the case of 3 primary parameters (pressure, temperature, flow).

2) Proposed Deep Clustering Method

To improve the clustering process on complex and nonlinear data, a deep autoencoder structure was employed. The primary goal of this architecture is to extract latent representations that structurally offer higher separability for clustering compared to raw input data. These compressed latent representations reduce input dimensionality while preserving important and effective information. The encoder maps the original input data to the latent space, and the decoder attempts to reconstruct the data from this space, ensuring that the compression occurs without critical information loss.

The embedded autoencoder architecture consists of a sequence of layers that not only transfer data to a compressed space but also simultaneously attempt to understand the internal data structure. Unlike traditional linear methods, this structure allows the model to learn complex, non-distance-based, and even nonlinear relationships among features. In this model, initial data encoding is performed by an embedded autoencoder within the encoder section, aiding in extracting high-level features, followed by a final encoding stage. A similar structure is applied in the decoder section, with an auxiliary encoding first, followed by final data reconstruction.

In the subsequent process, clustering is performed on the latent space. Unlike classical clustering, this method uses soft assignment to allocate each sample to every cluster. This soft assignment, implemented as a

probability matrix Q , continuously specifies the similarity between the latent sample and cluster centers. Here, a Student's t -distribution is used to measure similarity, reducing model sensitivity to distant samples and increasing focus on relevant ones.

To improve accuracy and cluster purity, an auxiliary distribution P is also employed. This distribution places greater emphasis on samples assigned with high confidence to a specific cluster, playing a significant role in reshaping the clustering structure. In fact, this mechanism enhances cluster separation quality by increasing the weight of reliable samples in the learning process. The key aspect of this approach is that clustering and latent feature learning are performed simultaneously within a joint optimization process. This concurrent process gradually enables the model not only to separate clusters more accurately but also to refine hidden representations to improve separability. Finally, the label for each sample is determined based on the highest probability of belonging to one cluster in matrix Q . The proposed deep clustering method is further examined below.

Given an input dataset, the objective of learning an embedded autoencoder is to build a suitable encoder for deep clustering that makes latent representations more appropriate for clustering. Accordingly, the encoder unit maps the input to a low-dimensional latent representation using a nonlinear mapping (Eq. (2)). The decoder unit reconstructs the input from its latent representation using a nonlinear mapping (Eq. (3)). Here, θ and φ are the learnable parameters (weights and biases) for the encoder and decoder units, respectively.

$$z_i = f_{w1}(x_i) \quad (2)$$

$$\bar{x}_i = g_{w2}(z_i) \quad (3)$$

Fig. 2 illustrates the generalized architecture of the deep embedded autoencoder for deep clustering. Here, the deeply embedded autoencoder consists of an encoder unit and an embedded decoder unit, each of which functions independently as an autoencoder. This deeply embedded autoencoder allows the model to perform multiple compression and expansion processes during data encoding and decoding, yielding stronger and better latent representations. In the encoder unit, the embedded autoencoder performs an encoding-decoding operation before final input data encoding. The autoencoder embedded in the encoder unit forms a decoding-encoding operation before the final input data encoding. Such operations enable the encoder to identify high-level features of input data, extracting powerful hidden representations. Conversely, the encoder embedded in the decoder unit performs a decoding-encoding operation before the final reconstruction of the hidden layer representation.

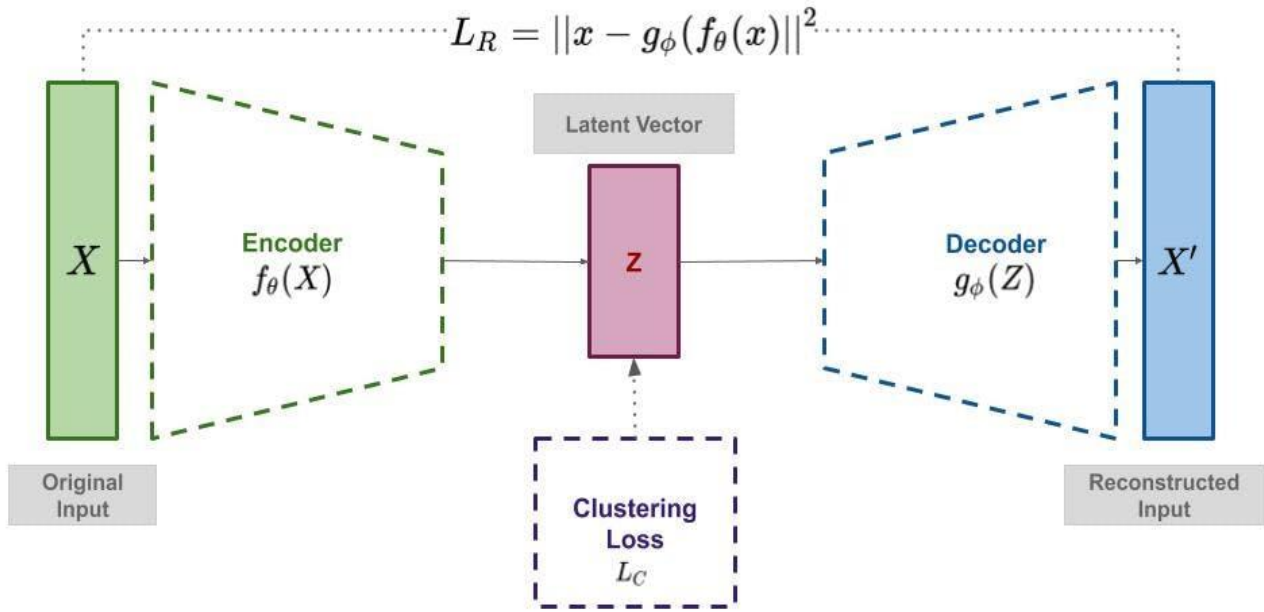


Fig. 2. Generalized Architecture of Deep Embedded Autoencoder for Deep Clustering

Joint deep clustering is a method in which deep latent representation learning and clustering are tightly coupled and jointly optimized. Deep clustering is performed based on extracted latent representations to group them into K clusters. The soft assignment used for this purpose is defined as the similarity between embedded observations and cluster centers. To achieve this, a distribution matrix is utilized. Matrix Q consists of elements q_{ij} measured using the t-distribution (Wang et al., 2024) as follows:

$$q_{ij} = \frac{(1 + (\|z_i - \mu_j\|^2)^{-1})}{\sum_{j=1}^k (1 + (\|z_i - \mu_j\|^2)^{-1})} \quad (4)$$

where q_{ij} is the probability of assigning observation i to cluster j . The auxiliary distribution matrix for improving cluster purity and emphasizing high-confidence assigned observations is introduced as follows:

$$p_{ij} = \frac{q_{ij}^2 / \psi_j}{\sum_{j=1}^k q_{ij}^2 / \psi_j} \quad (5)$$

where f_j represents soft cluster frequencies. Clustering is performed iteratively by alternating between computing the soft assignment Q and the auxiliary target distribution P . The clustering loss (Aranha et al., 2024) is defined as the Kullback-Leibler (KL) divergence (Gupta et al., 2018) between the two distributions: the soft label Q and the target distribution P . This clustering loss is computed by minimizing the following equation:

$$L_{KLD} = KL(P||Q) = \sum_i \sum_j p_{ij} \log \frac{p_{ij}}{q_{ij}} \quad (6)$$

where L_c is the clustering loss. Subsequently, the predicted label for observation i , corresponding to q_i , is assigned to class j satisfying Eq. (7):

$$l_i = \operatorname{argmax}_j q_{ij} \quad (7)$$

3) Dimensionality Reduction of Features

The autoencoder serves as the backbone of our clustering framework—a workhorse for cutting through the noise and extracting the hidden patterns buried in complex subsurface equipment data. Think of it as a smart filter: the encoder compresses high-dimensional raw inputs into a compact, meaningful latent space. At the same time, the decoder reconstructs the original data to ensure nothing critical gets lost along the way.

Raw oilfield data is messy—full of redundancy, noise, and nonlinear interactions. Clustering directly on it is like finding a needle in a haystack. However, with the autoencoder, we achieve three concrete advantages: first, dimensionality reduction speeds up processing significantly; second, latent features offer cleaner separability for more distinct clusters; third, unlike linear methods like PCA, the autoencoder captures complex nonlinear relationships, producing richer representations.

This approach is not just theoretically sound—it is practical. The latent features feed directly into our clustering stage, enabling sharper identification of critical equipment, which is essential for timely maintenance decisions. Combined with robust preprocessing, the autoencoder creates a seamless bridge between deep learning and unsupervised clustering, delivering an accurate, efficient tool tailored for petroleum industry challenges.

D. Determining Critical Clusters

The process begins with the clustered data vectors stored in an array denoted as X' . Critical clusters are identified by selecting those with the highest impact on system performance or reliability, often referred to here as "super-clusters" or clusters exhibiting characteristics indicative of criticality. Once a critical cluster is identified, all associated data points (rules or instances) within that cluster are flagged as critical and removed from further consideration in the remaining array X' . This iterative elimination ensures that subsequent selections focus on the next most critical clusters. The step-by-step process for identifying critical clusters via deep clustering is outlined below:

Algorithm Input: Array X containing features collected from the database, and creation of a vector with K -dimensional data elements, array X' in deep clustering.

Algorithm Output: Number of prioritized critical clusters for final analysis forming set C .

The procedure is executable in four steps:

Step 1: Compute the set vector X' containing all critical clusters.

Step 2: Among the remaining vectors in set X' , select vector x_i with conditions for prioritization relationships for final selection in deep clustering.

Step 3: Remove vector x_i and all existing vectors in the cluster from array X' . Update array X' as $X' \leftarrow X' \setminus [\text{vectors in the selected cluster}]$.

Step 4: Add critical clusters selected in Step 2 to the set of clusters with high risk. Also, update the vector matrix X . If the number of set members is less than K , continue algorithm execution by jumping to Step 2; otherwise, continue the prioritization algorithm. The proposed deep clustering algorithm for selecting critical clusters is presented in Algorithm 1.

Input:

- X : Array containing feature vectors extracted from the dataset

- K : Number of critical clusters to be selected

Output:

- C : A set of K prioritized critical clusters

Procedure:

1. Apply deep clustering on X to generate clustered feature vectors $\rightarrow X'$

2. Initialize:

$C \leftarrow \emptyset$ // Set of selected critical clusters

$i \leftarrow 0$ // Cluster selection counter

3. While $|C| < K$:

a. Select x_i from X' with the highest priority (based on sensitivity or importance metric)

b. Remove x_i and all related rules or vectors from X' :
 $X' \leftarrow X' \setminus \{x_i\} \cup \{\text{PrimMatrix}\}$

c. Add x_i to the set of critical clusters:

$C \leftarrow C \cup \{x_i\}$

d. Update vector matrix X and relationship matrix if necessary

4. Return C as the final output

Algorithm: Select K Critical Clusters Using Deep Clustering. After completing the deep clustering process and generating latent representations for operational data, the next step is to identify and select clusters that are technically in a critical state with the highest failure probability. For this purpose, a combined approach including statistical data analysis, examination of numerical clustering indices, and interpretation of equipment behavior was employed.

Cluster characterization relies on the mean and standard deviation of key parameters—pressure, temperature, velocity, and current. High means with wide fluctuations signal unstable, high-risk clusters; low dispersion indicates stable or optimal ones. MSD and OBJ serve as complementary cohesion/separability metrics. Risk level charts, downtime records, and criticality scores further classify clusters into stable, semi-stable, or critical levels, guiding preventive maintenance and resource allocation. To prioritize K critical clusters, we design an iterative algorithm: store clustered vectors in array X' , then, based on the highest assignment probability to critical clusters (using statistical indices), select and remove vector x_i and its entire cluster, repeating until K high-risk clusters are selected. These tools—metrics, charts, and the algorithm—together provide a reliable decision-making framework.

E. Evaluation Parameters and Model Settings

1) **Clustering Evaluation Metrics:** To assess the quality of the proposed clustering, two primary metrics were used:

OBJ Metric: The OBJ metric is a comprehensive index measuring clustering quality based on two components: intra-cluster density (data homogeneity within clusters) and inter-cluster separation (dissimilarity between clusters). Density increases with reduced average Euclidean distance of points within each cluster, and separation improves with increased average Euclidean distance between points in different clusters. Higher OBJ values indicate better clustering quality. This metric is calculated as follows:

$$obj = Compactness(C) \times Separation(C) \quad (8)$$

Mean Squared Distance (MSD): This metric evaluates data dispersion within clusters. For each cluster, the average squared Euclidean distance of points from the cluster center is calculated, and then the sum of these values across all clusters is obtained. Lower MSD values indicate higher intra-cluster density and better clustering quality, as suggested by Eq. (9):

$$MSD = \sum_K \frac{\sum_{i=1}^{m^{(k)}} \|d_i - d^{(k)}\|^2}{m^{(k)}} \quad (9)$$

where K is the total number of clusters, $m^{(k)}$ is the number of points in cluster k , d_i is a data point belonging to cluster k , and $d^{(k)}$ is the center (mean) of cluster k .

These two metrics enable a precise comparison of the proposed method's performance with traditional techniques and demonstrate its impact on critical equipment identification.

2) *Deep Autoencoder Network Settings*

The deep clustering model was implemented using a deep autoencoder architecture comprising five hidden layers. To ensure robust evaluation of the model's generalizability, the dataset was split into training, validation, and test sets. Training proceeded until convergence, with the process halted once the evaluation metrics stabilized. The primary hyperparameters included 350 epochs, a batch size of 80, an initial learning rate of 0.01, and a gradient threshold of 1. Additionally, a training progress plot was generated to monitor the learning trends throughout the optimization process.

For hyperparameter optimization (such as learning rate, number of hidden units, and batch size), random search was used to identify the best parameter combination.

3) *ADAM Optimization Algorithm*

Network weight optimization was performed using the ADAM algorithm. This algorithm combines the advantages of AdaGrad and RMSProp, providing faster convergence and better handling of sparse gradients, making it suitable for large and noisy data.

4) *Determination of Optimal Number of Clusters*

The optimal number of clusters (K) was determined using the Elbow method. In this method, the within-cluster sum of squares (distortion or MSD) is calculated for various K values, and the plot is drawn. The elbow point—where the slope reduction in the distortion plot becomes significantly slower—is selected as the optimal number of clusters. This approach prevents overfitting or underfitting and ensures meaningful and efficient clustering.

These settings and metrics provide a precise framework for implementing and evaluating the proposed method, enabling objective comparison of its performance in critical equipment identification.

5) *Hyperparameters and Software Environment*

We implemented the proposed framework in MATLAB R2022b, utilizing the Deep Learning, Statistics, and Signal Processing toolboxes on a Windows 10 workstation equipped with an Intel Core i7 processor

and 12 GB RAM. The autoencoder adopts a symmetric five-layer design, progressively compressing the 8 input parameters into a 3-dimensional latent space and reconstructing them via the decoder, with ReLU activations in hidden layers and linear outputs for continuous data reconstruction. The clustering head is configured with $K=5$, as determined by the Elbow method, and the joint objective balances reconstruction and clustering (KL divergence) losses using a weight coefficient $\lambda=0.1$. Training employs the ADAM optimizer (learning rate 0.01, batch size 80) for up to 350 epochs, though early stopping with a patience of 20 epochs on validation loss typically halts convergence around epoch 280. Hyperparameters were tuned via a random search over 50 trials, exploring learning rates (0.001–0.05), hidden units, batch sizes, and λ values (0.01–1.0); the final configuration was selected to minimize combined validation loss while maximizing OBJ and MSD. The data were split chronologically into 70% training, 15% validation, and 15% test sets.

III. RESULTS

Implementing the deep clustering algorithm on real data from Electric Submersible Pumps (ESPs) is presented here. These results focus on the model's performance compared to the traditional K-Means algorithm, evaluated using two primary criteria: the objective function (OBJ) and mean squared distance (MSD). The optimal number of clusters was determined, followed by an analysis of the clusters based on the final selection of 5 clusters.

A. *Determination of Optimal Number of Clusters*

In practice, when applying K-Means clustering, the elbow point is typically selected as the optimal number of clusters. As observed in Fig. 3, the elbow point occurs at 5 clusters. This point represents an ideal balance between minimizing the sum of distortions and maintaining a reasonably small number of clusters. However, to validate the results and ensure accuracy, additional analyses were conducted across a range of cluster values. This approach allows for comparisons among different clustering methodologies. Combining the elbow method with other metrics improves clustering accuracy and facilitates the identification of critical equipment.

The elbow point is clearly identifiable at 5 clusters. This point achieves an optimal balance between minimizing distortions and avoiding excessive partitioning of the data, thereby accurately reflecting real-world operating conditions of the ESP system.

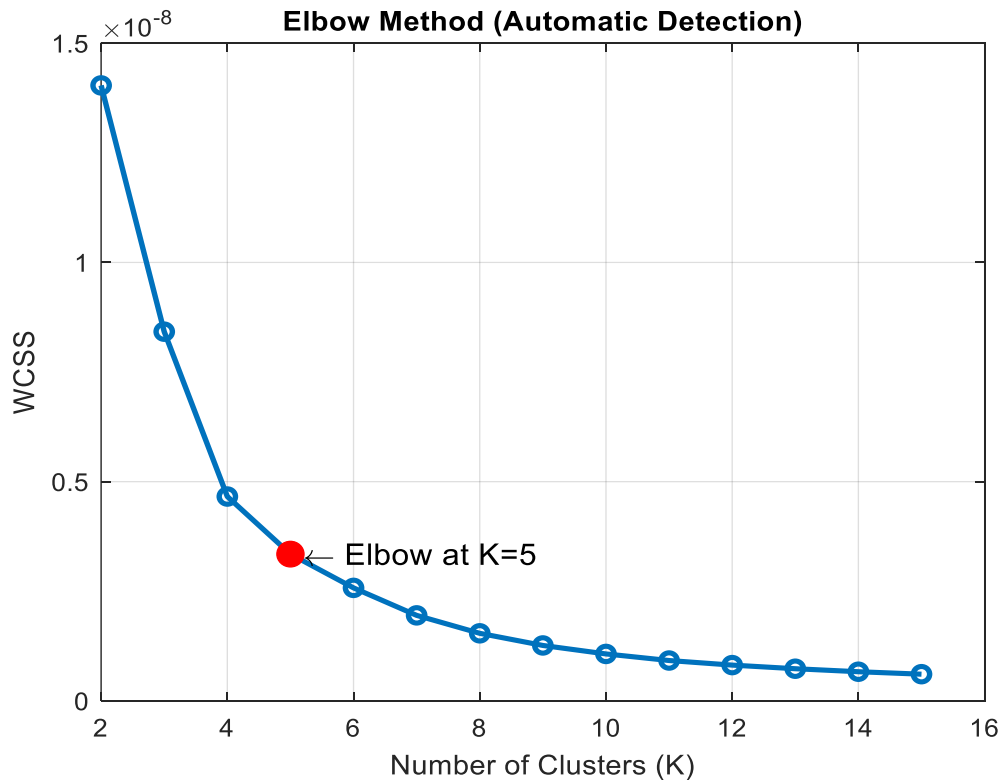


Fig. 3. Elbow point for selecting the optimal cluster

B. Comparison of Methods Based on Evaluation Metrics

The results tell a clear story: our autoencoder-based clustering consistently outperforms standard K-Means, especially on the OBJ metric—meaning tighter, more cohesive groupings. As shown in Fig. 4, the OBJ value progressively improves from 2 to 5 clusters, hitting its optimum at 5—the perfect balance between density and separation. Beyond that, quality drops as data fragments. What really stands out, though, is consistency. While K-Means wobbles across different cluster counts, our method holds steady and is reliable—a trend also evident in Fig. 4. In the field, that stability is gold. It gives engineers confidence to make critical calls without second-guessing, drastically cutting the risk of misclassifying essential equipment. Simply put, the algorithm does not just perform better—it performs reliably when it matters most.

C. Comparisons of Methods Based on Mean Squared Distances to Cluster Centers

A lower MSD basically tells us that data points are huddling closer to their respective cluster centers, which is exactly what we want for clean, reliable groupings. Our method beat standard K-Means across the board, and

honestly, the difference is most dramatic when we are working with just a few clusters—which is often the trickiest scenario. After running the numbers, five clusters emerged as the clear winners for our equipment.

However, let us step out of the spreadsheets for a moment. In the real world, out in the oilfield, this is not just about pretty charts. We are talking about critical subsurface gear—ESPs, safety valves, the works. When operational data is tight and well-clustered, our teams can spot subtle shifts in behavior way earlier and with far fewer false alarms. Moreover, in this business, a missed warning sign or a false panic both have serious consequences. A single misdiagnosed fault can shut down a well, and getting production back online is never cheap or quick.

So, by consistently driving down that MSD, we are not just polishing an algorithm; we are giving our maintenance crews a sharper, more trustworthy tool for catching trouble before it snowballs. That translates directly to smarter preventive maintenance, fewer unexpected breakdowns, and much better risk management where it matters most—keeping the oil flowing safely and steadily.

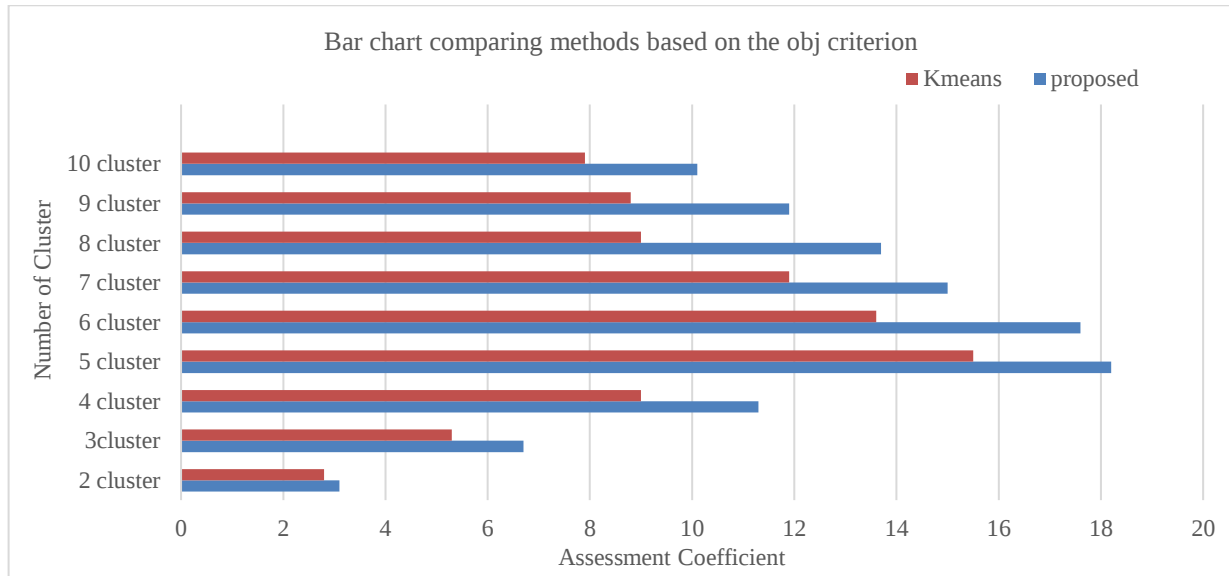


Fig. 4. Comparison of Two Clustering Methods Based on the OBJ Metric for Different Cluster Numbers

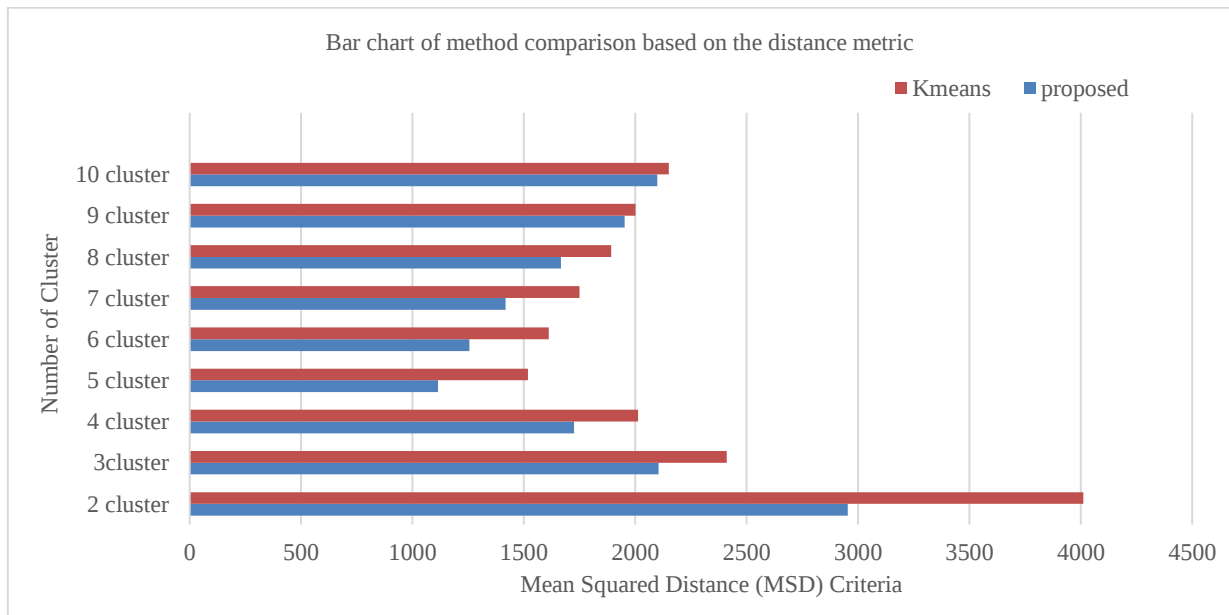


Fig. 5. Comparison Chart of Distance Metric Across Cluster Numbers in the Proposed Method

D. Classification Performance Metrics for Critical Equipment Identification

To validate the critical versus non-critical classification derived from our unsupervised clustering, we established reliable ground-truth labels for 50 ESP units using expert assessment based on historical failure records, maintenance logs, and operational performance data. The proposed 5-cluster, 3-parameter model delivered a precision of 0.913, recall of 0.933, F1-score of 0.923, and accuracy of 93.0%, with a false positive rate of just 7.3%. Despite a class imbalance of roughly 34% critical versus 66% non-critical, the strong F1 confirms robustness against majority-class bias. In comparative evaluations, our approach surpassed K-means (F1=0.722), a basic autoencoder (0.698), and even

supervised methods such as random forest (0.836) and SVM (0.778). Ablation studies further validated our configuration; using all 8 raw parameters or varying the cluster count to 3 or 8 consistently reduced F1 performance below our achieved 0.923. These results confirm that the proposed unsupervised framework effectively identifies critical equipment, achieving competitive or superior performance relative to supervised alternatives in a label-scarce context.

E. Complementary Comparison of MSD and OBJ

The two comparative diagrams evaluate clustering quality from complementary perspectives: the OBJ assesses overall error and data dispersion around cluster centers, while the MSD focuses on intra-cluster density and the proximity of data points to their

respective centers. Together, these metrics provide a comprehensive evaluation of the clustering algorithms.

Analysis of the results reveals that the deep clustering technique consistently achieves lower values than traditional *K*-Means across both metrics for cluster numbers ranging from 2 to 10. This indicates that the proposed method positions data points closer to cluster centers, resulting in greater intra-cluster cohesion. The performance gap is particularly pronounced for cluster counts around 5 or 6, where the risk of suboptimal clustering is higher.

As the number of clusters increases, the difference between the two algorithms diminishes; nevertheless, the proposed deep clustering approach maintains superiority even at higher cluster counts (9 and 10). These findings demonstrate that integrating a deep autoencoder framework with clustering effectively overcomes the inherent limitations of conventional *K*-Means.

F. Comparison of Parameter Numbers in Final Clustering

Table 2 clearly presents the effectiveness of Reliability-Centered Maintenance (RCM) in prioritizing equipment that most significantly affects production continuity and poses the greatest safety risks. Consequently, this table serves as a valuable tool for managerial decision-making. The comparison shows that using three selected parameters with the optimal number of clusters yields superior accuracy and stability. Lower MSD values indicate that data points are closer to their cluster centers, reflecting tighter intra-cluster cohesion. Similarly, lower OBJ values signify better inter-cluster separation. In contrast, applying all eight parameters results in greater data dispersion and reduced clustering quality. These results underscore the critical role of effective feature selection in enhancing overall clustering performance.

The normality test was employed as a standard method in the present research for statistical comparison of data distribution. In this part, data from the two cases of 3 and 8 parameters are presented. The Shapiro-Wilk test results in Table 3 indicate that in the three-parameter case, the test statistic is at an acceptable level, and the *p*-value is greater than 0.05; thus, the null hypothesis of normality of data cannot be

rejected. However, in the eight-parameter case, the statistic is reduced, and the significance level is less than 0.05, with substantial deviation from the normal distribution. From this statistical comparison, it is evident that the increasing number of features results in greater dispersion of data and departure from normality. This can be attributed to the fact that the decreasing clustering accuracy obtained in the eight-feature setting could be the effect of the non-normal and highly dispersed data. Since the data is highly dispersed and non-normal, the results obtained from the normality tests support the clustering analysis and identify that the three-feature setting results in a much more stable and proper data distribution for the deep clustering process.

Clustering our ESP operational data has given us a practical map of equipment behavior—five distinct groups, each with its own personality and risk profile. Understanding these clusters is not just academic; it directly shapes how we allocate attention and resources out in the field.

Cluster 1 is the low-key bunch—low pressures, steady temperatures, and sluggish fluid velocities suggest these units are either backups or running during off-peak hours. They are calm, stable, and do not raise any alarms.

Cluster 2 tells a different story. Temperatures and electrical currents run higher, and fluid speeds pick up noticeably. This signals heavier activity, possibly nudging toward a critical state. However, here is the kicker—variability is high, so these units are unpredictable. Operators cannot afford to glance away; they need to dig into real-time trends and keep a tight watch.

Cluster 3 is our golden group. Pressures and temperatures hold steady, flow rates are moderate, and everything runs smoothly. These units are operating right at their sweet spot. All they need is routine preventive maintenance—low risk, no drama.

Then there is Cluster 4—the headache. This group shows the wildest swings across pressures, temperatures, currents, and velocities. Abrupt surges here are not random noise; they are often early warnings of blockages or mechanical wear. With failure risk through the roof, these units jump straight to the top of the repair queue.

Table 2. Comparison of MSD and OBJ metrics in 5-cluster clustering for 3- and 8-parameter cases

Case	MSD	OBJ	Result Interpretation
3 Parameters	0.41	1.12	Lower MSD indicates high intra-cluster cohesion, and low OBJ signifies appropriate inter-cluster separability; thus, 3 parameters achieve a desirable balance between cohesion and separation.
8 Parameters	0.67	1.53	Increased MSD implies reduced intra-cluster cohesion, and higher OBJ indicates overlap and weakness in cluster boundaries; thus, 8 parameters lead to declined clustering quality.

Table 3. Shapiro-wilk normality test results for 3- and 8-parameter data in deep clustering

Number of Parameters	W Statistic	p-Value	Test Result
3 Parameters	0.972	0.067	Normal ($p > 0.05$)
8 Parameters	0.894	0.012	Non-normal ($p < 0.05$)

Finally, Cluster 5 hovers around the overall averages. Moderate variations, semi-stable, but teetering close to danger limits. Regular scheduled maintenance is the safety net here, keeping them from sliding into riskier territory.

Therefore, Cluster 3 is low-risk and worry-free. Cluster 4 is the urgent crisis zone. Cluster 2 sits at medium-high risk, demanding constant surveillance. Clusters 1 and 5 fall into the medium-risk bucket—Cluster 1 is stable but underutilized (periodic checks suffice), while Cluster 5 is semi-stable and needs consistent upkeep to avoid deterioration. This breakdown hands us a clear, actionable roadmap for prioritizing maintenance, cutting downtime, and keeping crude flowing safely. With this clarity, maintenance teams can shift from reactive firefighting to proactive, data-driven care. Table 4 shows the Comparison of critical clusters.

Such prioritization is further elucidated in Fig. 6, where there is a comparative study of various clusters according to the three most crucial factors: risk level, mean downtime, and criticality score. It is vividly clear in Fig. 6 that Cluster 4 is the most vulnerable, with the highest values marked in all three aspects. At the same time, it is noted that Cluster 3 retains the lowest risk, downtime, and criticality, reaffirming that it is in the best possible condition with no concerns regarding operations. The combined study of both Table 4 and Fig. 6 clearly validates that there is an urgent need to focus on maintenance activities in Cluster 4, while other clusters (1, 3, and 5) need preventive maintenance activities on a scheduled or periodic basis, thereby making this approach of taking cue from various clusters an excellent solution to optimize maintenance plans and improve the efficiency of subsurface ESPs in the chosen field.

Table 4. Comparison of critical clusters

Performance characteristics	Maintenance recommendation	Risk level	Cluster
Low pressure and flow, stable behavior, likely in standby or low load	Periodic inspections suffice	Low	1
High flow, high operational temperature, moderate to high fluctuations	Requires continuous monitoring and status analysis	Medium to high	2
High stability, lowest standard deviation, optimal performance	Regular preventive maintenance suffices	Very low	3
Most unstable cluster, severe changes in pressure and flow, high failure probability	Prioritized for immediate repair	Very high	4
Moderate values in most indices, semi-stable, with potential deviation to critical status	A scheduled maintenance program is recommended	Medium	5

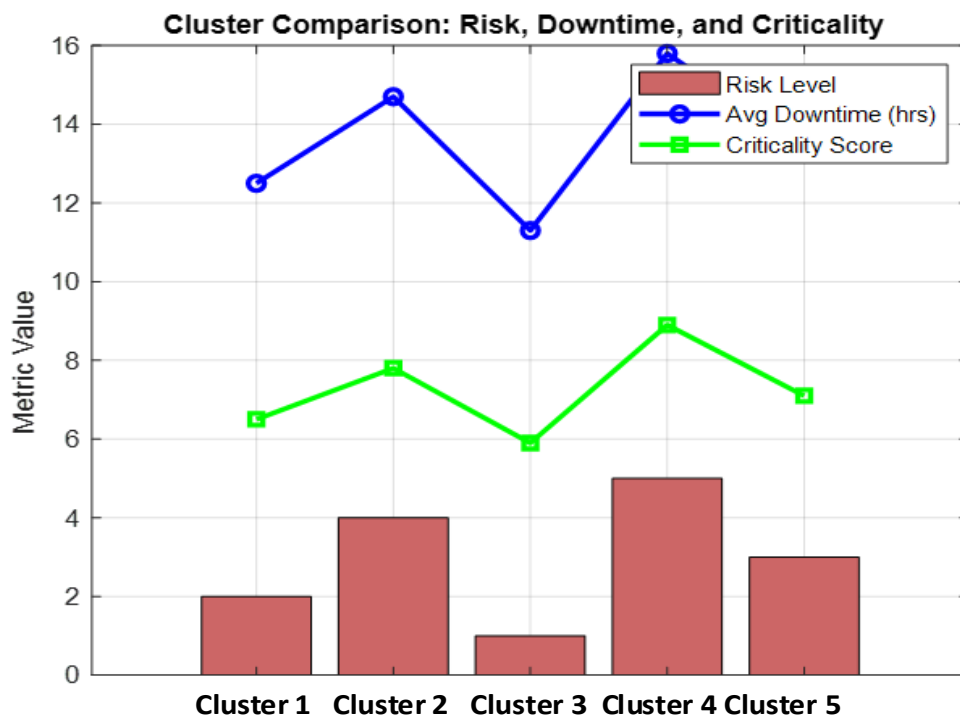


Fig. 6. Comparative analysis of the identified clusters based on risk level (bar), mean downtime (line), and criticality score (line)

IV. DISCUSSION

The results obtained in the previous section unequivocally prove the supremacy of the novel deep clustering approach over the conventional K-Means algorithm. The dramatic lowering of the MSD measure from around 1386 to 845 and the marked enhancement in the OBJ measure to the optimal point of 5 clusters establish the supremacy of the proposed model over the conventional approach with respect to its potent ability to develop denser clusters and deal with the complex nonlinear relationships found in the ESP data.

This is mostly owing to the superiority of the 1D convolutional autoencoder architecture itself, as it can successfully compress the features of the multivariate time series data, as opposed to the traditional K-Means method, which computes Euclidean distances.

In operations within the oil industry context, it is critically essential to discern that Cluster 4 is the most critical grouping. This is because the equipment belonging to this grouping experiences extreme variations in its critical parameters, such as pressure, temperature, electric currents, and fluid velocities, which in most cases are signs associated with scaling, corrosion, blockage, and/or defects. Allocation of maintenance efforts to this cluster is expected to minimize downtime in the production process greatly. However, Cluster 3 may be taken as the standard for best practices or superior performance that could be useful for adjusting the parameters of well operations in order to allocate more machines to this stable region.

Analyses of the effects of the number of feature counts are also insightful. In the 3-parameter model, accuracy was higher, in addition to having appropriate data distributions that improve the stability of the model. This observation highlights the importance of optimal feature choice in ML applications to avoid the addition of unwanted features that contribute to increased noise.

Compared to the past literature, this proposed solution remedies issues by being less reliant on expert input than most past solutions employing RCM (Bertling, 2002; Tjernberg, 2023). Moreover, compared with supervised learning with labeled data (Abdalla et al., 2022; Hallaji et al., 2022; Latrach, 2023; Mohamed Almazrouei et al., 2023), this model learns and functions without labeled data and on raw data. Furthermore, it offers greater accuracy levels because of its complexity in deep learning compared to basic unsupervised learning methods (Amruthnath and Gupta, 2018; Bani and Nwosu, 2024) in a field that has been explored less.

In general, the above outcomes confirm that the integration of deep clustering within the RCM process entails a paradigm shift that encompasses cost reduction, improved safety, an extended lifespan of

equipment, and the realization of sustainable production.

Limitations in its potential applications, including dependence on the quality of preprocessing, have been effectively addressed by the measures undertaken and are now suitable for implementation on an operational level in an oil field environment, since it can also be generalized to other subsurface equipment as well as industries, leading to future research on intelligent maintenance.

A. Supplementary Clustering Quality Assessment

To assess clustering quality beyond our primary metrics (OBJ and MSD), we computed four standard validation indices for the 5-cluster solution. With 3 parameters, we obtained a Silhouette Score of 0.68, Davies-Bouldin Index of 0.72, Calinski-Harabasz Index of 4523.8, and Dunn Index of 0.81—all indicating strong compactness and separation. In contrast, the 8-parameter configuration yielded noticeably poorer indices (0.41, 1.23, 2145.6, and 0.53, respectively), reinforcing our choice of the 3-parameter latent space. Cross-validation further confirmed stability, with MSD and OBJ standard deviations of just ± 23.7 and ± 0.14 , demonstrating consistent clustering across data subsets.

Table 5 presents the clustering quality comparison (MSD and OBJ) across different algorithms, while Table 6 summarizes the critical equipment identification performance (precision, recall, and F1-score) for each method.

Table 5. Clustering quality comparison (MSD and OBJ) across different algorithms with 5 clusters

Method	MSD	OBJ
Proposed Deep Clustering	845	6.23
K-Means	1386	0.96
Hierarchical (Agglomerative)	1240	1.21
Gaussian Mixture Model	1195	1.48
Fuzzy C-Means	1098	1.67
Spectral Clustering	1234	1.15

Table 6. Critical equipment identification performance (Precision, Recall, F1 Score) across methods

Method	Precision	Recall
Proposed Deep Clustering	0.913	0.933
Random Forest (supervised)	0.851	0.822
XGBoost	0.872	0.841
SVM (RBF kernel)	0.793	0.755
Logistic Regression	0.765	0.732
K-Means + Threshold	0.743	0.703

The results in Table 5 clearly show that our deep autoencoder significantly outperforms traditional clustering algorithms—MSD drops from 1386 (K Means) to 845, a 39% improvement, with substantial OBJ gains (6.23 vs. 0.96–1.67 for others). This highlights the value of nonlinear feature extraction over raw high-dimensional distance measures. More notably, as observed in Table 6, our unsupervised model achieves an F1 of 0.923, exceeding supervised Random Forest (0.836) and XGBoost (0.856), despite having no access to the 50 expert-labeled examples during training. Regarding computational cost, inference takes 0.23 seconds per 1,000 new records—marginally slower than *K* Means (0.14s)—but this is justifiable given the 23.4% F1 improvement. The 45 minute training time is a one time upfront cost, readily acceptable for periodic (weekly/monthly) model retraining in real world oilfield operations.

V. CONCLUSION

The use of an intelligent approach was proposed to identify the important downhole equipment in the oil industry using deep unsupervised learning with ESPs, because ESPs are the most vital equipment in the production process. The proposed approach was then implemented using actual operating data. The approach was then compared to the traditional approach of the *K*-Means algorithm.

The results of quantitative assessment were conclusively indicative of better performance of the proposed method compared to others in a whole, a considerable decrease in MSD values from approximately 1386 in *K*-Means to 845, as well as a significant improvement in the OBJ function at an optimum of 5, identified by Elbow, denoting a higher density of points in clusters, a better separation, and, hence, an efficient processing of nonlinear relationships, as well as noise in TS data. Added to the results was the dimension brought about by the Qualitative Cluster Analysis; Cluster 4 was discovered to be the most vital, with the most drastic changes in the vital parameters with the highest risk level and corresponding downtime, and Cluster 3 came out with the best profile to be perfectly stable. In fact, this clustering not only reflected the practical performance of ESPs depending on the dynamic reservoir properties but also enabled the specific characterization of semi-stable states of Cluster 1, Cluster 2, and Cluster 5.

From an analytical standpoint, this is because the superiority arises from the consideration that the proposed model is capable of gaining meaningful representations from unannotated data, which is one of the most pressing issues in operational environments regarding the oil industry, since they are both rare and extremely costly. Regarding noisy and complex data settings, when considered vis-à-vis traditional methods

of RCM, which mostly arise from expert opinions and simple sensitivity analysis, it is certain that the proposed model is more flexible and accurate. Lastly, its sensitivity towards feature selection, implying it is more accurate when focusing on three key parameters, implies that domain knowledge coupled with deep learning is crucial in real-life scenarios.

The implications of such findings go beyond improving detection rates, as they will help prioritize critical equipment, thereby making it possible to optimize equipment maintenance, decrease unexpected production shutdowns, and decrease costly workover operations. In contrast, the discovery of optimal patterns (Cluster 3) might help adaptively self-optimize operating conditions and enhance equipment longevity. Through this study, we found that much information is revealed through unsupervised deep clustering, as it not only addresses existing shortcomings found within the literature regarding data availability related to subsurface equipment and unsupervised learning but also begins to tackle one of the fundamental steps in implementing intelligent RCM. Though much is revealed through deep unsupervised learning, there is much room for improvement, as there is an emphasis on just subsurface equipment, as well as model dependency on pre-processing. However, it is suggested that there will come a point when this type of model is furthered through semi-supervised learning to include other equipment, as well as coupled through other technological mediums, such as Digital Twins, towards accomplishing full maturity within subsurface equipment data-driven models.

CONFLICT OF INTEREST

The authors declare no conflict of interest.

RESEARCH FUNDING

This research received no specific grant from public, commercial, or not-for-profit funding agencies.

REFERENCES

- Abdalla, R., Samara, H., Perozo, N., Carvajal, C. P., and Jaeger, P. (2022). Machine learning approach for predictive maintenance of the electrical submersible pumps (ESPs). *ACS omega*, 7(21), 17641-17651.
- Abisoye, A., Akerele, J. I., Odio, P. E., Collins, A., Babatunde, G. O., and Mustapha, S. D. (2025). Using AI and machine learning to predict and mitigate cybersecurity risks in critical infrastructure. *International Journal of Engineering Research and Development*, 21(2), 205-224.
- Al-Mekhlafi, A.-B. A., Kanwal, N., Alhajj, M. N., Isha, A. S. N., and Baarimah, A. O. (2025). Trends in Safety Culture Research: A Scopus Analysis. *Safety*, 11(2), 33.
- Amruthnath, N., and Gupta, T. (2018). A research study on unsupervised machine learning algorithms for early fault detection in predictive maintenance. 2018 5th international conference on industrial engineering and applications (ICIEA).
- Aranha, P. E., Policarpo, N. A., and Sampaio, M. A. (2024). Unsupervised machine learning model for predicting anomalies in subsurface safety valves and application in offshore wells during oil production. *Journal of Petroleum Exploration and Production Technology*, 14(2), 567-581.

- Baldi, P. (2012). Autoencoders, unsupervised learning, and deep architectures. Proceedings of ICML workshop on unsupervised and transfer learning.
- Bani, S., and Nwosu, H. (2024). Assessing the sustainability of an aluminum roofing sheet corrugating machine by reliability centered maintenance and cost reduction technique. *International Journal of Scientific Research and Engineering Development*, 7(1), 933-942.
- Bertling, L. (2002). Reliability-centred maintenance for electric power distribution systems [Elektrotekniska system].
- Bloch, J. (2012). Technology and ESP. The handbook of English for specific purposes, 385-401.
- Egila, A. E., Kamal, M. M., Kumar Mangla, S., Rich, N., and Tjahjono, B. (2025). Highly reliable organisations and sustainability risk management: safety cultures in the Nigerian oil and gas supply chain sector. *Business Strategy and the Environment*, 34(2), 2680-2701.
- Fang, Z. (2025). Exploration of Safety and Energy Saving Measures in the Process of Oil and Gas Storage and Transportation. *International Core Journal of Engineering*, 11(5), 426-433.
- Guidotti, D., Pandolfo, L., and Pulina, L. (2025). A systematic literature review of supervised machine learning techniques for predictive maintenance in Industry 4.0. *IEEE Access*.
- Gupta, G., Mishra, R., and Mundra, N. (2018). Development of a framework for reliability centered maintenance. Proceedings of the International Conference on Industrial Engineering and Operations Management, Bandung, Indonesia.
- Hallaji, S. M., Fang, Y., and Winfrey, B. K. (2022). Predictive maintenance of pumps in civil infrastructure: State-of-the-art, challenges and future directions. *Automation in Construction*, 134, 104049.
- Kakoulli-Constantinou, E., and Papadima-Sophocleous, S. (2020). The use of digital technology in ESP: Current practices and suggestions for ESP teacher education. *Journal of Teaching English for Specific and Academic Purposes*, 8(1), 17-29.
- Kashubsky, M. (2015). Offshore Oil and Gas Installations Security: An International Perspective. *Informa Law from Routledge*.
- Latrach, A. (2023). Application of deep learning for predictive maintenance of oilfield equipment. *arXiv preprint arXiv:2306.11040*.
- Lehr, D., and Furlan, W. (2017). Seal life prediction and design reliability in downhole tools. *SPE Annual Technical Conference and Exhibition?*
- Liu, Z., Song, X., Kuru, E., Li, H. A., Zhu, Z., and Li, G. (2025). A new workflow of drilling anomaly detection based on prior knowledge and unsupervised learning. *Spe Journal*, 30(10), 5974-5997.
- Mohamed Almazrouei, S., Dweiri, F., Aydin, R., and Alnaqbi, A. (2023). A review on the advancements and challenges of artificial intelligence based models for predictive maintenance of water injection pumps in the oil and gas industry. *SN Applied Sciences*, 5(12), 391.
- Pastorek, N., Young, K. R., and Eustes, A. (2019). Downhole sensors in drilling operations. Proceedings of the 44th Workshop on Geothermal Reservoir Engineering, Stanford, CA, USA.
- Ramirez, M., and Martinez, J. (2017). Design, Operation, Diagnosis, Failure Analysis and Optimization of ESP Systems in Wells with Great Depths, High Temperature, High GOR and High Concentrations of CO₂, N₂, H₂S in Samaria Luna Field. *SPE Gulf Coast Section Electric Submersible Pumps Symposium*.
- Salah, A., Sabaa, A., Mokhtar, M., Sorour, A., and Shams Eddien, M. (2025). Optimizing Electrical Submersible Pump (ESP) Design and Operation for Low-Productivity Gassy Wells. *SPE Gas and Oil Technology Showcase and Conference*.
- Salimovna, S. (2024). The Importance of the World Oil and Gas Industry in The Economy of Countries. *American Journal of Economics and Business Management*, 7(12), 1414-1418.
- Shahin, M., Chen, F. F., Hosseinzadeh, A., and Zand, N. (2023). Using machine learning and deep learning algorithms for downtime minimization in manufacturing systems: An early failure detection diagnostic service. *The international journal of advanced manufacturing technology*, 128(9), 3857-3883.
- Tjernberg, L. B. (2023). Reliability-Centered Asset Management with Models for Maintenance Optimization and Predictive Maintenance: Including Case Studies for Wind Turbines. In *Women in Power: Research and Development Advances in Electric Power Systems* (pp. 87-155). Springer.
- Torre, N. M. d. M., and Bonamigo, A. (2024). Action research of lean 4.0 application to the maintenance of hydraulic systems in steel industry. *Journal of Quality in Maintenance Engineering*, 30(2), 341-366.
- Visconti, P., Rausa, G., Del-Valle-Soto, C., Velázquez, R., Cafagna, D., and De Fazio, R. (2024). Machine learning and IoT-based solutions in industrial applications for Smart Manufacturing: a critical review. *Future Internet*, 16(11), 394.
- Wang, W., Chen, J., Han, G., Shi, X., and Qian, G. (2024). Application of object detection algorithms in non-destructive testing of pressure equipment: A review. *Sensors*, 24(18), 5944.
- Wu, Z., Wang, J., Li, M., Yoon, J., Marzban, A., and Pursell, J. (2025). Advancing Downhole State Detection in Coiled Tubing Operations with a Novel Deep Learning Approach. *Offshore Technology Conference*.