

**DETC2021-71333**

## MACHINE LEARNING TO PREDICT MEDICAL DEVICES REPAIR AND MAINTENANCE NEEDS

**Hao-yu Liao**

Graduate Research Assistant  
Environmental Engineering Sciences  
University of Florida, Gainesville, FL, 32611  
[haoyuliao@ufl.edu](mailto:haoyuliao@ufl.edu)

**Willie Cade**

ICR Management  
Chicago, IL, 60618  
[williecade@gmail.com](mailto:williecade@gmail.com)

**Karthik Boregowda**

Graduate Research Assistant  
Environmental Engineering Sciences  
University of Florida, Gainesville, FL, 32611  
[kboregowda@ufl.edu](mailto:kboregowda@ufl.edu)

**Sara Behdad\***

Associate Professor  
Environmental Engineering Sciences  
University of Florida, Gainesville, FL, 32611  
[sarabehdad@ufl.edu](mailto:sarabehdad@ufl.edu)

### ABSTRACT

Products often experience different failure and repair needs during their lifespan. Prediction of the type of failure is crucial to the maintenance team for various reasons, such as realizing the device performance, creating standard strategies for repair, and analyzing the trade-off between cost and profit of repair. This study aims to apply machine learning tools to forecast failure types of medical devices and help the maintenance team properly decides on repair strategies based on a limited dataset. Two types of medical devices are used as the case study. The main challenge resides in using the limited attributes of the dataset to forecast product failure type. First, a multilayer perceptron (MLP) algorithm is used as a regression model to forecast three attributes, including the time of next failure, repair time, and repair time z-scores. Then, eight classification models, including Naïve Bayes with Bernoulli (NB-Bernoulli), Gaussian (NB-Gaussian), Multinomial (NB-Multinomial) model, Support Vector Machine with linear (SVM-Linear), polynomial (SVM-Poly), sigmoid (SVM-Sigmoid), and radial basis (SVM-RBF) function, and K-Nearest Neighbors (KNN) are used to forecast the failure type. Finally, Gaussian Mixture Model (GMM) is used to identify maintenance conditions for each product. The results reveal that the classification models could forecast failure type with similar performance, although the attributes of the dataset were limited.

Keywords: Machine Learning, Repair, Maintenance, Medical Devices, Regression, Classification, Clustering

### NOMENCLATURE

|       |                                   |
|-------|-----------------------------------|
| $rt$  | Repair time                       |
| $mrt$ | Mean of repair time               |
| $rs$  | Standard deviation of repair time |

|                 |                                                        |
|-----------------|--------------------------------------------------------|
| $C_k$           | Class $k$                                              |
| $X$             | Input dataset                                          |
| $P(C_k X)$      | Posterior in Bayes' theorem                            |
| $P(X C_k)$      | Prior in Bayes' theorem                                |
| $P(C_k)$        | Likelihood in Bayes' theorem                           |
| $P(X)$          | Evidence in Bayes' theorem                             |
| $N_k$           | Number of class $k$                                    |
| $N$             | Number of all classes.                                 |
| $p_{k_i}^{x_i}$ | Probability when event $k_i$ occurs for sample $x_i$ . |
| $u_k$           | Mean of class $k$                                      |
| $\sigma_k^2$    | Variances of class $k$                                 |
| $\phi(x)$       | Nonlinear function                                     |
| $w$             | Weight of support vector machine                       |
| $b$             | Bias of support vector machine                         |
| $R$             | Structural risk function                               |
| $c^*$           | Penalty parameter                                      |
| $\varepsilon$   | Error tolerance                                        |
| $\xi, \xi^*$    | Slack variables                                        |
| $a, a^*$        | Lagrange multipliers                                   |
| $K(x, x_k)$     | Kernel function                                        |
| $x_k$           | Support vector                                         |
| $r$             | Scaling factor                                         |
| $d$             | Degree term                                            |
| RMSE            | Root mean square error                                 |
| MAE             | Mean absolute error                                    |
| CC              | Correlation coefficient                                |
| Accuracy        | Accuracy of classifiers                                |

### 1. INTRODUCTION

The US is known as a “throw-away” society due to limited repair and reuse practices [1]. One way to handle end-of-use products is to discard them in municipal solid waste (MSW). In 2000, the main components of MSW were

discarded products [2]. Although landfill is a convenient way to dispose of items, it creates significant environmental issues [2]–[4]. It threatens public health through groundwater contamination and air pollution [2], leading to land loss for other usages like housing, leisure, and agriculture [4]. Besides, landfill threatens the house values, especially in higher-priced neighborhoods [5]. Thus, product recovery options such as repair and reuse seem proper replacement for landfills.

Looking at electronic waste (e-waste) as an example reveals the extent of the problem. The US abandons 30 million computers each year, and Europe discards 100 million phones annually. The Environmental Protection Agency estimates that approximately 15 to 20% of e-waste is recycled, and the remainder is disposed of into landfills and incinerators [6]. A significant number of studies have been conducted on the relationship between e-waste and sustainability [7]–[10]. Reuse, recycle, refurbishment, recycling, and repair are several actions to handle e-waste. Among them, repair is a preferred option due to its environmental benefits and materials conservation; however, it is challenging to handle consumer behavior towards repair [4]. It is essential to encourage consumers to repair by showing the economic and environmental aspects of repair.

Previous studies have discussed the importance of investigating product lifecycle. To name a few studies, Gopalakrishnan and Behdad (2017) analyzed the failure factors of hard disk drives using the product lifecycle data [11]. Romero and Vieira (2014) analyzed the relationship between product lifecycle management and maintenance, repair, and overhaul (MRO) service to enhance services [12]. Zhang et al. (2017) also proposed a framework of product lifecycle management based on big data to improve MRO service quality [13]. Houssin and Coulibaly (2014) analyzed the entire product lifespan and used Markov models by considering operating time, maintenance time, and preparing time after failure to evaluate product performance [14]. Wu et al. (2017) proposed a framework by analyzing product lifecycle to enhance product design performance [15].

While previous studies have discussed the impact of analyzing the entire product lifecycle and particularly the repair and operation stage, detailed studies that make repair practices feasible are limited. Particularly, the use of data science techniques and machine learning to support maintenance services is very limited. This study aims to build a forecasting framework to predict the product failure type based on its repair and maintenance history log. We expect the proposed framework will improve the efficiency of services, reduce the cost of repairing and ultimately lead to waste reduction.

The proposed framework consists of two main phases. First, the MLP model was used as a regression model to forecast the next failure time, repair time, and repair time z-scores. Then, in the next phase, these three forecasted values

were fed into a classification model to forecast the failure type. Different classification techniques, including NB-Bernoulli, NB-Gaussian, NB-Multinomial, SVM-Linear, SVM-Poly, SVM-Sigmoid, and SVM-RBF, and KKN, will be compared to each other. Finally, a GMM model will identify the maintenance condition for each product. We demonstrated the application of the above-mentioned machine learning models in medical devices repair. In addition, the study discusses the challenge of limited attributes in the dataset for training classification models and further investigates the performance of different models.

## 2. BACKGROUND OF DATASET

### 2.1 Overview of Dataset

Two medical devices labeled as CA type 1 and CA type 2 are used in this paper. The primary function of these devices is the infusion pump. We obtained the dataset from one of the largest healthcare providers in the US. The above-mentioned devices were selected due to sufficient dataset. The two types of products have been used for over 5 years and have 22 different types of failure reasons, as shown in Figures 1 and 2.

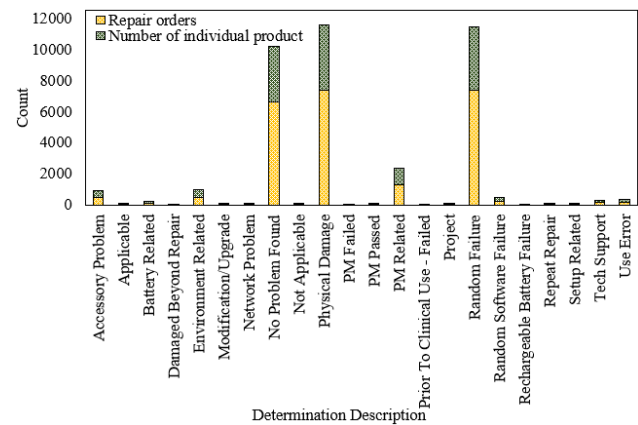


FIGURE 1: The failure reason counts for CA type 1.

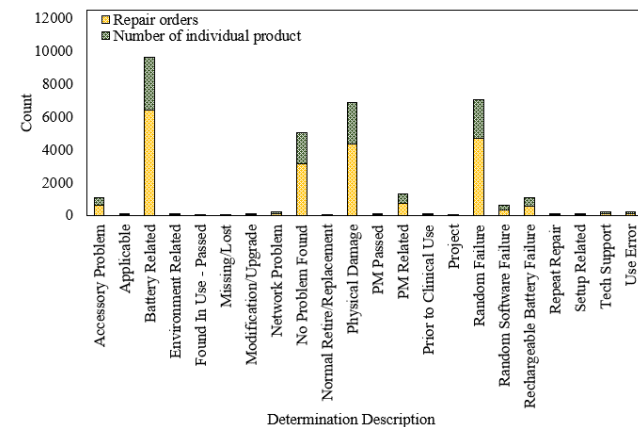


FIGURE 2: The failure reason counts for CA type 2.

The CA type 1 and CA type 2 have 24,516 repair records of 14,660 samples and 21,234 repair records of 12,496 samples respectively, and the average number of repair orders for CA type 1 and CA type 2 is 14 per day. The main maintenance reasons for CA type 1 are the No Problem Found, Physical Damage, and Random Failure, and for CA type 2 are the Battery Related, No Problem Found, Physical Damage, and Random Failure.

Table 1 shows the dataset attributes such as work order and product type. In this table, the Date Created, Completed Date, and Determination Description are the repair start time, repair completion time, and the reason for repair, respectively. The only information available in the dataset is the repair start time and completion time. The aim is to forecast the failure type for the next upcoming repair order based on the available data.

**TABLE 1:** The description of dataset fields.

| Data Attribute            | Description                   |
|---------------------------|-------------------------------|
| WO Number                 | The work order of repair      |
| Type Code                 | Product type                  |
| First Asset Number        | Property unique number        |
| First Asset Description   | Function of product           |
| First Asset Manufacturer  | Manufacturer Name             |
| First Asset Model Number  | Product name                  |
| First Asset Serial Number | Property unique serial number |
| Date Created              | Repair start time             |
| Completed Date            | Repair completion time        |
| Determination Description | The reason for repair         |

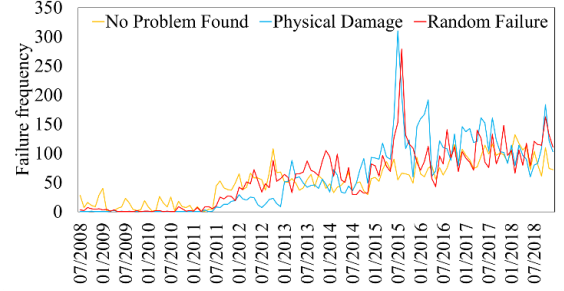
## 2.2 Data analysis

First, we calculated the current failure and repair time by using the data set, and then we used these data to predict the next failure time, repair time, and repair time z-scores. The regression model was trained by available data and used to forecast the next failure time, repair time, and repair time z-scores. The failure time of a product is assumed to be the same as the arrival time of a repair request. The repair time is the difference between Date Created (repair start time) and Completed Date (repair completion time). The repair time z-score (RTZS) is the normalized repair time. Using RTZS improves the accuracy of the classification models.

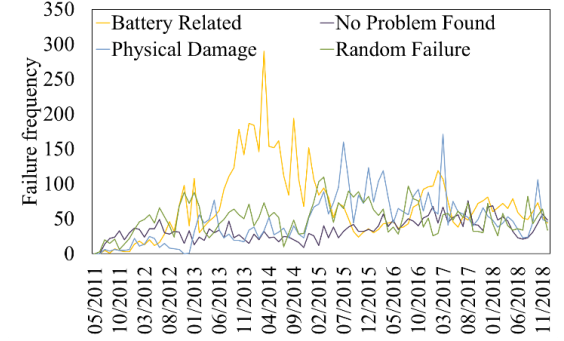
However, the model outputs are influenced by the limitation of the current dataset. Figures 3 and 4 show the frequency of main failure reasons in each year. The sequence of repair order arrivals is not smooth. For example, CA type 1 has peaks on 07/2015 for Physical Damage and Random Failure. CA type 2 has a peak around 04/2014 for Battery Related reason. Several reasons such as a change in repair and replacement policy and purchasing actions will unstable the system.

For example, CA type 1 has three main maintenance reasons including No Problem Found, Physical Damage, and Random Failure. The main maintenance reasons for CA type 2 are Battery Related, No Problem Found, Physical Damage, and Random Failure. Although each product might

have failed due to 22 different reasons, we have only used the data for the most frequent failure types. 87% of the dataset for CA type 1 includes main failure reasons.



**FIGURE 3:** The frequency of main failure reasons for CA type 1.

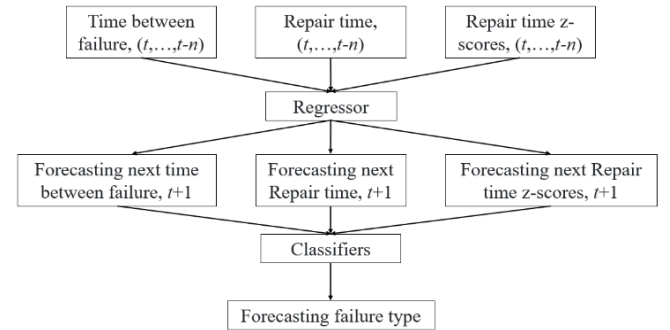


**FIGURE 4:** The frequency of main failure reasons for CA type 2.

## 3. METHOD

### 3.1 Data input and output

Figure 5 shows the proposed framework, along with the data input and output. The MLP as regressor was trained to forecast the next failure time, repair time, and RTZS before the classifiers is trained by forecasting results to predict failure type for the next repair order.



**FIGURE 5:** The proposed framework for data analysis.

The initial data exploration and analysis show that after adding RTZS, classifiers' accuracy will be better. The RTZS is calculated as:

$$RTZS = \frac{rt - mrt}{rs} \quad (1)$$

where  $rt$  is repair time.  $mrt$  is the mean of repair time, and  $rs$  is the standard deviation of repair time.

We normalized each repair type independently. We found that after applying RTZS, the accuracy of classifiers is improved. Although RTZS does not have any statistical significance in this study, RTZS can enhance accuracy by its normalization property. Therefore, we still used RTZS as one of the inputs to each classifier. In addition, the regression model is evaluated by root-mean-square error (RMSE), mean absolute error (MAE), and correlation coefficient. The classifiers will be assessed by accuracy.

We implemented regression, classification, and cluster learning on the two medical devices. In CA type 1, we have 21,312 records for main failure reasons. CA type 2 has 18,487 records for main failure reasons. The 70% and 30% of the dataset have been used for training and testing respectively. We defined Mean Time Between Failures (MTBF) as the average time between failure for a specific failure reason for cluster learning. For cluster learning as unsupervised learning, however, we do not have the true label for the dataset to our MTBF definition. The clustering learning will find the best group based on the characteristic of the dataset. We applied cluster learning to the MTBF of Random Failure and Physical Damage to analyze the properties of each clustering group.

### 3.2 Multilayer Perceptron (MLP)

MLP is popular for regression analyses. Liu et al. (2021) used stepwise regression analysis and combined other networks like MLP for air quality detectors [16]. Leszczynski and Jasinski (2020) used MLP to build an evaluation model for estimating product life cycle cost [17]. Liu et al. (1995) demonstrated MLP to analyze the reliability data and identify the underlying distribution of failure data [18]. Xu et al. (2003) used MLP to forecast the engine systems reliability [19]. When training MLP, several parameters will need to be considered, such as the number of hidden layers and learning rate. These parameters will influence the performance of forecasting results. The MLP structure is shown in Figure 6.

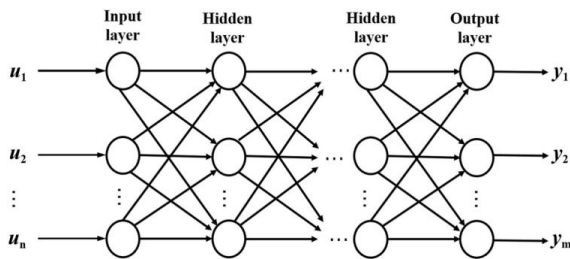


FIGURE 6: The MLP network structure [16].

We used MLP in the first phase of the framework for predicting time. In the second phase, we will use different classification techniques described in future subsections.

### 3.3 Naïve Bayes

The Naïve Bayes will consider the probability distribution of the training dataset based on Bayes' assumption.

The approach has been widely used in the literature. Let's review some formulas from the Bayes' theorem. The conditional probability can be expressed as:

$$P(C_k|X) = \frac{P(X|C_k)P(C_k)}{P(X)} \quad (2)$$

where  $C_k$  is class  $k$ ,  $X$  is the input dataset. In the Bayes' theorem, the  $P(C_k|X)$  is posterior,  $P(X|C_k)$  is prior multiplied by  $P(C_k)$  as the likelihood, and  $P(X)$  is the evidence. The  $P(C_k)$  and  $P(X)$  are indicated as:

$$P(C_k) = \frac{N_k}{N} \quad (3)$$

$$P(X) = \sum_k P(X|C_k)P(C_k) \quad (4)$$

where,  $N_k$  is the number of class  $k$  and  $N$  is the number of all classes. After considering different probability distributions, the Naïve Bayes classifiers with Bernoulli, Gaussian, and Multinomial are expressed as [20] [21]:

$$P(X|C_k) = \prod_{i=1}^n p_{k_i}^{x_i} (1 - p_{k_i})^{(1-x_i)} \quad (5)$$

$$P(X|C_k) = \frac{1}{\sqrt{2\pi\sigma_k^2}} e^{-\frac{(x-u_k)^2}{2\sigma_k^2}} \quad (6)$$

$$P(X|C_k) = \frac{(\sum_i x_i)!}{\prod_i x_i!} \prod_i p_{k_i}^{x_i} \quad (7)$$

where  $p_{k_i}^{x_i}$  is the probability when an event  $k_i$  occurs for sample  $x_i$ .  $u_k$  and  $\sigma_k^2$  are the mean and variances of class  $k$ . Equations (3), (4), and (5) are for NB-Bernoulli, NB-Gaussian, and NB-Multinomial.

### 3.4 Support Vector Machine (SVM)

In 1995, Vapnik proposed SVM to solve classification problems [22]. Previous researchers have confirmed the high performance of SVM in different applications [23]–[25]. According to [24], the  $w$  (weights) and  $b$  (bias) can be determined by minimizing the structural risk function as follow:

$$\text{Min } R(w, b, \xi, \xi^*) = \frac{1}{2} \|w\|^2 + c^* \sum_{i=1}^n L_\epsilon(\xi + \xi^*) \quad (8)$$

$$\text{Subject to } \begin{cases} y_i - \hat{y}_i = \hat{y}_i - (w^T \phi(x) + b) \leq \epsilon + \xi_i \\ y_i - \hat{y}_i = (w^T \phi(x) + b) - \hat{y}_i \leq \epsilon + \xi'_i \\ \xi_i \geq 0 \\ \xi'_i \geq 0 \\ i = 1, 2, \dots, n \end{cases} \quad (9)$$

where  $c^*$  is a penalty parameter for making a tradeoff between model complexity,  $\epsilon$  is error tolerance,  $y$  is the target,  $\hat{y}$  is evaluated output, and  $\xi$ ,  $\xi^*$  are slack variables. After finding the best  $w$  and  $b$ , the evaluated  $\hat{y}$  is expressed as:



$$\hat{y} = f(x) = \mathbf{w}^T \phi(x) + b \quad (10)$$

Where  $\phi(x)$  is a nonlinear kernel function, the SVM has four kernel functions as follow:

Linear function (SVM-LN):

$$K(x_i, x_j) = x_i^T \cdot x_j \quad (11)$$

Polynomial function (SVM-Poly):

$$K(x_i, x_j) = (r \cdot x_i^T \cdot x_j + c)^d \quad (12)$$

Radial basis function (SVM-RBF):

$$K(x_i, x_j) = \exp(-r \|x_i - x_j\|^2) \quad (13)$$

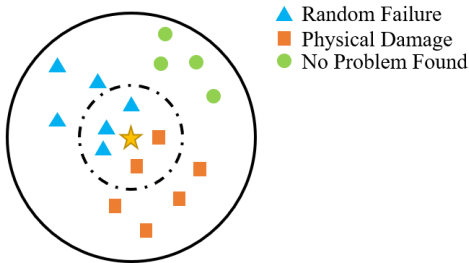
Sigmoid function (SVM-Sigmoid):

$$K(x_i, x_j) = \tanh(r \cdot x_i^T \cdot x_j + c) \quad (14)$$

Where  $r$  is a scaling factor,  $c$  is the bias term,  $d$  is the degree term. This study uses the grid search method [26] to find parameters ( $r$ ,  $c^*$ ,  $d$ , and  $\epsilon$ ).

### 3.5 K-Nearest Neighbors (KNN)

Another popular classification technique is KNN. The KNN classifiers will consider near samples to determine the class for input data. Figure 7 shows the process to determine the class of the star point. When  $k$  equal to 5 votes, the KNN will find the nearest 5 samples to the input data. After finding the nearest 5 samples, the most frequent class is chosen as the class of input data. The KNN model is useful in larger datasets and has a good performance in classification. Previous studies have confirmed the performance of KNN is various classification problems [27] [28].

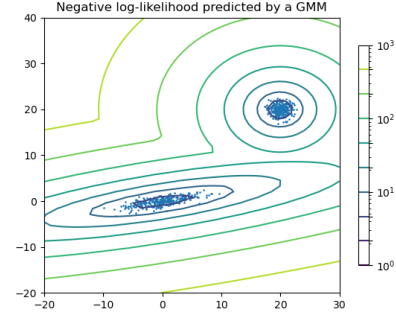


**FIGURE 7:** Classifying the yellow star to Random Failure when  $k$  is 5.

### 3.6 Gaussian Mixture Models (GMM)

We also applied the Gaussian Mixture Models to do cluster learning on analyzing products' performance based on failure reasons. The GMM is a probabilistic model. It assumes the dataset is created by a finite number of Gaussian distributions [29]. Previous researchers have confirmed that GMM is a useful tool for cluster learning. For example, McLachlan et al. (2014) investigated GMM on cluster learning and found that the dataset's components are constructed based on the probability distribution [30]. Sahbi (2008) also used GMM to retrieve information contained in images [31]. Glowacz et al. (2018) applied GMM to the early fault diagnostic of the single-phase induction motor

[32]. Wang and Zarader (2012) used GMM for analyzing aircraft fault diagnosis[33]. The optimal number of clusters is decided by the Dunn index. The Dunn index considers the distance between each point and each cluster. The larger is Dunn index, the better is the number of clusters [34]. The GMM model considers the distribution of the dataset and evaluates the equi-probability surfaces for the dataset as shown in Figure 8.



**FIGURE 8:** Example of two-component Gaussian Mixture Model [29].

## 4. THE RESULTS OF FORECASTING, CLASSIFICATION, AND CLUSTERING

In this section, we discuss the outcome of the prediction, classification, and clustering techniques. We applied scikit-learn python library for this study. The best-trained parameters are shown in Table 5. The explanation of each parameter can be obtained from scikit-learn documents [29].

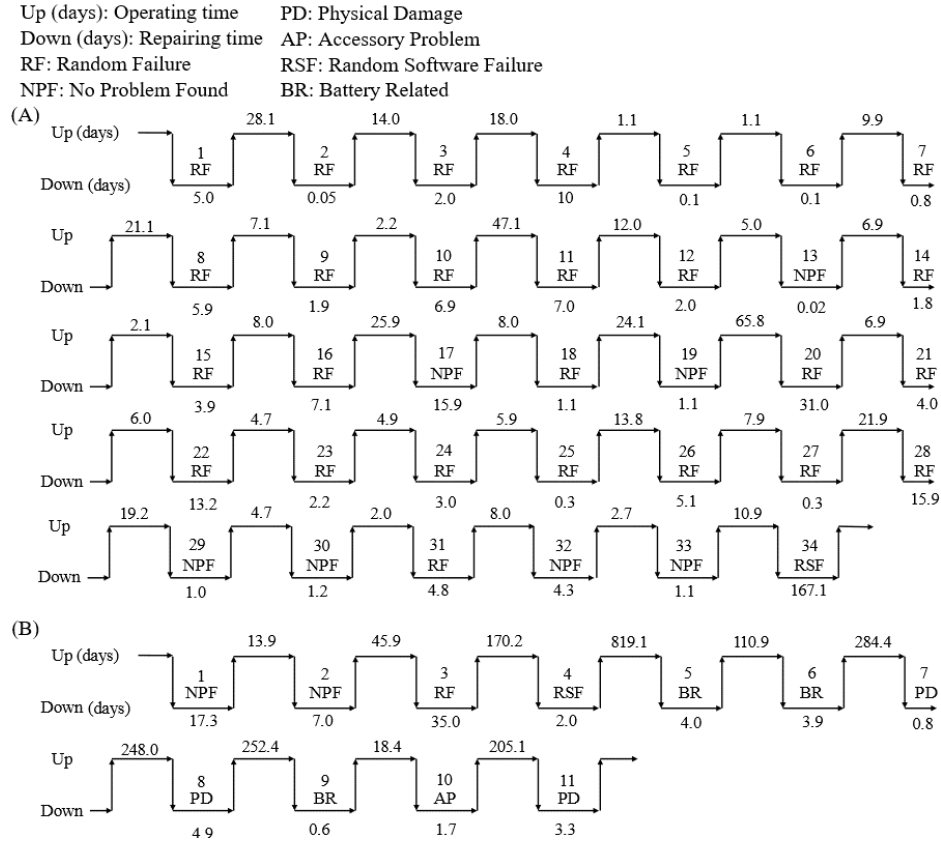
### 4.1 Product lifecycle analysis

To provide a better understanding of the analysis, the lifecycle of two individual products are mapped in Figure 9. Figure 9 (A) shows the product lifecycle of CA type 1 with ID 94012568, which has more maintenance records than others, and Figure 9 (B) demonstrates CA type 2 with ID 94027178, which is a randomly selected sample. The first product is an extreme case with 34 maintenance records with three types of failure including Random Failure, No Problem Found, and Random Software Failure. The most frequent maintenance record is Random Failure, and the longest maintenance time is 31 days as shown in the 20<sup>th</sup> records.

The longest maintenance time for No Problem Found is 15.9 days as shown in the 17<sup>th</sup> records. The longer time of maintenance could be due to the long diagnosis time or longer repair time. The CA type 2 with ID 94012568 has similar conditions on No Problem Found. The longest maintenance time is 17.3 days in the 1<sup>st</sup> record.

The longer maintenance time shows the necessity for predicting the type of failures. The prediction results help the maintenance team better manage labor cost and maintenance schedules.

Table 2 shows the summary of the two products. The information such as the number of failures, and % of up and down timetable is useful in deciding whether to repair or replace a medical device.



**FIGURE 9:** The product lifecycle of (A) CA type 1 with ID 94012568 and 34 maintenance records and (B) CA type 2 with ID 94027178 and 11 maintenance records, including Random Failure (RF), No Problem Found (NPF), Physical Damage (PD), Accessory Problem (AP), Random Software Failure (RSF), and Battery Related (BR) (Up: Operating Time and Down: Maintenance Time).

**TABLE 2:** Summary of the two selected individual products.

| Product ID         | 94012568 | 94027178 |
|--------------------|----------|----------|
| Number of Failures | 34       | 11       |
| % of up time       | 0.57     | 0.96     |
| % of down time     | 0.43     | 0.04     |
| Number of RF       | 26       | 1        |
| Number of PD       | 0        | 3        |
| Number of NPF      | 7        | 2        |
| Number of AP       | 0        | 1        |
| Number of RSF      | 1        | 1        |
| Number of BR       | 0        | 3        |

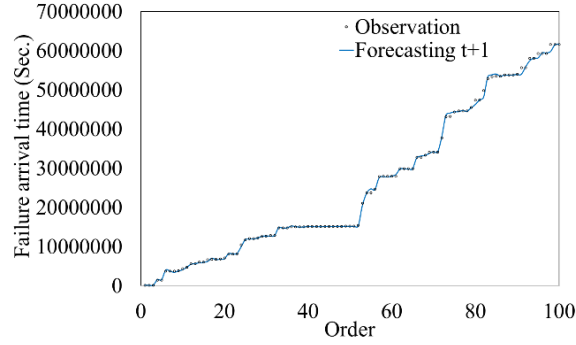
#### 4.2 Regressor results

After training MLP regressors, the best parameters were determined based on RMSE, MAE, and CC. Table 3 shows the training and testing results for MLP. The results reveal that MLP forecasts the next failure time better than two other variables including next repair time and z-score. The CC is almost one for CA type 1 and CA type 2. The performance of forecasting repair time and RTZS are similar for both CA type 1 and CA type 2. The results of the regression are shown in Figures 10 to 12. As seen, the forecasted failure time catches the observation, but the forecasted repair time and z-scores overestimate the actual values. As expected from the analysis, forecasting the next repair's timing from historical data is challenging due to the independence of the

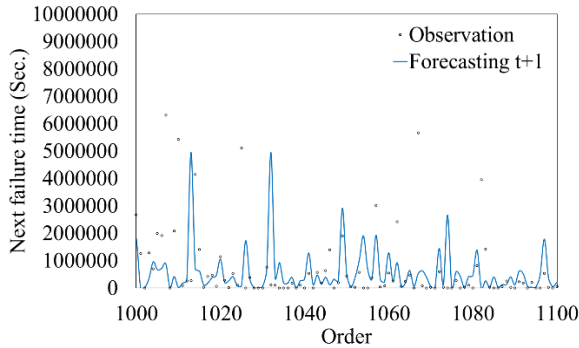
repair events. The previous repair event and the current repair event have different failure reasons and different repair time ranges. Also, we found the correlation coefficient between  $t$  and  $t-1$  is near 0.37. Therefore, it is hard to forecast the repair time for the next failure from the current dataset.

**TABLE 3:** The MLP training and testing results for CA type 1 and CA type 2.

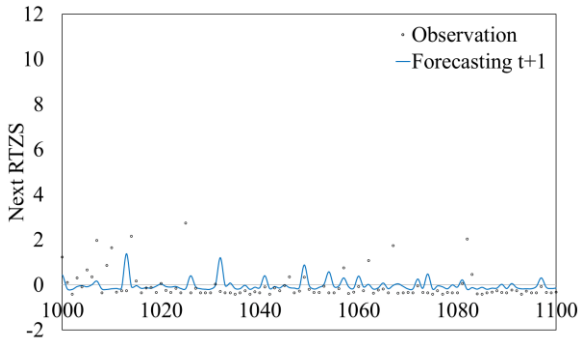
|                                      | Lead time<br>(t+1)  | RMSE<br>(Sec.)    | MAE<br>(Sec.) | CC |
|--------------------------------------|---------------------|-------------------|---------------|----|
| <b>Forecasting next failure time</b> |                     |                   |               |    |
| Train                                | 87422,<br>35926     | 24399,<br>19169   | 1, 1          |    |
| Test                                 | 84110,<br>34708     | 24491,<br>18719   | 1, 1          |    |
| <b>Forecasting next repair time</b>  |                     |                   |               |    |
| Train                                | 1983521,<br>1860649 | 777089,<br>640744 | 0.42, 0.39    |    |
| Test                                 | 1852907,<br>1759914 | 772451,<br>619292 | 0.47, 0.45    |    |
| <b>Forecasting next RTZS</b>         |                     |                   |               |    |
| Train                                | 0.92, 0.96          | 0.42, 0.38        | 0.41, 0.44    |    |
| Test                                 | 0.86, 0.88          | 0.42, 0.38        | 0.44, 0.47    |    |



**FIGURE 10:** The forecasted time for the next failure of CA type 1.



**FIGURE 11:** The forecasted repair time for the next failure in CA type 1.

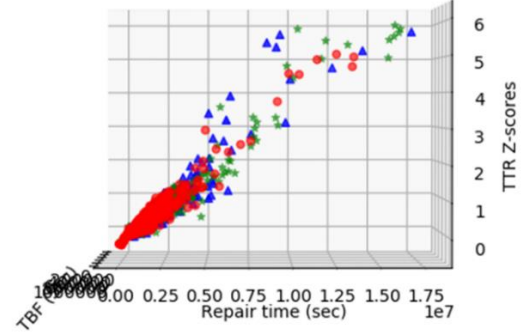


**FIGURE 12:** Forecasting the next RTZS for CA type 1.

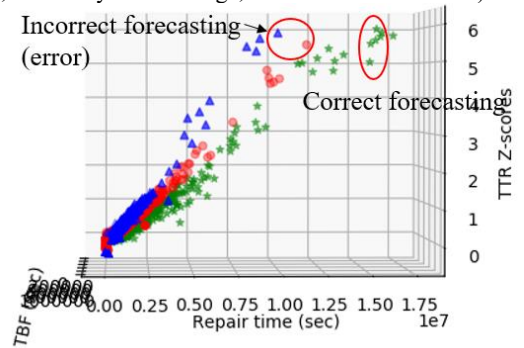
### 4.3 Classifiers results

Our main objective is to build classification models that can predict the next failure type. Once the maintenance services are aware of the type of failure that happens to a product, they can prepare strategies and resources needed for fixing the next product. After training the regressor, the regression results are fed into classifiers to find the future failure types. As a reminder, the main failure types for CA type 1 are Random Failure, Physical Damage, and No Problem Found, and for CA type 2 are Battery Related, Random Failure, Physical Damage, and No Problem Found. Table 4 shows the classification results for each classifier. Among them, KNN outperforms in both products. The testing accuracy is 0.363 for CA type 1 and 0.350 for CA type 2. The results have slight overfitting even though we have adjusted different parameters such as k votes and

weights. Figures 13 and 14 show the actual and forecasted labels for CA type 1 with three main failure reasons. For example, the green star shows the actual label for Random Failure in Figure 13. In Figure 14, if the model forecasts this point with a green star, it is the correct forecasting, otherwise, it is the incorrect forecasting. The results show that the forecasting model can predict the failure type correctly as shown in Figure 14.



**FIGURE 13:** The actual label of CA type 1 (Red: No Problem Found; Blue: Physical Damage; Green: Random Failure).



**FIGURE 14:** The forecasted label of CA type 1 by SVM-RBF models (Red: No Problem Found; Blue: Physical Damage; Green: Random Failure).

**TABLE 4:** The training and testing results of each classifier for CA type 1 and CA type 2.

| Model          | Training Accuracy | Testing Accuracy |
|----------------|-------------------|------------------|
| NB-Bernoulli   | 0.348, 0.344      | 0.348, 0.344     |
| NB-Gaussian    | 0.361, 0.245      | 0.361, 0.242     |
| NB-Multinomial | 0.350, 0.344      | 0.350, 0.344     |
| SVM-Linear     | 0.338, 0.331      | 0.340, 0.331     |
| SVM-Poly       | 0.349, 0.343      | 0.350, 0.343     |
| SVM-RBF        | 0.362, 0.354      | 0.350, 0.326     |
| SVM-Sigmoid    | 0.355, 0.343      | 0.355, 0.343     |
| KNN            | 0.381, 0.355      | 0.363, 0.350     |

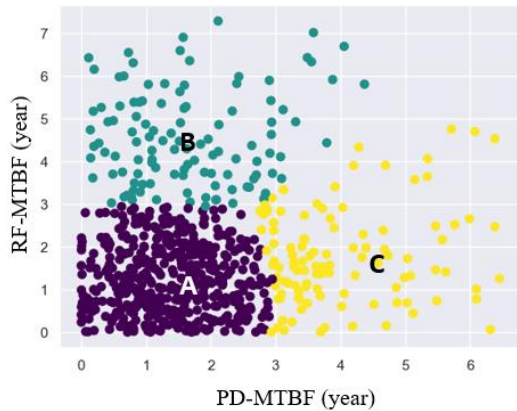
### 4.4 The results of clustering

The GMM clustering model is applied to the dataset. The GMM is an unsupervised learning clustering algorithm. The

clustering model will divide each group based on the dataset. After separating each group, we analyzed the properties of each group based on an unsupervised GMM model. Figures 15 and 16 show the clustering results for CA type 1 and CA type 2. Each group has its own meaning for maintenance purposes. For example, in Figure 9 (B), individual product 94027178 has three Physical Damage records (7<sup>th</sup>, 8<sup>th</sup>, and 11<sup>th</sup>). The MTBF of Physical Damage for 94027178 was computed based on the three records. The optimal number of clusters for both CA type 1 and CA type 2 is 3 based on the largest Dunn Index, as shown in Figure

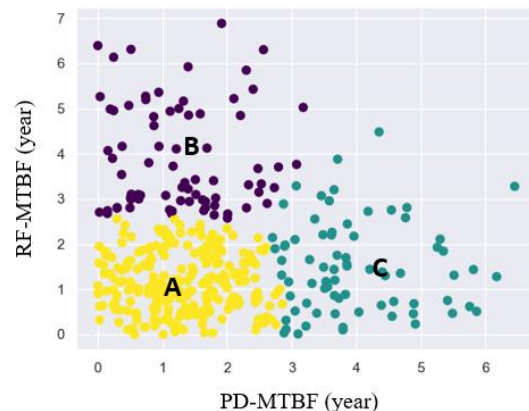
**TABLE 5:** Best trained parameters of CA type 1 and CA type 2.

| Item                                                                                                                        | Model                | Parameters          |                    |              |               |        |
|-----------------------------------------------------------------------------------------------------------------------------|----------------------|---------------------|--------------------|--------------|---------------|--------|
| Next failure time forecasting<br>Next repair time forecasting<br>Next RTZS forecasting<br><br>Next failure type forecasting | MLP                  | hidden_layer_sizes  | learning_rate_init | momentum     |               |        |
|                                                                                                                             |                      | (100, 20), (50, 20) | 1E-03, 1E-04       | 0.1, 0.1     |               |        |
|                                                                                                                             |                      | (4, 2), (4, 2)      | 1E-05, 1E-05       | 0.1, 0.1     |               |        |
|                                                                                                                             | (500, 20), (500, 20) | 1E-03, 1E-03        | 0.1, 0.1           |              |               |        |
|                                                                                                                             | NB                   | Type                | alpha              | binarize     | var_smoothing |        |
|                                                                                                                             |                      | Bernoulli           | 1E-09, 1E-09       | 1, 1         | None          |        |
|                                                                                                                             |                      | Gaussian            | None               | None         | 1E-09, 1E-09  |        |
|                                                                                                                             | Multinomial          | 1E-09, 1E+02        | None               | None         |               |        |
|                                                                                                                             | SVM                  | kernel              | C                  | coef0        | gamma         | degree |
| Linear                                                                                                                      |                      | 1, 500              | 0, 0               | scale, scale | None          |        |
| Poly                                                                                                                        |                      | 10, 1               | 1, 1               | scale, scale | 2, 1          |        |
| RBF                                                                                                                         |                      | 1, 1                | 0, 0               | auto, auto   | None          |        |
| Sigmoid                                                                                                                     |                      | 10, 1               | 1, 100             | scale, scale | None          |        |
| KNN                                                                                                                         | n_neighbors          | p                   | weights            |              |               |        |
|                                                                                                                             | 240, 180             | 2, 1                | uniform, uniform   |              |               |        |
| GMM                                                                                                                         | n_components         | covariance_type     |                    |              |               |        |
|                                                                                                                             | 3, 3                 | full                |                    |              |               |        |
| Clustering on MTBF-RF & PD                                                                                                  |                      |                     |                    |              |               |        |



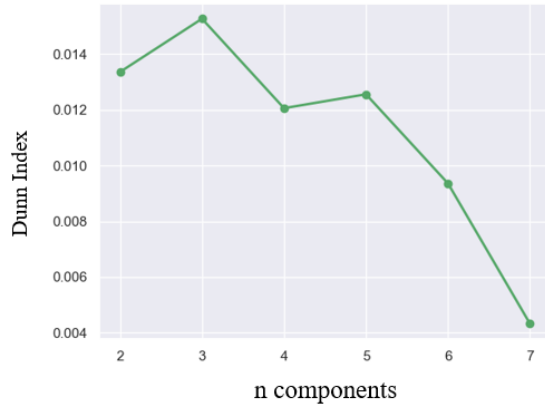
**FIGURE 15:** The results of clustering CA type1 (RF: Random Failure and PD: Physical Damage).

17. In Group A, the overall samples have less MTBF for Physical Damage and Random Failure reasons. It means that if a product is classified in group A, the product may need more care on maintenance to extend the product lifespan as the MTBF is low for any type of failure. Group B has more lifespan on Random Failure and less lifespan on Physical Damage. That means these samples need more care on Physical Damage maintenance. Group C is the inverse of Group B. The samples in this group need more attention to the maintenance of Random Failure



**FIGURE 16:** The results of clustering CA type 2 (RF: Random Failure and PD: Physical Damage).





**FIGURE 17:** The Dunn Index for CA type 1.

The clustering results provide the maintenance team with helpful insights for repair decisions when the next failure type happens. For example, if the classifiers predict that the next failure type is Physical Damage. By reviewing the product records, if the product is in Group B, the maintenance team needs to inspect the product for Physical Damage to increase product's lifespan based on the combination of both classifiers and clustering results.

## 5. CONCLUSION

This study aims to build a data analytics framework to forecast the time and type of the product's next failure toward proper repair and maintenance planning. A set of regression, classification, and clustering techniques have been used. The prediction of the time and type of failure helps repair shops with efficient resource management. Although it is challenging to build prediction models given limited or imperfect datasets, the results show that the classifier models can forecast the type of next failure type as a piece of supportive information for maintenance services. A dataset of medical equipment has been used to show the application of the proposed data analytics framework. Two types of products have been selected for product lifecycle analysis. Although the dataset is limited, the results of prediction, classification, and clustering provide helpful insights for the maintenance team.

This research can be extended in several ways to improve the accuracy of models. First, applying the proposed framework to other datasets with complete data attributes increases the prediction and classification accuracy. Second, extending the framework to other failure reasons. We only considered the most frequent failure reasons. Therefore, the classification methods can be extended to enhance the learning accuracy of all failure reasons together. Third, other product features such as frequency of failure, average operating time, and average repair time can be included in ML techniques. Fourth, deep learning techniques can be applied. For example, transferring product's records into signal images to identify the maintenance condition. Finally, developing an expert

system to decide whether to repair or replace based on the maintenance records.

## ACKNOWLEDGEMENT

This material is based upon work supported by the National Science Foundation–USA under grants #CMMI-2017968 and CBET-2017971. Any opinions, findings, and conclusions or recommendations expressed in this material are those of the authors and do not necessarily reflect the views of the National Science Foundation.

## REFERENCES

- [1] J. Mccollough and J. Mccollough, "Consumer Discount Rates and the Decision to Repair or Replace a Durable Product: A Sustainable Consumption Issue a Durable Product: A Sustainable Consumption Issue," vol. 3624, 2014, doi: 10.2753/JEI0021-3624440109.
- [2] B. Sheehan, F. D. Survey, and C. Science, "Managing product waste: The next frontier for MSW," no. February 2006, 2018.
- [3] C. Cole, T. Cooper, and A. Gnanapragasam, "Extending Product Lifetimes through WEEE Reuse and Repair : Opportunities and Challenges in the UK Producer Responsibility ( PR )," pp. 1–9, 2016.
- [4] A. M. King, S. C. Burgess, W. Ijomah, and C. A. McMahon, "Reducing waste: repair, recondition, remanufacture or recycle?," *Sustain. Dev.*, vol. 14, no. 4, pp. 257–267, Oct. 2006, doi: 10.1002/sd.271.
- [5] A. C. Nelson, J. Genereux, and M. M. Genereux, "Price effects of landfills on different house value strata," *J. Urban Plan. Dev.*, vol. 123, no. 3, pp. 59–67, 1997.
- [6] D. K. J. Ara Begum, "Electronic Waste (E-Waste) Management in India: A Review," *IOSR J. Humanit. Soc. Sci.*, vol. 10, no. 4, pp. 46–57, 2013, doi: 10.9790/0837-1044657.
- [7] R. G. de Souza, J. C. N. Clímaco, A. P. Sant'Anna, T. B. Rocha, R. de A. B. do Valle, and O. L. G. Quelhas, "Sustainability assessment and prioritisation of e-waste management options in Brazil," *Waste Manag.*, vol. 57, pp. 46–56, 2016, doi: 10.1016/j.wasman.2016.01.034.
- [8] Y. Xu and C. H. Yeh, "Sustainability-based selection decisions for e-waste recycling operations," *Ann. Oper. Res.*, vol. 248, no. 1–2, pp. 531–552, 2017, doi: 10.1007/s10479-016-2269-2.
- [9] B. Debnath, R. Chowdhury, and S. K. Ghosh, "Sustainability of metal recovery from E-waste," *Front. Environ. Sci. Eng.*, vol. 12, no. 6, pp. 1–12, 2018, doi: 10.1007/s11783-018-1044-9.
- [10] L. H. Xavier, E. C. Giese, A. C. Ribeiro-Duthie, and

- F. A. F. Lins, "Sustainability and the circular economy: A theoretical approach focused on e-waste urban mining," *Resour. Policy*, no. October 2017, p. 101467, 2019, doi: 10.1016/j.resourpol.2019.101467.
- [11] P. Kumare Gopalakrishnan and S. Behdad, "Usage of Product Lifecycle Data to Detect Hard Disk Drives Failure Factors." 06-Aug-2017, doi: 10.1115/DETC2017-67973.
- [12] A. Romero and D. R. Vieira, "Using the product lifecycle management systems to improve maintenance, repair and overhaul practices: the case of aeronautical industry," in *IFIP International Conference on Product Lifecycle Management*, 2014, pp. 159–168.
- [13] Y. Zhang, S. Ren, Y. Liu, T. Sakao, and D. Huisingh, "A framework for Big Data driven product lifecycle management," *J. Clean. Prod.*, vol. 159, pp. 229–240, 2017, doi: <https://doi.org/10.1016/j.jclepro.2017.04.172>.
- [14] R. Houssin and A. Coulibaly, "Safety-based availability assessment at design stage," *Comput. Ind. Eng.*, vol. 70, pp. 107–115, 2014, doi: <https://doi.org/10.1016/j.cie.2014.01.005>.
- [15] Z. Wu *et al.*, "Product lifecycle-oriented knowledge services: Status review, framework, and technology trends," *Concurr. Eng.*, vol. 25, no. 1, pp. 81–92, 2017.
- [16] B. Liu, Q. Zhao, Y. Jin, J. Shen, and C. Li, "Application of combined model of stepwise regression analysis and artificial neural network in data calibration of miniature air quality detector," *Sci. Rep.*, vol. 11, no. 1, pp. 1–12, 2021.
- [17] Z. Leszczyński and T. Jasiński, "Comparison of Product Life Cycle Cost Estimating Models Based on Neural Networks and Parametric Techniques—A Case Study for Induction Motors," *Sustainability*, vol. 12, no. 20, p. 8353, 2020.
- [18] M. C. Liu, W. Kuo, and T. Sastri, "An exploratory study of a neural network approach for reliability data analysis," *Qual. Reliab. Eng. Int.*, vol. 11, no. 2, pp. 107–112, Jan. 1995, doi: <https://doi.org/10.1002/qre.4680110206>.
- [19] K. Xu, M. Xie, L. C. Tang, and S. L. Ho, "Application of neural networks in forecasting engine systems reliability," *Appl. Soft Comput.*, vol. 2, no. 4, pp. 255–268, 2003, doi: [https://doi.org/10.1016/S1568-4946\(02\)00059-5](https://doi.org/10.1016/S1568-4946(02)00059-5).
- [20] Wikipedia contributors, "Naive Bayes classifier." 2021.
- [21] K. P. Murphy, "Naive bayes classifiers," *Univ. Br. Columbia*, vol. 18, no. 60, 2006.
- [22] V. N. Vapnik, *The nature of statistical learning theory*. Springer-Verlag New York, Inc., 1995.
- [23] R. Muhammad, X. Yuan, O. Kisi, and Y. Yuan, "Streamflow Forecasting Using Artificial Neural Network and Support Vector Machine Models," *Am. Sci. Res. J. Eng. Technol. Sci.*, vol. 29, no. 1, pp. 286–294, 2017.
- [24] M. J. Chang *et al.*, "A support vector machine forecasting model for typhoon flood inundation mapping and early flood warning systems," *Water (Switzerland)*, vol. 10, no. 12, 2018, doi: 10.3390/w10121734.
- [25] A. Zendeboudi, M. A. Baseer, and R. Saidur, "Application of support vector machine models for forecasting solar and wind energy resources: A review," *J. Clean. Prod.*, vol. 199, pp. 272–285, 2018, doi: 10.1016/j.jclepro.2018.07.164.
- [26] J. F. Bard, "An Efficient Point Algorithm for a Linear Two-Stage Optimization Problem," *Oper. Res.*, vol. 31, no. 4, pp. 670–684, Feb. 1983.
- [27] R. Damarta, A. Hidayat, and A. S. Abdullah, "The application of k-nearest neighbors classifier for sentiment analysis of PT PLN (Persero) twitter account service quality," in *Journal of Physics: Conference Series*, 2021, vol. 1722, no. 1, p. 12002.
- [28] L. A. Demidova and Y. S. Sokolova, "Approbation of the data classification method based on the SVM algorithm and the k nearest neighbors algorithm," in *IOP Conference Series: Materials Science and Engineering*, 2021, vol. 1027, no. 1, p. 12001.
- [29] F. Pedregosa *et al.*, "Scikit-learn: Machine learning in Python," *J. Mach. Learn. Res.*, vol. 12, pp. 2825–2830, 2011.
- [30] G. J. McLachlan and S. Rathnayake, "On the number of components in a Gaussian mixture model," *Wiley Interdiscip. Rev. Data Min. Knowl. Discov.*, vol. 4, no. 5, pp. 341–355, 2014.
- [31] H. Sahbi, "A particular Gaussian mixture model for clustering and its application to image retrieval," *Soft Comput.*, vol. 12, no. 7, pp. 667–676, 2008.
- [32] A. Glowacz, W. Glowacz, Z. Glowacz, and J. Kozik, "Early fault diagnosis of bearing and stator faults of the single-phase induction motor using acoustic signals," *Measurement*, vol. 113, pp. 1–9, 2018, doi: <https://doi.org/10.1016/j.measurement.2017.08.036>.
- [33] Z. Wang, J.-L. Zarader, and S. Argentieri, "A novel aircraft fault diagnosis and prognosis system based on Gaussian mixture models," in *2012 12th International Conference on Control Automation Robotics & Vision (ICARCV)*, 2012, pp. 1794–1799.
- [34] B. Desgraupes, "Clustering indices," *Univ. Paris Ouest-Lab Modal'X*, vol. 1, p. 34, 2013.