

CLIMA 2016 - proceedings of the 12th REHVA World Congress

volume 6

Heiselberg, Per Kvols

Publication date:
2016

Document Version
Publisher's PDF, also known as Version of record

[Link to publication from Aalborg University](#)

Citation for published version (APA):
Heiselberg, P. K. (Ed.) (2016). *CLIMA 2016 - proceedings of the 12th REHVA World Congress: volume 6*. Department of Civil Engineering, Aalborg University.

General rights

Copyright and moral rights for the publications made accessible in the public portal are retained by the authors and/or other copyright owners and it is a condition of accessing publications that users recognise and abide by the legal requirements associated with these rights.

- Users may download and print one copy of any publication from the public portal for the purpose of private study or research.
- You may not further distribute the material or use it for any profit-making activity or commercial gain
- You may freely distribute the URL identifying the publication in the public portal -

Take down policy

If you believe that this document breaches copyright please contact us at vbn@aub.aau.dk providing details, and we will remove access to the work immediately and investigate your claim.

Classifying Office Plug Load Appliance Events in the context of NILM using Time-series Data Mining

Balaji Kalluri^{#1}, Andreas Kamilaris^{^2}, Sekhar Kondepudi^{#3}, Kua Harn Wei^{#4}, Tham Kwok Wai^{#5}

[#]Department of Building, National University of Singapore, Singapore 119077

¹balajikalluri@nus.edu.sg

³sekhar.kondepudi@nus.edu.sg

⁴bdgkuahw@nus.edu.sg

⁵bdgtkw@nus.edu.sg

[^]Department of Computer Science, University of Cyprus, Cyprus 20537

²kami@cs.acy.ac.cy

Abstract

Smart building energy management requires knowledge of individual appliance operation from reduced metering points. The key purpose of this study is to present a classification framework for offices that can help discover individual appliances and its operational modes from single-point aggregate measurements. This approach to non-intrusive load monitoring is supervised through labeled Office Plug Load Dataset. The classification approach is based on short episodes (also called subsequences) from time-series dataset within which appliance events lie hidden. A popular technique for discretizing time-series data known as Symbolic Aggregate approXimation lies at the heart of this framework. Mining large time-series dataset, extracting characteristic appliance features and classifying them appropriately based on individual appliance events is facilitated through “Bag of Patterns” based Vector Space Model. This study focuses on classifying multiple events from three common aggregate appliance use-case scenarios in an office environment. The approach is promising at analyzing subsequence patterns from more than 1700 time-series episodes in the dataset. The results from classifying multi-functional device operations from aggregate signature show errors less than 22% in scenario where three appliances are in operation, whereas error is less than 37% when two appliances are in operation. The results also indicate that the approach is likely to work better as the dataset grows as in the case of big data. Additionally, the proposed approach enables visualizing subsequences of a time-series using color-coding scheme. Such visualization helps in understanding the relative specificity of an event to others in the time series.

Keywords - office appliances; time-series mining; classification; NILM; big data.

1. Introduction

In modern office buildings, Information and Communication Technology (ICT) based office plug loads account for a significant share (about 15%) comparable to heating and cooling loads and they are predicted to grow to 36% by 2030 [1]. The problem aims at improving the energy efficiency of office plug load appliances not at the design stage but when buildings are in operation. For instance, operating ICT appliances rationally in offices within university campus has shown potential energy savings of more than 130,000 kWh annually [2]. The knowledge of appropriate appliance operation together with actionable feedback has shown a potential of 311kWh of desktop PC energy savings in a small office [3]. However, it is not only challenging but also impractical to meter and measure literally every office appliance in a fully operational building. Such an approach is expensive, time consuming and obstructive to occupants. To alleviate the situation several research efforts towards measuring a group of devices from a single point with multiple appliances operating in a single electrical circuit in a building has been evolving since 90s [4, 5]. This has been identified as a classical approach in building energy measurement and verification termed as *Non-Intrusive Load Monitoring (NILM)* or simply *load disaggregation* [6].

1.1. Non-intrusive Load Monitoring in Offices

NILM is a generic term to address methods and techniques to decipher individual appliance energy consumption and their states from a single, aggregate measurement point in a circuit. NILM in offices poses surprisingly different set of challenges unlike homes such as physically large spaces, diverse appliance types, several identical instances and overlapping operational patterns [6]. Therefore, non-intrusive plug load audits in offices need to be treated differently using a different kind of dataset and approach. The challenge for energy audit of ICT appliances in an office is the presence of several identical appliances on each of several workstations connected to a circuit and operating concurrently. A heuristic approach to the dataset is previously proposed [7]. This approximates every office workstation into a (dis) aggregating point associated with three or four common office appliances. Thus, the problem of NILM in offices reduces the metering points (for example only 25 instead of 75 appliances). This results in a considerable saving in both metering (and in turn auditing) cost and time.

1.2. Office Plug Load Datasets

The plug load dataset is a repository of electricity consumption data of several individual appliances (e.g. microwave oven, refrigerator, and computers) and their combination measured across multiple levels (e.g. whole-building or appliance-level) using energy meters. Generally such

datasets vary based on building type (e.g. residences or offices) as well. A quick list of references to all public plug load datasets for NILM is available [8].

The OPLD consists of measurements of both aggregate and individual appliance data of four common office workstation appliances. The appliances are: desktop PC, laptop PC, monitor and multi-function printing device (MFD). There are 5 possible use case scenarios possibly with these 4 appliances combination. Typically every single appliance state combination is measured approximately for duration of 8 to 10 minutes at a rate of a sample per second. For more details on actual instrumentation scheme, experimental data collection plan and other data sanctity measures with respect to OPLD refer to [7].

1.3. Applicability of time series data mining to NILM

Typically the appliance load signatures collected by metering devices in buildings are temporal in nature. A time series dataset is simply a collection of several measurements made chronologically [9]. However, the application of time series data mining techniques to NILM, energy disaggregation and building energy analytics is limited. The methodology and results presented in this study ascertain the applicability of using time series subsequence data mining to study office appliances' transient operations and classify them in aggregate data. The idea of discovering individual appliance events from aggregate signature is introduced as a classification problem based on supervised bag-of-rules approach. Visualizing the hidden subsequence patterns from several aggregate time series measurements in OPLD present the potential for individual load identification.

1.4. Problem Statement

The problem of NILM in the above context of offices therefore scales down to disambiguating single time series measurement obtained from disaggregating points into individual appliance states. The measurements are both time-stamped and labeled for individual appliance events. Therefore, it suits supervised approaches to time series classification.

2. Methodology

In this paper, NILM is treated as a classification problem. The goal is to classify episodes within aggregate time series data appropriately into individual hidden appliance events. The current scope of this paper is limited to disambiguate several transient events of multi-functional devices (MFD) that lie hidden in the aggregate measurements from several workstations. The reason for choice of MFD is that showed repetitive subsequence patterns also called *characteristic motifs* within several episodes of time series measurements. The relative class specificity of subsequent patterns (motifs)

within time series episodes to be discussed later in section 3 is also a source of motivation.

2.1. Dataset for classification

This study employs measurements of three use case scenarios: (a) Desktop PC + Monitor + MFD (b) Laptop PC + Monitor + MFD and (c) Laptop PC + MFD from OPLD. For the classification presented in this paper, episodes from aggregate data involving MFD transient events such as copy, scan and print are retrieved. The data is not time stamped but temporally ordered. More than 1700 time series episodes from among all three scenarios are extracted. The length of each time series episodes is 480 samples. This forms the classification dataset for analysis in this study. Fig. 2 presents a simplistic view of such dataset for one such scenario.

		Length of Time series									
Length of Classification dataset ↓		1	2	3	4	5	6	7	-----	479	480
	1	79.993	81.02	80.36	79.92	86.488	96.248	83.039	-----	84.396	90.047
	2	85.167	80.653	86.598	82.158	82.011	79.809	80.617	-----	82.378	81.901
	3	81.497	83.736	83.772	81.644	80.433	80.763	80.837	-----	86.268	85.754
	1	81.571	82.525	81.938	80.727	84.947	80.433	83.369	-----	87.038	397.653
	2	77.718	82.305	85.864	85.314	83.919	83.075	82.782	-----	88.726	94.817
	3	81.681	81.644	84.249	81.791	81.571	82.855	81.681	-----	81.681	82.562
		↓		↓		↓			-----	↓	
	1	87.699	87.735	81.938	81.351	83.626	81.057	81.02	-----	1105.40	1033.376

Fig. 2 A simplistic representation of an example dataset used in classification

The rows in the dataset represent the individual episodes and columns represent every data point in the time series. The first column in the dataset represents the class corresponding to appliance event whereas the rest of the columns represent the time series data points. The following are some labels used in this study: (a) 1: MFD-COPY, (b) 2: MFD-SCAN and (c) 3: MFD-PRINT. Each row in the dataset has multiple characteristic subsequence patterns (motifs) corresponding to the labeled classes. The role of proposed classification framework is to identify and exploit such characteristic subsequence pattern from several rows within the large dataset in classifying them appropriately.

2.2. Classification Framework

The proposed approach to classifying episodes of unknown time series data into appropriate hidden appliance events is supervised. It is based on two techniques namely SAX [10] and VSM [11] as introduced in [12]. This approach is chosen for analysis for the few reasons. Firstly, it has shown promising results in classifying several UCR time series datasets across domains [12]. Secondly, it exploits popular time series subsequence mining

approach called SAX. Thirdly, entire classification framework is open source¹. The outline flow of the sequences of steps involved in the proposed classification framework is presented in Fig. 3.

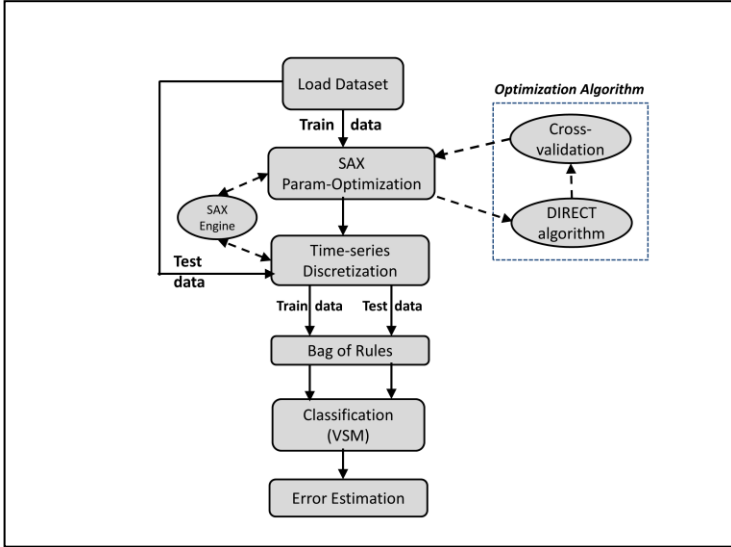


Fig. 3 Conceptual clasification framework based on time-series subsequence mining

The dataset for classification typically is composed of few hundreds of time series measurements. The first step is to prepare the dataset and load them into the framework. The plug load energy measurements obtained from OPLD is in CSV (Comma Separated Values) format as previously presented Fig. 2. This step also involves separating the entire dataset into train and test data. This is followed by discretizing and symbolizing each data set separately. A popular technique called Symbolic Aggregate approXimation (SAX) [10] is implemented for this process of time series transformation. This treatment is carried out to transform the data patterns within short subsequences of time series into a collection of words. Such discretization combined with symbolizing time series data aid in representing appliance events using specific Context Free Grammar (CFG) rules. This transformation enables application of text data mining approaches in classifying subsequence patterns [11]. This technique requires three user defined parameters namely *sliding window size* (w), *PAA size* (p) and *alphabet size* (a). Several characteristic motifs specific to various appliance events are obtained for different combinations of w , p and a . Since the

¹ <https://github.com/jMotif/jmotif-R>

classification dataset is large manual selection and tuning of SAX parameters for each and every time series data is impossible. Therefore at the next step, optimum selection of these parameters is performed with the help of Dividing RECTangles (DIRECT) algorithm and common cross-validation scheme. Such a SAX parameter optimization approach has been previously found applicable for wide range of time series datasets [12].

The next step is to transform the time series subsequences into large collection of SAX rules and apply Vector Space Model for classifying them into appropriate appliance events. This involves discretizing the entire dataset and creating *Bag of Rules (BoR)* repository very similar to the one described in [13]. However, other approaches to extract subsequent features from time series data for classification are also available [14-16]. Both parameter optimization and time series discretization steps in the framework employs SAX mechanism. Individual *bags* of SAX words are created for every time series in both train and test data as shown in Fig. 4. They are combined on the basis of *classes (labels)* to form *corpus* of SAX words. Representing CFG rules in vector space and techniques for classification them are derived from Information theory [11]. The vector space model builds a *term frequency (tf)* and *inverse document frequency (idf)* matrices for the entire dataset. The *tf* for every SAX word captures the number of times a particular subsequence pattern appears in a time series. On the other hand, the *idf* for every SAX word indicates a measure of its relative presence across entire corpus. The *tf*idf* based *weighting scheme* for every subsequence (word) is determined. Finally a cosine similarity measure between the *tf*idf* matrix for the corpus and the bag of rules for test dataset is computed. The classification of test data to an appliance event (i.e. *class*) is based on the highest cosine similarity score. By combining SAX based Bag of Rules and Vector Space Model (VSM), the proposed approach transforms time series dataset into vectors in space to help classify based on class-specific subsequence. Finally, a simple scheme for estimating the error in classification is implemented to determine the classification performance. This is done by comparing the percentage of false prediction of the number of labels of time series in test dataset.

The proposed SAX based VSM approach based on open source implementation [12] to supervised appliance event classification from aggregate signature has the following characteristics. Firstly, it considers every possible subsequence pattern (an appliance event) in the aggregate time series. Secondly, in classifying an unknown aggregate time series the cosine similarity measures considers relative weights of every SAX word (representative of appliance event) in the labeled (known) time series. The complete implementation of this classification framework is done in CRAN-R.

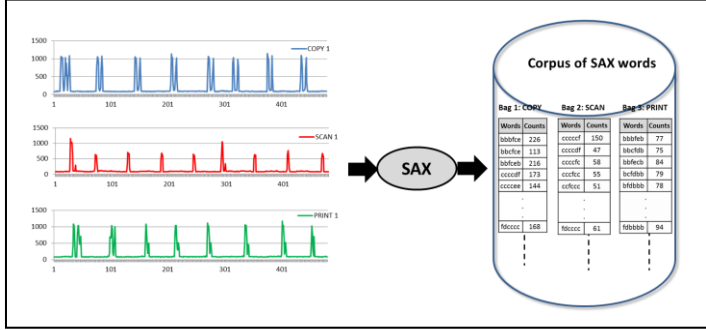


Fig. 4 Illustration of SAX discretization step to create corpus of SAX words

3. Results and Discussion

The application of time series classification framework presented in section 2.2 in the context of NILM is the main purpose of this study. The analysis of appropriately identifying the individual hidden appliance and its events from approximately 1700 disjoint aggregate time series measurements across multiple office appliance use case scenario is discussed. A brief discussion on the possible visualization of subsequence pattern specificity with the help of VSM is also presented.

A. Classification performance of three use case scenarios in offices

The three distinct office appliance use case scenarios from OPLD as discussed in section 2.1 are analyzed. The result of classification of aggregate time series measurements for each scenario is presented in Table 1. Note that for each scenario the target appliance to be classified from the aggregate time series measurement in this analysis is MFDs. Three different instances of MFD (i.e. P1, P2 and P3) from OPLD are analyzed. Some key measures of classification are summarized in Table 1.

They are the following: SAX parameters, count of total number of training and testing time series measurements, and %error in classification. The classification results obtained are with optimum selection of SAX parameters (w, p & a) as listed in the Table 1. The split strategy similar to UCR dataset [17] is considered for breaking the entire dataset into training and testing sets. The following are some observation drawn from each scenario in Table 1.

- **Desktop PC + Monitor + MFD:** For the scenario 1, a total of about 705 time series measurements are analyzed. The overall error in classifying MFDs across all events from aggregate time series measurements ranges from 0% to 22%.

Table 1. Summary of classification results of Vector Space Model applied to OPLD

S. No.	Appliance Use-case Scenarios	Classification Measures	Classification Output		
			P1 in aggregate	P2 in aggregate	P3 in aggregate
1.	D + M + P	SAX parameters [w, p, a]	[25, 6, 6]	[23, 7, 6]	[34, 6, 7]
		# Training Set	79	75	92
		# Testing Set	154	155	150
		% Error	18.18%	0%	22%
2.	L + M + P	SAX parameters [w, p, a]	[25, 6, 6]	[18, 7, 4]	[25, 6, 7]
		# Training Set	100	100	100
		# Testing Set	190	184	175
		% Error	11.05%	17.39%	13.14%
3.	L + P	SAX parameters [w, p, a]	[25, 6, 4]	[25, 6, 6]	[32, 6, 7]
		# Training Set	30	45	30
		# Testing Set	42	25	41
		% Error	35.71%	36%	36.59%

- Laptop PC + Monitor + MFD: In the scenario 2, a total of 849 time series measurements are analyzed and classification error ranges between 11% and 18%. The improved classification is possibly a result of increased number of time series instances in scenario 2 against scenario 1.
- Laptop PC + MFD: The scenario 3 is quite different from other two scenarios. A total of 213 time series measurements are analyzed. The overall misclassification of appliances ranges between 35% and 37%. The results produced show a slightly poor classification performance over others possibly due to the size of its input dataset. The size of input dataset ranges around 70 time series for each MFD instance unlike 250 plus in other scenarios.

However, there are other observations that can be drawn from across all three scenarios. They are summarized as follows.

- From among all three instances of MFD only appliance instance P1 strongly exhibits its characteristic events consistently across all scenarios. This is reflected in the optimized SAX parameters selected for classification. Additionally it also reflects a hidden appliance property that 3 distinct events such as *copy*, *scan* and *print* for P1 is likely to be similar due to identical optimal sliding window (w) and PAA size (p).
- It can be seen that although the MFD appliance instances (i.e. P1, P2 and P3) are same across all three scenarios, the SAX

parameters that yield optimal classification results are indeed different. This strongly suggests that appropriate parameter optimization step is necessary to make right choice of w , p and a .

In summary, the proposed classification framework helps to discover multiple hidden appliance operation from single-point, aggregate energy data. The analysis presented in this study is focused on discovering several multi-functional imaging devices (MFD) and their operation modes from aggregate appliance energy data. The overall result of classification is reasonably good (with error <22%) at disambiguating three aggregate office appliance scenario in contrast to two aggregate office appliance scenario (with error ~37%). The latter is possibly due to the limited dataset size. This indicates that larger the size of the dataset, accurate classification is possible with such a framework. Additionally the approach can be promising for big data NILM studies because SAX can help reduce the dimensionality of the data.

B. Visualizing hidden appliance events in aggregate data

Visualizing subsequence patterns (also called *motifs*) corresponding to individual hidden appliance events can be more powerful. A simple illustration of the proposed vector space model for visualizing appliance events such as MFD-COPY, SCAN and PRINT from one such aggregate time series data is presented in Fig. 5. A sample aggregate measurement from OPLD when a desktop PC, monitor and MFD is working together in an office circuit is presented in Fig. 5(a). This is described as a time series data with several labeled episodes corresponding to distinct MFD operation.

As discussed in section 2.2, the vector space model based on bag of rules help to build the $tf*idf$ matrix. The columns in the matrix indicate the measure of relative specificity of every subsequence pattern to the appliance class. This measure is color coded to present effective insights into the data patterns through visualization. An example episode for each class i.e. *copy*, *scan* and *print* from the aggregate time series is presented in Fig. 5(b), (c) and (d) respectively. From the color coded time series episodes, it can be seen that most of the characteristic appliance subsequent patterns are specific to its class represented in either blue or green colors. This serves an indication that $tf*idf$ matrix based classification of aggregate time series is appropriate in using hidden appliance subsequence patterns as features.

4. Conclusion

Classification as an approach to NILM in offices is demonstrated in this study with the help of a time series subsequence data mining framework. It is one of the first studies to transform office appliance plug load dataset into an approximate symbolic rules using SAX. An existing approach using vector space model as an extension to previously introduced Naïve classification of

OPLD is considered. The proposed framework can help discover multiple hidden appliance operation from single-point, aggregate energy data. The analysis presented in this study so far is focused on discovering several multi-functional imaging devices (MFD) and their operation modes from aggregate appliance energy data. The result of classification is promising and shows potential for application in big data NILM scenarios.

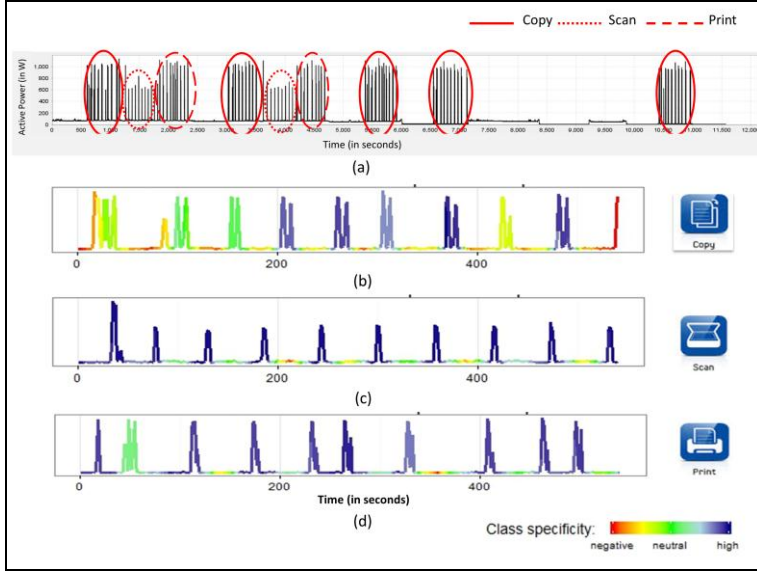


Fig. 5 (a) An example aggregate measurement of scenario-1: Desktop PC+Monitor+MFD with labeled episodes of MFD events. One sample episode for (b) MFD-COPY (c) MFD-SCAN and (d) MFD-PRINT using colour coding to capture relative specificity of time series subsequence obtained using SAX-VSM framework is presented.

5. Future Work

The analysis presented in this study is so far limited to only identifying one appliance and its operating state from a combination of two or three appliances operating together in an office sub-circuit. The next step would be to extend the implementation to identifying every individual appliance and its multiple operational states. The future work will also address slightly large errors resulting due to limited dataset size using alternate similarity metrics for classifying CFG rules.

Acknowledgment

The authors would like to thank the funding agency Singapore's Ministry of Education under AcRF grant R-296-000-145-112. A special thanks to

Jyothsna Devi in helping collect, clean and prepare the dataset used for classification in this paper.

References

- [1] U. DoE, "Buildings energy databook," *Energy Efficiency & Renewable Energy Department*, 2011.
- [2] A. Kamilaris, D. T. H. Ngan, A. Pantazaras, B. Kalluri, S. Kondepudi, and T. Kwok Wai, "Good practices in the use of ICT equipment for electricity savings at a university campus," in *Proceedings of the 5th International Green Computing Conference (IGCC)*, Dallas, 2014, pp. 1-11.
- [3] A. Kamilaris, J. Neovino, S. Kondepudi, and B. Kalluri, "A case study on the individual energy use of personal computers in an office setting and assessment of various feedback types towards energy savings," *Energy and Buildings*, vol. 104, 2015, pp. 73-86.
- [4] G. W. Hart, "Nonintrusive appliance load monitoring," *Proceedings of the IEEE*, vol. 80, 1992, pp. 1870-1891.
- [5] L. K. Norford and S. B. Leeb, "Non-intrusive electrical load monitoring in commercial buildings based on steady-state and transient load-detection algorithms," *Energy and Buildings*, vol. 24, 1996, pp. 51-64.
- [6] A. Kamilaris, B. Kalluri, S. Kondepudi, and T. Kwok Wai, "A literature survey on measuring energy usage for miscellaneous electric loads in offices and commercial buildings," *Renewable and Sustainable Energy Reviews*, vol. 34, 2014, pp. 536-550.
- [7] B. Kalluri, S. Kondepudi, K. Ham Wei, T. Kwok Wai, and A. Kamilaris, "OPLD: Towards improved non-intrusive office plug load disaggregation," in *Proceedings of the 1st IEEE International Conference on Building Energy Efficiency and Sustainable Technologies (ICBEST)*, Singapore, 2015, pp. (in-press).
- [8] O. Parson, "Public Datasets for NIALM," Available: <http://blog.oliverparson.co.uk/2012/06/public-data-sets-for-nialm.html>
- [9] T.-c. Fu, "A review on time series data mining," *Engineering Applications of Artificial Intelligence*, vol. 24, 2011, pp. 164-181.
- [10] J. Lin, E. Keogh, L. Wei, and S. Lonardi, "Experiencing SAX: a novel symbolic representation of time series," *Data Mining and Knowledge Discovery*, vol. 15, 2007, pp. 107-144.
- [11] C. D. Manning, P. Raghavan, and H. Schutze, "Introduction to Information Retrieval," *Cambridge University Press*, vol. 1, 2008.
- [12] P. Senin and S. Malinchik, "SAX-VSM: Interpretable Time Series Classification Using SAX and Vector Space Model," in *IEEE International Conference on Data Mining*, Dallas, Texas, 2013, pp. 1175-1180.
- [13] J. Lin and Y. Li, "Finding Structural Similarity in Time Series Data Using Bag-of-Patterns Representation," in *Proceedings of the 21st International Conference on Scientific and Statistical Database Management*, New Orleans, 2009, pp. 461-477.
- [14] U. Kamath, J. Lin, and K. De Jong, "SAX-EFG: An evolutionary feature generation framework for time series classification," in *Proceedings of the 2014 Genetic and Evolutionary Computation Conference (GECCO 2014)*, Vancouver, 2014, pp. 533-540.
- [15] J. Wang, P. Liu, M. F. She, S. Nahavandi, and A. Kouzani, "Bag-of-words representation for biomedical time series classification," *Biomedical Signal Processing and Control*, vol. 8, 2013, pp. 634-644.
- [16] K. Basu, V. Debusschere, S. Bacha, U. Maulik, and S. Bondyopadhyay, "Nonintrusive Load Monitoring: A Temporal Multilabel Classification Approach," *IEEE Transactions on Industrial Informatics*, vol. 11, 2015, pp. 262-270.
- [17] C. Yanping, E. Keogh, H. Bing, B. Nurjahan, B. Anthony, M. Abdullah, and B. Gustavo, "The UCR Time Series Classification Archive," Available: www.cs.ucr.edu/~eamonn/time_series_data/