

Robust Fault Detection with One-Class Training

Ark Ifeanyi¹, and Jamie Coble²

^{1,2} *University of Tennessee-Knoxville, Knoxville, TN, 37996, USA*

aifeanyi@vols.utk.edu

jamie@utk.edu

ABSTRACT

Fault detection is essential for maintaining the safety and reliability of industrial systems, particularly in critical infrastructure such as nuclear power plants, where undetected anomalies can lead to costly downtime or safety risks. This paper evaluates three one-class learning methods—Principal Component Analysis (PCA), Autoencoder (AE), and Deep Centered Embedding (DCE)—for detecting pump screen fouling in the circulating water system of the Hope Creek nuclear power plant. Using real sensor data from two distinct operational periods (2014 and 2018), we assess each model’s ability to identify faults in the presence of significant distributional changes. While all methods performed well on in-sample data, only the DCE model successfully generalized to previously unseen operational conditions. Feature importance analyses further revealed alignment in key fault-related signals across models, while DCE’s treatment of less salient features contributed to its generalization strength. These results demonstrate the potential of embedding-based one-class models for robust, data-efficient fault detection in safety-critical environments. Future work will explore domain adaptation strategies to enhance model resilience under changing operational regimes.

1. INTRODUCTION

Fault detection plays a critical role in maintaining the reliability, efficiency, and safety of industrial systems, particularly in the energy sector, where unplanned downtime or undetected faults can result in substantial economic loss, environmental hazards, and safety risks. From nuclear power plants to renewable energy facilities, real-time monitoring systems are increasingly tasked with identifying deviations from normal operating behavior before they escalate into failures. This underscores the growing need for robust and generalizable anomaly detection frameworks that can

operate effectively under varying conditions and system configurations.

A key challenge in deploying such systems lies in the scarcity of labeled fault data. In many industrial settings, faults are rare, diverse, and context-dependent, making it difficult to gather sufficient examples to train supervised machine learning models. To address this, anomaly detection approaches that rely solely on healthy data, commonly referred to as one-class training, have gained significant attention (Xu et al., 2024). These methods learn the distribution of normal behavior and subsequently flag deviations as potential anomalies. This approach is not only practical but also better aligned with the realities of industrial monitoring, where capturing the full spectrum of potential faults during training is infeasible.

This paper explores the performance and generalization capability of three data-driven one-class training methods for fault detection. The first is Principal Component Analysis (PCA), a linear dimensionality reduction method that captures dominant modes of variation in healthy data. The second is a classical Autoencoder (AE), which learns to reconstruct healthy input data and uses reconstruction error to signal anomalies. The third, referred to in this paper as Deep Centered Embedding (DCE), leverages a neural network-based encoder trained to map healthy data into a compact latent space, where the distance from a learned center serves as an anomaly score.

A central focus of this research is to evaluate how well these one-class models generalize to previously unseen distributions of healthy and faulty samples. In real-world deployments, industrial systems often undergo operating condition changes, upgrades, or environmental shifts. Therefore, it is essential not only to detect known anomalies but also to maintain high performance in the face of such distributional changes. By systematically testing each approach across both known and new operating scenarios, this study aims to assess their robustness and suitability for real-world applications.

In nuclear power plants with natural water sources, pump

Ark Ifeanyi et al. This is an open-access article distributed under the terms of the Creative Commons Attribution 3.0 United States License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.

screen fouling is a significant operational concern due to the intake of debris and particulate matter from the cooling water source. Fouling commonly arises from biological materials—including algae, seaweed, jellyfish, and fish—as well as organic debris such as leaves and twigs. These are especially prevalent in open-loop systems or those drawing from natural water bodies or open basins. In addition, inorganic materials like sediment, shells, and anthropogenic waste may enter the system, particularly following heavy rainfall or increased water turbulence. Internal sources, such as corrosion products from aging infrastructure, can also contribute to screen blockage. Over time, biofilm growth and mineral scaling may further reduce the effective open area of the screens (Venkatesan & Sriyutha Murthy, 2008).

Early detection of pump screen fouling is critical for maintaining system performance and reliability. As screens become obstructed, the reduced flow impairs the circulating water system's ability to effectively cool the condenser, leading to degraded thermal efficiency and increased turbine backpressure. This can result in reduced power output and potential derating of the plant (Agarwal et al., 2022, 2021). Furthermore, fouling increases the risk of pump cavitation due to inadequate suction flow, which can damage pump components and elevate maintenance costs. Excessive screen blockage also raises the hydraulic load on the pumps, accelerating wear and energy consumption. In severe cases, undetected fouling may trigger unplanned maintenance outages or violate environmental discharge constraints due to altered flow or temperature profiles at the outlet (Venkatesan & Sriyutha Murthy, 2008). Therefore, timely identification and mitigation of screen fouling are essential to ensure efficient, safe, and compliant operation of the circulating water system.

2. SYSTEM AND DATA

The Circulating Water System (CWS) plays a vital role in nuclear power plant operations, particularly as a non-safety-related system that supports thermal efficiency and environmental compliance. At facilities like Hope Creek, the CWS is essential for dissipating heat from the main steam turbine and its auxiliary systems (Agarwal et al., 2021). Its design focuses on maximizing the efficiency of the steam power cycle while ensuring that environmental impacts on the Delaware River remain within regulatory limits established by the state of New Jersey. Functionally, the system is responsible for filtering intake water before it flows through the condenser and for cooling the exhaust steam from the turbine.

The CWS is composed of several key components, including vertical motor-driven pumps (referred to as circulators), intake screens to remove debris and aquatic organisms, the main condenser's tube-side components, and support systems

for air removal, sampling, and screen washing. These elements are integrated with control instrumentation, piping, and valves that enable stable and efficient operation. At the Hope Creek nuclear power plant, which operates a single boiling water reactor, the CWS uses four circulators. Unlike typical river-fed systems, Hope Creek draws its cooling water from a cooling tower basin. As a result, it does not rely on traveling screens; instead, each pump is equipped with its own screen to prevent debris from entering the condenser. The design features a shared header that distributes water from the four pumps to six condenser waterboxes—two for each of the plant's three condensers (Agarwal et al., 2021). This arrangement, illustrated in Fig. 1, reflects a configuration optimized for both performance and environmental compliance.

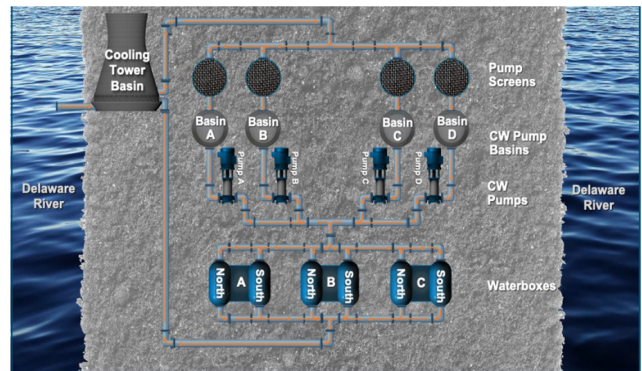


Figure 6. Schematic representation of the Hope Creek CWS.

Figure 1. The Hope Creek Circulating Water System

The Hope Creek CWS data comprises hourly measurements spanning from January 1, 2010, to May 18, 2021, but for this research, we extracted sample sequences of ten-day periods (240 hours) for the first five months of 2014 and 2018 only. This was because the exploration of the work order report of the plant revealed that these periods of these two years had the most occurrences of 'fouling' in 'pump screen A', the fault mode of interest (Agarwal et al., 2021). The ten days represent five days before and five days after the report of the maintenance activity to clear the fouling. In other words, the sequences are centered around the maintenance dates for the faulty samples. The healthy samples, on the other hand, are periods that do not have reports of any type of maintenance activity within the dates of interest. In 2018, there were 4 faulty samples and 18 healthy samples, whereas in 2014, there were 7 faulty samples and 23 healthy samples. The scarcity of the faulty samples for both years is typical of a nuclear power plant, driving the one-class training approach employed in this work.

Overall plant status data include gross load in megawatts and a 15-minute average of the ambient outside temperature. Each of the four circulators or circulating water pumps

(CWPs) measurements includes the ‘basin level’, ‘discharge pressure (DP)’, ‘motor winding temperature’, and ‘outbound and thrust bearing temperatures’. The temperatures at the north and south ends of each condenser were also recorded, with the total monitored signals being sixty-nine (69) for this research. This means each sample has 240 timesteps and 69 features (shape 240×69).

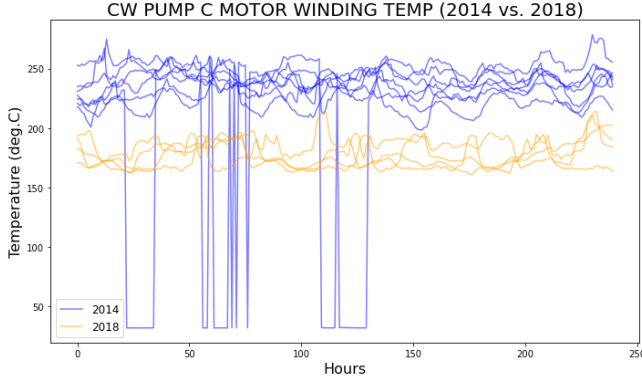


Figure 2. System Condition for the Different Years

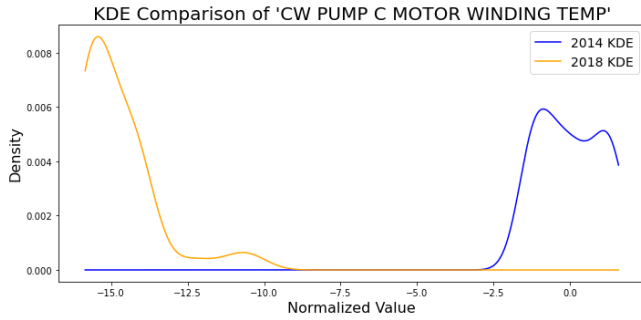


Figure 3. System Condition for the Different Years

However, variations in operating conditions across different years lead to changes in system response, even under similar fault modes. This is illustrated in Fig.2, which highlights differences in the winding temperature of pump C across several ten-day periods during similar seasons in different years. To characterize these differences more formally, kernel density estimation (KDE) is employed to estimate the probability density function (PDF) of the variable in a non-parametric manner (Chen, 2017). Fig.3 shows KDE representations using Gaussian kernels for one monitored system variable, further emphasizing the divergence in data distributions between the two years. These distributional shifts present challenges for developing detection models that generalize well to similar faults across multiple years.

3. METHODOLOGY

Fig. 4 highlights our overall approach, where we train different models to detect pump screen fouling faults by

processing the raw signals of the sensors in the plant. We train the model with healthy and faulty samples from a particular year and test the predictive performance of the trained models for the same year, along with their abilities to generalize to a different year. The three tested models are PCA, AE, and DCE as mentioned in section 1. We also attempted to explain the observed performances of the neural network-based approaches using different suitable techniques for each model. The AE reconstructs the multivariate timeseries as explained in section 3.2. Since the highly recognized explainability tool, SHAP, is not currently applicable to timeseries outputs, the highest contributors to the high reconstruction errors of the faulty samples can be regarded as the most useful features for detection (A. O. Ifeanyi et al., 2024). Shapley values calculated with SHAP (A. Ifeanyi et al., 2025; Ghorbani & Zou, 2019) are then used to rank feature importance for the DCE since its outputs are vectors of embeddings. The subsequent subsections dive into workings of the tested models.

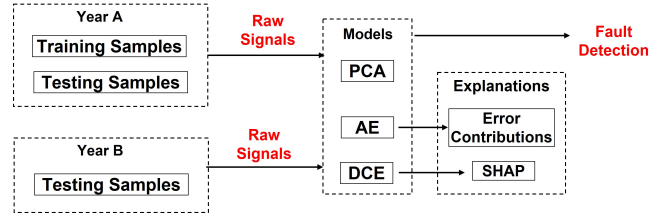


Figure 4. Summary of Approach

3.1. Principal Component Analysis (PCA)

Principal Component Analysis (PCA) is a linear dimensionality reduction technique that identifies the directions of maximum variance in high-dimensional data. In the context of one-class anomaly detection, PCA is trained using only healthy (non-faulty) data to learn a compact subspace. Anomalies are detected by measuring reconstruction error—how far a test sample deviates from this learned subspace (A. Ifeanyi et al., 2023).

3.1.1. Training Phase (One-Class Learning)

Let the input dataset be represented by the matrix $\mathbf{X} \in \mathbb{R}^{n \times d}$, where n is the number of healthy training samples and d is the number of features (e.g., flattened sensor sequences).

The data is mean-centered as:

$$\bar{\mathbf{X}} = \mathbf{X} - \mathbf{1}_n \boldsymbol{\mu}^\top, \quad (1)$$

where the mean vector $\boldsymbol{\mu}$ is given by:

$$\boldsymbol{\mu} = \frac{1}{n} \sum_{i=1}^n \mathbf{x}_i. \quad (2)$$

Next, PCA computes the empirical covariance matrix of the centered data:

$$\Sigma = \frac{1}{n} \bar{\mathbf{X}}^\top \bar{\mathbf{X}}. \quad (3)$$

We then compute the top- k eigenvectors of Σ , which form the matrix $\mathbf{V}_k \in \mathbb{R}^{d \times k}$. These eigenvectors span the principal subspace with number of principal components which retain 95% of the variance in the original data (Maćkiewicz & Ratajczak, 1993).

3.1.2. Reconstruction and Anomaly Scoring

A data point $\mathbf{x}_i \in \mathbb{R}^d$ is projected onto the principal subspace as:

$$\mathbf{z}_i = \mathbf{V}_k^\top (\mathbf{x}_i - \boldsymbol{\mu}), \quad (4)$$

and reconstructed as:

$$\hat{\mathbf{x}}_i = \mathbf{V}_k \mathbf{z}_i + \boldsymbol{\mu} = \mathbf{V}_k \mathbf{V}_k^\top (\mathbf{x}_i - \boldsymbol{\mu}) + \boldsymbol{\mu}. \quad (5)$$

The reconstruction error, used as the anomaly score, is computed as:

$$e_i = \|\mathbf{x}_i - \hat{\mathbf{x}}_i\|_2^2. \quad (6)$$

A high reconstruction error indicates that the test sample lies far from the healthy subspace and is likely anomalous (A. Ifeanyi et al., 2023).

3.1.3. Thresholding

To classify a test sample $\mathbf{x}_i$, the anomaly score e_i is compared against a threshold θ . Our approach was to set θ based on a high percentile (98%) of the reconstruction errors from healthy validation data:

$$\text{Anomaly}(\mathbf{x}_i) = \begin{cases} \text{normal}, & \text{if } e_i \leq \theta \\ \text{anomalous}, & \text{if } e_i > \theta \end{cases} \quad (7)$$

PCA-based reconstruction provides a lightweight, interpretable method for modeling healthy system behavior. Its primary limitation lies in its assumption of linearity, which may not capture nonlinear dependencies common in complex industrial time-series data. Nevertheless, it serves as a strong baseline for one-class training, especially in applications where fault examples are unavailable or rare. Since temporal dependencies are not captured by PCA, the following approaches were investigated.

3.2. 1D-CNN Autoencoder for One-Class Anomaly Detection

A 1D-Convolutional Neural Network (1D-CNN) based autoencoder is a powerful architecture for learning compressed representations of time-series data. In our context of one-class anomaly detection, the model is trained solely on healthy samples to reconstruct the input sequences as accurately as possible. The underlying assumption is that the model will struggle to reconstruct anomalous data, leading to larger reconstruction errors, which can be used as indicators of faults (A. O. Ifeanyi et al., 2024).

3.2.1. Model Architecture

The 1D-CNN autoencoder consists of two primary components:

- **Encoder:** Applies 1D convolutional layers to the input sequence to extract high-level temporal features and compress the sequence into a lower-dimensional latent representation.
- **Decoder:** Uses upsampling followed by convolution to reconstruct the original input from the latent representation.

Given an input sequence

$$\mathbf{X} = \{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_T\}, \quad \mathbf{x}_t \in \mathbb{R}^d,$$

the autoencoder learns a function $\mathcal{F}_\theta$ that attempts to reproduce the input sequence:

$$\hat{\mathbf{X}} = \mathcal{F}_\theta(\mathbf{X}), \quad (8)$$

where $\hat{\mathbf{X}}$ is the reconstructed sequence, and θ represents the learnable parameters of the encoder and decoder networks.

3.2.2. Training Objective

The model is trained using a reconstruction loss, typically the mean squared error (MSE) between the original input and the reconstruction:

$$\mathcal{L}_{\text{rec}} = \frac{1}{T} \sum_{t=1}^T \|\mathbf{x}_t - \hat{\mathbf{x}}_t\|^2. \quad (9)$$

The goal is to minimize $\mathcal{L}_{\text{rec}}$ over healthy data samples so that the model effectively captures the normal patterns present in the time-series.

3.2.3. Anomaly Detection

After training, the reconstruction error for unseen sequences is used as an anomaly score. Sequences with high reconstruction errors are flagged as potentially anomalous

(A. O. Ifeanyi et al., 2024). The anomaly score for a sequence $\mathbf{X}$ is computed as:

$$\text{Score}(\mathbf{X}) = \frac{1}{T} \sum_{t=1}^T \|\mathbf{x}_t - \hat{\mathbf{x}}_t\|^2. \quad (10)$$

A detection threshold τ is determined as $> 98\text{th percentile}$ using the distribution of scores on the healthy validation set. If $\text{Score}(\mathbf{X}) > \tau$, then $\mathbf{X}$ is considered anomalous.

This method is particularly effective in capturing spatial (or temporal) local patterns in time-series data due to the convolutional nature of the encoder. Moreover, its reliance only on healthy data makes it attractive for industrial applications where fault examples are rare or unavailable. The convolutional layers also enable efficient processing and robustness to local shifts and noise in the signal (A. O. Ifeanyi et al., 2024).

3.3. Deep Center Encoding (DCE) for One-Class Anomaly Detection

The Deep Center Encoding (DCE) method is a neural network-based approach for one-class anomaly detection that learns a compact representation of healthy data in a latent space. Unlike traditional reconstruction-based methods, DCE explicitly optimizes a loss function that minimizes the distance between the learned embeddings of healthy samples and a central reference point (or center) in the latent space. This method is inspired by the Deep Support Vector Data Description (Deep SVDD) paradigm (Kim et al., 2015; Han et al., 2025), but adapted to time-series data through 1-dimensional convolutional encoders.

3.3.1. Model Overview

Given an input sequence

$$\mathbf{X} = \{\mathbf{x}_1, \mathbf{x}_2, \dots, \mathbf{x}_T\}, \quad \mathbf{x}_t \in \mathbb{R}^d,$$

a neural network encoder $\mathcal{E}_\theta$ maps the sequence to a latent representation $\mathbf{z}$:

$$\mathbf{z} = \mathcal{E}_\theta(\mathbf{X}), \quad \mathbf{z} \in \mathbb{R}^k. \quad (11)$$

The goal of the DCE approach is to ensure that all embeddings $\mathbf{z}$ corresponding to healthy sequences lie close to a central point $\mathbf{c} \in \mathbb{R}^k$ in the latent space similar to the subspace of the PCA.

3.3.2. Center Initialization and Loss Function

To ensure stable training, the center $\mathbf{c}$ is computed after an initial warm-up phase. After training the encoder for a

few epochs, the embeddings of the healthy data are used to estimate the center:

$$\mathbf{c} = \frac{1}{N} \sum_{i=1}^N \mathcal{E}_\theta(\mathbf{X}^{(i)}). \quad (12)$$

This computed center is then fixed, and training continues using the loss in Equation (13).

The network is trained by minimizing the average squared Euclidean distance between the embeddings of the healthy samples and the center $\mathbf{c}$:

$$\mathcal{L}_{\text{DCE}} = \frac{1}{N} \sum_{i=1}^N \left\| \mathcal{E}_\theta(\mathbf{X}^{(i)}) - \mathbf{c} \right\|^2, \quad (13)$$

where $\mathbf{X}^{(i)}$ denotes the i -th healthy input sequence and N is the number of training samples. This loss encourages the encoder to map all healthy samples to a compact region around $\mathbf{c}$ in the latent space.

3.3.3. Anomaly Detection

At inference time, sequences are encoded using the trained encoder, and their distance to the center $\mathbf{c}$ is used as the anomaly score:

$$\text{Score}(\mathbf{X}) = \left\| \mathcal{E}_\theta(\mathbf{X}) - \mathbf{c} \right\|^2. \quad (14)$$

A threshold τ is set using the 98th percentile of the error from a validation set. A sample is flagged as anomalous if $\text{Score}(\mathbf{X}) > \tau$.

DCE combines the strengths of both PCA and AE. Like PCA, it maps raw input data into a lower-dimensional embedding space that captures the essential characteristics of the samples. Like AE, it can model complex nonlinear relationships, but it avoids the added complexity of a decoder by directly using the embedding distance from a learned center for anomaly detection. This makes DCE both expressive and computationally efficient.

4. RESULTS

4.1. Predictive Performance

This section presents the performances of the investigated models when tested on unseen samples from the same year of training. As seen in Fig. 5, Fig. 6, and Fig. 7, all models performed excellently when trained and tested on only 2014 and 2018 samples, respectively. When samples from both years are combined for training and testing of the models, PCA misclassified one sample in each class whereas the neural network-based models maintained their excellence.

This superiority of the neural network-based models is not unexpected as they have been shown to capture complex relationships better than PCA in different application areas including anomaly detection for nuclear systems.

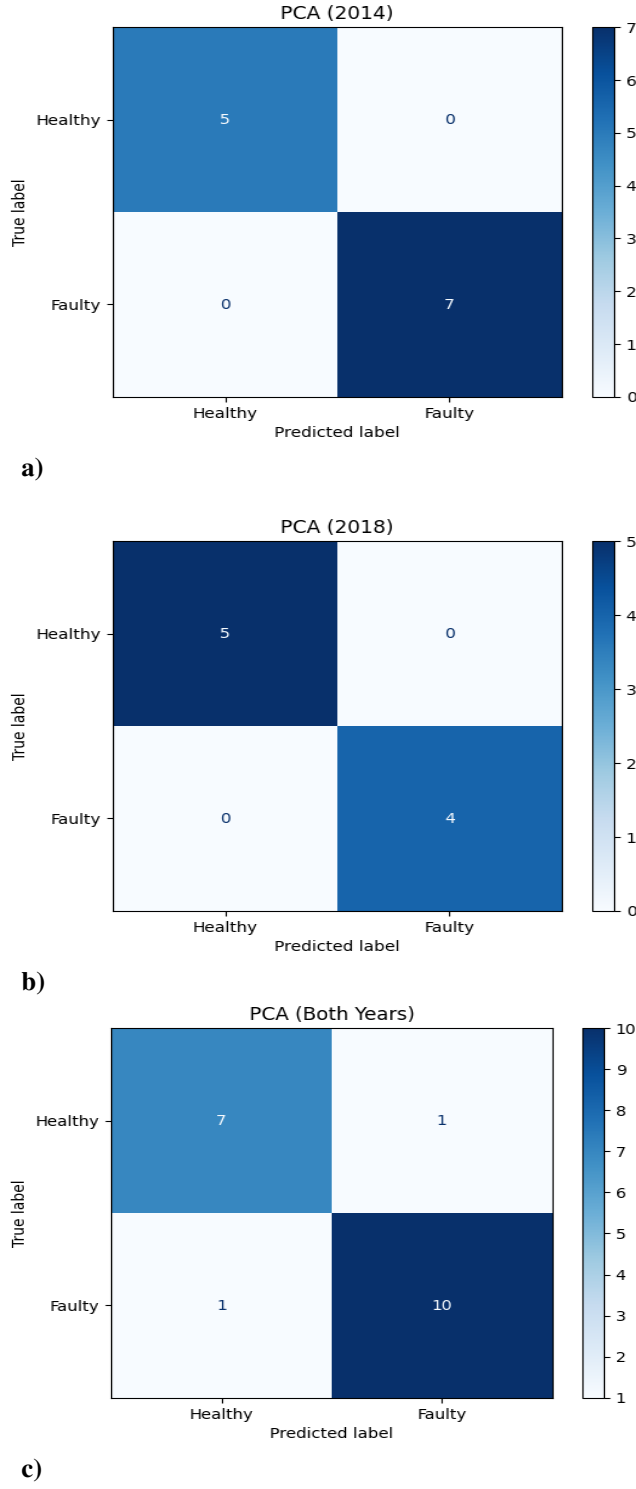


Figure 5. PCA Predictions. a) - 2014, b) - 2018, and c) - Both Years.

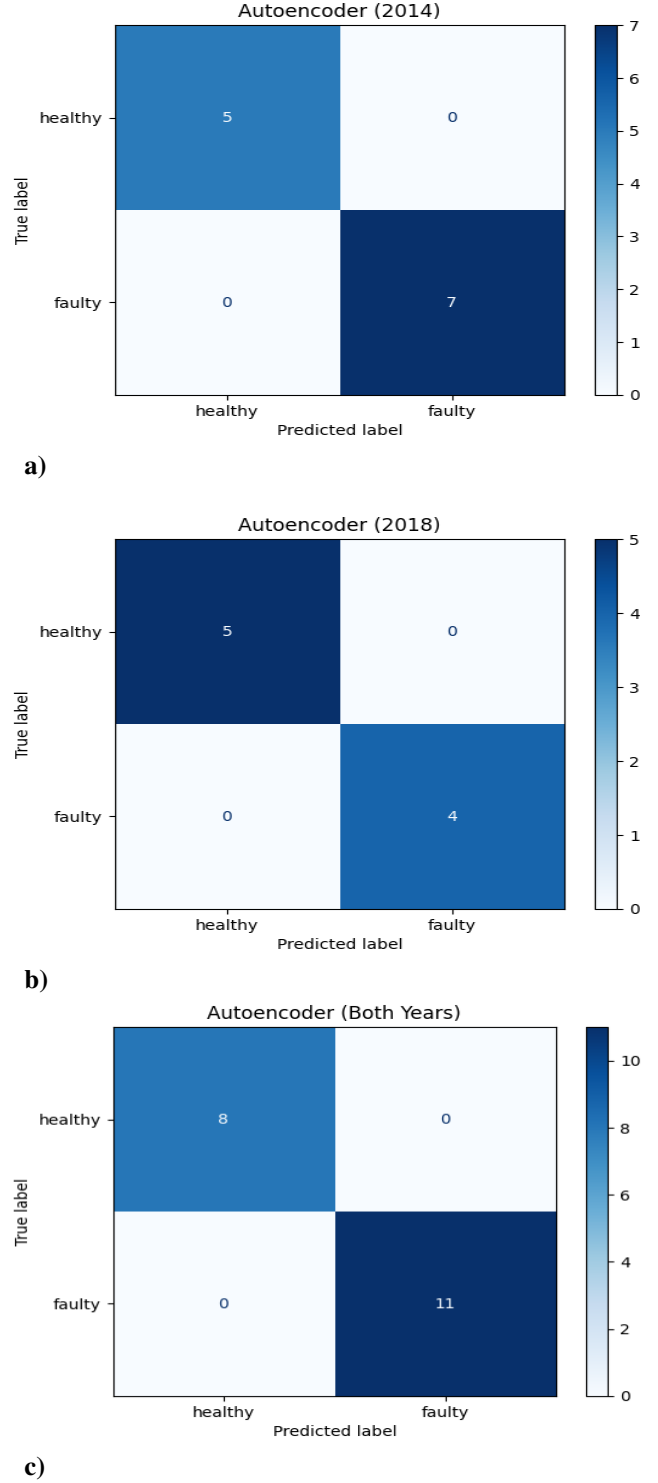
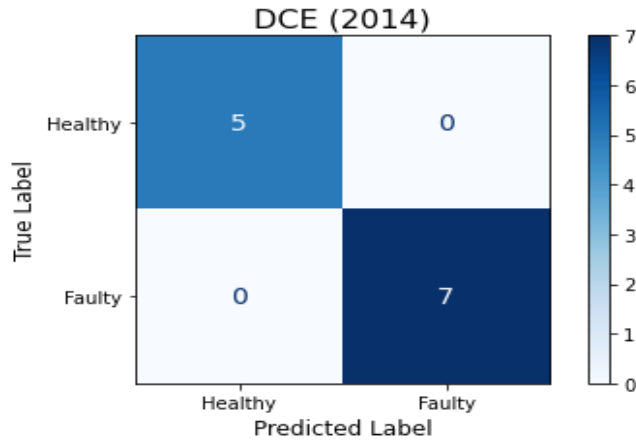


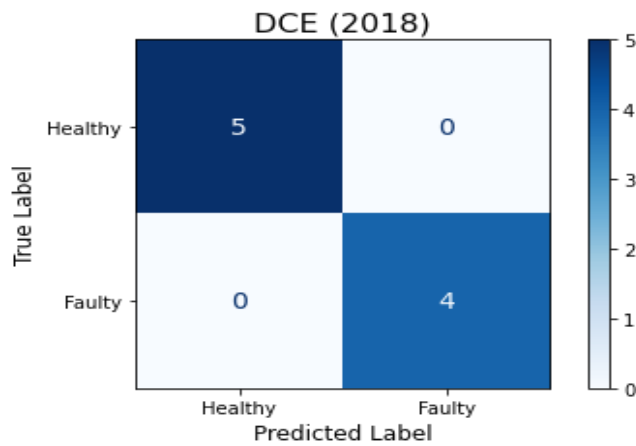
Figure 6. Autoencoder Predictions. a) - 2014, b) - 2018, and c) - Both Years.

4.2. Generalization

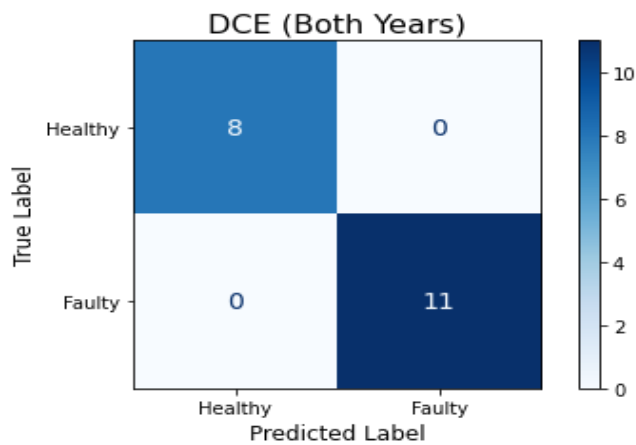
Here, we test how well the models can detect the same fault type in a similar season but in a different operational



a)



b)

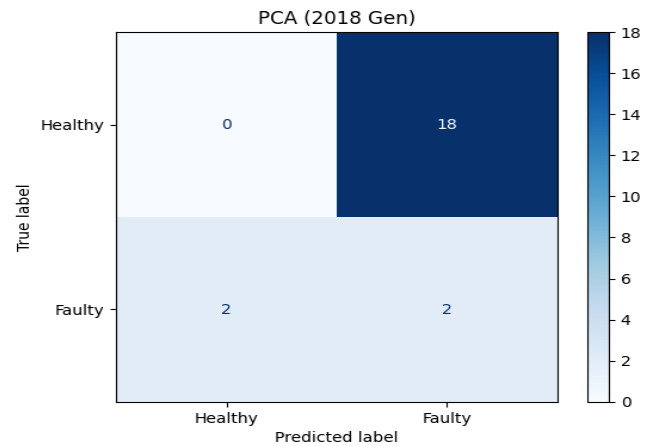


c)

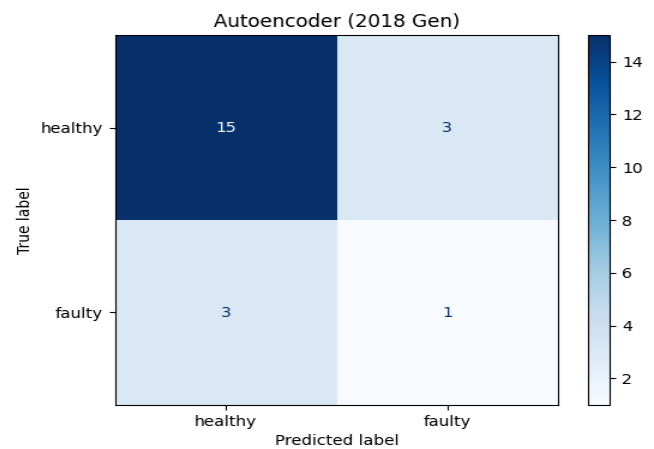
Figure 7. Deep-Centered Encoder Predictions. a) - 2014, b) - 2018, and c) - Both Years.

environment four years later. As seen in Fig. 8, DCE generalized perfectly after adjusting the threshold of the samples' distance from the center required for faulty

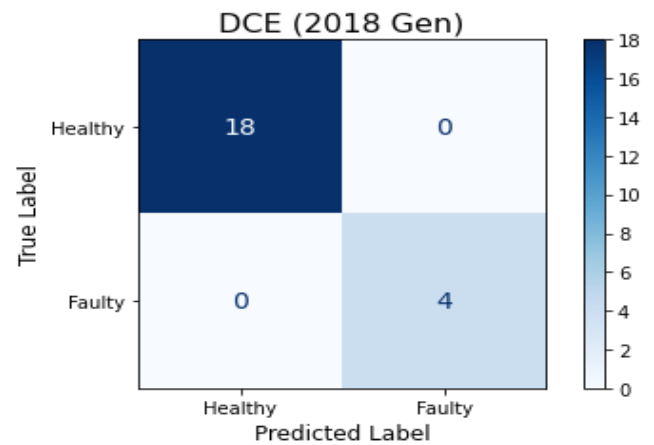
classification. The AE model correctly classified most of the healthy samples but struggled to detect the faulty ones. PCA was the worst overall performer for this test of generalization.



a)



b)



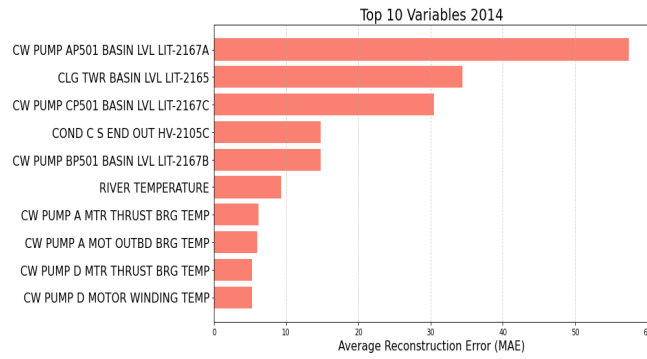
c)

Figure 8. Generalization Test. a) - PCA, b) - Autoencoder, and c) - Deep-Centered Encoder.

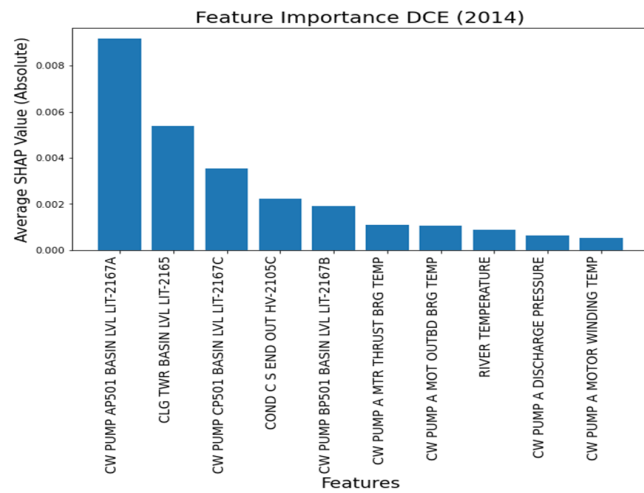
As can be deduced from the explanations in section 3.3 and the results in Fig. 8, DCE provides several potential advantages:

- **Compact Latent Representations:** By directly enforcing proximity to a center, DCE avoids the need for reconstruction, reducing overfitting to healthy patterns.
- **Robust to Unseen Faults:** Since it only learns the compact structure of normal data, it can generalize well to a wide range of faults without prior exposure.
- **Low Computational Overhead:** The approach requires only forward passes through the encoder and a simple distance computation for detection.

Overall, DCE is an efficient and potentially generalizable method for fault detection in settings where faulty examples are scarce or unavailable, which is often the case in critical industrial domains such as power generation and process monitoring.



a)



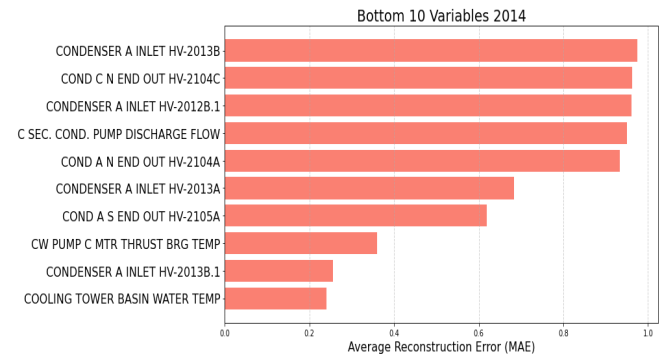
b)

Figure 9. Top Features. a) - Autoencoder, and b) - Deep-Centered Encoder.

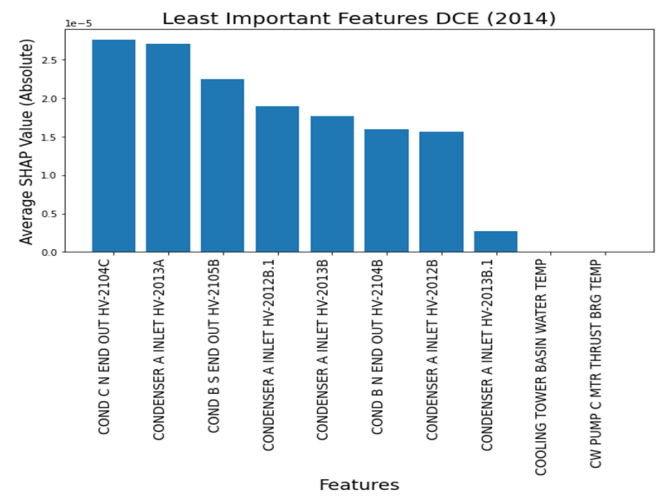
4.3. Explanations for Autoencoder and DCE

This section compares the performance of the AE and DCE models trained on 2014 data, as these models were used to evaluate generalization capability. Despite employing different explainability techniques, both models identified eight of the same features among their top ten most important, with complete agreement in the rankings of the top five features (see Fig. 9). This consistency suggests that these top eight features are likely central to fault detection within the 2014 dataset, particularly since both models demonstrated similar performance in this test scenario.

However, this also implies that the features enabling the DCE model to generalize to data from different distributions may lie beyond these top eight. The divergence between the models becomes more apparent when examining the least important features (see Fig. 10), where there is no overlap in feature rankings. This highlights a significant difference in how each model interprets and utilizes the remaining input features beyond those critical for in-sample fault detection.



a)



b)

Figure 10. Bottom Features. a) - Autoencoder, and b) - Deep-Centered Encoder.

5. CONCLUSION

This study evaluated the performance and generalization capability of three one-class fault detection models—PCA, Autoencoder (AE), and Deep Centered Embedding (DCE)—applied to pump screen fouling detection in a nuclear power plant’s circulating water system. Using real operational data from Hope Creek, we demonstrated that while all models performed well when tested on data from the same year as training, only DCE maintained high accuracy when generalizing to samples from a different year with distinct operating conditions. DCE’s ability to embed healthy samples in a compact latent space and classify faults based on distance from a learned center allowed it to overcome the challenges posed by distributional shifts across years, making it particularly suitable for fault detection in settings with limited or unavailable faulty data.

Analysis of feature importance further revealed that both AE and DCE models consistently identified the most critical features for fault detection within the training year, but diverged in how they treated less important features. This divergence likely contributed to DCE’s superior ability to generalize across different data distributions. These findings highlight the advantages of embedding-based one-class models like DCE in industrial environments where labeled fault data are scarce or evolving conditions are common.

As a direction for future work, we propose integrating domain adaptation techniques or adversarial training to further enhance the model’s robustness to distributional shifts caused by changes in system configuration, environmental conditions, or sensor calibration. Such methods could improve performance in long-term deployments by enabling the model to adapt continuously to new operating regimes without requiring labeled fault data.

ACKNOWLEDGMENT

This material is based upon work supported by the Department of Energy under Award Number DE-NE0009278.

This report was prepared as an account of work sponsored by an agency of the United States Government. Neither the United States Government nor any agency thereof, nor any of their employees, makes any warranty, express or implied, or assumes any legal liability or responsibility for the accuracy, completeness, or usefulness of any information, apparatus, product, or process disclosed, or represents that its use would not infringe privately owned rights. Reference herein to any specific commercial product, process, or service by trade name, trademark, manufacturer, or otherwise does not necessarily constitute or imply its endorsement, recommendation, or favoring by the United States Government or any agency thereof. The views and

opinions of authors expressed herein do not necessarily state or reflect those of the United States Government or any agency thereof.

REFERENCES

- Agarwal, V., Manjunatha, K. A., Gribok, A. V., Mortenson, T. J., Bao, H., Reese, R. D., ... others (2021). *Scalable technologies achieving risk-informed condition-based predictive maintenance enhancing the economic performance of operating nuclear power plants* (Tech. Rep.). Idaho National Laboratory (INL), Idaho Falls, ID (United States).
- Agarwal, V., Smith, J. A., Gribok, A. V., Joe, J. C., Oxstrand, J. H., & Primer, C. A. (2022). *Data architecture and analytics requirements for artificial intelligence and machine learning applications to achieve condition-based maintenance* (Tech. Rep.). Idaho National Laboratory (INL), Idaho Falls, ID (United States).
- Chen, Y.-C. (2017). A tutorial on kernel density estimation and recent advances. *Biostatistics & Epidemiology*, 1(1), 161–187.
- Ghorbani, A., & Zou, J. (2019). Data shapley: Equitable valuation of data for machine learning. In *International conference on machine learning* (pp. 2242–2251).
- Han, T., Wang, X., Guo, J., Chang, Z., & Chen, Y. (2025). Health-aware joint learning of scale distribution and compact representation for unsupervised anomaly detection in photovoltaic systems. *IEEE Transactions on Instrumentation and Measurement*.
- Ifeanyi, A., Saxena, A., & Coble, J. (2023). Within-bank condition monitoring and fault detection of fine motion control rod drives. 13th nuclear plant instrumentation. *Control & Human-Machine Interface Technologies (NPIC&HMIT 2023)*, 66(8), 1053–1062.
- Ifeanyi, A., Zanotelli, M., & Coble, J. (2025). System-level prognostics with explainable artificial intelligence. *Control & Human-Machine Interface Technologies (NPIC&HMIT 2025)*.
- Ifeanyi, A. O., Dos Santos, D., Saxena, A., & Coble, J. (2024). Fault detection and isolation in simulated batch operation of fine motion control rod drives. *Nuclear Technology*, 210(12), 2387–2403.
- Kim, S., Choi, Y., & Lee, M. (2015). Deep learning with support vector data description. *Neurocomputing*, 165, 111–117.
- Maćkiewicz, A., & Ratajczak, W. (1993). Principal components analysis (pca). *Computers & Geosciences*, 19(3), 303–342.
- Venkatesan, R., & Sriyutha Murthy, P. (2008). *Macrofouling control in power plants*. Springer.
- Xu, H., Wang, Y., Jian, S., Liao, Q., Wang, Y., & Pang, G. (2024). Calibrated one-class classification for

unsupervised time series anomaly detection. *IEEE Transactions on Knowledge and Data Engineering*.

BIOGRAPHIES

Ark Ifeanyi graduated from the University of Benin, Nigeria with a B.Eng. in Electrical and Electronic Engineering before obtaining an MSc. in Renewable Energy Systems Technology from Loughborough University, Leicestershire, UK. He is currently pursuing a Ph.D. in Energy Science and Engineering at the Bredesen Centre for Interdisciplinary Research and Graduate Education at the University of Tennessee. His current research is focused on the application of machine learning and artificial intelligence to the Prognostics and Health Management (PHM) of nuclear plants including small modular reactors (SMRs). Specifically, he is developing advanced machine learning algorithms to detect

and diagnose faults in nuclear power plants, with the ultimate goal of improving the safety, reliability, and efficiency of nuclear energy production.

Jamie Coble is an Associate Professor of Nuclear Engineering at the University of Tennessee, Knoxville (UTK). She earned her Ph.D. in Nuclear Engineering in 2010 from UTK. Prior to joining the faculty, she was a scientist in the Applied Physics group at Pacific Northwest National Laboratory. Her research interests lie mainly in applications of data analytics and machine learning in operations and maintenance of nuclear power plants. She is a member of American Nuclear Society and U.S. Women in Nuclear, senior member of IEEE, and fellow of the Prognostics and Health Management Society and of the International Society of Engineering Asset Management.