

Degradation index-based prediction for remaining useful life using multivariate sensor data

Kang, Wenda; Jongbloed, Geurt; Tian, Yubin; Chen, Piao

DOI

[10.1002/qre.3615](https://doi.org/10.1002/qre.3615)

Publication date

2024

Document Version

Final published version

Published in

Quality and Reliability Engineering International

Citation (APA)

Kang, W., Jongbloed, G., Tian, Y., & Chen, P. (2024). Degradation index-based prediction for remaining useful life using multivariate sensor data. *Quality and Reliability Engineering International*, 40(7), 3709-3728. <https://doi.org/10.1002/qre.3615>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

Green Open Access added to TU Delft Institutional Repository

'You share, we take care!' - Taverne project

<https://www.openaccess.nl/en/you-share-we-take-care>

Otherwise as indicated in the copyright section: the publisher is the copyright holder of this work and the author uses the Dutch legislation to make this work public.

Degradation index-based prediction for remaining useful life using multivariate sensor data

Wenda Kang¹  | Geurt Jongbloed¹ | Yubin Tian² | Piao Chen³

¹Department of Applied Mathematics,
Delft University of Technology, Delft,
Netherlands

²School of Mathematics and Statistics,
Beijing Institute of Technology, Beijing,
China

³ZJUI Institute, Zhejiang University,
Haining, China

Correspondence

Piao Chen, ZJUI Institute, Zhejiang
University, Haining, China.
Email: piaochen@intl.zju.edu.cn

Funding information

National Natural Science Foundation of
China, Grant/Award Number: 12131001

Abstract

The prediction of remaining useful life (RUL) is a critical component of prognostic and health management for industrial systems. In recent decades, there has been a surge of interest in RUL prediction based on degradation data of a well-defined degradation index (DI). However, in many real-world applications, the DI may not be readily available and must be constructed from complex source data, rendering many existing methods inapplicable. Motivated by multivariate sensor data from industrial induction motors, this paper proposes a novel prognostic framework that develops a nonlinear DI, serving as an ensemble of representative features, and employs a similarity-based method for RUL prediction. The proposed framework enables online prediction of RUL by dynamically updating information from the in-service unit. Simulation studies and a case study on three-phase industrial induction motors demonstrate that the proposed framework can effectively extract reliability information from various channels and predict RUL with high accuracy.

KEYWORDS

degradation index, electrical motors, multivariate sensor data, prognostics, remaining useful life

1 | INTRODUCTION

1.1 | Background and motivation

The prediction of remaining useful life (RUL) is crucial for complex systems such as electrical systems and has become an increasingly popular research topic in recent years.¹ The RUL refers to the time remaining until a system can no longer perform its intended function, and accurate RUL prediction is essential for ensuring system safety and reliability, minimizing maintenance costs, and maximizing its lifespan. Modern sensor technology has facilitated the collection of multivariate sensor data, allowing real-time monitoring of a system's health status. These multivariate data provide valuable information on the system's performance, which can be used to predict the RUL. Therefore, developing effective methods for analyzing and processing multivariate sensor data is critical for accurate RUL prediction.

The major challenge in RUL prediction based on multivariate sensor data is the inability of any one-dimensional signal to fully capture the variation of RUL.² One example of such data is the multivariate three-phase industrial induction motor data used in our case study. Figure 1 shows an example of a motor with 11 channels of raw signals including current, voltage, and temperature presented as a function of the experimental period (i.e., the cycles shown in Figure 1), where the true value of RUL is also available. As seen, most signals do not exhibit a significant trend during the experiment, making

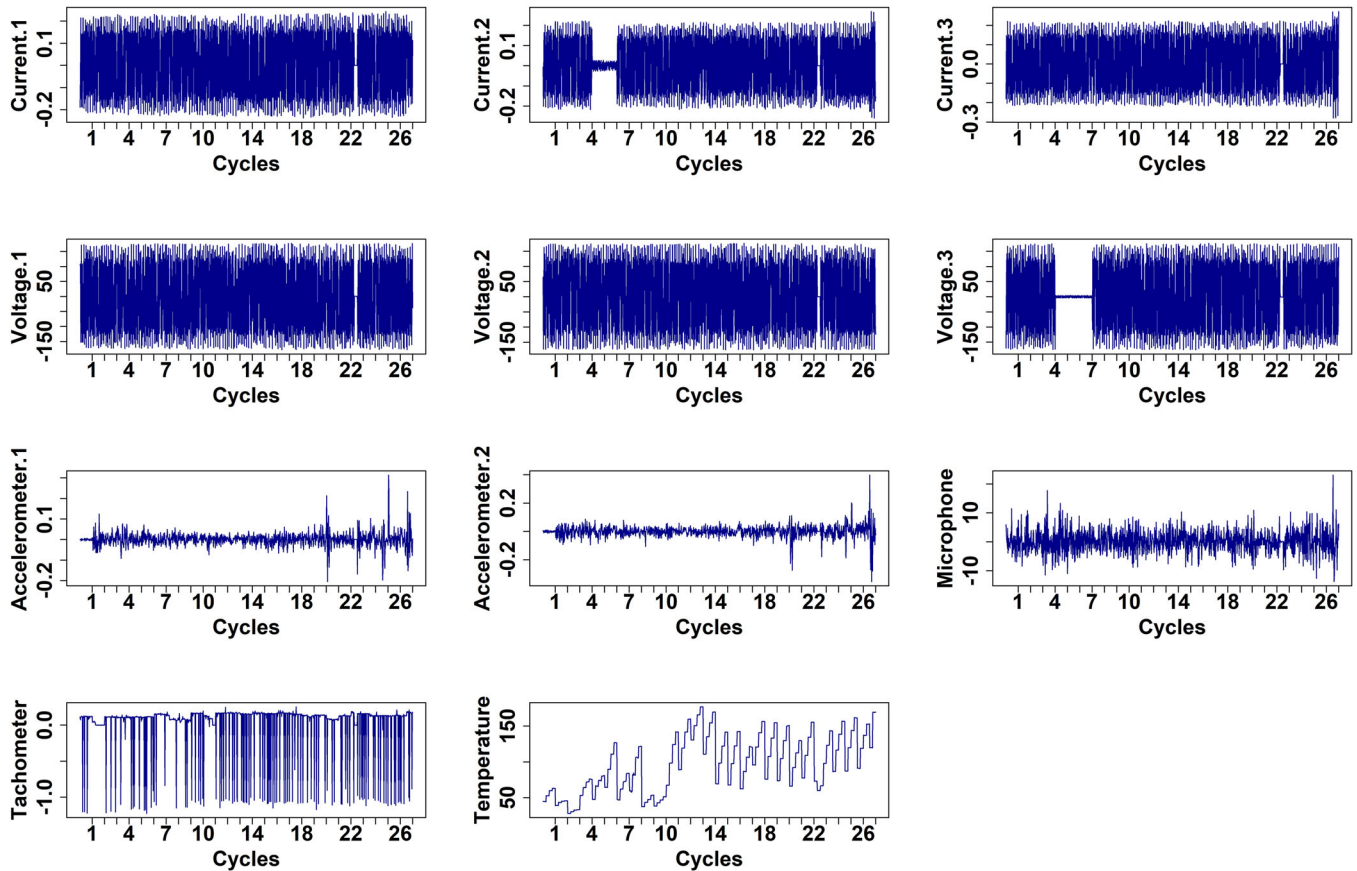


FIGURE 1 Example of multivariate sensor data for a three-phase induction motor.

them unsuitable for direct use in RUL prediction. This challenge demonstrates the need for effective techniques to extract relevant information from multiple sensor data for accurate RUL prediction. In the following subsection, we provide a comprehensive review of the existing literature on RUL prediction based on multivariate sensor data, covering the three main areas of research: sensor fusion techniques, degradation index (DI) methods, and RUL prediction models.

1.2 | Literature review

1.2.1 | Sensor fusion

In the field of degradation modeling for complex systems that collect multivariate sensor data, an effective fusion of the sensor data is a critical task. Existing fusion methods for multivariate sensor data can be broadly categorized into three groups: signal-level, feature-level, and decision-level.^{3–5}

Signal-level fusion involves the direct integration of all raw sensor signals. For instance,⁶ directly fuses multisensor data using a 2-D convolutional neural network and applies several artificial intelligence (AI) methods to the fused data for fault detection and diagnosis of gearboxes. However, signal-level fusion requires caution since sensor recordings may have different acquisition, pre-filtering, and amplification settings, and raw data fusion often requires commensurate data as input.³

Feature-level fusion predicts the health status by combining extracted features from the data of each raw sensor. This approach has been widely used due to its simplicity and effectiveness. For example, ref. [7] proposes an RUL prediction method by performing a gated recurrent unit network on the extracted nonlinear features generated using kernel principal component analysis.⁸ proposes an integrated deep multiscale feature fusion network for aero engine RUL prediction using multisensor data, and they integrate features extracted from the convolutional neural network and gated recurrent unit network.

The third category, decision-level fusion, involves integrating the decisions made from independent analyses of multivariate sensor data, such as fault diagnosis, RUL prediction, or other types of analysis tasks. For example,⁹ develops a decision-level fusion method by combining the high-dimensional decisions transformed from low-dimensional decisions made based on individual sensor data.¹⁰ proposes a decision-level method for multisensor fusion for collaborative fault diagnosis by using an enhanced voting fusion strategy. However, this approach is a post-processing technique that heavily depends on the quality of the raw data and is highly sensitive to the decision fusion rules, limiting its practical applications.¹¹

1.2.2 | Degradation index construction

The reviewed multivariate sensor data fusion methods in Section 1.2.1 share a common drawback: the absence of a univariate index that credibly reflects the underlying degradation process. While some of these methods employ the raw sensor data or extracted features as input to different AI models, there is still no satisfactory fused indicator that meets requirements such as monotonicity, smoothness, and maximum range information.¹² As a result, these approaches tend to be less interpretable with respect to the underlying degradation process, making many existing statistical methods inapplicable. Consequently, the construction of an informative univariate index, or the DI as referred to in this paper, is a crucial step towards describing the underlying degradation process based on multivariate sensor data.^{13,14}

Several methods have been proposed for constructing the DI. For example,¹³ proposes a method to construct the DI by fusing multi-sensor data at the signal level and using the resulting DI for the degradation modeling of an aircraft gas turbine engine. Subsequent work has been done by refs. [12, 15–18]. In particular, ref. [19] presents a DI building method for multivariate sensor data with censoring, which can automatically select informative sensor signals using the group least absolute shrinkage and selection operator penalty. However, these methods may not be suitable for all practical cases as they assume there should be a trend in some raw signals, which may not be the case where only extracted features show such trends. Furthermore, these methods are all focused on raw sensor data and may be time-consuming for high-dimensional feature spaces. In addition,¹⁹ also notes that existing methods cannot perform automatic variable selection, and the DI and variable selection procedure in their own work lacked an explanation for the contribution of each sensor.

1.2.3 | RUL prediction

To accurately predict the RUL, it is necessary to establish a precise correlation between the constructed DI and RUL. Univariate DI-based RUL prediction methods typically fall into three categories: physical-based, data-driven, and hybrid approaches.^{20,21} Physical-based methods require a thorough understanding of the degradation behavior based on failure mechanisms, which can be challenging to obtain for complex systems. Conversely, data-driven methods have gained attention in recent years due to their mechanism-agnostic approach, which infers the health status of products from monitored degradation signals. Hybrid methods combine both physical-based and data-driven methods, but their effectiveness may be limited by the difficulty of obtaining accurate failure mechanisms for complex systems.

Data-driven methods can be further classified into statistical and AI methods.²² Statistical methods based on the Wiener process, Gamma process, and inverse Gaussian process have been widely used. Examples and applications can be found in refs. [23–28] and references therein. However, these stochastic process methods have some strong assumptions such as the Markov property, and are also prone to model misspecification problems, which limit their application in engineering.²⁹ In contrast, AI methods are not affected by these limitations.³⁰ Among them, the similarity-based method is widely used for DI-based RUL prediction due to its intuitive and interpretable nature.³¹ Further examples can be found in the review paper.³²

1.3 | Objective and overview

Based on the literature review, the issues of existing methods can be summarized as follows. Although direct mapping of multivariate sensor data and RUL is possible, DI-based methods are often more intuitive and explainable. However, existing DI-based methods mainly focus on cases where the raw sensor data have significant trends, which is not suitable

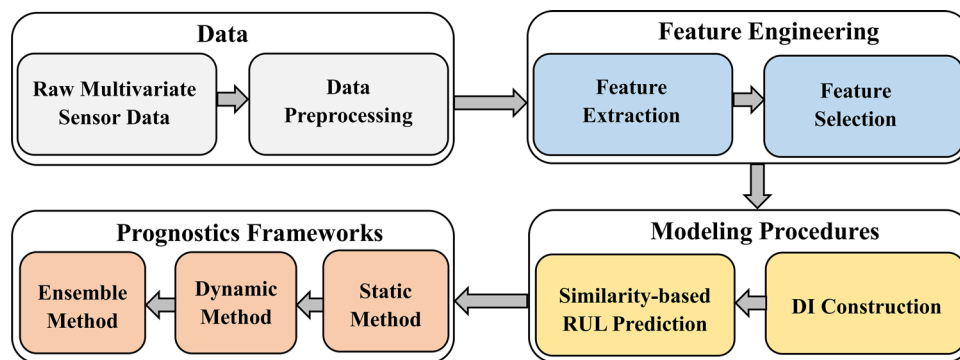


FIGURE 2 The basic procedures of the developed prognostics frameworks.

for many applications, as demonstrated in Figure 1. Despite the inclusion of feature engineering procedures, existing DI-based methods may still lack efficiency and applicability in high-dimensional feature spaces. Additionally, the accuracy of existing DI-based methods for RUL prediction heavily relies on the form of DI and the sample size of the training dataset, limiting their usefulness in certain applications.

Motivated by the above-mentioned issues, this paper proposes DI-based prognostic frameworks for predicting RUL in complex systems. In contrast to existing DI-based methods, our constructed DI is feature-based and performs automatic feature selection, which is essential for accurately capturing the underlying degradation trend of the system. Furthermore, our proposed DI incorporates a nonlinear relationship between the selected features and the degradation process, better reflecting the complex and nonlinear nature of practical engineering applications. Based on the constructed DI and similarity matching method, we have developed three frameworks for prognostic RUL prediction of complex systems that collect multivariate sensor data. These frameworks are designed to overcome the challenges of accurately predicting RUL. The basic principles of these frameworks are illustrated in Figure 2, highlighting the importance of data preprocessing, feature extraction and selection, DI construction, and RUL prediction.

The main contributions of this paper are summarized as follows:

- (1) Introduction of feature-level prognostics frameworks for DI-based RUL prediction, which can also automatically select informative features.
- (2) Development of a nonlinear form of DI to amalgamate representative features extracted from collected multivariate sensor data.
- (3) Proposal of an ensemble approach that stably integrates common and individual features.

The remainder of this paper is organized as follows. In Section 2, we provide details of the feature engineering process for multivariate sensor data and the method used to construct a DI. The procedure for deriving similarity-based RUL and quantifying the uncertainty of predictions is presented in Section 3. We then illustrate the developed prognostics frameworks for DI-based RUL in Section 4. In Section 5, we conduct a simulation study to investigate the performance of the proposed frameworks. In Section 6, we provide a case study based on real-world induction motors degradation data. Finally, we give some concluding remarks and discussions in Section 7.

2 | DEGRADATION INDEX CONSTRUCTION

2.1 | Feature engineering

Raw sensor data typically consists of time series data with a fixed sampling frequency. However, analyzing the data at each time point can be computationally expensive and may not yield useful information. To address this challenge, feature extraction is commonly used to generate features from the raw time series that accurately describe the data while reducing computational costs.^{1,20} Additionally, feature selection can be employed to select the most informative subset of features, as not all extracted features may be useful. Therefore, feature extraction and feature selection techniques are crucial for exploring useful information and reducing computational costs.

TABLE 1 Extracted features from raw sensor signals.

Feature	Description	Equation
p_1	Average amplitude	$\frac{1}{k} \sum_{i=1}^k h(i)$
p_2	Standard deviation	$\left(\frac{\sum_{i=1}^k (h(i)-p_1)^2}{k-1} \right)^{1/2}$
p_3	Root mean square amplitude	$\left(\frac{1}{k} \sum_{i=1}^k h(i)^2 \right)^{1/2}$
p_4	Squared mean rooted absolute amplitude	$\left(\frac{1}{k} \sum_{i=1}^k h(i) ^{1/2} \right)^2$
p_5	Kurtosis coefficient	$\frac{\sum_{i=1}^k (h(i)-p_1)^4}{(k-1)p_2^4}$
p_6	Skewness coefficient	$\frac{\sum_{i=1}^k (h(i)-p_1)^3}{(k-1)p_2^3}$
p_7	Peak value	$\max h(i) $
p_8	Peak factor	$\frac{p_7}{p_3}$
p_9	Margin factor	$\frac{p_7}{p_4}$
p_{10}	Waveform factor	$\frac{p_3}{\frac{1}{k} \sum_{i=1}^k h(i) }$
p_{11}	Impulse factor	$\frac{p_7}{\frac{1}{k} \sum_{i=1}^k h(i) }$

In this study, we focus on investigating time-domain feature extraction techniques. Specifically, we employ the time domain features used in previous works such as refs. [1] and [20]. The details of the extracted features are presented in Table 1, where h represents a time series with a length of k . Among all the features, (p_1, p_3, p_4, p_7) are used to capture the amplitude and energy of each signal, while the remaining ones reflect the distribution of the signal over the time domain. Note that the features listed in Table 1 differ for each signal and k denotes the total length of the signal.

Since noisy features can impede modeling accuracy, feature selection is often utilized to retain the most important subset of features. In many existing works on DI construction, Fisher's discriminant ratio is used as a criterion for feature selection.²⁹ This ratio can be formulated as follows:

$$S_F(X_j) = \frac{(\mu_{j,1} - \mu_{j,2})^2}{\sigma_{j,1}^2 + \sigma_{j,2}^2}, \quad (1)$$

where $\mu_{j,c}$ and $\sigma_{j,c}^2$ are the mean and variance of feature X_j within the healthy ($c = 1$) or unhealthy ($c = 2$) class. To determine these two classes, we follow the approach used in refs. [1] and [20], which assumes that the first few cycles are relatively healthy and the last few cycles become faulty.

Fisher's discriminant ratio method only identifies unit-specific informative features rather than general informative features across multiple reference units. To address this limitation, we propose a new feature selection method that selects the most common informative features across reference units. First, we calculate Fisher's discriminant ratios for all features in Table 1 of each unit and sort them in descending order. We then select a certain percentage of features with the highest ratios. In practice, this percentage can be chosen by engineering background or cross-validation. In this paper, the top 50% is used for a fair comparison, which is consistent with refs. [1] and [20]. Next, we generate the most frequent features by selecting a certain percentage of the reference units (e.g., 5 out of 7 units in our case study) and taking the intersection of features from each subset. Finally, we select the union of these intersections as the set of selected features. This approach improves the robustness and generalization of our feature selection method. The entire feature engineering process is presented in Algorithm 1.

2.2 | Constructing degradation index

After feature engineering, the next step is to construct a suitable DI based on these selected features. Let p be the number of selected features from the raw sensor data, and $\mathbf{x}(s) = [x_1(s), x_2(s), \dots, x_p(s)]$ be the corresponding features generated at operational time s .

ALGORITHM 1 The overall process of feature engineering.

Input : Multivariate sensor signal data for all reference units after preprocessing.

Output: The general informative features of most reference units.

```

1 for each unit do
2   Calculate the Fisher's discriminant ratios for all features according to Table 1;
3   Sort these ratios in descending order;
4   Take out a certain percentage of the top-ranked features;
5 Take a certain percentage of the reference units as subsets;
6 Find the intersection of features for each subset;
7 return The union of these intersections.
```

To construct the DI, we employ the cumulative damage model,^{2,19} which assumes that the degradation of a system accumulates over time and is widely used in many engineering systems. Specifically, the DI $Z(t)$ is defined as follows:

$$Z(t) = \int_0^t u(s)ds, \quad (2)$$

where $u(s)$ is constructed by the selected features following the additive model

$$u(s) = \sum_{j=1}^p \beta_j f_j[x_j(s), \phi_j], \quad (3)$$

where β_j is the parameter reflecting the contribution of the j th feature, and $f_j[x_j(s), \phi_j]$ is the corresponding effect function. To make $f_j[x_j(s), \phi_j]$ flexible to capture potential nonlinear patterns of the features, we adopt a linear combination of spline basis

$$f_j[x_j(s), \phi_j] = \sum_{l=1}^L \phi_{jl} bs_{jl}(s), \quad (4)$$

where j is the feature index, $bs_{jl}(s)$, $l = 1, \dots, L$ are the B-spline basis functions, L is the number of degrees of freedom, and ϕ_{jl} are the corresponding weight coefficients which can be derived by fitting the j th feature. The B-spline is used due to its simplicity of computation and wide applicability. More details can be found in ref. [33].

Regarding the properties to construct the DI, we employ the following three widely used properties^{13,15–19}:

- Monotonic degradation trend*: The trend of a constructed DI is assumed to be monotonic, showing a clear increasing or decreasing degradation trend. Without loss of generality, we assume that the DI is monotonically increasing in this work.
- Consistent initial status*: In practice, the initial states of different units are often assumed to be the same, which can be achieved by setting the initial value of the constructed DI to 0.
- Maximized range information*: The range information of the constructed DI starting from the initial to the failure time point should be maximum to guarantee a clear degradation trend.

The constructed $Z(t)$ naturally satisfies the first two properties, that is, $Z(0) = 0$ and $Z(t)$ is monotonically increasing. As pointed out in refs. [12] and [18], the maximum range information ensures that the range of the constructed DI (i.e., initial degradation level to the level of failure) is maximized, providing a clear degradation trend. To achieve this, we propose the following unconstrained optimization problem to determine the parameters $\beta = (\beta_1, \dots, \beta_p)$,

$$\min_{\beta} (1 - \lambda)R(\beta) + \lambda \|\beta\|_1, \quad (5)$$

where $R(\beta) = 1/\min_i Z_i$ with Z_i being the DI value of the i th reference unit at the failed time, $\lambda \in [0, 1]$ is a tuning parameter which can be determined by cross-validation, and $\|\cdot\|_1$ is the L_1 norm. Note that minimizing $R(\beta)$ maximizes the

overall range of the reference units. Moreover, by incorporating the lasso penalty, the proposed method automatically performs feature selection, which is a critical step often overlooked in DI-related studies. Because the objective function (5) is highly complex, it is recommended to use heuristic optimization algorithms³⁴ such as simulated annealing for optimization.

3 | RUL PREDICTION

3.1 | Similarity-based RUL

With the constructed DI, we propose a similarity-based method for RUL prediction. The similarity-based prediction method is a popular data-driven approach that is widely used in RUL prediction because it does not require pre-knowledge of failure mechanisms or specific degradation models.^{31,32} The basic principle is to compare the DI of a test unit with those of reference units at specific time points. If they are similar, the RULs of the test and reference units should also be similar. The commonly used Euclidean distance is adopted in this paper to measure the similarity between the DI of the test unit and the DIs of the reference units.^{31,32}

Because the length of DIs for the test and reference units are often different, the distance cannot be calculated directly. Therefore, a reconstruction of the DIs is necessary to ensure that the test and reference units share the same length of DI at the reconstructed segments. Assuming there are n failed reference units, the DI of the i th reference unit is $\mathbf{Z}_i = (Z_{i,1}, \dots, Z_{i,n_i})$, where $i = 1, 2, \dots, n$ and $Z_{i,j}$ is the value of the DI for the i th reference unit at time j . Note that the n_i values are integers, representing the cycle number at which the i th reference unit failed. Let the DI of the test unit be $\mathbf{Z}_{T,t} = (Z_{T,1}, \dots, Z_{T,t})$, where t is the current operating time, which is also an integer. Then, the DI of the i th reference unit is reconstructed into $n_i - t + 1$ segments $\{\mathbf{Z}_{i,1}, \dots, \mathbf{Z}_{i,j}, \dots, \mathbf{Z}_{i,n_i-t+1}\}$, where $\mathbf{Z}_{i,j} = (Z_{i,j}, \dots, Z_{i,j+t-1})$, $j = 1, \dots, n_i - t + 1$. The Euclidean distance between $\mathbf{Z}_{T,t}$ and $\mathbf{Z}_{i,j}$ is calculated by

$$d_{i,j,t} = \|\mathbf{Z}_{T,t} - \mathbf{Z}_{i,j}\|_2, \quad (6)$$

where $\|\cdot\|_2$ is the L_2 norm.

The most similar segment from the i th reference unit is then selected by using the smallest distance $d_{i,t} = \min_j d_{i,j,t}$. Let k be the index number of the most similar segment, and $\mathbf{Z}_{i,k} = (Z_{i,k}, \dots, Z_{i,k+t-1})$. The corresponding predicted RUL at time t based on the i th reference unit is derived as

$$\text{RUL}_{i,t} = n_i - (k + t - 1). \quad (7)$$

Consequently, the similarity-based RUL of the test unit at time t is estimated by a weighted sum of the RULs derived from all the reference units, which can be expressed as

$$\text{RUL}_t = \sum_{i=1}^n \omega_{i,t} \text{RUL}_{i,t}, \quad (8)$$

where $\omega_{i,t} := \frac{1/d_{i,t}}{\sum_{i=1}^n 1/d_{i,t}}$ is the weight for the i th reference unit at time t . Note that it is reasonable to use $\omega_{i,t}$ as the weight since a smaller value of $d_{i,t}$ indicates stronger similarity and $\sum_{i=1}^n \omega_{i,t} = 1$.³²

3.2 | Prediction interval for RUL

In practice, the prediction interval is often more valuable and can be used to measure the uncertainty in prediction. To calculate the prediction interval of similarity-based RUL, we propose a method that applies the bootstrap method after constructing the DIs of all units. The basic idea is to resample with replacement from the DIs of the reference units and then derive the prediction for the similarity-based RUL for the test unit using the procedures described in Section 3.1 based on the resampled data. This procedure is repeated M times to derive M predicted RULs. Finally, the corresponding prediction interval can be obtained by using the empirical percentiles of the M predicted RULs. For example, the 95% prediction interval can be calculated using the algorithm described in Algorithm 2 by setting $\gamma = 0.05$.

ALGORITHM 2 Procedures of constructing the $100(1 - \gamma)\%$ prediction interval for RUL.

- Input** : Constructed DIs of reference units and the DI of the test unit at time t .
Output: The $100(1 - \gamma)\%$ prediction interval for the RUL of the test unit.
- 1 **for** $m \leftarrow 1$ **to** M **do**
 - 2 Resample with replacement from the construct DIs of reference units, and the sampling size is consistent with the number of reference units;
 - 3 Derive the most similar segments for the DI of the test unit based on the resampled DIs and (6);
 - 4 Determine the similarity-based RUL of the test unit at time t using (7) and (8);
 - 5 Calculate the $\gamma/2$ and $1 - \gamma/2$ quantiles of the M predicted RULs;
 - 6 **return** The $\gamma/2$ and $1 - \gamma/2$ quantiles of the M predicted RULs.

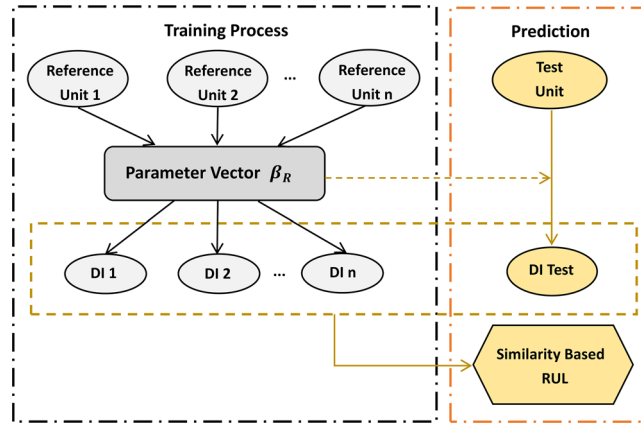


FIGURE 3 Flowchart of the static prognostics.

4 | FRAMEWORKS FOR RUL PREDICTION

With the identified features, we can compute the similarity-based RUL for the test unit using the DIs constructed in Section 2 and the prediction procedure outlined in Section 3. In this section, we introduce three frameworks for RUL prediction. The first framework considers only the information from the reference units, the second framework dynamically incorporates information from both the test and reference units, and the third framework is an ensemble of the first two methods.

4.1 | The static framework

For the collected multivariate sensor signal data, a straightforward approach is to train the parameter vector β in (5) based on the selected common features from the reference units and then use this parameter vector to predict the RUL for the test unit. We refer to this method as a static method since it relies solely on the information from reference units to derive the parameter vector β .

To determine the DIs for the reference and test units, the common features are first selected from the raw sensor data of the reference units using the feature engineering method outlined in Section 2.1. Then, the contribution parameters β_j , $j = 1, \dots, p$ in (3) can be derived by solving (5) using these data, denoted as β_R . Thus, the DIs of reference and test units can be determined by substituting β_R and the corresponding feature data into (2). Using these DIs and the prediction method described in Section 3, the similarity-based RUL for the test unit can be derived. The flowchart of this method is illustrated in Figure 3.

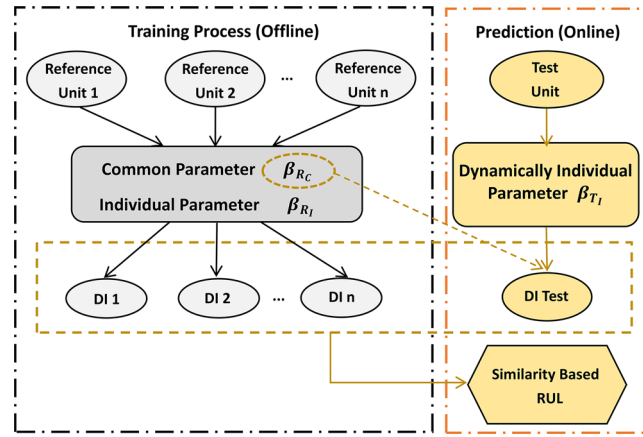


FIGURE 4 Flowchart of the dynamically updating prognostics.

4.2 | The dynamic framework

While the static method presented relies solely on the common information shared by the reference units, it is important to note that each unit also has unique individual features that affect its degradation process. In order to account for both the common and individual features, we propose a dynamic prognostic method that integrates both. We divide the DI at time t in (2) into two parts, the common part $Z^{(1)}(t)$ and the individual part $Z^{(2)}(t)$, which can be expressed as

$$Z(t) = Z^{(1)}(t) + Z^{(2)}(t) = \int_0^t u_1(s)ds + \int_0^t u_2(s)ds, \quad (9)$$

and,

$$u_1(s) = \sum_{j=1}^{p_1} \beta_{1j} f_{1j}[x_{1j}(s), \phi_{1j}],$$

$$u_2(s) = \sum_{j=1}^{p_2} \beta_{2j} f_{2j}[x_{2j}(s), \phi_{2j}], \quad (10)$$

where p_1 and p_2 are respectively the numbers of common and individual features, β_{1j} and β_{2j} are the corresponding contribution parameters, and $f_{1j}[x_{1j}(s), \phi_{1j}]$ and $f_{2j}[x_{2j}(s), \phi_{2j}]$ capture the effects of the corresponding features.

To obtain the DI of the test unit, the contribution parameters of the common part $\beta_{1j}, j = 1, \dots, p_1$ are assumed to be consistent with the reference units, while the individual parameters $\beta_{2j}, j = 1, \dots, p_2$ are allowed to vary based on the data collected from the test unit at operating time t . Using the common features of reference units and (5), $\beta_{1j}, j = 1, \dots, p_1$ can be derived by the static framework, denoted as β_{RC} . Each reference unit's individual contribution parameter, denoted as β_{RI} , is independently computed by solving (5). This computation focuses on the top informative individual features exclusive to each reference unit, excluding common features. It is important to note that there is no overlap between the sets of common and individual features, and the values of β_{RI} vary among distinct reference units.

Up until this point, the proposed framework only utilizes the degradation data from the reference units and can be performed offline: the DIs of reference units can be determined by β_{RC} and β_{RI} . As for the test unit, the individual contribution parameter β_{TI} can be dynamically updated by solving (5) using the selected individual features at operating time t . Combining with the common parameter β_{RC} , we can dynamically obtain the DI of the test unit and calculate the corresponding similarity-based RUL at each operation time point. The flowchart of the dynamic method is illustrated in Figure 4.

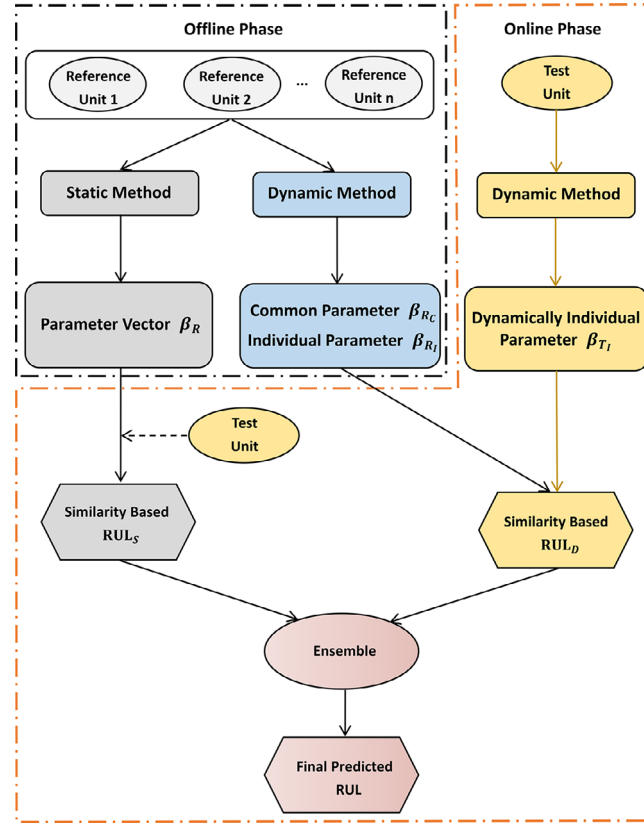


FIGURE 5 Flowchart of the ensemble prognostics.

4.3 | The ensemble framework

The effectiveness of the proposed dynamic framework largely depends on the individual features selected for the test unit. This is because the data size of the test unit is typically smaller than that of the reference units, making it more susceptible to random errors during the dynamic updating process, particularly when the operating time t is short. To address this issue, the following ensemble framework is proposed to achieve a more stable prediction. The basic idea is to ensemble the predicted RULs from the static and dynamic frameworks based on the fact that the RUL of a unit will not improve over time and will not experience a sudden big drop.^{1,20} Let $RUL_{S,t}$ denote the predicted RUL at the operating time t from the static framework, $RUL_{D,t+1}$ denote the predicted RUL at the operating time $t + 1$ from the dynamic framework. Then $RUL_{S,t} - RUL_{D,t+1}$ can be defined as the ensemble condition. Specifically, if $RUL_{S,t} - RUL_{D,t+1}$ is smaller than a preset minimum drop, $\min_d$, then the prediction at time $t + 1$ is $RUL_{S,t} - \min_d$. If $RUL_{S,t} - RUL_{D,t+1}$ is larger than a preset maximum drop, $\max_d$, then the prediction at time $t + 1$ is $RUL_{S,t} - \max_d$. Otherwise, the prediction at time $t + 1$ is $RUL_{D,t+1}$. The values of $\min_d$ and $\max_d$ can be determined through cross-validation. The flowchart of the ensemble method is illustrated in Figure 5.

5 | SIMULATION STUDY

5.1 | Simulated dataset

In this section, a simulation study is conducted to evaluate the performance of the proposed frameworks. The settings are designed to be similar to those in the case study in Section 6. Specifically, there are 10 units, and for each unit, 10 signal data are collected. The failed cycles for the units are (24, 26, 25, 23, 28, 22, 25, 24, 21, 19). Similar to ref. [19], we assume each signal is a trend function of the experimental cycle with some noise, which can be formulated as $X_m(t) = g_m(t) + \varepsilon_m(t)$, $m = 1, \dots, 10$. Without loss of generality, the trend function $g_m(t)$ can be a constant, linear, power, or trigonometric

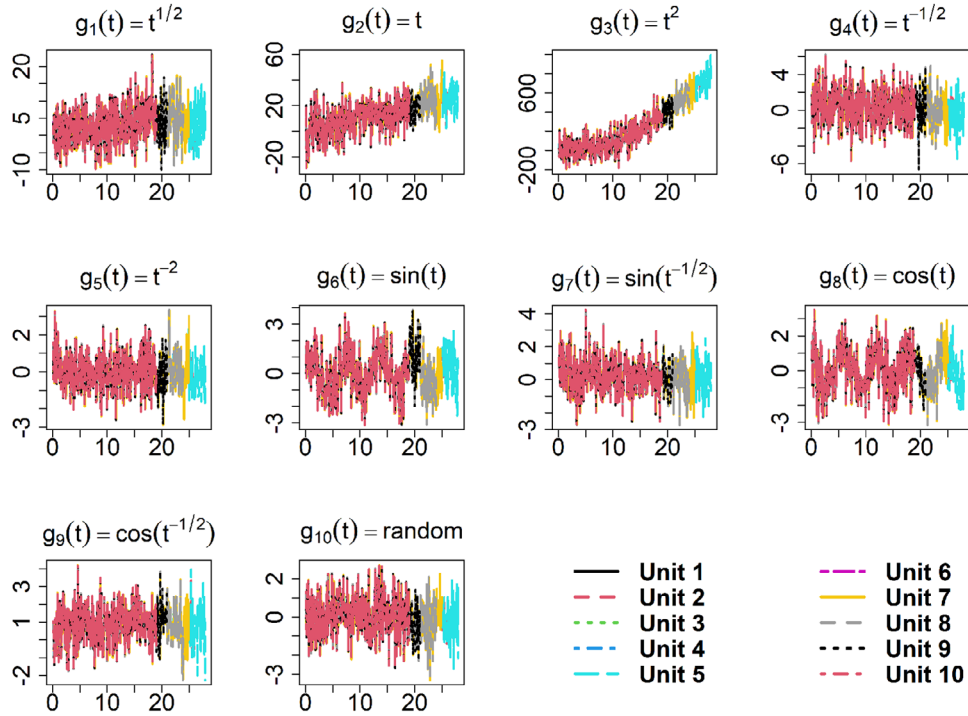


FIGURE 6 Simulated data for 10 units with 10 signals. Each panel represents a single signal. The horizontal axis shows experimental cycles and the vertical axis shows signal measurement. Different colors and line types represent different units.

function, and the noise $\varepsilon_m(t)$ follows a zero-mean normal distribution. The simulated dataset and corresponding functions $g_m(t)$ are shown in Figure 6.

5.2 | Feature engineering and data normalization

Prior to applying the proposed frameworks, feature engineering is necessary as discussed in Section 2.1. Using the data in Figure 6, the features in Table 1 can be calculated. Each unit has a total of 110 features, with 10 signals and 11 features per signal. In the feature selection stage, we employ Fisher's discriminant ratio and Algorithm 1 to select useful features. In the simulated dataset, the number of subsets is 36, as we consider 7 out of 9 as the percentage mentioned in Algorithm 1. The number of selected common features for different units are as follows: (43, 43, 43, 48, 44, 48, 43, 43, 48, 45). Figure 7 illustrates the selected features and their corresponding coefficients estimated by the static framework using simulated data. The visual representation demonstrates the successful generation and selection of informative features by the proposed framework. Additionally, it effectively captures both increasing and decreasing trends.

To reduce the impact of varying data magnitudes, the min-max approach is used to normalize the selected features before model training.³⁵ For a selected feature, it can be formulated as

$$X' = \frac{X - \min(X)}{\max(X) - \min(X)}, \quad (11)$$

where X denotes the original feature, $\max(X)$, $\min(X)$ are calculated over cycles, and X' refers to the normalized version.

5.3 | Performance evaluation

To evaluate the performance of the RUL prediction, the commonly used metric, the root mean square error (RMSE), is employed.¹ Let $\mathbf{RUL}$ and $\widehat{\mathbf{RUL}}$ be the vector of true and predicted RULs of one specific unit, respectively, and n_T be the

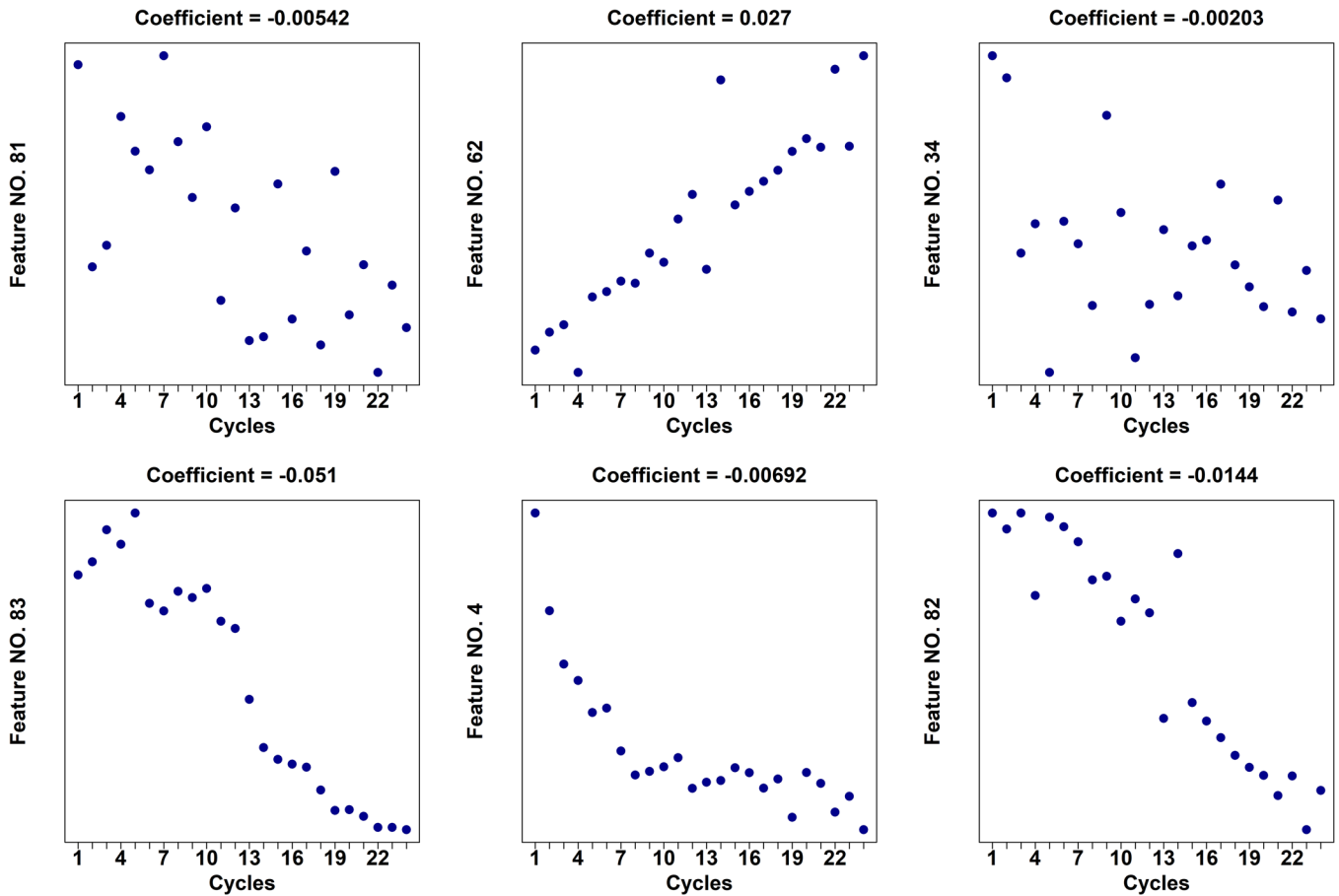


FIGURE 7 Example of the selected features and the corresponding coefficients in the simulated dataset.

corresponding failed cycle number. Then, the RMSE can be formulated as,

$$\text{RMSE} = \|\mathbf{RUL} - \widehat{\mathbf{RUL}}\|_2 / \sqrt{n_T}. \quad (12)$$

In addition to RMSE, it is also important to quantify the uncertainty associated with RUL predictions. To do this, a 95% prediction interval is widely used, which can be calculated using Algorithm 2 by setting $\gamma = 0.05$.

5.4 | Simulation performance

To verify the feasibility and effectiveness of the proposed frameworks in this paper, we utilize two additional methods from refs. [1] and [20] as benchmarks since they also concentrate on the same dataset as in our case study.

These studies assume that the DI and RUL have a fixed linear¹ or nonlinear²⁰ relationship. To predict the RUL, they first build the model from the input features to the DIs using the feed-forward neural network with one-hidden layer, and then they smooth the DI dynamically to improve the quality of the DI. Finally, with the fixed linear or nonlinear relationship, they predict the RUL based on the constructed DI. More details can be found in refs. [1] and [20].

The leave-one-out approach is utilized to validate the performance and determine the reference units. For instance, if unit 1 is chosen as the test unit, the remaining 9 units are treated as reference units. RMSEs based on different methods are reported in Table 2, where M1 is the method in ref. [1], M2 is the method in ref. [20], M3 is the static method in Section 4.1, M4 is the dynamic method in Section 4.2, and M5 is the ensemble method in Section 4.3. The average value of the RMSE in predicting all the units is reported in the last row in the table, and the minimum RMSE value for each unit is highlighted in bold. For uncertainty quantification, the results of 4 representative units are shown in Figure 8, where the solid black line is the true value of RUL.

TABLE 2 RMSE of RUL prediction for the simulated dataset.

Test Unit	M1	M2	M3	M4	M5
Unit 1	2.44	3.14	0.58	1.51	0.42
Unit 2	3.20	6.32	1.84	3.76	1.90
Unit 3	2.60	4.90	1.00	2.55	0.87
Unit 4	2.71	1.83	1.43	2.76	1.28
Unit 5	4.61	9.02	4.61	4.99	4.40
Unit 6	3.28	1.39	3.25	4.69	3.26
Unit 7	2.61	4.90	0.93	1.52	0.80
Unit 8	2.44	3.22	1.43	3.59	1.17
Unit 9	3.99	2.53	0.83	2.66	1.43
Unit 10	5.84	5.91	1.96	4.36	2.20
Mean	3.37	4.32	1.79	3.24	1.77

The bold values represent the minimum RMSE for each unit or motor across different methods. These values are highlighted for easier comparison.

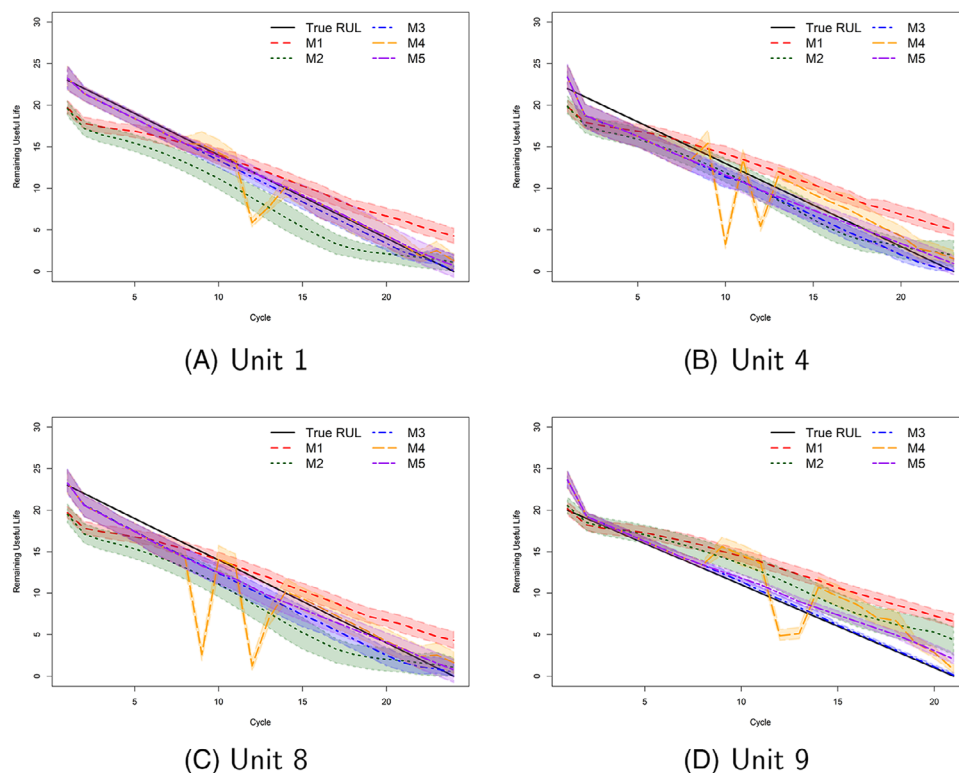
**FIGURE 8** RUL predictions and the 95% prediction interval for the simulated dataset. RUL, remaining useful life.

Table 2 and Figure 8 suggest that the proposed methods (M3 and M5) generally outperform the existing methods (M1 and M2), as indicated by their smaller RMSE and narrower prediction intervals. In particular, the ensemble method (M5) outperforms the static method (M3) in most cases (6 out of 10 and the mean RMSE case), highlighting the effectiveness of integrating static and dynamic methods using the proposed ensemble framework. The dynamic method (M4) yields comparable results per unit to the existing methods, suggesting that it is able to extract useful information using dynamic updating. However, the inconsistent performance also underscores the necessity of the ensemble method (M5).

Note that M1 and M2 demonstrate similar performances, as evidenced by their comparable RMSEs. This could be due to the simulated dataset having a predominantly linear relationship between DI and RUL. As a result, M1 and M2 may have similar capabilities in capturing the underlying trend.

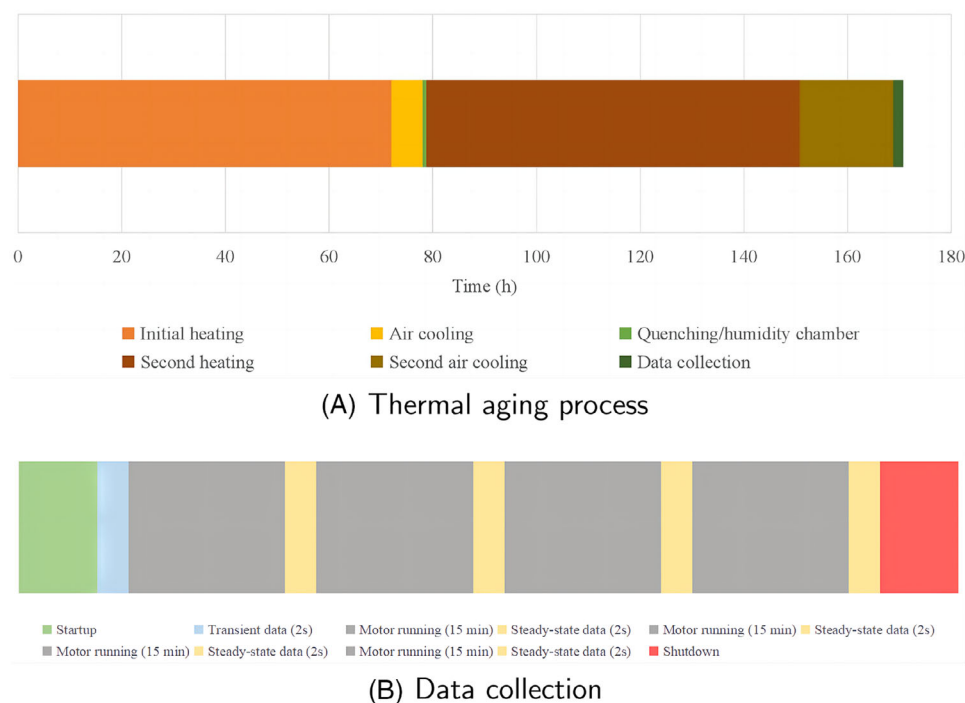


FIGURE 9 Timeline of the accelerated thermal aging and data collection process.

6 | CASE STUDY

In this section, a case study on the degradation data of 8 three-phase industrial induction motors is presented to demonstrate the implementation of the proposed frameworks.

6.1 | Data overview

The data were reported by ref. [36], where ten 5-horsepower motors were used for the accelerated thermal aging process. As depicted in Figure 9A, each thermal aging cycle lasts approximately one week. Further details of each cycle are provided below.

- (1) Initial heating: Heat the motor in one of 3 identical EW-52402-91 ovens at 160°C (or 140°C) for 72 h.
- (2) Air cooling: Remove and allow to air cool for 6 h.
- (3) Quenching/humidity chamber: Quench in an enclosed shallow water pool for 15 minutes.
- (4) Second heating: Immediately place back in the oven and heat again for 72 h.
- (5) Second air cooling: Air cool for 18 h before data collection.

In the data collection stage, as shown in Figure 9B, each motor was connected to a Winco generator through an elastomeric coupling and instrumented with a data collection system. The steady-state data were collected for 2 s every 15 min at 10 kHz and 4 times per cycle for each motor.^{1,20,36} The processes of thermal aging and data collection were repeated until the motor fails to startup normally. The experimental device setup for accelerated thermal aging motor experiments is illustrated in Figure 10.³⁶

In general, this experiment collected 13 channels of key signals including three-phase current (Current 1, 2, 3), three-phase voltage (Voltage 1, 2, 3), two directions of vibration (Accelerometer 1, 2), acoustic (Microphone), speed (Tachometer), temperature (Temperature), and load (output current and voltage) signals. The two channels of load signals were excluded because they were measured by connecting a motor to specific load equipment which is unavailable in practical systems.^{1,20} Since 2 out of the 10 motors experienced abnormal faults during the experiment, the data from 11 signal

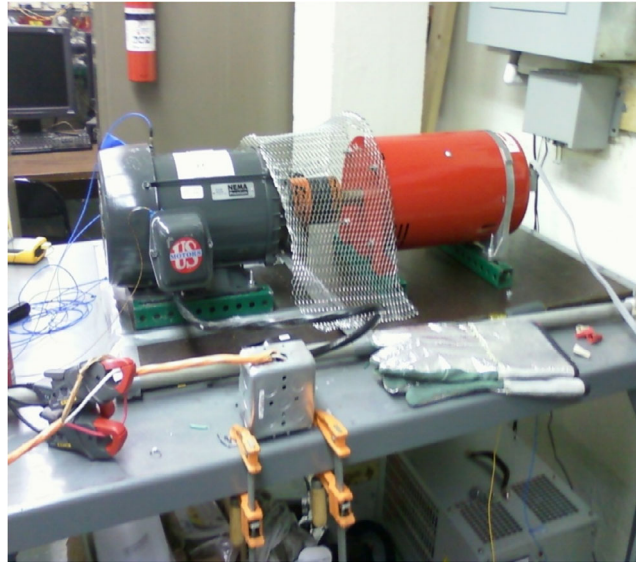


FIGURE 10 Experimental device for accelerated thermal aging of electric motor.

TABLE 3 Details of time to failure for each motor and the corresponding missing values.

Motor number	Failed cycle	Missing values
1	18	Current 2: cycle 5 and 6; Voltage 3: cycle 5, 6, and 7
2	27	Current 2: cycle 5 and 6; Voltage 3: cycle 5, 6 and 7
3	26	Current 1: cycle 7; Current 2: cycle 5, 6, and 7; Current 3: cycle 7; Voltage 3: cycle 5, 6, and 7
4	29	Current 2: cycle 5 and 6; Voltage 3: cycle 5, 6, and 7
5	28	Current 2: cycle 2; Voltage 3: cycle 2 and 3
6	27	Current 2: cycle 5; Voltage 3: cycle 5, 6, and 7
7	27	Current 2: cycle 5 and 6; Voltage 3: cycle 5, 6, and 7
8	25	Current 2: cycle 5; Voltage 3: cycle 5 and 6

channels of the rest 8 motors were used in this paper. The details of time to failure for each motor and the corresponding missing values are given in Table 3.

6.2 | Data pre-processing and feature engineering

As shown in Table 3, some cycle signals have missing values due to human error in data acquisition or short-term damage to the data collection system. Following refs. [1, 20], these missing values are replaced by the nearest historical values. Specifically, when a cycle of signals at a channel is missing, it is replaced with values from its previous cycle. For instance, for the missing values of Current 1 in cycle 7 for Motor 3, they are replaced by signals from cycle 6. Then, the features listed in Table 1 can be calculated. Each motor has 121 features in total, as there are 11 channel signals with 11 features per signal. The leave-one-out approach is employed to validate the performance and determine the reference motors, which is consistent with the simulation study. As highlighted in Section 2.1, we first select useful features using Fisher's discriminant ratio and Algorithm 1. For the proposed frameworks, the number of subsets is 21, as we consider 5 out of 7 as the percentage mentioned in Algorithm 1, and the number of selected common features for different test motors are (50, 50, 50, 48, 47, 48, 50, 51).

Figure 11 displays an illustration of the selected features and their corresponding coefficients estimated by the static framework. The visual representation reveals that the proposed framework successfully generated and selected informative features, while also effectively capturing both increasing and decreasing trends.

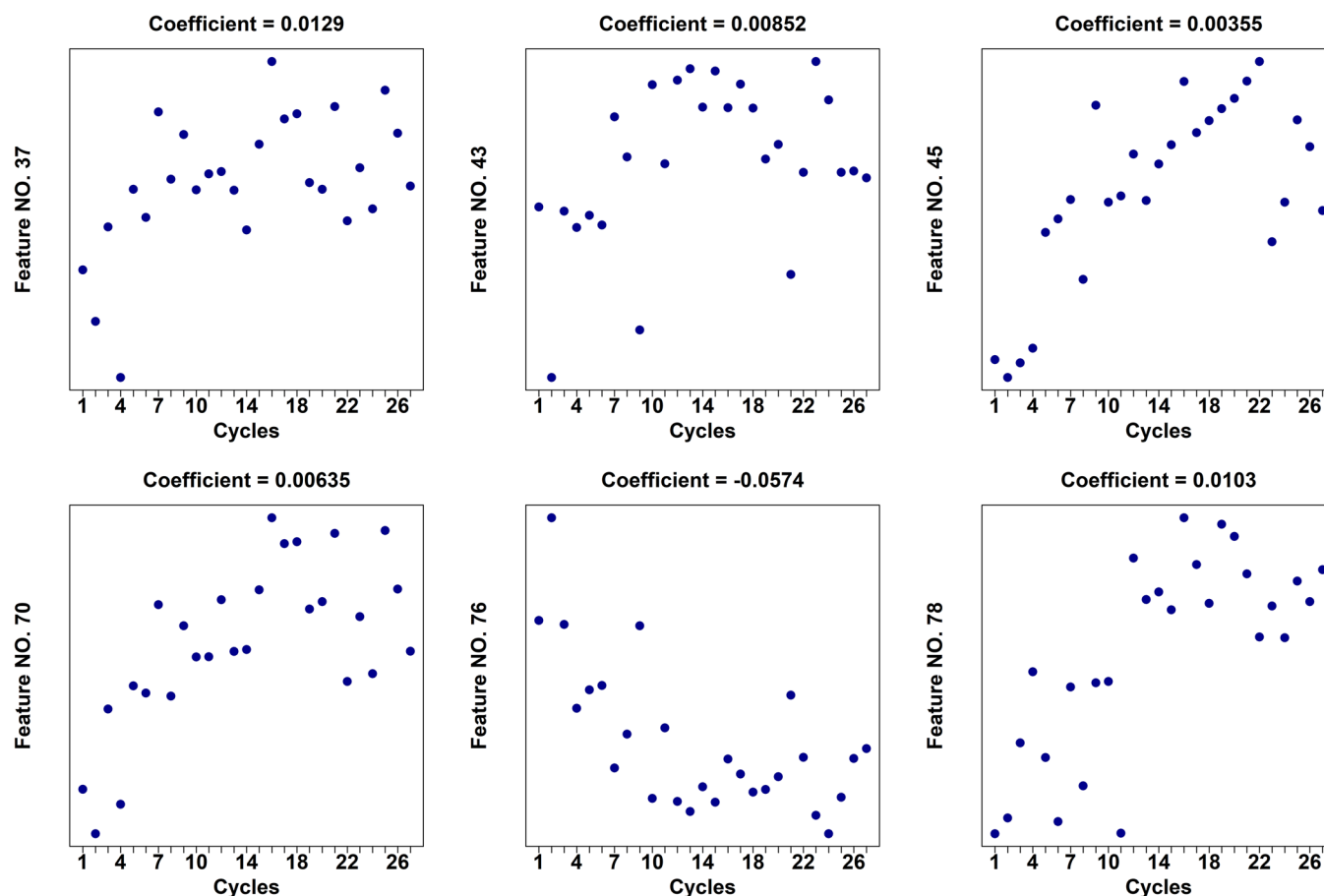


FIGURE 11 Example of the selected features and the corresponding coefficients in the case study.

TABLE 4 RMSE of RUL prediction for the case study.

Test motor	M1	M2	M3	M4	M5
Motor 1	5.38	10.70	8.44	6.83	8.50
Motor 2	3.50	4.77	1.18	2.57	0.30
Motor 3	4.00	4.05	1.01	1.17	1.17
Motor 4	5.53	6.71	2.18	2.49	2.13
Motor 5	6.17	5.50	1.23	2.54	1.09
Motor 6	3.45	5.73	0.96	2.88	0.33
Motor 7	4.81	4.58	3.45	4.09	1.50
Motor 8	2.95	3.72	1.06	1.75	1.76
Mean	4.47	5.72	2.44	3.04	2.10

The bold values represent the minimum RMSE for each unit or motor across different methods. These values are highlighted for easier comparison.

6.3 | Performance of the proposed frameworks

Similar to the simulation study, the RMSE in RUL prediction based on different methods of each motor is reported in Table 4. Recall that M1 refers to the method proposed in ref. [1], M2 corresponds to the method presented in ref. [20], M3 denotes the static method detailed in Section 4.1, M4 represents the dynamic method shown in Section 4.2, and finally, M5 indicates the ensemble method elaborated in Section 4.3. The last row in the table shows the mean value of the RMSE in predicting all motors, and the minimum RMSE value in each row is highlighted in bold for easy comparison of these methods. For uncertainty quantification, the results of 4 representative motors are shown in Figure 12, where the solid black line is the true value of RUL.

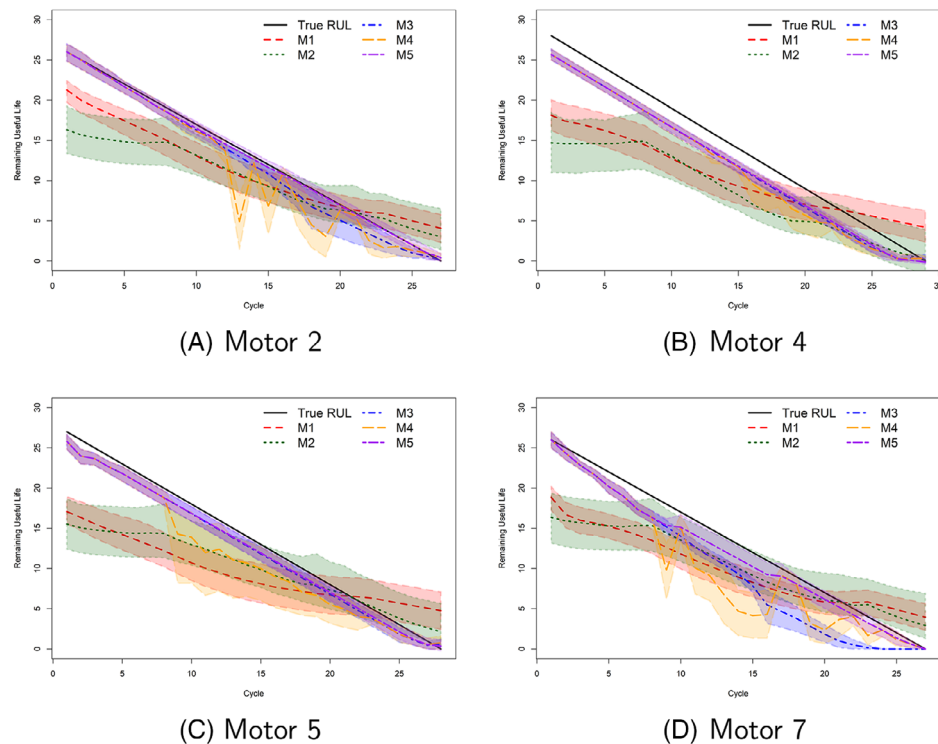


FIGURE 12 RUL predictions and the 95% prediction interval for the case study.

Table 4 and Figure 12 indicate that, in general, the proposed methods (M3 and M5) outperform the existing methods (M1 and M2), as evidenced by their lower RMSE values and narrower prediction intervals. While the dynamic method (M4) does not consistently improve prediction performance compared to the static method (M3), it still provides valuable information, underscoring the importance of the ensemble method (M5). Notably, Table 4 and Figure 12 demonstrate that M5 achieves the best performance in terms of RMSE in most cases, indicating that the proposed ensemble framework effectively integrates the static and dynamic methods. Further comparisons of the RMSE performance of the proposed methods show that Motor 1 always got the worst RMSE performance compared to other motors, and this may be due to the much shorter failure time of Motor 1 (18 cycles vs. 25 ~ 29 cycles).

7 | CONCLUSIONS

This paper presented novel DI-based prognostic frameworks for predicting RUL in complex systems that collect multivariate sensor data. The proposed DI, $Z(t)$ in (2), is feature-based and performs automatic feature selection, which is essential for capturing the underlying degradation trend of the system. Moreover, $Z(t)$ incorporates a nonlinear relationship between the selected features and the degradation process, which reflects the complex and nonlinear nature of the practical engineering applications. Based on the constructed DI, three frameworks were developed for prognostic RUL prediction utilizing various degradation sources. The proposed frameworks do not require prior knowledge of failure mechanisms or specific degradation models, making them applicable to a wide range of engineering systems. The performances of the proposed frameworks were validated through both simulation studies and a case study on the degradation data of 8 three-phase industrial induction motors. The numerical results demonstrate that the proposed prognostic frameworks outperform existing methods by a large margin in terms of predictive accuracy.

One potential direction for future research is to explore more efficient optimization algorithms to solve (5). The computational efficiency of the proposed frameworks is highly dependent on the number of reference units and experimental cycles. Given the relatively small number of motors in our case study, the proposed frameworks can be efficiently implemented on a personal laptop. Nevertheless, more efficient optimization algorithms or parallel computation techniques have to be invoked in the presence of a large number of reference units or experimental cycles. Another possible future research direction is to leverage information from additional sources, such as data collected from different experimental

settings. This may be achieved by developing other types of objective functions and exploring more advanced techniques such as transfer learning.

ACKNOWLEDGMENTS

The authors would like to thank the editor and two anonymous referees for their valuable comments and constructive suggestions, which have significantly improved the quality and presentation of this work. This research was supported by the National Science Foundation of China (Grant No. 12131001).

DATA AVAILABILITY STATEMENT

The data that support the findings of this study are available from the corresponding author upon reasonable request.

ORCID

Wenda Kang  <https://orcid.org/0000-0002-9656-2678>

REFERENCES

- Yang F, Habibullah MS, Zhang T, Xu Z, Lim P, Nadarajan S. Health index-based prognostics for remaining useful life predictions in electrical machines. *IEEE Trans Ind Electron*. 2016;63(4):2633-2644. doi:10.1109/TIE.2016.2515054
- Hong Y, Zhang M, Meeker WQ. Big data and reliability applications: the complexity dimension. *J Qual Technol*. 2018;50(2):135-149. doi:10.1080/00224065.2018.1438007
- Eleftheroglou N, Zarouchas D, Loutas T, Alderliesten R, Benedictus R. Structural health monitoring data fusion for in-situ life prognosis of composite structures. *Reliab Eng Syst Saf*. 2018;178:40-54. doi:10.1016/j.ress.2018.04.031
- Niu G, Yang B-S, Pecht M. Development of an optimized condition-based maintenance system by data fusion and reliability-centered maintenance. *Reliab Eng Syst Saf*. 2010;95(7):786-796. doi:10.1016/j.ress.2010.02.016
- Broer AA, Benedictus R, Zarouchas D. The need for multi-sensor data fusion in structural health monitoring of composite aircraft structures. *Aerospace*. 2022;9(4):183. doi:10.3390/aerospace9040183
- Azamfar M, Singh J, Bravo-Imaz I, Lee J. Multisensor data fusion for gearbox fault diagnosis using 2-D convolutional neural network and motor current signature analysis. *Mech Syst Sig Process*. 2020;144:106861. doi:10.1016/j.ymssp.2020.106861
- Chen J, Jing H, Chang Y, Liu Q. Gated recurrent unit based recurrent neural network for remaining useful life prediction of nonlinear deterioration process. *Reliab Eng Syst Saf*. 2019;185:372-382. doi:10.1016/j.ress.2019.01.006
- Li X, Jiang H, Liu Y, Wang T, Li Z. An integrated deep multiscale feature fusion network for aeroengine remaining useful life prediction with multisensor data. *Knowledge Based Syst*. 2022;235:107652. doi:10.1016/j.knosys.2021.107652
- Wei Y, Wu D, Terpenney J. Decision-level data fusion in quality control and predictive maintenance. *IEEE Trans Autom Sci Eng*. 2020;18(1):184-194. doi:10.1109/TASE.2020.2964998
- Shao H, Lin J, Zhang L, Galar D, Kumar U. A novel approach of multisensory fusion to collaborative fault diagnosis in maintenance. *Inf Fusion*. 2021;74:65-76. doi:10.1016/j.inffus.2021.03.008
- Kim M, Song C, Liu K. A generic health index approach for multisensor degradation modeling and sensor selection. *IEEE Trans Autom Sci Eng*. 2019;16(3):1426-1437. doi:10.1109/TASE.2018.2890608
- Liu K, Chehade A, Song C. Optimize the signal quality of the composite health index via data fusion for degradation modeling and prognostic analysis. *IEEE Trans Autom Sci Eng*. 2015;14(3):1504-1514. doi:10.1109/TASE.2015.2446752
- Liu K, Gebrael NZ, Shi J. A data-level fusion model for developing composite health indices for degradation modeling and prognostic analysis. *IEEE Trans Autom Sci Eng*. 2013;10(3):652-664. doi:10.1109/TASE.2013.2250282
- Gu L, Zheng R, Zhou Y, Zhang Z, Zhao K. Remaining useful life prediction using composite health index and hybrid LSTM-SVR model. *Qual Reliab Eng Int*. 2022;38(7):3559-3578. doi:10.1002/qre.3151
- Song C, Liu K, Zhang X. Integration of data-level fusion model and kernel methods for degradation modeling and prognostic analysis. *IEEE Trans Reliab*. 2017;67(2):640-650. doi:10.1109/TR.2017.2715180
- Chehade A, Song C, Liu K, Saxena A, Zhang X. A data-level fusion approach for degradation modeling and prognostic analysis under multiple failure modes. *J Qual Technol*. 2018;50(2):150-165. doi:10.1080/00224065.2018.1436829
- Song C, Liu K. Statistical degradation modeling and prognostics of multiple sensor signals via data fusion: a composite health index approach. *IIEE Trans*. 2018;50(10):853-867. doi:10.1080/24725854.2018.1440673
- Wang D, Liu K. An integrated deep learning-based data fusion and degradation modeling method for improving prognostics. *IEEE Trans Autom Sci Eng*. 2024;21:1713-1726. doi:10.1109/TASE.2023.3242355
- Wang Y, Lee IC, Hong Y, Deng X. Building degradation index with variable selection for multivariate sensory data. *Reliab Eng Syst Saf*. 2022;227:108704. doi:10.1016/j.ress.2022.108704
- Yang F, Habibullah MS, Shen Y. Remaining useful life prediction of induction motors using nonlinear degradation of health index. *Mech Syst Sig Process*. 2021;148:107183. doi:10.1016/j.ymssp.2020.107183
- Peng W, Wei Z, Huang CG, Feng G, Li J. A hybrid health prognostics method for proton exchange membrane fuel cells with internal health recovery. *IEEE Trans Transp Electr*. 2023;9(3):4406-4417. doi:10.1109/TTE.2023.3243788

22. Zhu J, Chen N, Peng W. Estimation of bearing remaining useful life based on multiscale convolutional neural network. *IEEE Trans Ind Electron*. 2018;66(4):3208-3216. doi:[10.1109/TIE.2018.2844856](https://doi.org/10.1109/TIE.2018.2844856)
23. Zhai Q, Chen P, Hong L, Shen L. A random-effects Wiener degradation model based on accelerated failure time. *Reliab Eng Syst Saf*. 2018;180:94-103. doi:[10.1016/j.ress.2018.07.003](https://doi.org/10.1016/j.ress.2018.07.003)
24. Wang Z, Zhai Q, Chen P. Degradation modeling considering unit-to-unit heterogeneity-a general model and comparative study. *Reliab Eng Syst Saf*. 2021;216:107897. doi:[10.1016/j.ress.2021.107897](https://doi.org/10.1016/j.ress.2021.107897)
25. Luo C, Shen L, Xu A. Modelling and estimation of system reliability under dynamic operating environments and lifetime ordering constraints. *Reliab Eng Syst Saf*. 2022;218:108136. doi:[10.1016/j.ress.2021.108136](https://doi.org/10.1016/j.ress.2021.108136)
26. Xu A, Wang B, Zhu D, Pang J, Lian X. Bayesian reliability assessment of permanent magnet brake under small sample size. *IEEE Trans Reliab*. 2024. doi:[10.1109/TR.2024.3381072](https://doi.org/10.1109/TR.2024.3381072)
27. Guan Q, Wei X, Zhang H, Jia L. Remaining useful life prediction for degradation processes based on the Wiener process considering parameter dependence. *Qual Reliab Eng Int*. 2024;40(3):1221-1245. doi:[10.1002/qre.3461](https://doi.org/10.1002/qre.3461)
28. Zhai Q, Ye Z, Li C, Revie M, Dunson DB. Modeling recurrent failures on large directed networks. *J Am Stat Assoc*. 2024. doi:[10.1080/01621459.2024.2319897](https://doi.org/10.1080/01621459.2024.2319897)
29. Lei Y, Li N, Guo L, Li N, Yan T, Lin J. Machinery health prognostics: a systematic review from data acquisition to RUL prediction. *Mech Syst Sig Process*. 2018;104:799-834. doi:[10.1016/j.ymssp.2017.11.016](https://doi.org/10.1016/j.ymssp.2017.11.016)
30. Huang C-G, Huang H-Z, Li Y-F. A bidirectional LSTM prognostics method under multiple operational conditions. *IEEE Trans Ind Electron*. 2019;66(11):8792-8802. doi:[10.1109/TIE.2019.2891463](https://doi.org/10.1109/TIE.2019.2891463)
31. Liu Y, Hu X, Zhang W. Remaining useful life prediction based on health index similarity. *Reliab Eng Syst Saf*. 2019;185:502-510. doi:[10.1016/j.ress.2019.02.002](https://doi.org/10.1016/j.ress.2019.02.002)
32. Xue B, Xu H, Huang X, Zhu K, Xu Z, Pei H. Similarity-based prediction method for machinery remaining useful life: a review. *Int J Adv Manuf Technol*. 2022;121(3):1501-1531. doi:[10.1007/s00170-022-09280-3](https://doi.org/10.1007/s00170-022-09280-3)
33. Jahani S, Kontar R, Zhou S, Veeramani D. Remaining useful life prediction based on degradation signals using monotonic B-splines with infinite support. *IIEE Trans*. 2020;52(5):537-554. doi:[10.1080/24725854.2019.1630868](https://doi.org/10.1080/24725854.2019.1630868)
34. Luo C, Chen P, Jaillet P. Portfolio optimization based on almost second-degree stochastic dominance. *Manage Sci*. 2024. doi:[10.1287/mnsc.2022.01092](https://doi.org/10.1287/mnsc.2022.01092)
35. Zhuang L, Xu A, Wang XL. A prognostic driven predictive maintenance framework based on Bayesian deep learning. *Reliab Eng Syst Saf*. 2023;234:109181. doi:[10.1016/j.ress.2023.109181](https://doi.org/10.1016/j.ress.2023.109181)
36. Sharp ME. *Prognostic Approaches Using Transient Monitoring Methods*. PhD Thesis. University of Tennessee; 2012. https://trace.tennessee.edu/utk_graddiss/1431/

How to cite this article: Kang W, Jongbloed G, Tian Y, Chen P. Degradation index-based prediction for remaining useful life using multivariate sensor data. *Qual Reliab Eng Int*. 2024;40:3709-3728. <https://doi.org/10.1002/qre.3615>

AUTHOR BIOGRAPHIES



Wenda Kang received his bachelor's degree in mathematics and applied mathematics from Ocean University of China in 2018. Currently, he is pursuing his PhD degree in statistics with the Delft Institute of Applied Mathematics, Delft University of Technology. His research interests focus on data fusion for reliability engineering and transfer learning.



Geurt Jongbloed received his master's degree in applied mathematics in 1991 and his PhD degree in statistics in 1995, both from Delft University of Technology. He is currently a full professor with the Delft Institute of Applied Mathematics, Delft University of Technology. His research interests include inverse problems, statistical inference, nonparametric estimation, and computational statistics.



Yubin Tian received her bachelor's degree in 1987, and master's degree in 1990, both in mathematics from Beijing Normal University. She obtained her PhD degree in weapon science and technology from Beijing Institute of Technology in 2000. She is currently a full professor with the School of Mathematics and Statistics, Beijing Institute of Technology. Her research interests include mathematical statistics, experimental design, and reliability optimization.



Piao Chen received his bachelor's degree in industrial engineering from Shanghai Jiao Tong University in 2013 and his PhD degree in industrial and systems engineering management from the National University of Singapore in 2017. He is currently an associate professor with the ZJUI Institute, Zhejiang University. His research interests include reliability engineering, statistical learning, and operations research.