



Universiteit
Leiden
The Netherlands

Trustworthy anomaly detection for smart manufacturing

Li, Z.

Citation

Li, Z. (2025, May 1). *Trustworthy anomaly detection for smart manufacturing*. *SIKS Dissertation Series*. Retrieved from <https://hdl.handle.net/1887/4239055>

Version: Publisher's Version

License: [Licence agreement concerning inclusion of doctoral thesis in the Institutional Repository of the University of Leiden](#)

Downloaded from: <https://hdl.handle.net/1887/4239055>

Note: To cite this publication please use the final published version (if applicable).

Chapter 3

Feature Selection for Fault Detection and Prediction based on Event Log Analysis

Authors: **Zhong Li**, Matthijs van Leeuwen

Published in ACM SIGKDD Explorations Newsletter 24.2 (2022): 96-104.

Abstract

Event logs are widely used for anomaly detection and prediction in complex systems. Existing log-based anomaly detection methods usually consist of four main steps: log collection, log parsing, feature extraction, and anomaly detection, wherein the feature extraction step extracts useful features for anomaly detection by counting log events. For a complex system, such as a lithography machine consisting of a large number of subsystems, its log may contain thousands of different events, resulting in abounding extracted features. However, when anomaly detection is performed at the subsystem level, analyzing all features becomes expensive and unnecessary. To mitigate this problem, we develop a feature selection method for log-based anomaly detection and prediction. Specifically, our method consists of three main modules: the *Log Event Vectorization* module that converts semi-structured log texts into time series; the *Selection of Relevant Features* module that leverages Kendall rank correlation and Granger causality test to select log events for fault detection and prediction; and the *Removal of Redundant Features* module that utilizes Kendall rank correlation to reduce redundant log events. Results on 25 real-world datasets show that our method can detect and predict faults more accurately by selecting a small proportion of log events, thereby improving the effectiveness and efficiency.

3.1 Introduction

A lithography machine is a piece of complex structural equipment used to manufacture chips. Typically, it consists of the following main subsystems: the light source subsystem, the objective lens subsystem, the table subsystem, the mask table subsystem, the mask transfer subsystem, the wafer transfer subsystem, and the exposure subsystem [111]. Particularly, the wafer transfer subsystem serves to transfer silicon wafers between the track and wafer stage, having a great impact on the precision of chip fabrication. A wafer transfer subsystem usually contains two robots, namely a load robot and an unload robot. When the lithography machine goes into production, these two robots may encounter faults. We assume there are $L + M$ types of faults, viz. $GF_1, GF_2, \dots, GF_L$ and $SF_1, SF_2, \dots, SF_M$. Specifically, $GF_1, GF_2, \dots, GF_L$ represent faults that occur gradually and thus can be detected in an early stage (i.e., they are typically predictable). In contrast, $SF_1, SF_2, \dots, SF_M$ denote faults that generally occur suddenly and are often hard to predict.

To minimize machine downtime and thus maximize productivity, the possible faults of load and unload robots should be detected and predicted (if possible) in an automated way. To this end, the wafer transfer subsystem usually uses sensors to measure the position of the two robots in real time, collecting time series data that can be used for data-driven fault detection and prediction. However, due to the limited information contained in sensor data, it is challenging to detect all possible faults based on time series data alone. Meanwhile, as shown in Figure 3.1 and Table 3.1, a lithography machine also has an information system that records all triggered events—in the form of logs—when the machine is working. Since the different subsystems in a lithography machine are interconnected, a fault incurred in one subsystem (e.g., the wafer transfer system) may trigger events not only in that subsystem, but also in other subsystems. Besides, the components in the same subsystem are usually closely interconnected. Therefore, the fault of one component is very likely to cause faults of other components. Therefore, the event logs contain important information for fault detection and prediction. Traditional log-based anomaly detection methods can be used to detect such faults [113].

Due to the complexity of the lithography machine, there can be thousands of unique log events, resulting in millions of log events in a relatively short working time of the machine. Moreover, when attempting to detect faults of certain components in a specific subsystem (e.g., the load and unload robots in the wafer transfer subsystem), many of these log events are irrelevant or abundant. Hence, a direct application of

3.1. Introduction

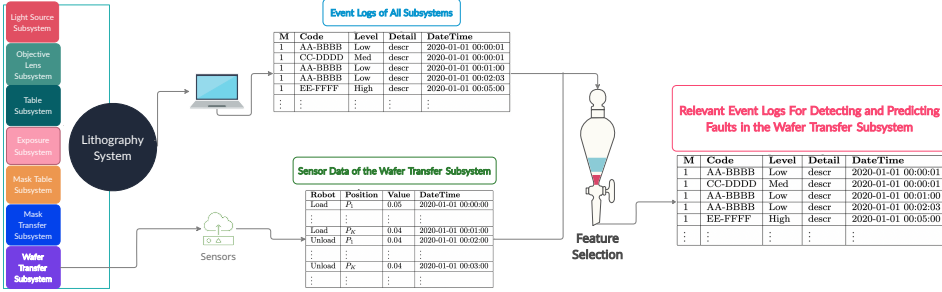


Figure 3.1: The composition of the lithography system, and an illustration of the feature selection problem addressed in this work.

existing log-based anomaly detection methods on all log events can be computationally prohibitive and may also produce misleading detection results due to the inclusion of irrelevant log events. To mitigate this problem, we regard each log event as a feature and develop a feature selection method that aims to select relevant features for log-based fault detection and prediction.

In brief, our method consists of three main modules, namely *Log Event Vectorization*, *Selection of Relevant Features* and *Removal of Redundant Features*. Specifically, the *Log Event Vectorization* module aims at converting unstructured log events into time series data; the *Selection of Relevant Features* module attempts to select relevant features for fault detection and prediction by using the variables measured by sensors as target; and the *Removal of Redundant Features* module focuses on eliminating redundant features to further reduce the number of selected features. The main contributions of this paper can be summarized as follows:

- 1) We formalize the problem of feature selection for fault detection and prediction based on event log analysis. To our knowledge, this problem is not uncommon in the industry, but has not yet been studied.
- 2) To address this problem, we further propose a novel approach using two types of feature selection.
- 3) We demonstrate the effectiveness of our approach by conducting extensive experiments on 25 real-world datasets gathered from lithography systems.

The remainder of this paper is organized as follows. Chapter 3.2 introduces related work. Next, Chapter 3.3 introduces terminology, formalizes the problem, and presents our proposed method in detail. Chapter 3.4 empirically evaluates the performance of

Table 3.1: An example event log. All values are fictional.

M	Code	Level	Detail	DateTime
1	AA-BBBB	Low	descr	2020-01-01 00:00:01
1	CC-DDDD	Med	descr	2020-01-01 00:00:01
1	AA-BBBB	Low	descr	2020-01-01 00:01:00
1	AA-BBBB	Low	descr	2020-01-01 00:02:03
1	EE-FFFF	High	descr	2020-01-01 00:05:00
⋮	⋮	⋮	⋮	⋮

our method. Finally, Chapter 3.5 concludes the paper and discusses possible future directions.

3.2 Related Work

Existing log-based anomaly detection methods usually consist of four main steps: *Log Collection*, *Log Parsing*, *Feature Extraction* and *Anomaly Detection* [114]. First, the *Log Collection* step is responsible for recording triggered events in the form of logs. A record is called a *log message*, which usually contains the date and time of occurrence and the detailed description of event. More concretely, detailed descriptions are typically presented in predefined templates, and may also include parameters. Second, the *Log Parsing* step aims at converting each log message into a specific *log event* template [93]. Usually, a *log event* corresponds to a unique template. Third, based on derived log events, the *Feature Extraction* step attempts to convert each *log sequence* into a *log count vector*. Specifically, a *log sequence* is composed of multiple *log events*. In general, a *log count vector* is a vector with each entry indicating the number of times that the corresponding *log event* was triggered. Note that the entries of the *log count vector* can be computed in another refined way [115]. Finally, the *Anomaly Detection* step performs anomaly detection based on the extracted *log count vectors*.

There exist many log parsing methods based on clustering [116], frequent pattern mining [117], or heuristic techniques [92]. Since our work centers around feature selection (e.g., log event selection) for fault detection and prediction, we will only consider the *Feature Extraction* and *Anomaly Detection* steps.

He et al. [118] have performed a systematic comparison of five state-of-the-art log-based anomaly detection methods, including DeepLog [65], LogAnomaly [66], PLELog [71], LogRobust [119], CNN [120]. Their findings show that log-based anomaly detec-

3.3. Method

tion methods are often not as good as expected (i.e., as claimed by the original authors of each method) in real-world datasets. Importantly, they found that some log-based anomaly detectors, such as LogAnomaly, perform poorly on datasets with numerous log events. Therefore, feature selection can be exploited to improve the performance of these anomaly detection methods. However, we are not aware of any existing publications that attempt to address this problem.

3.3 Method

In the following section, we first introduce the terminology used in this work and then formalize the problem. Next, we elucidate our proposed method that consists of three modules: *Log Event Vectorization*, *Selection of Relevant Features*, and *Removal of Redundant Features*.

3.3.1 Terminology and Problem Formulation

We assume that we have access to two types of data. First, as shown in Table 3.1, we assume that the log data in a lithography machine, denoted by $\mathbf{X}$, has been collected and accurately parsed. Without loss of generality, we suppose there is a **Code** as the unique identifier for each *log event*, a **Level** roughly indicating the severity level of triggered *log event*, a **Detail** describing the detail of each *log message* that is an instantiation of a *log event* using a predefined template, and a **DateTime** containing the corresponding date and time. Hereinafter, we also call each *log event* a *log feature*. In addition, we may have log data for multiple lithography machines, and we use $\mathbf{M}$ to represent the corresponding name of the machine.

Ideally, applying an existing log-based anomaly detection method on $\mathbf{X}$ can detect most faults related to load and unload robots. However, due to the large number of *log features*, it is computationally prohibitive to directly use most existing anomaly detection methods. Furthermore, the presence of irrelevant *log features* may significantly degrade detection performance and even lead to misleading detection results.

Second, we also assume the availability of sensor data, denoted by $\mathbf{Y}$, that measures the positions of robots. As shown in Table 3.2, there are measurements and corresponding timestamps from K different positions for the load robot and the unload robot, respectively. By using the *Log Event Vectorization* module in our proposed method (as will be explained in the sequel), $\mathbf{Y}$ can be rewritten as $(\mathbf{LP}_1, \dots, \mathbf{LP}_K, \mathbf{UP}_1, \dots, \mathbf{UP}_K)$. Specifically, for $k \in \{1, \dots, K\}$, $\mathbf{LP}_k = \{(Value_t, DateTime_t)\}_{t \in \mathbf{T}}$ de-

Table 3.2: Example sensor data. All values are fictional.

Robot	Position	Value	DateTime
Load	P_1	0.05	2020-01-01 00:00:00
$\vdots$	$\vdots$	$\vdots$	$\vdots$
Load	P_K	0.04	2020-01-01 00:01:00
Unload	P_1	0.04	2020-01-01 00:02:00
$\vdots$	$\vdots$	$\vdots$	$\vdots$
Unload	P_K	0.04	2020-01-01 00:03:00
$\vdots$	$\vdots$	$\vdots$	$\vdots$

notes the corresponding time series of load robot from position k , and at the same time $\mathbf{UP}_k = \{(Value_t, DateTime_t)\}_{t \in \mathbf{T}}$ denotes the corresponding time series of unload robot from position k , where $\mathbf{T}$ denotes the timestamp when the time series was sampled.

By applying an appropriate time series anomaly detection method on $(\mathbf{LP}_1, \dots, \mathbf{LP}_K)$ and $(\mathbf{UP}_1, \dots, \mathbf{UP}_K)$, we can detect certain faults (especially gradual faults) of the load and unload robot, respectively. However, due to the limited fault information contained in $\mathbf{Y}$, these faults are difficult to predict using sensor data only.

Therefore, we aim to address the following problem: *Suppose there is a complex system Δ that is composed of several interconnected subsystems $\{\Gamma, \Lambda, \dots, \Theta\}$. Given a database of logs $\mathbf{X}_\Delta$ generated by the entire system Δ and a database $\mathbf{Y}_\Theta$ consisting of time series measured by sensors from a certain subsystem Θ , we assume the set of unique log events in $\mathbf{X}_\Delta$ is $\mathbf{C}_\Delta = \{C_1, C_2, \dots, C_N\}$. Based on $\mathbf{Y}_\Theta$, we attempt to select a subset of $\mathbf{C}_\Delta$, denoted by $\mathbf{C}_\Theta = \{C'_1, C'_2, \dots, C'_L\}$ with $L \ll N$ and resulting in a new database $\mathbf{X}_\Theta \subset \mathbf{X}_\Delta$, that mainly keeps relevant log events for detecting and predicting faults that occurred in subsystem Θ .*

Specifically, as shown in Figure 3.1, system Δ is a lithography machine and subsystem Θ is the wafer transfer subsystem in this work. To solve the above problem, we propose a method consisting of three modules, namely *Log Event Vectorization*, *Selection of Relevant Features* and *Remove of Redundant Features*, each of which will be separately described next.

3.3.2 Log Event Vectorization (Module 1)

The first module, namely *Log Event Vectorization*, mainly aims at converting unstructured or semi-structured log events into time series data. Considering the log messages

3.3. Method

given in Table 3.1, it is straightforward to generate a time series for each **Code** per machine by keeping only the corresponding records. Without loss of generality, we assume that the smallest unit of time in the original data is seconds. Accordingly, for each log event in a certain machine, we can obtain a time series in the form of $\{(Value_t, DateTime_t)\}_{t \in \mathbf{T}}$, meaning that this log event is triggered $Value_t$ times at the timestamp $DateTime_t$ for $t \in \mathbf{T}$, where $\mathbf{T}$ represent all time points (an ordered list) when this log event was triggered.

After preliminary results and based on domain knowledge, we have decided to take each day as an interval to count the number of times a log event is triggered. Besides, we take the start point of this time interval to represent the time point when this log event is triggered.

3.3.3 Selection of Relevant Features (Module 2)

Given that the load robot and unload robots are very similar, for simplicity, we only consider the load robot when elucidating the proposed method. That is, we assume a multivariate time series database $\mathbf{Y} = (\mathbf{Y}_1, \dots, \mathbf{Y}_K)$. Note that all these time series are of equal length. Therefore, we assume $\mathbf{Y}_k = \{y_{kt}\}_{t \in \mathbf{T}}$ for $k \in \{1, \dots, K\}$, with $\mathbf{T}$ denoting the timestamp when $\mathbf{Y}_k$ was sampled. Meanwhile, after applying the *Log Event Vectorization* module on the log event database $\mathbf{X}$, we can obtain another multivariate time series database $\mathbf{Z} = (\mathbf{Z}_1, \dots, \mathbf{Z}_n, \dots, \mathbf{Z}_N)$, where N represents the number of unique log features (i.e., log events). Although different log events are usually triggered at different time points, we have taken each day as an interval to count the number of times that each log event is triggered. As a result, for $n \in \{1, 2, \dots, N\}$, $\mathbf{Z}_n$ has a fixed length and thus we assume $\mathbf{Z}_n = \{z_{ns}\}_{s \in \mathbf{S}}$, with $\mathbf{S}$ denoting the timestamp when $\mathbf{Z}_n$ was sampled.

To detect faults, for $k \in \{1, \dots, K\}$, we can apply a univariate time series anomaly detector $\phi(\cdot)$ on $\mathbf{Y}_k$, resulting in $\phi(\mathbf{Y}_k)$. Alternatively, we can apply a multivariate time series anomaly detector $\Phi(\cdot)$ on $\mathbf{Y}$ by jointly considering all $\mathbf{Y}_k$ for $k \in \{1, \dots, K\}$, leading to $\Phi(\mathbf{Y}_1, \dots, \mathbf{Y}_K)$. As shown in Figure 3.2 (the two subplots at the bottom), the identification of faults is feasible by applying time series anomaly detectors on $\mathbf{Y}$. However, for gradual faults, just identifying them is not enough. It is also necessary to predict them accurately. Since there is only limited fault information in $\mathbf{Y}$, it is difficult to predict these faults based on $\mathbf{Y}$ alone. Therefore, we attempt to select relevant log features from $\mathbf{Z}$ to better detect and predict faults. For $k \in \{1, \dots, K\}$, by considering $\mathbf{Y}_k$ as the target variable and $\mathbf{Z}_1, \dots, \mathbf{Z}_n, \dots, \mathbf{Z}_N$ as the prediction variables, it

Chapter 3. Feature Selection for Fault Detection and Prediction based on Event Log Analysis

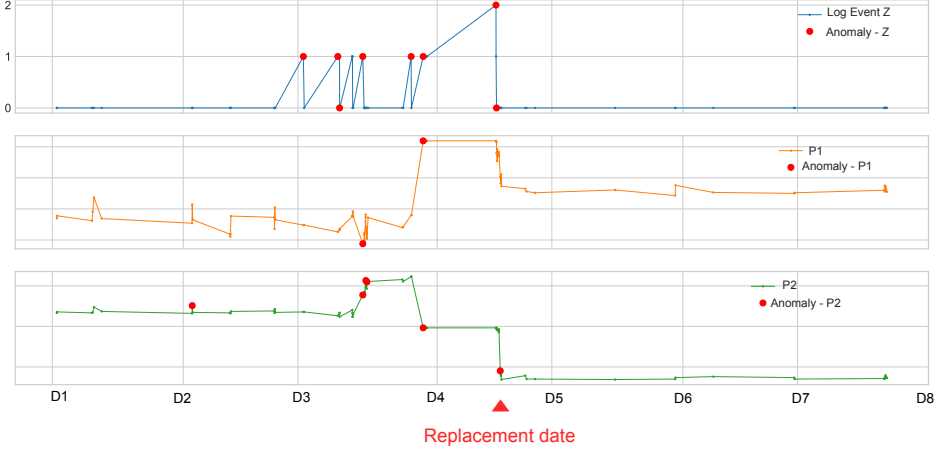


Figure 3.2: An example showing the relevance of a specific log event (upper plot) to detecting and predicting faults of a load robot based on sensor data (lower two plots). The shared x-axis represents the timestamp and the y-axes represent the measured values of each feature. Note that the y-axes of P_1 and P_2 are intentionally hidden and their timestamps are anonymized.

becomes a supervised feature selection problem. Compared to traditional supervised feature selection problems, however, there are three novel challenges:

- 1) The features considered are time series rather than numeric tabular data;
- 2) The target and prediction variables are not of equal length, and their timestamps are also different;
- 3) Traditional similarity metrics (e.g. Euclidean distance, dynamic time warping) do not give meaningful results when measuring the relevance/similarity of the predictor variable to the target variable.

We now detail why traditional similarity metrics fail to provide meaningful results when trying to find relevant log features. As shown in Figure 3.2, we can see that from ‘D3’ some faults started to appear in the robot and became detectable after a certain time based on P_1 and P_2 . These faults disappeared after the replacement of specific components on ‘Replacement Date’. Meanwhile, we can observe that log event Z was triggered several times before the replacement date. More importantly, it was triggered several times even before the faults became detectable based on P_1 and P_2 . At other times, this log event was not triggered. In other words, the log event Z could

3.3. Method

potentially be used to detect and predict these faults. However, if we consider their shapes (wrapped or not) or values (normalized or not), we can see that log feature $\mathbf{Z}$ is not similar to $\mathbf{P1}$ or $\mathbf{P2}$. To address this problem, we propose a novel similarity metric that consists of three steps, as follows.

As summarized in Algorithm 2, for $k \in \{1, \dots, K\}$, we first apply an appropriate univariate time series anomaly detector $\phi(\cdot)$ on $\mathbf{Y}_k$, resulting in a time series of anomaly scores $\mathbf{U}_k = \{u_{kt}\}_{t \in \mathbf{T}}$ (Lines 5-8). Second, for $n \in \{1, 2, \dots, N\}$, we apply an appropriate univariate time series anomaly detector $\varphi(\cdot)$ on $\mathbf{Z}_n$, resulting in a time series of anomaly scores $\mathbf{V}_n = \{v_{ns}\}_{s \in \mathbf{S}}$ (Lines 9-12). Note that $\phi(\cdot)$ and $\varphi(\cdot)$ can be different considering that $\mathbf{Z}_n$ is sampled at a regular frequency (i.e., an observation per day) but $\mathbf{Y}_k$ is sampled at an irregular frequency (e.g., multiple observations in day A but no observation in day B). Third, we can select relevant log events by comparing $\mathbf{U}_k$ with $\mathbf{V}_n$ (Lines 13-23). Note that the lengths and scales of $\mathbf{U}_k$ and $\mathbf{V}_n$ may be different, but their peaks should overlap (for detection) or preferably the peak of $\mathbf{V}_n$ precedes the corresponding peak of $\mathbf{U}_k$ (for prediction) if we compare them using the same timeline. A peak here means a relatively high degree of outlyingness.

Time series anomaly detector $\phi(\cdot)$

Since $\mathbf{Y}_k$ is sampled at an irregular frequency, it may have multiple observations on a given day, but no observations in the following days. However, traditional time series anomaly detection methods usually assume that the input time series is regularly sampled. To circumvent this limitation, we adapt a simple yet effective anomaly detection strategy, which is called *persistence checking*. Specifically, for each time series value Y_{kt} in $\mathbf{Y}_k$, we compare this value with its previous value to obtain an anomaly score, defined as $U_{kt} = \phi(Y_{kt}) = |Y_{kt} - Y_{kh}|$ with $t = h + 1$. Particularly, this anomaly detector can deal with time series sampled at irregular frequencies.

Time series anomaly detector $\varphi(\cdot)$

Although $\mathbf{Z}_n$ is in the form of a time series, we are not concerned about the temporal order of observations when detecting anomalies. This is because if the lithography machine is working under a normal production environment, the occurrence of each log event should remain stable. Therefore, we can apply a traditional anomaly detector designed for tabular data on it. Specifically, we define $V_{ns} = \varphi(Z_{ns}) = \frac{|Z_{ns} - \text{med}(\mathbf{Z}_n)|}{\text{std}(\mathbf{Z}_n)}$ as the anomaly score for the sample point Z_{ns} in the time series $\mathbf{Z}_n$, where *med* and *std* denote the median and standard deviation of all sample points in $\mathbf{Z}_n$, respectively.

Feature selection for fault detection

We perform feature selection for fault detection by comparing obtained anomaly scores $\mathbf{U}_k$ and $\mathbf{V}_n$. Assuming that the robot works continuously for T days, for $\mathbf{V}_n$ we have a value per day. However, for $\mathbf{U}_k$, we may have multiple values on some days, but no values on most days. To make $\mathbf{U}_k$ and $\mathbf{V}_n$ comparable, we modify $\mathbf{U}_k$ as follows: for each day, if there are multiple anomaly scores, we take the maximum of these scores as the final anomaly score; if there is no anomaly score, we set the final anomaly score as zero. We denote the modified $\mathbf{U}_k$ as $\hat{\mathbf{U}}_k$, which has the same length as $\mathbf{V}_n$.

On this basis, we compute the Kendall's τ coefficient [121] between $\hat{\mathbf{U}}_k$ and $\mathbf{V}_n$ for feature selection. We prefer Kendall's rank correlation measure over other correlation measures for several reasons:

- it measures monotonicity relationships and has a straightforward interpretation;
- it does not require the tested features to follow a normal distribution, as anomaly scores usually do not follow a normal distribution;
- it is robust to noise and can be calculated easily.

Suppose $\mathbf{V}_n = (\alpha_1, \dots, \alpha_T)$ and $\hat{\mathbf{U}}_k = (\beta_1, \dots, \beta_T)$, then Kendall's τ coefficient measures the similarity of orderings of elements in $\hat{\mathbf{U}}_k$ and $\mathbf{V}_n$. Specifically, we let $(\alpha_1, \beta_1), \dots, (\alpha_T, \beta_T)$ be paired observations. Furthermore, (α_i, β_i) and (α_j, β_j) are regarded as concordant if and only if one of the following conditions are fulfilled:

- $\alpha_i < \alpha_j$ and $\beta_i < \beta_j$;
- $\alpha_i > \alpha_j$ and $\beta_i > \beta_j$.

Otherwise, (α_j, β_j) are considered discordant. Kendall's τ coefficient is then defined as:

$$\tau = \frac{\#concordant - \#discordant}{\#concordant + \#discordant}, \quad (3.1)$$

where $\#concordant$ and $\#discordant$ represent the number of concordant pairs and the number of discordant pairs, respectively. Hence, τ can take a value in $[-1, 1]$, with a high absolute value indicating a high relevancy of log event $\mathbf{Z}_n$ to $\mathbf{Y}_k$ for detecting faults. By setting a threshold τ_0 , we can select a subset of log events that are considered to be the most relevant for fault detection.

3.3. Method

Feature selection for fault prediction based on Kendall rank correlation test

For $k \in \{1, \dots, K\}$ and $n \in \{1, \dots, N\}$, after obtaining $\mathbf{V}_n = (\alpha_1, \dots, \alpha_T)$ and $\hat{\mathbf{U}}_k = (\beta_1, \dots, \beta_T)$ as described in Chapter 3.3.3, we achieve feature selection for fault prediction by comparing $\mathbf{V}_n$ and the lagged values of $\hat{\mathbf{U}}_k$. The underlying rational is that if $\mathbf{Z}_n$ is useful for predicting faults in $\mathbf{Y}_k$, then the anomaly score vector $\mathbf{V}_n$ should be correlated with the lagged values of anomaly score vector $\hat{\mathbf{U}}_k$. In other words, the value of Kendall's τ coefficient between $\mathbf{V}_n$ and $\text{Lag}(\hat{\mathbf{U}}_k, l)$ should be high for some $l \in \{1, 2, \dots, T\}$, where $\text{Lag}(\cdot, l)$ denotes the lag operator and l represents the the number of lagged times. For instance,

$$\text{Lag}((\beta_1, \beta_2, \beta_3, \dots, \beta_{T-2}, \beta_{T-1}, \beta_T), 2) = (\beta_3, \beta_4, \beta_5, \dots, \beta_T).$$

Hence, the resulting vector is shorter than the original vector.

For a specific l , we calculate the Kendall's τ coefficient between $\mathbf{V}_n$ and $\text{Lag}(\hat{\mathbf{U}}_k, l)$ as follows. We let $(\alpha_1, \beta_{1+l}), \dots, (\alpha_{T-l}, \beta_T)$ be paired observations, and then count the number of concordant pairs and discordant pairs, followed by using Equation (3.1) to obtain the corresponding value of τ . A high value of τ indicates a high relevancy of log event $\mathbf{Z}_n$ to predict faults in $\mathbf{Y}_k$. Particularly, although l can take a value in $\{1, 2, \dots, T\}$, an overly large value will cause few observations, and may produce misleading results. Therefore, we only consider $l \in \{1, 2, 3, 4, 5\}$ in this work. Similarly, this coefficient can take a value in $[-1, 1]$, with a high absolute value indicating a high relevancy of log event $\mathbf{Z}_n$ to predict $\mathbf{Y}_k$. By setting a threshold τ_1 , we can select a subset of log events that are considered to be the most relevant for fault prediction.

Our preliminary results suggest that only applying Kendall rank correlation to select features for fault prediction may not be sufficient. Given a value for l , we only consider a single and equal number of lag operations at each time point. Therefore, it may result in few features selected for prediction. To mitigate this problem, we utilize Granger causality test to select more features for fault prediction, as follows.

Feature selection for fault prediction based on Granger causality test

The Granger causality test is widely used to determine whether a time series is useful for predicting another time series [122]. Formally, given an information set $\Omega_t = (\alpha, \beta)$ with $\alpha = (\alpha_1, \dots, \alpha_t)$ and $\beta = (\beta_1, \dots, \beta_t)$, α is said to Granger cause β if and only if $\delta_f^2(\beta_t - P(\beta_t | \beta_{j:j < t}, \alpha_{i:i < t})) < \delta_r^2(\beta_t - P(\beta_t | \beta_{j:j < t}))$. Specifically, δ_f^2

represents the variance of prediction errors generated by the optimal linear predictor $P(\beta_t | \beta_{j:j < t}, \alpha_{i:i < t})$ based on full information $(\alpha_{i:i < t}, \beta_{j:j < t})$. Meanwhile, δ_r^2 denotes the variance of prediction errors generated by the optimal linear predictor $P(\beta_t | \beta_{j:j < t})$ based on reduced information $(\beta_{j:j < t})$. In other words, the ‘usefulness’ for prediction in Granger causality test means the ability to increase the prediction accuracy. Therefore, the Granger causality test is suitable for selecting log features that can be used to more accurately predict faults.

More concretely, we specify the following linear prediction equation:

$$\beta_t = c + \sum_{i=1}^I a_i \alpha_{t-i} + \sum_{j=1}^J b_j \beta_{t-j} + \mu_t, \quad (3.2)$$

where I and J represent the number of lagged operations considered for α and β , respectively, while μ_t denotes the corresponding residual. Intuitively, if $\sum_{i=1}^I |a_i| \neq 0$, we can say α is useful for predicting β . Accordingly, we call Equation (3.2) the full or unrestricted regression model. Meanwhile, we call the following equation the reduced or restricted regression model:

$$\beta_t = c + \sum_{j=1}^J b_j \beta_{t-j} + \mu_t, \quad (3.3)$$

Specifically, given two time series $\alpha = (\alpha_1, \dots, \alpha_t)$ and $\beta = (\beta_1, \dots, \beta_t)$, the Granger causality test calculates the following F -statistic:

$$F = \frac{\frac{SSR_r - SSR_f}{J}}{\frac{SSR_f}{D - (I + J + 1)}}, \quad (3.4)$$

where SSR_r is the sum of squared residuals on the reduced regression model and SSR_f is the sum of squared residuals on the full regression model. Besides, D denotes the degrees of freedom (i.e., the number of observations), while I and J have the same meaning as in Equation (3.2). Consequently, a higher value of F indicates that α is more useful for predicting β . As a result, α is declared Granger causal for β if the F -statistic is larger than the $(1 - Sig)\%$ quantile of an $F(J, D - (I + J + 1))$ distribution, where Sig denotes the significance level.

We should be aware that the Equation (3.4) requires several explicit and implicit assumptions to effectively identify Granger causality [122]:

- **Linearity:** the generating processes of the two time series are linear, and their

3.4. Experiments and Results

causal effects are also linear;

- Continuous-valued series: the two time series are supposed to have continuous-valued observations;
- Discrete-time: the two time series are sampled on a discrete and regular manner;
- Stationarity: the statistical properties of these two time series do not change over time;
- Known lag: the linear dependency on past values are assumed to have a known order;
- Perfectly observed: the two time series do not have measure errors;
- Complete system: there are no unmeasured confounders.

Overall, this subsection is summarized on Lines 5-23 in Algorithm 2.

3.3.4 Removal of Redundant Features (Module 3)

After selecting a subset of relevant log events for detecting and predicting faults in Module 2, we can further reduce the number of selected log events by computing the correlation between them. Specifically, we compute the pairwise Kendall’s τ coefficient between selected log events and remove the redundant ones by setting a threshold τ_2 for the τ coefficient. More concretely, given two time series α and β , if the absolute value of their pairwise τ coefficient exceeds the threshold, we randomly remove one feature. This subsection is summarized on Lines 24-27 in Algorithm 2.

3.4 Experiments and Results

To demonstrate the effectiveness and efficiency of our method, we perform a series of experiments. The experimental setup, results, and corresponding analysis are described as follows.

3.4.1 Data Description

Given the novelty of the problems faced by this work, we are not aware of any publicly available datasets that can be used. The traditional benchmark datasets for

Chapter 3. Feature Selection for Fault Detection and Prediction based on Event Log Analysis

Table 3.3: Summary of datasets and experiment results. *Replacement* indicates the date when some components of the robot are replaced (faults always happen on or earlier than this date). *#Messages* represents the number of log messages. *#Raw* denotes the number of log features in the original dataset and *#Selected* denotes the number of selected log features. Besides, AD indicates whether the fault is detected based on the corresponding log features, where ‘Yes’ means the failure has been successfully detected or predicted and ‘No’ otherwise. If a fault is detected early (i.e., predicted), we further use ‘P’ to highlight it. Moreover, if a fault is only detected or predicted with the top 3 points but not the top 1 point, we denote it with an asterisk. Note that the value of *Fault* is anonymized, where *GF*, *SF* and *Unknown* represent gradual faults, sudden faults or unknown types of faults, respectively.

Machine	Robot	Fault	Replacement	#Messages	#Raw(AD)	#Selected(AD)
1	Unload	<i>GF</i> ₁	2021-01-05	103294	301(Yes)*	34(Yes)
2	Unload	<i>GF</i> ₁	2021-01-06	90974	246(No)	47(Yes/P)
3	Unknown (Load)	<i>GF</i> ₁	2021-01-10	90783	276(Yes)*	37(Yes)*
4	Unload	<i>GF</i> ₁	2021-02-08	93729	437(No)	48(Yes)*
5	Unload	<i>GF</i> ₁	2021-02-08	76797	245(No)	16(Yes)*
6	Load	<i>SF</i> ₁	2019-06-24	86988	303(Yes)	38(Yes)
7	Unload	<i>SF</i> ₁	2020-03-10	95513	376(Yes/P)	49(Yes/P)
8	Load	<i>SF</i> ₁	2020-07-21	23262	404(Yes)	47(Yes)
9	Unload	<i>SF</i> ₁	2020-07-24	34158	366(Yes)	69(Yes)
11	Unknown (Load)	<i>SF</i> ₂	2019-01-24	13173	434(No)	27(Yes)
12 (1)	Load	<i>SF</i> ₁	2019-03-22	51877	358(Yes/P)	95(Yes/P)
12 (2)	Unload	<i>GF</i> ₁	2020-10-06	100223	337(Yes)	56(Yes)
13	Load	<i>SF</i> ₃	2019-12-04	94426	294(Yes)	51(Yes)
14 (1)	Unload	<i>Unknown</i>	2020-01-28	116177	246(Yes)	15(Yes)
14 (2)	Unload	<i>GF</i> ₁	2021-03-11	117146	327(Yes)*	72(Yes/P)
15 (1)	Load	<i>Unknown</i>	2020-03-24	101482	238(Yes)	9(Yes)
15 (2)	Unload	<i>Unknown</i>	2020-11-18	96719	293(Yes)*	39(Yes)
17	Unload	<i>GF</i> ₁	2020-06-03	107414	217(No)	39(Yes)
18	Load	<i>Unknown</i>	2020-06-09	93875	402(Yes/P)	21(Yes/P)
19 (1)	Load	<i>GF</i> ₁	2020-06-14	94125	259(No)	14(Yes)*
19 (2)	Load	<i>GF</i> ₁	2020-10-19	106073	310(No)	19(No)
19 (3)	Load	<i>Unknown</i>	2021-03-29	105027	227(Yes)*	12(Yes)*
20	Unload	<i>GF</i> ₁	2020-06-22	106070	215(Yes)*	21(Yes)*
21	Load	<i>Unknown</i>	2020-08-23	28414	345(Yes/P)	76(Yes/P)
22	Load	<i>GF</i> ₁	2021-03-29	89350	187(No)	15(Yes)*

log anomaly detection, including HDFS [64], BGL [123], Spirit [123] and Thunderbird [123], are mainly produced by complex systems such as cloud service providers or supercomputers. As a result, sensor data for monitoring hardware is often unavailable.

Therefore, we use 25 real-world datasets provided by our industrial partners to test our method. Specifically, the 25 datasets were generated by 20 different test machines, of which 16 machines each generated one dataset and the remaining 4 machines each generated at least 2 datasets. A summary of the datasets is given in Table 3.3. These datasets contain multiple types of faults, including gradual faults, sudden faults, and some unknown types of faults. Moreover, these faults may occur in the loading and/or unloading robot. In particular, for some machines it is unknown whether the fault

3.5. Conclusion and Future Work

occurred in the unloading or loading robot. Therefore, we examined the datasets of both robots to find the robot most likely to have failed.

3.4.2 Baseline, Hyperparameter settings and Performance metrics

Based on domain knowledge and preliminary experiments, in module 2 we set the threshold $\tau_0 = 0.6$ for the similarity coefficient to select log events for fault detection. Furthermore, we set $\tau_1 = 0.6$ in the Kendall rank correlation and $Sig = 0.05$ in the Granger causality test to select features for fault prediction. In module 3, we further remove some highly correlated log features by setting $\tau_2 = 0.95$.

After constructing the event count matrix based on the selected features, we apply a commonly used unsupervised anomaly detector on it: KNN [124]. To demonstrate the necessity of feature selection, we also build an event count matrix from all original features and then apply KNN on it as a baseline. Specifically, a fault is considered detected if at least one of the points with the highest anomaly scores (we consider the top 3 points in our experiments) is on the date of the known fault. Moreover, a fault is considered predicted if one of the points with the highest anomaly scores is located slightly earlier (we consider 1-7 days) than the date when the fault is known to occur. On this basis, as a performance metric, we count the number of datasets in which the faults are accurately detected and/or predicted.

3.4.3 Results and Analysis

As shown in Table 3.3, our proposed feature selection method can help improve log-based anomaly detection performance. Specifically, based on the selected log features, KNN was able to accurately detect or predict faults in 24 out of 25 machines. In contrast, KNN can only accurately detect or predict faults in 17 out of 25 datasets based on all log features. One possible reason is that logs from all subsystems are entangled and the inclusion of many irrelevant log events renders the detection/prediction of gradual faults difficult.

3.5 Conclusion and Future Work

In this work we have proposed a simple yet effective feature selection method for log based anomaly detection. This method has been empirically proven to be effective on 25 real-world datasets. In the future, we plan to include more datasets for testing,

evaluating whether our approach generalizes to other settings, and investigating hyperparameter tuning. More importantly, we will try more anomaly detection methods when defining $\phi(\cdot)$ and $\varphi(\cdot)$. Particularly, the Equation (3.4) used in Granger causality requires several explicit and implicit assumptions to effectively identify Granger causal effects. Some of these assumptions may not be fulfilled by our use-case though. Therefore, we will further explore and improve Granger causality test [125] or other similar techniques to find log events that can be used to predict sensor time series anomalies. Furthermore, we have not fully explored the causal relationships between different log events. In the future, by constructing a causality graph using techniques such as the PC-algorithm [126] on log events, we can investigate the causal relationships between log events. As a result, it might be possible to pinpoint the root causes of anomalies.

3.5. Conclusion and Future Work

Algorithm 2 Feature Selection for Fault Detection and Prediction (FS4FDP)

Input: Event log database $\mathbf{Z} = (\mathbf{Z}_1, \dots, \mathbf{Z}_n, \dots, \mathbf{Z}_N)$; Sensor database $\mathbf{Y} = (\mathbf{Y}_1, \dots, \mathbf{Y}_k, \dots, \mathbf{Y}_K)$; Anomaly detector $\varphi(\cdot)$ for $\mathbf{Z}_n$; Anomaly detector $\phi(\cdot)$ for $\mathbf{Y}_k$; Threshold parameters τ_0, τ_1, τ_2 and Sig with $\tau_0 = \tau_1 = 0.6, \tau_2 = 0.95$ and $Sig = 0.05$ by default.

Output: Subset of log features $\mathbf{F}$

```

1: procedure FS4FDP( $\mathbf{Z}, \mathbf{Y}, \varphi(\cdot), \phi(\cdot), \tau_0, \tau_1, \tau_2, Sig$ )
2:    $\mathbf{F} \leftarrow \{\}$ 
3:    $\mathbf{U} \leftarrow \{\}$ 
4:    $\mathbf{V} \leftarrow \{\}$ 
5:   for  $k \in \{1, \dots, K\}$  do  $\triangleright$  Obtaining anomaly score for each sensor time series
6:      $\mathbf{U}_k = \varphi(\mathbf{Y}_k)$ 
7:      $\mathbf{U}.append(\mathbf{U}_k)$ 
8:   end for
9:   for  $n \in \{1, \dots, N\}$  do  $\triangleright$  Obtaining anomaly score for each log event time series
10:     $\mathbf{V}_n = \varphi(\mathbf{Z}_n)$ 
11:     $\mathbf{V}.append(\mathbf{V}_n)$ 
12:  end for
13:  for  $\mathbf{U}_k \in \mathbf{U}$  and  $\mathbf{V}_n \in \mathbf{V}$  do
14:     $\hat{\mathbf{U}}_k \leftarrow Modify(\mathbf{U}_k)$   $\triangleright$  Modifying the length of anomaly score vector of
    each sensor time series
15:    if  $|\tau(\hat{\mathbf{U}}_k, \mathbf{V}_n)| > \tau_0$  then  $\mathbf{F}.append(\mathbf{Z}_n)$   $\triangleright$  Feature selection for fault
    detection using Kendall rank correlation
16:    end if
17:    for  $l \in \{1, 2, 3, 4, 5\}$  do  $\triangleright$  Feature selection for fault prediction using
    Kendall rank correlation on lagged values
18:      if  $|\tau(Lag(\hat{\mathbf{U}}_k, l), \mathbf{V}_n)| > \tau_1$  then  $\mathbf{F}.append(\mathbf{Z}_n)$ 
19:      end if
20:    end for
21:    if  $Pvalue(GrangerCausalityTest(\mathbf{V}_n, \hat{\mathbf{U}}_k)) < Sig$  then  $\mathbf{F}.append(\mathbf{Z}_n)$   $\triangleright$ 
    Feature selection for fault prediction using Granger causality test
22:    end if
23:  end for
24:  for  $\mathbf{F}_i, \mathbf{F}_j \in \mathbf{F}$  with  $i \neq j$  do  $\triangleright$  Removing redundant features using Kendall
    rank correlation
25:    if  $|\tau(\mathbf{F}_i, \mathbf{F}_j)| > \tau_2$  then  $\mathbf{F}.delete(\mathbf{F}_i)$ 
26:    end if
27:  end for
28:  return  $\mathbf{F}$ 
29: end procedure

```
