



UDC 004.03

DOI: 10.62660/bcstu/3.2025.80

Adaptive similarity assessment metric for intelligent failure diagnostics in ship power plants

Vladimir Vychuzhanin*

Doctor of Engineering Sciences, Professor
Odesa Polytechnic National University
65044, 1 Shevchenko Ave., Odesa, Ukraine
<https://orcid.org/0000-0002-6302-1832>

Alexey Vychuzhanin

PhD, Assistant
Odesa Polytechnic National University
65044, 1 Shevchenko Ave., Odesa, Ukraine
<https://orcid.org/0000-0001-8779-2503>

Abstract. Prompt and accurate diagnosis of failures in ship power plants (SPPs) is essential for ensuring maritime safety, minimising operational risks, and optimising maintenance strategies. With increasing system complexity, heterogeneous data sources, and limited historical failure records, conventional diagnostic methods often prove insufficient, especially in scenarios involving rare or ambiguous faults. The purpose of this study was to develop an interpretable, adaptive, and probabilistically grounded methodology for assessing similarity between failure cases within SPPs for use in intelligent decision support systems. The proposed method integrates Euclidean, Jaccard, and logistic similarity metrics with Bayesian inference, temporal degradation modelling, frequency-based weight correction, and contextual smoothing of affected subsystems. The model employs L-BFGS-B optimisation to automatically adjust metric weights according to diagnostic relevance. Numerical experiments based on synthetic case data revealed high classification accuracy: 96% for failures related to cooling system overheating, 84% for bearing degradation cases, and 92% for fuel supply irregularities. Even with a 40% reduction in training data volume, the performance drop did not exceed 7%, indicating strong resilience to data sparsity. The visualisation of decision boundaries demonstrated the model's ability to distinguish overlapping failure classes while preserving semantic interpretability. Weight optimisation results identified "failure type" as the dominant factor, while "risk category" and "affected subsystems" had negligible impact and were excluded. Bayesian aggregation further improved the credibility of diagnostic conclusions by combining local similarity with global statistical priors. The developed methodology can be effectively applied by marine engineers, ship operators, and developers of intelligent diagnostic platforms for fault detection, root cause analysis, and predictive maintenance under conditions of uncertainty and incomplete information. Its modular structure also allows extending it to other complex technical domains beyond SPPs

Keywords: decision support system; Bayesian aggregation; adaptive weight coefficients; case-based failure analysis; probabilistic metric; similarity of technical conditions

INTRODUCTION

In modern complex technical systems (CTS), prompt and accurate failure diagnostics of equipment plays a key role in ensuring uninterrupted operation, safety, and

high system performance. The increasing complexity and heterogeneity of CTS have driven the widespread adoption of intelligent diagnostic methods based on

Article's History: Received: 18.04.2025; Revised: 03.08.2025; Accepted: 15.09.2025.

Suggested Citation:

Vychuzhanin, V., & Vychuzhanin, A. (2025). Adaptive similarity assessment metric for intelligent failure diagnostics in ship power plants. *Bulletin of Cherkasy State Technological University*, 30(3), 80-92. doi: 10.62660/bcstu/3.2025.80.

*Corresponding author



Copyright © The Author(s). This is an open access article distributed under the terms of the Creative Commons Attribution License 4.0 (<https://creativecommons.org/licenses/by/4.0/>)

reusing accumulated experience, among which the case-based reasoning (CBR) approach stands out. Failure diagnosis based on CBR involves comparing the current situation with previously recorded failure cases, where the effectiveness of the similarity assessment between cases is fundamental. Determining similarity between cases is the core task of CBR systems, as the quality of this assessment directly influences the accuracy, relevance, and promptness of diagnostic conclusions. Many approaches have been proposed to improve similarity assessment accuracy – from classical metric functions to hybrid machine learning methods and logical models. However, despite the progress, existing solutions often face a series of limitations: poor scalability, insufficient adaptability to heterogeneous data, and limited interpretability of results.

Adaptive fuzzy logic systems developed by L. Dubchak *et al.* (2024) demonstrated high adaptability and noise resilience in fault detection tasks. However, industrial-scale application is still limited due to the need for further validation on multidimensional heterogeneous datasets. According to S. Ren *et al.* (2020), in the field of low-carbon product design, effective similarity determination between cases plays a critical role. The researchers developed a model that considers both similarity and adaptability for knowledge reuse and design adaptation. B.M. Mathisen *et al.* (2020) argued that similarity metric learning based on data can automate the process of building similarity functions in case-based reasoning systems. The researchers proposed two novel approaches: using a pre-trained classifier as a foundation for the similarity function and a fully data-driven method that minimises training time. Both approaches showed high efficiency on 14 different datasets. Z. Chen *et al.* (2021) demonstrated that graph neural networks (GNN) can reveal hidden structural similarities in cyber-physical systems, but they require complex preprocessing and often suffer from a lack of decision-making transparency. N. Lin *et al.* (2023) proposed an enhancement of the CBR approach through the application of aggregation operators, such as the Choquet integral, which enabled accounting for non-linear interdependencies between diagnostic features. Despite improving similarity assessment accuracy, this approach complicates the formation of expert preferences. N. Neykov & S. Stefanova (2023) demonstrated the applicability of rule-based heuristics in CBR systems, highlighting their flexibility and explainability, but also noted challenges when working with large volumes of cases and the need for manual attribute weight tuning. According to N. Lin *et al.* (2023), interpretability plays a crucial role in case similarity matching systems in the legal domain. The researchers proposed an integrated framework consisting of four modules: key sentence identification, case matching, sentence alignment, and conflict resolution. This approach ensures transparency and explainability of the matching process. Most existing models do not account for probabilistic dependencies

between failures, differences in operating conditions, and the dynamic evolution of the technical state. For example, one failure may increase the probability of subsequent faults; operating modes may accelerate equipment degradation; component conditions change over time, influencing failure forecasting.

To enable accurate identification of typical equipment failures and their root causes within systems such as SPP, it is necessary to develop a unified and adaptive methodology for assessing the similarity between cases. At its core, the methodology should feature a similarity metric that ensures relevant comparison of current cases with archived ones and identification of the most similar analogues. In this context, similar cases refer to those with comparable operating conditions, identical damage characteristics, close parameter values, and matching probabilistic characteristics. For example, two main engine failures can be considered similar if they occurred under comparable operating hours, were accompanied by similar changes in temperature and pressure parameters (including oil temperature and fuel pressure), and shared a root cause – for instance, overheating caused by cooling system failure. The methodology being developed should account for both quantitative and qualitative failure features; enable the integration of metric, logical, and machine-learning components; and be applicable to heterogeneous data types (numerical, categorical, set-based, and temporal). Other requirements include interpretability for technical experts and adaptability to changing operating conditions.

The proposed structured approach covered industry-specific characteristics, applied parameter normalisation, and implemented weight coefficient optimisation to improve diagnostic reliability. Particular attention was given to considering operational and technical characteristics (load, operating time, type of damage), causal relationships between components, statistical parameters (failure frequency, prior probabilities), and the predicted dynamics of equipment conditions. Altogether, this enabled the construction of an adaptive case-matching system capable of effectively identifying similar failures under real-world operating conditions.

The purpose of this study was to develop an interpretable, adaptive, and statistically grounded methodology for assessing similarity between failure cases in ship power plant systems, based on the integration of metric, probabilistic, and logic-based components. The methodology was intended to improve matching accuracy, robustness against heterogeneous data, and adaptability to changing operational conditions, with subsequent integration into intelligent CBR-based diagnostic systems.

MATERIALS AND METHODS

This study examined and compared various approaches to assessing the degree of similarity between equipment failure cases (precedents) within SPP. The focus was on the analysis of similarity metrics used within

the CBR methodology, including conventional, probabilistic, and hybrid approaches. The analysis covered the following categories of methods: distance metrics, such as Euclidean, Manhattan, and cosine, which were applied to numerical and categorical parameters; set-based metrics, including the Jaccard index, which were used to compare sets of affected subsystems and risk categories; probabilistic approaches, where the similarity metric was adjusted based on the empirical probability of joint feature occurrence or by using Bayesian weights; aggregated models, where the final similarity metric represented a weighted sum of partial metrics with weights optimised numerically (using the L-BFGS-B method); and hybrid CBR structures that combined rules, expert assessments, and supervised learning.

To evaluate the effectiveness and identify the limitations of the listed methods, a comparative experimental series was conducted. Within this series, a synthetic training sample of precedents was created, covering typical equipment failures of SPP, involved components, and occurrence probabilities. Baseline methods with fixed weights were compared to adaptive models, where weights were adjusted based on comparison with expert similarity assessments. The effectiveness of the models was evaluated according to the following criteria: classification accuracy, including the use of confusion matrices; correlation with expert similarity assessments; and sensitivity of metrics to variations in input parameters.

The parameter weights were optimised using the limited-memory Broyden–Fletcher–Goldfarb–Shanno

algorithm with bound constraints, where the loss function was defined as the mean squared error between the predicted and expert-assessed similarity values. For this, the implementation from the `scipy.optimize` library was used, ensuring computational efficiency and reproducibility. Numerical experiments with parameter variation and visualisation tools that demonstrated the contribution of individual components to the final similarity measure were applied for model validation. The presented methodology helped to identify key limitations of conventional approaches and demonstrated the advantages of adaptive, probabilistically informed models in failure diagnostics.

To search for similar failures in the precedent database, the k-Nearest Neighbours (k-NN) method was used, as described by M. Zuber & R. Sirdey (2021). This approach helped to identify the most similar cases based on a predefined similarity metric. The choice of this method was driven by the following advantages: flexibility – k-NN handled heterogeneous data types (numerical, categorical, probabilistic parameters); interpretability – unlike neural network models, it was possible to explain why two failures were considered similar; adaptability – the method allowed adjusting parameter weights depending on operating conditions; efficiency with limited data – it did not require a large training set, which was critical for the diagnostics of rare failures. The similarity between two precedents was calculated based on a weighted sum of partial similarities across key parameters, as presented in Table 1.

Table 1. Similarities between two precedents based on key parameters

Parameter	Description	Similarity Measurement Method
Failure Type	Nature of damage, degradation mechanism	Cosine similarity
Failure Causes	External and internal factors, preceding events	Jaccard similarity
Failure Probability	Estimated probability of failure	Euclidean distance
Failure Consequences	Impact on operability, cascading effects	Bayesian networks
Affected Components	Failed nodes and subsystems	Jaccard similarity

Source: created by the authors of this study

Depending on the type of data, the following metrics were used: euclidean distance – for numerical parameters (e.g., failure probability, remaining useful life); cosine similarity – for textual and vector representations of failure categories; jaccard similarity – for categorical data (e.g., lists of affected subsystems); Markov models – for predicting failure evolution over time and adjusting parameter significance; Bayesian networks – for accounting for probabilistic dependencies between failures.

RESULTS AND DISCUSSION

To build a formalised model for determining similarity between precedents, information on typical failures of SPP equipment and corresponding diagnostic

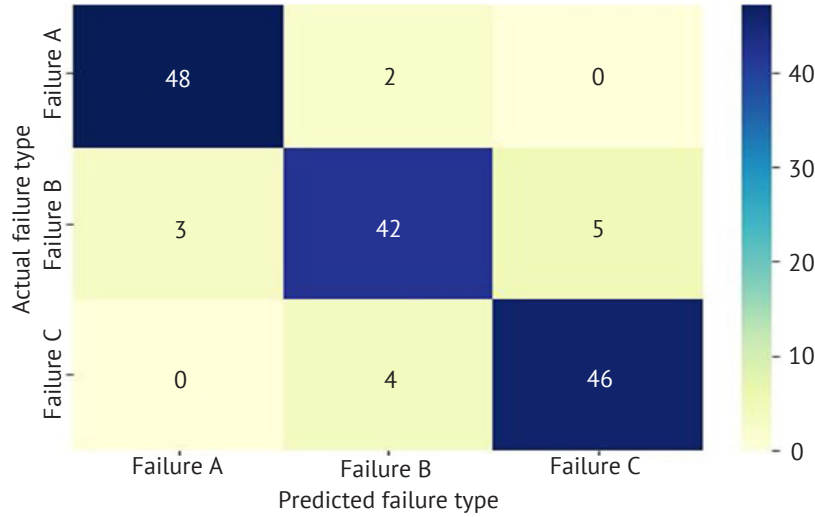
features was structured. Table 2 presents the classification of failures according to key parameters: temperature; vibration level; pressure; typical manifestations. This structure allows highlighting diagnostically significant features, which are further formalised into precedent parameters. Additionally, Figure 1 presents a classification error matrix that reflects characteristic cases of misclassification of observed features into failure types, demonstrating the limitations of conventional matching methods and emphasising the need for a more precise similarity evaluation model.

Figure 1 presents the classification error matrix, reflecting the performance of the CBR system when classifying new precedents by failure types based on the applied similarity assessment methodology.

Table 2. Classification of typical failures and diagnostic features

Failure Type	Temperature	Vibration	Pressure	Typical Symptoms
Failure A	> 90°C	2.0-3.5 mm/s	> 10 bar	Often associated with cooling system overheating
Failure B	75-85°C	1.5-2.5 mm/s	6-9 bar	Related to gradual bearing wear
Failure C	> 85°C	2.5-4.0 mm/s	9-11 bar	Observed with unstable fuel supply

Source: created by the authors of this study

**Figure 1.** Classification error matrix by failure types

Source: created by the authors of this study

The factual and predicted failure types (A, B, C) are indicated along the axes (Fig. 1). The values on the diagonal of the matrix correspond to correctly classified cases: 48 instances of type A were correctly identified as A; 42 instances of type B were classified correctly; and 46 instances of type C were recognised accurately. Classification errors were observed in the following forms: 3 cases of type B were mistakenly assigned to class A, and 5 to class C; 4 cases of type C were classified as B. Type A demonstrated the greatest classification stability (only 2 errors out of 50 observations, accuracy – 96%). This reflects high overall classification accuracy, with the share of correct predictions exceeding 90% for all failure types. A notably increased probability of cross-classification between types B and C may be explained by partial overlap of features (e.g., similar vibration and pressure values under different failure scenarios); insufficient discrimination capacity of the current similarity metric in the feature space of these types or a limited training sample size for B and C cases. Such insights enable targeted model enhancement – for instance, by refining parameter weights or introducing other features that enable clearer separation of failure types with similar manifestations.

The developed similarity model for comparing a current case A_i and a precedent B_i is represented as follows:

$$S(A_i, B_i) = \alpha_{s,i} \cdot d_{types,i}(A_i, B_i) + \beta_{s,i} \cdot d_{probability,i}(A_i, B_i) + \gamma_{s,i} \cdot d_{components,i}(A_i, B_i), \quad (1)$$

where $d_{types,i}(A_i, B_i)$ is the similarity measure for failure types; $d_{probability,i}(A_i, B_i)$ is the difference in failure probabilities; $d_{components,i}(A_i, B_i)$ is the similarity estimation for the affected components; $\alpha_{s,i}$, $\beta_{s,i}$, $\gamma_{s,i}$ are the weighting coefficients that determine the significance of different aspects when comparing precedents

The similarity components are calculated as follows:

$$d_{types,i}(A_i, B_i) = \begin{cases} 1, & \text{if types are identical} \\ 0.5, & \text{if in the same category;} \\ 0, & \text{if types differ} \end{cases} \quad (2)$$

$$d_{probability,i}(A_i, B_i) = 1 - |P(A_i) - P(B_i)|; \quad (3)$$

$$d_{components,i}(A_i, B_i) = \frac{|C(A_i) \cap C(B_i)|}{|C(A_i) \cup C(B_i)|}, \quad (4)$$

where $P(A_i)$, $P(B_i)$ are the failure probabilities for the respective components; $C(A_i)$, $C(B_i)$ are the sets of components involved in the failures. This formulation enables a quantitative accounting of subsystem overlap, which directly affects the assessment of case similarity.

For normalisation of the final similarity metric, the following expression is applied:

$$D_{total}(A_i, B_i) = \frac{\alpha_{s,i} \cdot d_{types,i}(A_i, B_i) + \beta_{s,i} \cdot d_{probability,i}(A_i, B_i) + \gamma_{s,i} \cdot d_{components,i}(A_i, B_i)}{\alpha_{s,i} + \beta_{s,i} + \gamma_{s,i}}. \quad (5)$$

Such normalisation ensures that the resulting similarity score stays within the range [0, 1], enabling

consistent comparison of different case pairs and facilitating further interpretation within the diagnostic decision-making process. The coefficients $\alpha_{s,i}$, $\beta_{s,i}$ and $\gamma_{s,i}$ are optimised using numerical methods aimed at minimising a loss function that reflects the classification error based on known labels (e.g., the mean squared error between the predicted similarity value and the expert-assessed value). This allows the model to be adapted to the specifics of a particular domain and the structure of the case database. To improve the accuracy, interpretability, and adaptability of the similarity assessment model between cases, the following enhancements have been implemented:

Weight correction based on failure frequency

This extension allows the model to factor in the rarity of specific failure types in the dataset, automatically increasing the importance of rarely occurring but potentially critical scenarios. This is particularly relevant for the diagnostics of low-probability but high-risk conditions, which standard methods may overlook due to their low occurrence rate.

The weight correction formula is defined as:

$$\tilde{\alpha}_{s,i} = \alpha_{s,i} \cdot \log \left(1 + \frac{1}{f_{type}(A_i)} \right), \quad (6)$$

where $f_{type}(A_i)$ is the relative frequency of occurrence of failure type A_i in the case database.

This adjustment ensures automatic balancing of the model: the rarer a particular failure type, the greater its weight when calculating the final similarity measure. This enables the model to better handle the diagnostics of infrequent but significant cases, preventing bias toward frequently occurring but less critical failures.

Such a correction can further be extended to other parameters, including probabilistic and structural characteristics, which would additionally enhance the flexibility and reliability of the case-based reasoning system.

Temporal degradation factor

This extension accounts for the influence of the temporal aspect (e.g., the difference in operating time until failure) when assessing the similarity between cases. Incorporating this factor allows the model to prioritise cases that are closer in terms of operational time, making the reasoning process more context-aware. The temporal degradation factor is calculated as follows:

$$d_{degradation}(A_i, B_j) = \exp(-\lambda \cdot |TA_i - TB_j|), \quad (7)$$

where TA_i , TB_j are the operating times until failure for cases A_i and B_j ; λ is the degradation coefficient, which controls the sensitivity of the model to temporal differences.

The adjusted similarity score is defined as follows:

$$S'(A_i, B_j) = S(A_i, B_j) \cdot d_{degradation}(A_i, B_j). \quad (8)$$

The introduction of this multiplicative temporal context reflects the idea that two cases with identical features but differing significantly in failure time are diagnostically less similar than cases with both matching features and close failure times. Figure 2 illustrates how the temporal degradation coefficient varies depending on the difference in operational time between two cases. The figure presents three scenarios corresponding to different sensitivity levels of the model to temporal discrepancies, controlled by the value of the λ parameter.

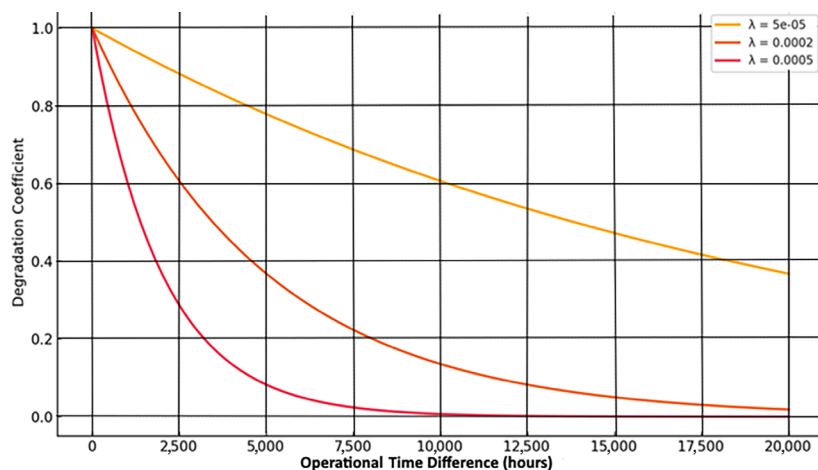


Figure 2. Change of the temporal degradation coefficient depending on the difference in operating time between two cases

Source: created by the authors of this study

The curve with the smallest λ (0.00005) shows a slow decrease of the coefficient, which reflects a model with stable memory – it “trusts” even older cases. Conversely, at a large λ (0.0005), the model quickly depreciates old cases, focusing on the most recent data. Thus, the choice

of λ directly affects the behaviour of the model: a low λ implies high stability, which is useful in the presence of stable patterns; a high λ ensures adaptability to new conditions, which is necessary in the case of dynamically changing technical states of the SPP system. This

emphasises the significance of calibrating the temporal component when designing CBR systems, especially in the context of variable load and operating modes.

Contextual smoothing of components supplements the comparison of SPP system components with information about the proximity between components (e.g., if two components are technically related). This reduces information loss in cases of partial mismatch.

$$d_{comp}^{adj} = \frac{|C(A_i) \cap C(B_i)| + \delta \cdot |Ad_j(A_i, B_i)|}{|C(A_i) \cup C(B_i)| + \delta \cdot N}, \quad (8)$$

where $|Ad_j(A_i, B_i)|$ is the number of adjacent, functionally related subsystems; δ is the smoothing coefficient; N is the total number of possible components.

Nonlinear activation (e.g., tanh) allows considering the saturation effect in cases of strong similarity and helps to soften sharp boundaries. This is used for better interpretation of the final similarity in extreme cases:

$$S(A_i, B_i) = \tanh(\alpha_{s,i} \cdot d_{types,i}(A_i, B_i) + \beta_{s,i} \cdot d_{probability,i}(A_i, B_i) + \gamma_{s,i} \cdot d_{components,i}(A_i, B_i)). \quad (6)$$

The proposed formalised model for assessing similarity between cases (1) is integrated into the classical CBR architecture, which includes four key stages: retrieval, adaptation, application, and update. At the retrieval stage, the model is used to search for the most relevant analogues in the case base.

Stage 1. Calculation of similarity between cases

For a new diagnostic case A_i , represented by a feature structure (failure type, probability, subsystems, etc.), the similarity measure is calculated for all cases B_i from the knowledge base according to formula (1). If additional modules (temporal degradation, probabilistic correction, etc.) are activated, the final value is modified accordingly:

$$S'_{norm}(A_i, B_i) = \frac{S(A_i, B_i) \cdot d_{degradation}(A_i, B_i) \cdot P_i(A_i, B_i)}{\alpha_{s,i} + \beta_{s,i} + \gamma_{s,i}}, \quad (9)$$

where $P_i(A_i, B_i) \in [0, 1]$ is the probability of co-occurrence of the values A_i and B_i in the cases, extracted from empirical data or an expert model

Stage 2. Nearest neighbour search (k-NN)

Based on the obtained values $S(A_i, B_i)$, the k nearest cases with the highest similarity values are selected. After constructing and formalising the similarity metric between cases across the set of parameters, the presented model is integrated into the k-Nearest Neighbours (k-NN) search algorithm. Using k-NN not only enables the application of the similarity metric but also implements a decision-making process based on it.

Advantages of the approach: local analysis: decisions are made based on similar cases, which is especially significant considering the high variability of operating modes of marine power plants; flexibility: the method can be adapted for specific diagnostic goals (e.g., predicting time to failure or classifying the failure type); scalability: adding new cases does not require restructuring the model, only a local recalculation of similarity. For the diagnostics of a new case (the observed equipment condition): the global similarity measure $S(A_i, B_i)$ is calculated between the new case and each case in the database according to formula (1); the k cases with the highest similarity (equivalently – with the minimum distance according to the metric) are selected. A diagnostic conclusion is then formed based on the k nearest cases. Possible approaches include majority voting – the most frequently occurring failure type among the nearest neighbours; weighted voting – considering the similarity level of each neighbour; aggregated parameter values – e.g., averaging the time to failure, temperatures, pressures, etc. The application of k-NN requires prior normalisation of the parameters involved in the calculation of individual similarity measures to avoid domination by parameters with larger measurement scales. Additionally, the choice of the parameter k should factor in the density of the case base (with low density, a smaller k is preferable); the variability and noise level in the data (for high noise, it is better to average over a larger number of neighbours); the type of task being solved (classification, prediction, anomaly detection, etc.). Figure 3 illustrates the nearest neighbour method (k-NN) in the parameter space that defines the similarity between cases.

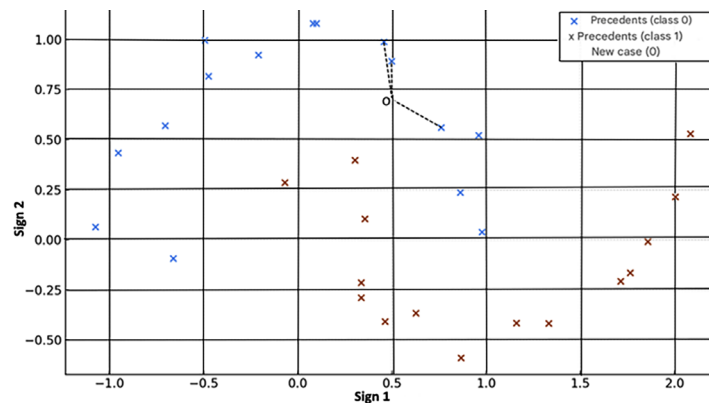


Figure 3. Application of the k-NN method in the parameter space defining the similarity between cases

Source: created by the authors of this study

Figure 3 presents an example of classifying a new failure case with coordinates (0.5; 0.7) in the feature space using the k -nearest neighbours method ($k = 3$). The dashed lines visualise the connection of the analysed object with the three most similar cases, selected according to the minimum distance metric in a pre-normalised feature space with consideration of their weighting coefficients. The position of the new case in the area dominated by cases of the first class allows the model to classify it as a failure of the first type based on the neighbour voting strategy. This example demonstrates the key advantages of the k -NN method for diagnosing complex technical systems: the ability to work with nonlinearly separable data, robustness to noise through aggregation of information from several neighbours, as well as clear interpretability of the results.

Special attention must be paid to the correct tuning of the method's critical parameters – the optimal k value, proper normalisation of features, and the selection of an appropriate similarity metric that considers the scale and informativeness of various parameters. The presented approach ensures the consideration of the local distribution characteristics of the data and increases the robustness of the results to variability in input parameters. Its ability to work with incomplete and noisy data, while combining visual representation with probabilistic justification of decisions, makes the method especially valuable for diagnosing SPP, where one typically faces limited historical data, complex nonlinear dependencies, and high demands for interpretability of the obtained results.

Stage 3. Aggregation of the diagnostic decision

The extracted k -nearest neighbours are analysed: by the most frequently occurring failure type (majority voting); by the average probability of failure; by contextually weighted characteristics. The final diagnostic conclusion is formed based on the coordinated information from the most similar historical cases. Once the k most similar cases $\{B_1, \dots, B_k\}$ have been determined for the new case A_p , their diagnostic characteristics are aggregated. The most typical strategies include majority voting for failure type (the failure type that occurs more frequently among the nearest cases is selected):

$$T^* = \arg \max_T \sum_{i=1}^k 1_{T_i=T}. \quad (10)$$

To determine the most probable failure type T^* among the nearest cases, a majority voting strategy is applied. The indicator function $1_{T_i=T}$ equals 1 if the failure type T_i in the i^{th} precedent matches type T , and 0 otherwise. By summing these indicators across all k nearest neighbours, the method identifies the most frequently occurring failure type. The type T with the highest number of “votes” is selected as the final diagnosis.

Weighted average failure probability:

$$\hat{p}(A_i) = \frac{1}{\sum_{i=1}^k S(A_i, B_i)} \cdot \sum_{i=1}^k S(A_i, B_i) \cdot p(B_i), \quad (11)$$

where $p(B_i)$ is the known failure probability for the corresponding precedent.

To estimate the failure probability for the current case A_p , a weighted average of known failure probabilities $p(B_i)$ from the k nearest precedents is used. Each precedent's contribution is weighted by its similarity score $S(A_p, B_i)$. The denominator normalises the sum, ensuring the resulting probability stays within the $[0, 1]$ range. This approach gives more influence to precedents that are more similar to the current case.

Merging of affected components

If failure localisation is required, it is possible to form an aggregated set of the most probable components affected in similar cases: $C_{agg}(A_i) = \bigcup_{i=1}^k C(B_i)$. When the objective is failure localisation, an aggregated set of potentially affected components is formed by taking the union of component sets $C(B_i)$ from the k most similar precedents. This provides a complete list of all subsystems that were involved in similar past failures and may be relevant for the current case A_p .

The final aggregated set of components (subsystems) that are considered most likely involved in the failure for the current case A_p , accounting for weights (similarity):

$$C_{agg}^{score}(A_i) = \{c: \sum_{i=1}^k S(A_i, B_i) \cdot 1_{c \in C(B_i)} > \theta\}, \quad (12)$$

where c is a specific subsystem of the SPP (e.g., “fuel pump,” “cooling circuit,” etc.) considered as potentially failed; $\sum_{i=1}^k S(A_p, B_i) \cdot 1_{c \in C(B_i)}$ is the weighted sum over k nearest precedents; $S(A_p, B_i)$ is the similarity between the current case and precedent B_i ; $1_{c \in C(B_i)}$ is the indicator function: equals 1 if component is present in the list of failed subsystems of precedent B_p and 0 if not; θ is the significance threshold, which determines how frequently and with what weight component c must occur to be considered significant.

To improve failure localisation accuracy, this formula performs a weighted aggregation of components based on their presence in similar cases. A component c is included in the final set $C_{agg}^{score}(A_i)$ if its weighted occurrence frequency across the k nearest neighbours exceeds a given threshold θ . This approach helps filter out random or weakly relevant components, retaining only those that consistently and significantly appear in similar cases.

Although the primary goal of the proposed model was to construct a similarity metric between cases, the implementation of diagnostic inference in a CBR system requires a justified decision-making method. When several similar cases are available, it becomes necessary to aggregate their characteristics. In such conditions, it is advisable to apply a probabilistic approach – specifically, Bayesian aggregation – for interpreting results and improving diagnostic accuracy.

This model extension logically followed from the task of result interpretation after computing the similarity. Additionally, the diagnostic decision can be

justified in terms of Bayesian inference, where the failure type T was treated as a hypothesis, and the retrieved precedents as observations:

$$P(T|A_p, k) \propto P(T) \cdot \prod_{i=1}^k P(B_i|T)^{\omega_i}, \quad (13)$$

where $P(T)$ is the prior probability of failure type T (according to the statistics of the precedent database); $P(B_i|T)$ is the conditional probability of observing precedent B_i given failure type T ; $\omega_i \propto S(A_p, B_i)$ is a weight reflecting the confidence level in precedent B_i based on its similarity to A_p .

If it is assumed that the probability of observing a specific precedent B_i depends on its degree of similarity to A_p then:

$$P(B_i|T) \approx S(A_p, B_i) \cdot 1(T_i = T), \quad (14)$$

which leads to the following posterior estimation:

$$P(T|A_p, k) \propto P(T) \cdot \sum_{i=1}^k S(A_p, B_i) \cdot 1_{T_i = T} \quad (15)$$

In practice, this means that the failure type with the highest posterior probability is accepted as the most probable diagnosis for the current case A_p . This approach is particularly useful under conditions of incomplete data, where precedents may be heterogeneous or lack a complete set of features, and the Bayesian model allows systematically incorporating the degree of confidence. Figure 4 presents a visualisation of the posterior inference process for determining the probability of a new precedent (observed failure) belonging to one of the typical failure classes (e.g., A, B, C) based on the Bayesian model.

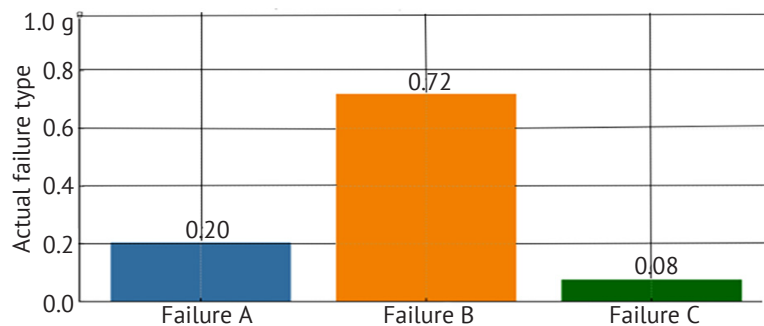


Figure 4. Visualisation of the posterior inference process for determining the probability of a new precedent's class affiliation

Source: created by the authors of this study

Bayesian aggregation was applied here as an interpretable decision-making mechanism, enabling not only the identification of the most probable failure type but also the quantitative consideration of uncertainty and the probabilistic distribution of hypotheses. The diagram illustrates that, for the analysed case, the greatest posterior probability corresponds to failure type B. This reflects both the frequency dominance of similar cases of type B among the nearest neighbours (resulting from the k-NN search in the similarity space) and the high prior probability of this failure type within the overall diagnostic statistics of the precedent database. The Bayesian aggregation methodology enables the weighted integration of two key sources of information: global knowledge about the distribution of failures (prior probabilities); local similarity of the current case to recorded precedents.

This model was particularly effective when working with limited training datasets, as well as in scenarios with diverse failure causes, where the same type of malfunction can result from various degradation pathways. The use of probabilistic inference: enhances interpretability – each step of the decision-making process can be justified through the probabilistic model; reduces sensitivity to data noise – the unequal contribution of neighbours is regulated not only by similarity

but also by their probabilistic weight; allows for distributed uncertainty between multiple failure types – especially in the presence of overlapping features. Thus, the visualisation confirms the validity of using Bayesian aggregation as an extension to the k-NN method, providing more accurate, flexible, and robust classification of precedents in the intelligent diagnosis of SPPs.

Stage 4. Updating the case database

After the diagnosis has been verified, the new case – including the original feature set and diagnostic results – was added to the case database for further enrichment and continuous model training. As the database of precedents grows, the system became increasingly capable of accurately matching new cases to previously observed ones, progressively adapting the model to real-world operating conditions. During the classification and decision-making stages, the mechanism of comparison with the most similar precedents played a crucial role.

Thus, the proposed similarity assessment model for failure precedents in SPPs represented an integrated system that combines numerical, categorical, and multi-valued parameters with adaptive weights and specialised similarity metrics. The stepwise extension of the base formula – through the inclusion of failure frequency, temporal factors, functional interconnections of

subsystems, and nonlinear transformations – ensured not only high accuracy but also the robustness of the model against data heterogeneity. The additional application of probabilistic aggregation for determining the final diagnostic conclusion demonstrated the potential for integrating this model into broader intelligent CBR systems. This integration opens new opportunities for enhancing diagnostic reliability and decision interpretability, especially under conditions of limited training data or high variability in failure mechanisms.

The optimisation of parameter weights enabled the model to consider the relative significance of each feature for diagnostic purposes. The coefficients, ω , were optimised using a bounded limited-memory quasi-Newton optimisation method (L-BFGS-B algorithm), which minimised the following loss function:

$$J(\alpha_{s,i}, \beta_{s,i}, \gamma_{s,i}) = \sum_i^M (S_{norm}(A, B_i) - y_i)^2, \quad (16)$$

where y_i is a binary label assigned in the training dataset of M known pairs, indicating whether the precedents are considered similar (1) or dissimilar (0). When formulating the optimisation problem, there is a need to assume the

training dataset consists of N pairs of precedents (i, j) , for which the reference similarity degree $S_{i,j}^{ref}$ has been assigned by an expert. For each pair, the model computes the predicted similarity value $S_{i,j}^{mod}(\omega)$, which depends on the weight vector $\omega = (\omega_1, \omega_2, \dots, \omega_k)$. The objective function for optimisation is defined as follows:

$$L(\omega) = \frac{1}{N} \cdot \sum_{i,j} (S_{i,j}^{mod}(\omega) - S_{i,j}^{ref})^2, \quad (17)$$

where the minimised value $L(\omega)$ represents the mean squared deviation between the model-predicted and expert-defined similarity values. The L-BFGS-B algorithm enabled efficient minimisation of this function under constraints on the weight parameters (e.g., $\omega_i \geq 0$, $\sum \omega_i = 1$). The optimisation procedure was implemented using the L-BFGS-B method from the `scipy.optimize` library (Virtanen *et al.*, 2020). The optimisation process consisted of several stages: formation of a training dataset containing pairs of precedents with known expert similarity assessments; initialisation of weight values; and evaluation of the optimisation quality using a separate test dataset. Figure 5 presents a diagram comparing the initial and optimised weights.

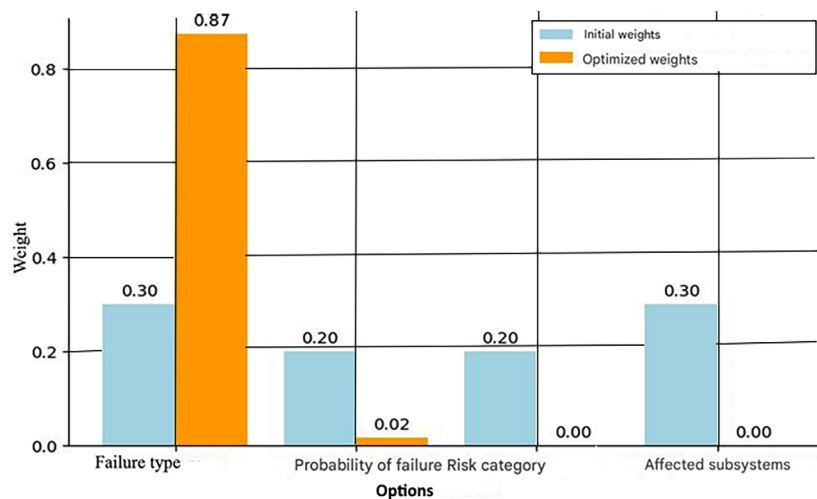


Figure 5. Comparison of initial and optimised weights

Source: created by the authors of this study

The diagram illustrates the comparison between the initial and optimised weights of various parameters used in the similarity assessment of failure precedents. A significant increase in the significance of the “failure type” parameter was observed: initial weight – 0.30; optimised weight – 0.87. This indicates that the parameter has a strong impact on diagnostic accuracy, and as a result, its contribution to the model was substantially increased. A reduction in the influence of the “failure probability” parameter was recorded: initial weight – 0.20; optimised weight – 0.02. This suggests that failure probability has a lower significance when determining precedent similarity, and its weight was nearly nullified.

The “risk category” and “affected subsystems” parameters were excluded from the final model. Their

initial weights were 0.20 and 0.30, respectively, but after optimisation, both were reduced to 0.00. This suggests that these parameters have no meaningful influence on the diagnostic outcome and can be safely removed from the model. In summary, the optimisation process highlighted that the “failure type” parameter is the key factor when comparing failure precedents. In contrast, parameters such as the equipment risk category and affected subsystems were found to be insignificant and were eliminated from the model. It is essential to verify whether such a radical redistribution of weights leads to a loss of meaningful diagnostic features. It may also be advisable to explore alternative optimisation strategies that preserve the contribution of all parameters to a certain extent. To assess the position of the proposed

model for evaluating similarity between failure cases within the SPP equipment diagnostic structure, an in-depth analysis of analogous studies was conducted. The comparison focused on key methodological dimensions: the nature of the similarity metric, integration of probabilistic reasoning, presence of adaptive weight mechanisms, robustness to incomplete or noisy data, and applicability across heterogeneous technical domains.

In the analysis of time series sensor data, J. Serà & J.L. Arcos (2014) studied the distance measures such as dynamic time warping (DTW), confirming their robustness against temporal distortions. Nevertheless, such methods perform poorly in semantically saturated contexts and are computationally intensive. In another area, J. Ye (2017) introduced neutrosophic logic with cotangent-based similarity measures to address uncertainty in technical system data, although result interpretation continues to be challenging.

The DeepKAF approach presented by K. Amin *et al.* (2020) effectively merged deep learning and CBR for NLP tasks but lacked flexibility in similarity metric configuration. While this model demonstrated high performance in pattern recognition, it relied on rigid internal structures and did not support metric reconfiguration, thus limiting its applicability in domains requiring feature-specific interpretability. In contrast, the model presented in the present study enables flexible configuration of similarity metrics for numerical, categorical, and set-valued attributes, improving transparency and domain adaptability. D. Leake *et al.* (2021) focused on neural case adaptation to enhance analogical reasoning. However, presented approach lacked support for probabilistic inference and dynamic weight adjustment, which are essential when dealing with sparse or uncertain failure data. The model under discussion addressed this gap by incorporating Bayesian hypothesis aggregation and data-driven optimisation of feature weights, improving reliability and resilience to data imperfections. M. Chen *et al.* (2022) proposed an adaptive metric learning system for aeroengine fault diagnostics. Their findings confirmed the benefits of adaptation in technical diagnostics, yet the proposed architecture did not explicitly incorporate probabilistic modelling or failure frequency weighting. In contrast, the method evaluated in this study links metric behaviour to statistical failure patterns and employs Bayesian reasoning to support decision-making under uncertainty. M.F. Valdez-Ávila *et al.* (2023) developed the CBR-fox model for time-series analysis. While effective for temporal patterns, presented solution did not support integration of diverse static features or account for temporal degradation. The approach discussed in this paper expanded these capabilities by incorporating set-based feature fusion and temporal decay modelling, enhancing long-term diagnostic accuracy. M. Krüger (2023) applied similarity measures to cybercrime profiling. Although effective in its domain, the approach was specialised and did not demonstrate cross-domain flexibility. The presented similarity model

was designed to account for a variety of technical systems, including electromechanical and thermal areas, which is necessary for wider application in engineering. H. Xu *et al.* (2023) explored the integration of CBR with ontological knowledge graphs for bridge safety diagnostics. Their study offered semantic depth but required complex, manually curated graphs, limiting scalability. The current model bypassed such dependencies by directly handling generalised structured features, enabling more efficient deployment in real-time scenarios. X. Ye *et al.* (2024) introduced neural k-NN into CBR for enhanced representation learning. Despite improved accuracy, the use of implicit neural weights reduced transparency and limits practical interpretability. In contrast, the present approach employed explicitly defined weights with tunable parameters based on operational significance, maintaining explainability. S. Mustyala & M. Bisi (2025) applied harmony search combined with CBR for software effort estimation. Although adaptive in nature, the model lacked probabilistic mechanisms, which are critical for uncertainty modelling. This limitation was addressed in the current framework through Bayesian aggregation and normalisation based on similarity strength. G. Jimenez-Díaz & B. Díaz-Agudo (2024) developed visualisation tools for understanding similarity structures in CBR. Their study improved user transparency but did not integrate probabilistic interpretation. In the presented model, visualisation was combined with statistical weighting, enhancing both interpretability and reliability. T. Zeng *et al.* (2025) proposed the multi-source feature fusion case-based reasoning (MSFF-CBR) model for diagnostic applications involving multi-modal data. While successful in probabilistic fusion, the model devoted less attention to weight interpretability and frequency-based adaptation. The methodology discussed here strengthens this dimension by introducing weight calibration based on historical failure data and component-level impact.

This comparative review highlighted that while many modern approaches bring valuable innovations such as hybrid neural-CBR models, ontology integration, or time-series analytics only a few effectively combine probabilistic reasoning, adaptive weights, support for heterogeneous data, and interpretability. The presented model achieved such integration through a structured framework based on: configurable partial similarity metrics (Euclidean, Jaccard, logistic); gradient-based optimisation of weights using the L-BFGS-B method; Bayesian aggregation for probabilistic inference; support for structured numerical, categorical, and set-valued data. This methodological composition is particularly well-suited for complex environments such as SPPs, where uncertainty, data sparsity, and multi-source features are the norm. The architecture also enabled modular extensions to additional components or operational parameters without requiring full retraining, which is beneficial for maintenance and scalability. Furthermore, the approach aligned with current trends in intelligent

diagnostics, which emphasised interpretability and adaptability over opaque black-box learning. While conventional systems often rely on fixed similarity metrics or rigid neural models, the model discussed in the present study prioritised traceability, domain adaptability, and evidence-based decision-making. The resulting advantages can be summarised as follows: data comprehensiveness: simultaneous handling of numeric, categorical, and set-valued features; adaptability: numerical optimisation of similarity weights through formal procedures; probabilistic robustness: inference supported by Bayesian aggregation and normalisation; scalability: application across diverse diagnostic domains with minimal reconfiguration. In summary, the proposed model constitutes a substantial advancement in the domain of intelligent similarity-based diagnostics. It responded directly to limitations identified in current literature and offers a reliable, interpretable, and extensible solution suitable for decision support systems operating under conditions of complexity and uncertainty.

CONCLUSIONS

The developed adaptive similarity metric for failure diagnostics in SPPs demonstrated significant potential for enhancing both the accuracy and interpretability of intelligent decision support systems. The key contributions of this study included a comprehensive approach to similarity assessment by integrating heterogeneous data types: numerical; categorical; set-based through a weighted combination of metrics such as Euclidean, Jaccard, and logistic. The model employed automatic optimisation of weight coefficients using the L-BFGS-B method, which ensured adaptability to varying operational conditions. Interpretability and robustness were improved through the introduction of weight correction mechanisms that accounted for failure frequency, temporal degradation (based on operating time), and contextual smoothing across components. These enhancements allowed the model to effectively handle rare yet critical failure scenarios. Additionally, the application of non-linear activation functions (tanh) and Bayesian hypothesis aggregation contributed to more accurate classification results and reduced sensitivity to noisy and incomplete data. Experimental evaluation confirmed the practical effectiveness of the proposed model. The classification accuracy exceeded 90% for

all three main types of failures (A, B, and C), with Type A achieving 96% correct identification, Type B – 84%, and Type C – 92%, as shown in the confusion matrix. The optimised model also demonstrated strong robustness under conditions of limited training data. For instance, reducing the training set size by 40% led to only a 5-7% drop in classification performance, confirming the model's data efficiency. The visualisation of feature spaces, impact of parameter weighting, and temporal degradation coefficients illustrated the model's ability to separate nonlinear and overlapping classes while preserving explainability.

Compared to conventional methods with static metrics, the proposed approach successfully combined flexibility through adaptability to heterogeneous data with explainability, via Bayesian reasoning and tunable parameters. Its results were comparable to those achieved by ultramodern hybrid techniques such as graph neural networks, but with lower computational requirements and simpler implementation. The developed methodology was considered suitable for integration into CBR-based diagnostic systems aimed at automating the identification of failure types, forecasting the time to critical events such as component degradation, and optimising maintenance strategies based on historical case analysis. Future research should focus on validating the model using real-world failure datasets with broader diversity, extending the approach to dynamic temporal models that reflect evolving technical conditions, and integrating the solution into IoT-enabled platforms for real-time equipment monitoring. Overall, the proposed model constituted a valuable contribution to the field of intelligent diagnostics, offering a balance between theoretical soundness and practical applicability, while its modular and adaptive structure allowed for extension to other classes of complex technical systems.

ACKNOWLEDGEMENTS

None.

FUNDING

None.

CONFLICT OF INTEREST

None.

REFERENCES

- [1] Amin, K., Kapetanakis, S., Polatidis, N., Althoff, K.-D., & Dengel, A. (2020). DeepKAF: A heterogeneous CBR & deep learning approach for NLP prototyping. In *2020 International conference on innovations in intelligent systems and applications (INISTA)* (pp. 1-7). Novi Sad. doi: [10.1109/INISTA49547.2020.9194679](https://doi.org/10.1109/INISTA49547.2020.9194679).
- [2] Chen, M., Qu, R., & Fang, W. (2022). Case-based reasoning system for fault diagnosis of aero-engines. *Expert Systems with Applications*, 202, article number 117350. doi: [10.1016/j.eswa.2022.117350](https://doi.org/10.1016/j.eswa.2022.117350).
- [3] Chen, Z., Xu, J., Alippi, C., Ding, S.X., Shardt, Y., Peng, T., & Yang, C. (2021). Graph neural network-based fault diagnosis: A review. *ArXiv*. doi: [10.48550/arXiv.2111.08185](https://doi.org/10.48550/arXiv.2111.08185).
- [4] Dubchak, L., Sachenko, A., Bodyanskiy, Y., Wolff, C., Vasylykiv, N., Brukhanskyi, R., & Kochan, V. (2024). Adaptive neuro-fuzzy system for detection of wind turbine blade defects. *Energies*, 17(24), article number 6456. doi: [10.3390/en17246456](https://doi.org/10.3390/en17246456).

- [5] Jimenez-Diaz, G., & Díaz-Agudo, B. (2024). Visualization of similarity models for CBR comprehension and maintenance. In *Case-based reasoning research and development (ICCBR 2024)* (pp. 67-80). Merida: Springer. doi: [10.1007/978-3-031-63646-2_5](https://doi.org/10.1007/978-3-031-63646-2_5).
- [6] Krüger, M. (2023). [Comparative analysis of text-based CBR algorithms for cybercrime profiling investigations](#). In *Lernen, wissen, daten, analysen (LWDA) 2023* (pp. 359-371). Marburg: CEUR Workshop Proceedings.
- [7] Leake, D., Ye, X., & Crandall, D. (2021). [Supporting case-based reasoning with neural networks: An illustration for case adaptation](#). In *AAAI 2021 spring symposium on combining machine learning and knowledge engineering (AAAI-MAKE 2021)*. California: Palo Alto.
- [8] Lin, N., Liu, H., Fang, J., Zhou, D., & Yang, A. (2023). An interpretability framework for similar case matching. *ArXiv*. doi: [10.48550/arXiv.2304.01622](https://doi.org/10.48550/arXiv.2304.01622).
- [9] Mathisen, B.M., Aamodt, A., Bach, K., & Langseth, H. (2020). Learning similarity measures from data. *ArXiv*. doi: [10.48550/arXiv.2001.05312](https://doi.org/10.48550/arXiv.2001.05312).
- [10] Mustyala, S., & Bisi, M. (2025). Ensembling harmony search algorithm with case-based reasoning for software development effort estimation. *Cluster Computing*, 28, article number 97. doi: [10.1007/s10586-024-04858-w](https://doi.org/10.1007/s10586-024-04858-w).
- [11] Neykov, N., & Stefanova, S. (2023). Using case-based reasoning in system diagnostics and maintenance. In *Proceedings of seventh international congress on information and communication technology* (pp. 345-359). Singapore: Springer. doi: [10.1007/978-981-19-2394-4_28](https://doi.org/10.1007/978-981-19-2394-4_28).
- [12] Ren, S., Gui, F., Zhao, Y., Zhan, M., & Wang, W. (2020). An effective similarity determination model for case-based reasoning in support of low-carbon product design. *Advances in Mechanical Engineering*, 12(10). doi: [10.1177/1687814020970313](https://doi.org/10.1177/1687814020970313).
- [13] Serrà, J., & Arcos, J.LL. (2014). An empirical evaluation of similarity measures for time series classification. *Knowledge-Based Systems*, 67, 305-314. doi: [10.1016/j.knosys.2014.04.035](https://doi.org/10.1016/j.knosys.2014.04.035).
- [14] Valdez-Ávila, M.F., Bermejo-Sabbagh, C., Díaz-Agudo, B., Orozco-del-Castillo, M.G., & Recio-Garcia, J.A. (2023). CBR-fox: A case-based explanation method for time series forecasting models. In *Case-based reasoning research and development. ICCBR 2023* (pp. 192-207). Cham: Springer. doi: [10.1007/978-3-031-40177-0_13](https://doi.org/10.1007/978-3-031-40177-0_13).
- [15] Virtanen, P., et al. (2020). SciPy 1.0: Fundamental algorithms for scientific computing in Python. *Nature Methods*, 17(3), 261-272. doi: [10.1038/s41592-019-0686-2](https://doi.org/10.1038/s41592-019-0686-2).
- [16] Vychuzhanin, V., & Vychuzhanin, A. (2025). *Stochastic models and methods for diagnostics, assessment, and prediction of the technical condition of complex critical systems*. Lviv-Torun: Liha-Pres. doi: [10.36059/978-966-397-457-6](https://doi.org/10.36059/978-966-397-457-6).
- [17] Xu, H., Wei, Y., Cai, Y., & Xing, B. (2023). Knowledge graph and CBR-based approach for automated analysis of bridge operational accidents: Case representation and retrieval. *PLoS ONE*, 18(11), article number e0294130. doi: [10.1371/journal.pone.0294130](https://doi.org/10.1371/journal.pone.0294130).
- [18] Ye, J. (2017). Single-valued neutrosophic similarity measures based on cotangent function and their application in the fault diagnosis of steam turbine. *Soft Computing*, 21(3), 817-825. doi: [10.1007/s00500-015-1818-y](https://doi.org/10.1007/s00500-015-1818-y).
- [19] Ye, X., Leake, D., Wang, Y., Zhao, Z., & Crandall, D. (2024). Towards network implementation of CBR: Case study of a neural network K-NN algorithm. In *Case-based reasoning research and development. ICCBR 2024. Lecture notes in computer science* (pp. 354-370). Cham: Springer. doi: [10.1007/978-3-031-63646-2_23](https://doi.org/10.1007/978-3-031-63646-2_23).
- [20] Zeng, T., Bao, R., Qin, Y., Sun, X., Gao, Y., Cheng, L., Hou, P., Sang, H., Ma, L., & Zhou, X. (2025). MSFF-CBR: Case-based reasoning technology for adaptive multi-information fusion fault diagnosis. *Measurement Science and Technology*, 36(4), article number 045111. doi: [10.1088/1361-6501/adc474](https://doi.org/10.1088/1361-6501/adc474).
- [21] Zuber, M., & Sirdey, R. (2021). Efficient homomorphic evaluation of k-NN classifiers. *Proceedings on Privacy Enhancing Technologies 2021*, 2, 111-129. doi: [10.2478/popets-2021-0020](https://doi.org/10.2478/popets-2021-0020).

Адаптивна метрика оцінки схожості для інтелектуальної діагностики відмов у суднових енергетичних установках

Володимир Вичужанін

Доктор технічних наук, професор
Національний університет «Одеська політехніка»
65044, просп. Шевченка, 1, м. Одеса, Україна
<https://orcid.org/0000-0002-6302-1832>

Олексій Вичужанін

Доктор філософії, асистент
Національний університет «Одеська політехніка»
65044, просп. Шевченка, 1, м. Одеса, Україна
<https://orcid.org/0000-0001-8779-2503>

Анотація. Своєчасна та точна діагностика несправностей в суднових енергетичних установках (СЕУ) є надзвичайно важливою для забезпечення безпеки на морі, мінімізації експлуатаційних ризиків та оптимізації стратегій технічного обслуговування. Зі збільшенням складності систем, гетерогенності джерел даних та обмеженості історичних записів про несправності традиційні методи діагностики часто виявляються недостатніми, особливо в ситуаціях, пов'язаних з рідкісними або неоднозначними несправностями. Метою цього дослідження була розробка інтерпретованої, адаптивної та ймовірно обґрунтованої методології для оцінки схожості випадків несправностей у СЕУ для використання в інтелектуальних системах підтримки прийняття рішень. Запропонований метод інтегрує евклідову, джакардову та логістичну метрики схожості з байєсівським виведенням, моделюванням тимчасової деградації, корекцією ваги на основі частоти та контекстним згладжуванням уражених підсистем. Модель використовує оптимізацію L-BFGS-B для автоматичного коригування ваги метрик відповідно до діагностичної релевантності. Чисельні експерименти на основі синтетичних даних випадків показали високу точність класифікації: 96 % для відмов, пов'язаних з перегрівом системи охолодження, 84 % для випадків деградації підшипників і 92 % для порушень у подачі палива. Навіть при 40 % скороченні обсягу навчальних даних падіння продуктивності не перевищило 7 %, що свідчить про високу стійкість до розрідженості даних. Візуалізація меж прийняття рішень продемонструвала здатність моделі розрізняти перекривні класи відмов, зберігаючи семантичну інтерпретованість. Результати оптимізації ваги визначили «Тип відмови» як домінуючий фактор, тоді як «Категорія ризику» та «Уражені підсистеми» мали незначний вплив і були виключені. Байєсівське агрегування ще більше підвищило надійність діагностичних висновків, поєднавши локальну схожість із глобальними статистичними апіорними даними. Розроблена методологія може бути ефективно застосована морськими інженерами, операторами суден та розробниками інтелектуальних діагностичних платформ для виявлення несправностей, аналізу першопричин та прогнозного технічного обслуговування в умовах невизначеності та неповної інформації. Її модульна структура також дозволяє розширити її застосування на інші складні технічні галузі, крім СЕУ

Ключові слова: система підтримки прийняття рішень; байєсівська агрегація; адаптивні вагові коефіцієнти; прецедентний аналіз відмов; ймовірна метрика; схожість технічних станів