

Feature Selection from Normal Operation Data in Nuclear Power Plants Using Shape-Based Clustering

Jae Min Kim

Advanced I&C Research Division, Korea Atomic Energy Research Institute, Republic of Korea. E-mail:
jaemink@kaeri.re.kr

Ji Hyeon Shin

Advanced I&C Research Division, Korea Atomic Energy Research Institute, Republic of Korea. E-mail:
jhshin0127@kaeri.re.kr

Seo Ryong Koo

Advanced I&C Research Division, Korea Atomic Energy Research Institute, Republic of Korea. E-mail:
srkoo@kaeri.re.kr

The application of artificial intelligence in nuclear power plants is often hindered by the high dimensionality of sensor data and the scarcity of labeled records for abnormal events. While most existing diagnostic models rely on supervised learning with simulated fault data, this study proposes an unsupervised feature selection framework that leverages the abundant data available from normal operations. We hypothesize that the tight physical coupling within nuclear power plant systems causes redundant variables to exhibit synchronized time-series shapes, even during normal load-following transients. To identify these redundancies, we introduce a shape-based clustering approach using dynamic time warping, which robustly groups variables with similar dynamic behaviors regardless of temporal delays. This method is evaluated against traditional trend-based and agglomerative clustering techniques using data from a large-scale generic pressurized water reactor simulator. From an initial set of thousands of variables, the proposed method successfully identified a minimal subset of representative variables, achieving an extensive dimensionality reduction. The validity of the selected features was verified by training a simplified multi-layer perceptron to diagnose several distinct abnormal scenarios. The model trained on the reduced feature set demonstrated diagnostic accuracy comparable to a reference model using all variables, confirming that the selected representatives preserve the essential dynamic characteristics of the plant. This study demonstrates that unsupervised analysis of normal operation data can effectively filter redundant information, enabling the development of lightweight, interpretable, and computationally efficient diagnostic systems suitable for real-time monitoring.

Keywords: Normal operation, clustering, feature selection, dynamic time warping.

1. Introduction

The integration of digital twin technologies and Artificial Intelligence (AI) has become a pivotal element in enhancing the operational safety and monitoring efficiency of modern Nuclear Power Plants (NPPs). These facilities are equipped with thousands of high-fidelity sensors that continuously monitor critical physical parameters, such as reactor coolant temperature, pressure, flow rate, and secondary side water levels, generating massive volumes of time-series data (Cobb 2012). While this abundance

of data offers a significant opportunity for automated status assessment and anomaly detection, it also presents a "curse of dimensionality." Although modern high-capacity deep learning models can achieve high accuracy even with thousands of inputs, this approach introduces significant practical challenges: prohibitive computational costs for real-time edge deployment, increased vulnerability to individual sensor failures, and a "black-box" nature that hinders operator trust and

interpretability (Bellman 1957; Hinton and Salakhutdinov 2006).

Recent research in the nuclear domain has achieved notable success by applying deep learning architectures to the diagnosis of abnormal events. For instance, Shin and Kim (2022) demonstrated that interpretable convolutional neural networks could effectively identify fault-specific signal patterns, while Kim and Lee (2019) utilized long short-term memory networks combined with novelty detection to diagnose both trained and untrained accident scenarios. Furthermore, Kim et al. (2023) proposed a consistency check algorithm to validate these diagnoses against physical correlations. Despite these advancements, most existing diagnostic models rely on supervised learning, facing a bottleneck due to the scarcity of labeled abnormal data. Since real-world accident data is rare, models are often over-tuned to specific simulation datasets. Conversely, normal operation data is abundant and continuously generated, yet its potential for identifying the essential "skeleton" of plant dynamics remains underutilized (Ma and Jiang 2011).

This study is motivated by the observation that NPP subsystems are characterized by strong physical coupling governed by conservation laws. Consequently, many variables do not move independently but exhibit synchronized response patterns—or shape-based redundancy—even during normal transients. We hypothesize that identifying these physically coupled groups using only normal operation data allows for the selection of a small set of "representative variables" that preserve the full diagnostic capability of the original system.

The goal of this research is not merely to improve diagnostic accuracy, which is already inherently high for most major trip scenarios, but to achieve massive dimensionality reduction without performance degradation. By doing so, we aim to enable lightweight, robust, and explainable diagnostic systems suitable for real-time monitoring and reduced sensor maintenance. In this paper, we propose a shape-based clustering framework using Dynamic Time Warping (DTW) to capture these redundancies, comparing it against trend-based and agglomerative approaches. Finally, we verify that a simplified diagnostic model using only the

selected representatives maintains the same level of performance as a full-variable model, confirming the effectiveness of our feature selection strategy.

2. Methodology

The core premise of this study is that the operational dynamics of an NPP are governed by tight thermal-hydraulic and neutronic coupling. Components are interconnected via closed loops where the transfer of mass and energy strictly adheres to conservation laws (Upadhyaya et al. 2019). Consequently, a perturbation in a single parameter typically manifests as synchronized response patterns across multiple sensors. We define this as "shape-based redundancy." By identifying and grouping variables that exhibit these consistent behaviors during normal load-following operations, we aim to extract the plant's essential dynamic structure. This section details the three clustering approaches evaluated for this purpose and the validation strategy using a simplified neural network.

2.1. Baseline Approach: Trend-based Grouping

As a computationally efficient baseline, we first categorize variables based on their macro-level trends. A linear regression model is applied to each normalized time-series window to calculate the correlation coefficient (r) with respect to time. Variables are then classified into broad categories: 'Increasing' ($r > 0.8$), 'Decreasing' ($r < -0.8$), 'Steady' (negligible variance), or 'Mixed'. While this method offers extreme speed ($O(N)$), it groups variables solely on their general direction of change, potentially merging physically unrelated parameters that coincidentally share a trend direction during a specific maneuver.

2.2. Agglomerative Hierarchical Clustering

To capture linear dependencies beyond simple trends, we utilize Agglomerative Hierarchical Clustering with Pearson correlation as the distance metric ($d = 1 - r$). This bottom-up approach iteratively merges variables that move in linear synchronization. While effective for stable correlations, this method is highly

sensitive to transport delays. In large-scale fluid systems, temperature or pressure waves take finite time to travel between sensors. Correlation-based metrics often interpret this lag as dissimilarity, failing to group physically redundant variables and resulting in a redundant feature set (Warren Liao 2005).

2.3. Shape-Based Clustering using DTW

To robustly identify redundancy even in the presence of temporal lags, we propose using DTW. DTW measures similarity by finding an optimal non-linear alignment between two time series, stretching or shrinking the time axis to match their shapes regardless of phase shifts. For two normalized sequences X and Y , the cumulative DTW distance $\gamma(i, j)$ is calculated recursively:

$$\gamma(i, j) = |x_i - y_j| + \min\{\gamma(i-1, j), \gamma(i, j-1), \gamma(i-1, j-1)\} \quad (1)$$

The final distance $\gamma(n, m)$ quantifies the shape similarity independent of phase shifts (Berndt and Clifford, 1994). This capability is crucial for NPPs, as it allows for the grouping of sensors that measure the same physical event propagating through the system with a delay. We apply hierarchical clustering on the computed DTW distance matrix to identify clusters representing unique physical response shapes.

2.4. Feature selection and validation model

From each identified cluster, a single representative variable is selected to serve as an input feature. We choose the variable located closest to the cluster center (centroid/medoid), ensuring it best reflects the collective behavior of the group. This preserves the physical interpretability of the features, as they remain actual raw sensor measurements.

To validate the representativeness of these features, we employ a simplified Multi-Layer Perceptron (MLP) architecture. The underlying logic is that if a minimal, unoptimized neural network can maintain the same high accuracy as a reference model using only a fraction of the variables, it provides definitive proof that the selection process captured the essential information while successfully filtering

redundant noise. This validation emphasizes robustness and information density over the internal complexity of the classifier.

3. Experimental Setup

3.1. Data acquisition

The dataset for this study was generated using the generic pressurized water reactor simulator, which models an NPP with a rated output of 1,400 MWe. The simulator is implemented using the 3KEYMASTER tool a high-fidelity software platform used for engineering analysis and operator training (Western Services Corporation 2010). To establish a baseline for variable selection, a reference set of 2,831 variables was identified. This set comprises all process parameters and system indicators that are directly displayable on the simulator's human-machine interface screens, representing the most comprehensive observation space available to an operator.

3.2. Data configuration

Two distinct types of operational data were utilized: normal operation data with time-series records obtained during stable power operations. This data served as the input for the three clustering algorithms described in Section 2 to identify physical redundancies. To test the validation model, abnormal operation data were used in 45 independent simulation runs covering 15 distinct abnormal scenarios. These scenarios include steam generator tube leakage, loss of coolant accident, main feedwater pump failure, and reactor coolant pump issues. Those were simulated as low-level leakage events rather than ruptures as abnormal events. Each scenario is represented by multiple files to ensure statistical robustness during validation.

3.3. Preprocessing and input dimension

To simulate a real-time monitoring environment, a sliding window technique was applied to the time-series data. Each input sample was defined by a window size of 10 seconds with a 1-second stride. For each variable within the window, 10 sequential data points were flattened into a single vector. Consequently, for the reference model using all 2,831 variables, the input dimension for each time step was 28,310 (2,831 variables $\times$ 10 seconds). The objective of the proposed

framework is to reduce this input dimension significantly while maintaining the diagnostic signal.

3.4. Diagnostic Model Architecture

To verify the representativeness of the selected features, a simplified MLP was employed as the classification model. The network architecture consists of an input layer corresponding to the number of selected variables, two hidden layers with 100 and 50 neurons respectively (using ReLU activation), and an output layer with 15 neurons for scenario classification. The training was conducted using the Adam optimizer with a learning rate of 0.001 and a maximum of 1,000 iterations. The dataset was split into training (67%) and testing (33%) sets using a stratified approach based on scenario labels to ensure balanced representation in both sets.

4. Results and Discussion

The effectiveness of the proposed shape-based feature selection was evaluated by comparing its performance against conventional methods and a full-variable reference model. The evaluation focused on three primary metrics: the degree of dimensionality reduction, classification accuracy, and computational efficiency (training time). The three clustering methods yielded significantly different numbers of representative variables from the original set of 2,831 parameters. The Agglomerative method resulted in 101 clusters, indicating a high degree of granularity. In contrast, the proposed DTW-based clustering successfully identified 27 unique response shapes, achieving a 99.05% reduction in the number of variables. The trend-based grouping was the most aggressive, selecting only 9 variables. The results confirm that the DTW method occupies a balanced position, capturing more nuanced dynamics than simple trends while being significantly more efficient than standard correlation-based agglomeration.

4.1. Analysis of Cluster Quality

To verify the physical meaningfulness of the clusters, we visualized the time-series behavior of variables within the groups identified by the DTW method. As illustrated in Figure 1, variables within the same DTW cluster exhibit

highly synchronized dynamic behaviors despite having different physical units and scales.

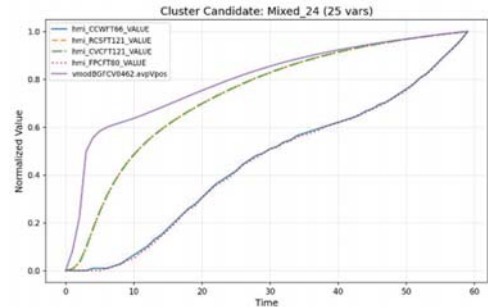


Fig. 1. Normalized time-series of representative variables in a DTW cluster.

For instance, in the ‘Mixed’ trend group, where variables show complex transient behaviors, DTW successfully grouped sensors that share identical inflection points and response curves. This confirms that the method effectively captures the “shape-based redundancy” resulting from the physical coupling between the primary and secondary loops. Unlike simple trend-based grouping which might merge unrelated increasing variables, DTW ensures that only variables with matching dynamic profiles are clustered, thereby preserving the precise signature of the plant’s state.

4.2. Model Performance

All models demonstrated exceptionally high diagnostic accuracy, as summarized in Table 1. Both the Reference model (using all variables) and the Agglomerative model achieved 100% accuracy on the test set. The Proposed DTW model and the baseline Trend-based model achieved 99.87% accuracy, with only minor misclassifications at the boundaries of specific windows.

Table 1. Comparison of Feature Selection Methods on Diagnostic Performance and Training Time.

Method	No. of Var.	Input Dim (10s)	Accuracy	Training Time (s)
Reference (All)	2831	28310	100%	32.48
Agglomeration	101	1010	100%	3.94

DTW	27	270	99.87%	1.6
Trend-based	9	90	99.87%	3.71

It is important to note that the primary objective was not to exceed the reference accuracy but to demonstrate that diagnostic performance does not degrade when input dimensionality is slashed by over 99%. The nearly identical performance of the DTW model (27 variables) and the Reference model (2,831 variables) suggests that normal operation data provides a sufficient basis for identifying the core physical couplings that define the plant’s response to failures.

The most significant impact of the feature selection framework was observed in computational efficiency. The training time for the Reference model using 28,310 input features was approximately 32.48 seconds. By utilizing the DTW-selected variables, the training time was reduced to 1.60 seconds, representing a 20-fold speedup. These results have profound implications for the practical deployment of AI in NPPs. A model with 27 variables is substantially easier to maintain, as it requires fewer sensors to be calibrated and functional at all times. Furthermore, the 95% reduction in input dimension compared to the Agglomerative method shows that DTW is superior at consolidating physically redundant but temporally lagged signals, making it ideal for edge computing environments where memory and processing power are limited.

4.3.Discussion

Visual analysis of the clusters reveals that the Agglomerative method tends to separate variables that are physically coupled but temporally shifted due to transport delays such as a temperature pulse moving along a pipe. From an engineering perspective, these lags are not mere noise; they represent the physical distance and response time between components. Therefore, the 101 clusters found by the Agglomerative method offer a high-fidelity map of the plant’s response network. In contrast, DTW-based clustering focuses on the “root phenomenon” by aligning these lags. By consolidating signals with the same underlying shape into 27 groups, it captures the fundamental dynamic skeleton of the plant.

This comparison highlights a valuable trade-off. If the diagnostic goal involves detailed root-cause analysis where the sequence of events is critical, the Agglomerative method may be preferred despite the larger feature set. However, for rapid status assessment and lightweight monitoring, the DTW method provides superior compression efficiency. The proposed method aims to identify physically consistent variable groups under nominal conditions, not to claim that the selected representatives are optimal across all stages of accident progression. As reactor dynamics deviate from the normal operating envelope, the relative importance of certain variables may shift. Therefore, the primary value of the proposed clustering lies in the consistency itself. Monitoring whether intra-cluster coherence is maintained or disrupted during plant operation can serve as a meaningful indicator for detecting deviations from normal behavior.

5. Conclusion

This study explored an unsupervised feature selection framework for nuclear power plant diagnosis based on shape-based clustering of normal operation data. By leveraging DTW to identify physically coupled variable groups, we demonstrated that the massive dimensionality of NPP sensor data can be effectively compressed without loss of diagnostic capability.

The experimental results revealed that a subset of only 27 representative variables could classify 15 different trip scenarios with 99.87% accuracy. This represented a 99% reduction in input variables and a 20-fold acceleration in model training compared to a full-variable reference model. The proposed method outperformed conventional correlation-based grouping by more effectively consolidating signals separated by transport lags.

Our findings suggest that the physical interdependencies observed during normal transients contain sufficient relational information to support robust fault diagnosis. This approach offers a clear path toward developing lightweight, interpretable, and maintainable AI systems for nuclear safety monitoring, reducing the reliance on high-performance hardware and extensive labeled

accident datasets. Future work will investigate the robustness of these clusters across different power levels and long-term aging effects, as well as the potential of intra-cluster consistency as an independent indicator for plant condition monitoring.

Acknowledgement

This work was supported by the Korea Institute of Energy Technology Evaluation and Planning (KETEP) and the Ministry of Climate, Energy & Environment (MCEE) of the Republic of Korea (No. 20224B10100130) and supported by the National Research Council of Science & Technology (NST) grant by the Korea government (MSIT) (No. GTL24031-000).

References

- Bellman, R. (1957). *Dynamic Programming*. Princeton University Press.
- Berndt, D. J. and Clifford, J. (1994). Using Dynamic Time Warping to Find Patterns in Time Series. *KDD Workshop*, 10, 359–370.
- Cobb, J. (2012). Nuclear Power Plant Diagnostics and Monitoring. *Nuclear Energy Enabling Technologies (NEET)*.
- Hinton, G. E. and Salakhutdinov, R. R. (2006). Reducing the dimensionality of data with neural networks. *Science*, 313(5786), 504–507.
- Kim, G., Kim, J., Shin, J., and Lee, S. J. (2023). Consistency Check Algorithm for Validation and Re-diagnosis to Improve the Accuracy of Abnormality Diagnosis in Nuclear Power Plants. *Nuclear Engineering and Technology*, 55, 1234–1245.
- Kim, J. and Lee, D. (2019). Nuclear Power Plant Accident Diagnosis Algorithm Including Novelty Detection Function Using LSTM. *Advances in Intelligent Systems and Computing*, 782, 123–130.
- Ma, J. and Jiang, J. (2011). Applications of Fault Detection and Diagnosis Methods in Nuclear Power Plants: A Review. *Progress in Nuclear Energy*, 53(3), 255–266.
- Na, M. G. et al. (2008). A method for sensor grouping and variable selection for nuclear power plants. *Nuclear Engineering and Technology*, 40(1).
- Shin, J. and Kim, J. (2022). An interpretable convolutional neural network for nuclear power plant abnormal events. *Annals of Nuclear Energy*, 178, 109345.
- Upadhyaya, B. R. et al. (2019). Sensor Fault Analysis and Isolation in Nuclear Power Plants. *Control and Dynamics of Nuclear Systems*.
- Warren Liao, T. (2005). Clustering of time series data—a survey. *Pattern Recognition*, 38(11), 1857–1874.
- Western Service Corporation (2010). *3KEYMASTER: Next Generation Simulation Technology* (Brochure). Frederick, MD: Western Services Corporation. Retrieved from <http://www.ws-corp.com>