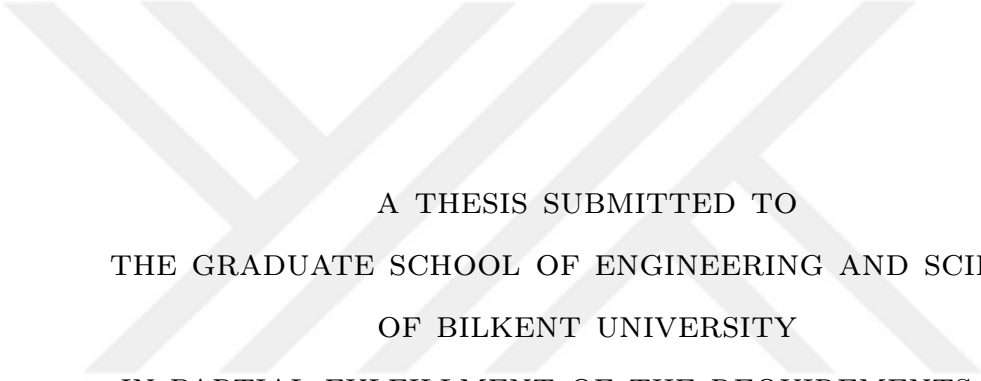


IMPLICIT CONCEPT DRIFT DETECTION FOR MULTI-LABEL DATA STREAMS



A THESIS SUBMITTED TO
THE GRADUATE SCHOOL OF ENGINEERING AND SCIENCE
OF BILKENT UNIVERSITY
IN PARTIAL FULFILLMENT OF THE REQUIREMENTS FOR
THE DEGREE OF
MASTER OF SCIENCE
IN
COMPUTER ENGINEERING

By
Ege Berkay Gülcan
January 2022

Implicit Concept Drift Detection for Multi-label Data Streams

By Ege Berkay Gülcan

January 2022

We certify that we have read this thesis and that in our opinion it is fully adequate, in scope and in quality, as a thesis for the degree of Master of Science.



Fazlı Can(Advisor)

Özgür Ulusoy

İsmail Sengör Altingövde

Approved for the Graduate School of Engineering and Science:

Ezhan Karaşan
Director of the Graduate School

ABSTRACT

IMPLICIT CONCEPT DRIFT DETECTION FOR MULTI-LABEL DATA STREAMS

Ege Berkay Gülcan

M.S. in Computer Engineering

Advisor: Fazlı Can

January 2022

Many real-world applications adopt multi-label data streams as the need for algorithms to deal with rapidly generated data increases. For such streams, changes in data distribution, also known as concept drift, cause the existing classification models to rapidly lose their effectiveness. To assist the classifiers, we propose a novel algorithm called Label Dependency Drift Detector (LD3), an implicit (unsupervised) concept drift detector using label dependencies within the data for multi-label data streams.

Our study exploits the dynamic temporal dependencies between labels using a label influence ranking method, which leverages a data fusion algorithm and uses the produced ranking to detect concept drift. LD3 is the first unsupervised concept drift detection algorithm in the multi-label classification problem area.

In this study, we perform an extensive evaluation of LD3 by comparing it with 14 prevalent supervised concept drift detection algorithms that we adapt to the problem area using 12 datasets and a baseline classifier. The results show that LD3 provides between 19.8% and 68.6% better predictive performance than comparable detectors on both real-world and synthetic data streams.

Keywords: Big data, multi-label data stream, multi-label classification, concept drift, drift detection, data fusion.

ÖZET

ÇOK ETİKETLİ VERİ AKIŞLARI İÇİN DENETİMSİZ KAVRAM KAYMA TESPİTİ

Ege Berkay Gülcan
Bilgisayar Mühendisliği, Yüksek Lisans
Tez Danışmanı: Fazlı Can
Ocak 2022

Bir çok gerçek dünya uygulaması, hızla oluşturulan verilerle başa çıkabilmek için çok etiketli veri akışlarını benimsemektedir. Bu tür akışlar için, kavram kayması olarak da bilinen veri dağılımındaki değişiklikler, mevcut sınıflandırma modellerinin etkinliğini hızla kaybetmesine neden olur. Bu tezde sınıflandırıcılara yardımcı olmak için, çok etiketli veri akışları için verilerdeki etiket bağımlılıklarını kullanan denetimsiz bir kavram sapma detektörü olan “Label Dependency Drift Detector” (LD3) adlı yeni bir algoritma öneriyoruz.

Çalışmamız, bir veri birleştirme algoritmasından faydalanır ve üretilen sıralamayı kavram kaymasını tespit etmek için kullanan bir etiket etkisi sıralama yöntemini kullanarak etiketler arasındaki dinamik zamansal bağımlılıklardan yararlanır. LD3, çok etiketli sınıflandırma problem alanındaki ilk denetimsiz kavram kayması algılama algoritmasıdır.

Bu çalışmada, LD3’ü, 12 veri seti ve bir temel sınıflandırıcı kullanarak problem alanına uyarladığımız 14 popüler denetimli kavram kayması algılama algoritması ile karşılaştırarak kapsamlı bir değerlendirme yapıyoruz. Sonuçlar, LD3’ün hem gerçek dünya hem de sentetik veri akışlarında karşılaştırılabilir dedektörlerden %19.8 ile %68.6 arasında daha iyi tahmin performansı sağladığını göstermektedir.

Anahtar sözcükler: Büyük veri, çok etiketli veri akışı, çok etiketli sınıflandırma, kavram kayması, sapma algılama, veri birleştirme.

Acknowledgement

First and foremost, I would like to thank my advisor Fazlı Can for his constant contributions and support regarding during this thesis and my education. Moreover, I would like to thank my dear family members and friends for their endless support, insights and inspirations.

January 2022, Ankara

Contents

1	Introduction	1
2	Problem Domain and Aim of the Study	4
3	Related Works	8
3.1	Supervised Concept Drift Detection Algorithms	8
3.1.1	Error Rate-based Algorithms	9
3.1.2	Algorithms for Multi-label Concept Drift Detection	11
3.2	Unsupervised Concept Drift Detection	12
4	Label Dependency Drift Detector	14
4.1	Evaluating the Changes in the Label Dependencies	16
4.2	Measuring Similarity Between Two Rankings	17
4.3	Detailed Calculation Example	19
5	Evaluation and Experimental Setup	22

5.1	Datasets	22
5.2	Experimental Setup and Evaluation Metrics	23
6	Experimental Results and Discussion	26
6.1	Effectiveness Analysis	26
6.2	Analysis with Different Effectiveness Measures	28
6.3	Statistical Significance Evaluation of Effectiveness Results	29
6.4	Time Change Analysis of Effectiveness	31
6.5	Efficiency Concerns and Advantages of Unsupervised Detection	33
6.6	Hyperparameter Analysis	34
6.7	Setting the Standard Deviation Multiplier t	34
6.8	Setting the Parameters Window Size w and Threshold for Number of Anomalies L	35
6.9	Analysis of Different Fusion Methods	36
6.10	Influence of Drift Types and Speed	37
7	Conclusion	39

List of Figures

2.1	Sources of concept drift based on how they occur probabilistically. Different colors represent different classes and the dotted line is the decision boundary.	5
2.2	Types of concept drift based on what happens to the concept over time (outlier is not a drift). Different colors represent different concepts.	7
4.1	Workflow of LD3. Red boxes represent false labels (0) and green boxes represent true labels (1).	18
6.1	Nemenyi critical distance diagrams for all metrics. The number given within parentheses after the method name indicates rank position of the method.	30
6.2	Example-based accuracy results of the top detectors and ND on four datasets. Y-axis of the plots is given with different scales. . .	32
6.3	Distribution of the rank correlations within W_{corr} of four datasets after a drift is detected with different t values (σ multipliers). X-axis of the plots is given with different scales.	34

List of Tables

3.1	Error rate-based drift detection algorithm characteristics (They are all our baselines; for those with more than one version, the number of variants is given in parentheses after the detector name). The symbols w_{hist} and w_{new} indicate the windows for old and new data in respective order. Landmark uses the entire set of data before drift and flushes this data when drift is detected to collect the new distribution.	9
4.1	Symbols table for Algorithm 1.	14
5.1	Table of multi-label datasets used in the experiments. The upper group contains the real-world datasets. N represents the number of samples, D is the number of features, n is the number of labels and $LC(\mathcal{D})$ and $LD(\mathcal{D})$ represents label cardinality and label density which are the average number of true labels of the samples in dataset $\mathcal{D}$ and $LC(\mathcal{D})$ divided by the number of labels, indicating the average label cardinality per label [1].	23

6.1	Experiment results for the detectors. ND is a classifier without any detection methods for the baseline. Average result is calculated from the average result of each algorithm, except LD3, for a dataset. <i>LD3 Imp.</i> represents LD3's improvement over the average. <i>Avg. Imp.</i> illustrates the average improvement for each metric. The best results are highlighted in bold.	27
6.1	Continued.	28
6.2	Example-based accuracy of LD3 with varying parameter settings for each real-world dataset The best scores are highlighted in bold. Datasets are listed in ascending order of the number of data items they contain.	36
6.3	Example-based accuracy results of LD3 using four different data fusion algorithms. Best results are highlighted in bold.	37
6.4	Example-based accuracy results of the top detectors with different drift types. Best results are highlighted in bold.	38

Chapter 1

Introduction

Many organizations generate temporal data in the form of data streams with high variety, volume, and velocity [2]. Data is created continuously on a massive scale and can be assigned multiple labels, which is termed multi-labeled. However, because of the scale of this data, they need to be processed immediately since the cost associated with storage and retrieval is high, which explains the current popularity of data stream mining [3].

Real-world applications are constantly evolving and over time, a change in data distribution may occur. This change in data is called concept drift [4] which is one of the most prevalent problems in data stream mining. Many of the traditional classification algorithms assume that the data is static, which causes them to lose effectiveness when faced with concept drift. An example area for concept drift is the energy sector where the drifted streams may cause instabilities for the learning models designed to predict energy consumption, production and distribution [5].

In this thesis, we propose a novel unsupervised (implicit) concept drift detection algorithm that exploits label dependencies between class labels for multi-label data streams through data fusion methods. In multi-label data stream mining, it is common for labels to have correlations and dependencies [6]. Several studies [7, 8, 9] show that incorporating label dependencies into multi-label classifiers

boosts their effectiveness. We aim to utilize these correlations among labels to demonstrate that label dependencies can be used for detecting concept drift. For this purpose, we model the label dependencies by creating a co-occurrence matrix of labels [10], and we use the generated matrix to represent the current and past stream representation through label ranking and then use data fusion to detect concept drift.

In this study our main contributions are the following.

1. To the best of our knowledge, for the first time in literature, we introduce the concept of label dependency ranking and use it for concept drift detection in multi-label classification.
2. We perform an extensive evaluation of our method by comparing it with 14 prevalent concept drift detection algorithms using 12 data streams and a base classifier. We compare drift detection algorithms with their influence on the predictive performance of the baseline classifier. In all cases, LD3 is the number one algorithm and in most cases, it enables performance that is statistically significantly higher than most of the baselines in terms of almost all effectiveness measures utilized in the experiments. Furthermore, our work provides the first study on the use of several different concept drift detection algorithms which are developed for multi-class environments, for concept drift detection in the multi-label classification problem domain.
3. Our method LD3 is the first unsupervised concept drift detection algorithm in the problem domain we focus on. The baseline drift detection algorithms used in the experiments are supervised and require true labels; on the other hand, our method uses only the predicted labels: In many streaming environments, true class labels needed by supervised methods are not always available; in some cases, only a percentage of them are available or arrive late or are potentially unavailable [11, 12]. These factors show the practical importance of our approach.

In this thesis, all algorithms are used with the default parameters provided by their respective publications. In the experiments, the baseline classifier starts

a new prediction model when a concept drift is detected. We should also note that, since we use a supervised classifier in our experiments, LD3 indirectly uses the ground truth labels of the data stream as it processes the predicted labels of the classifier. The reason behind this is the lack of online unsupervised multi-label classifiers. Although there are previously developed unsupervised multi-label classifiers, they are not designed for online learning, such as Wang and Zhang’s work [13]. Therefore, we chose to use a supervised classifier.

In the following chapters, we first describe the problem domain, the aim of the study, and introduce the previous work in Chapters 2 and 3. Then, we propose our solution in Chapter 4. Lastly, in Chapters 5 and 6, we discuss the evaluation methodology and present our results and discussions. Chapter 7 concludes the thesis.

Chapter 2

Problem Domain and Aim of the Study

Traditional multi-label data stream classifiers learn by performing the interleaved test-then-train method [14], i.e, prequential training, on a stream of incoming data. However, if there is a change in data distribution, a significant decrease in predictive performance is experienced since the classifier still uses the previously learned distribution in its predictions. A concept drift detector detects the change in data distribution and alerts the classifier so that it can start working on adaptation strategies, which increases the robustness of the classifier.

Concept drift is the change in the data distribution that occurs over time. Given a data stream $S_{0,t} = d_0, \dots, d_t$, in a time window $[0, t]$, where $d_i = (X_i, y_i)$ with X_i being the features and y_i being the labels of the i -th data instance, concept drift is [15]:

$$\exists t : P_t(X, y) \neq P_{t+1}(X, y) \quad (2.1)$$

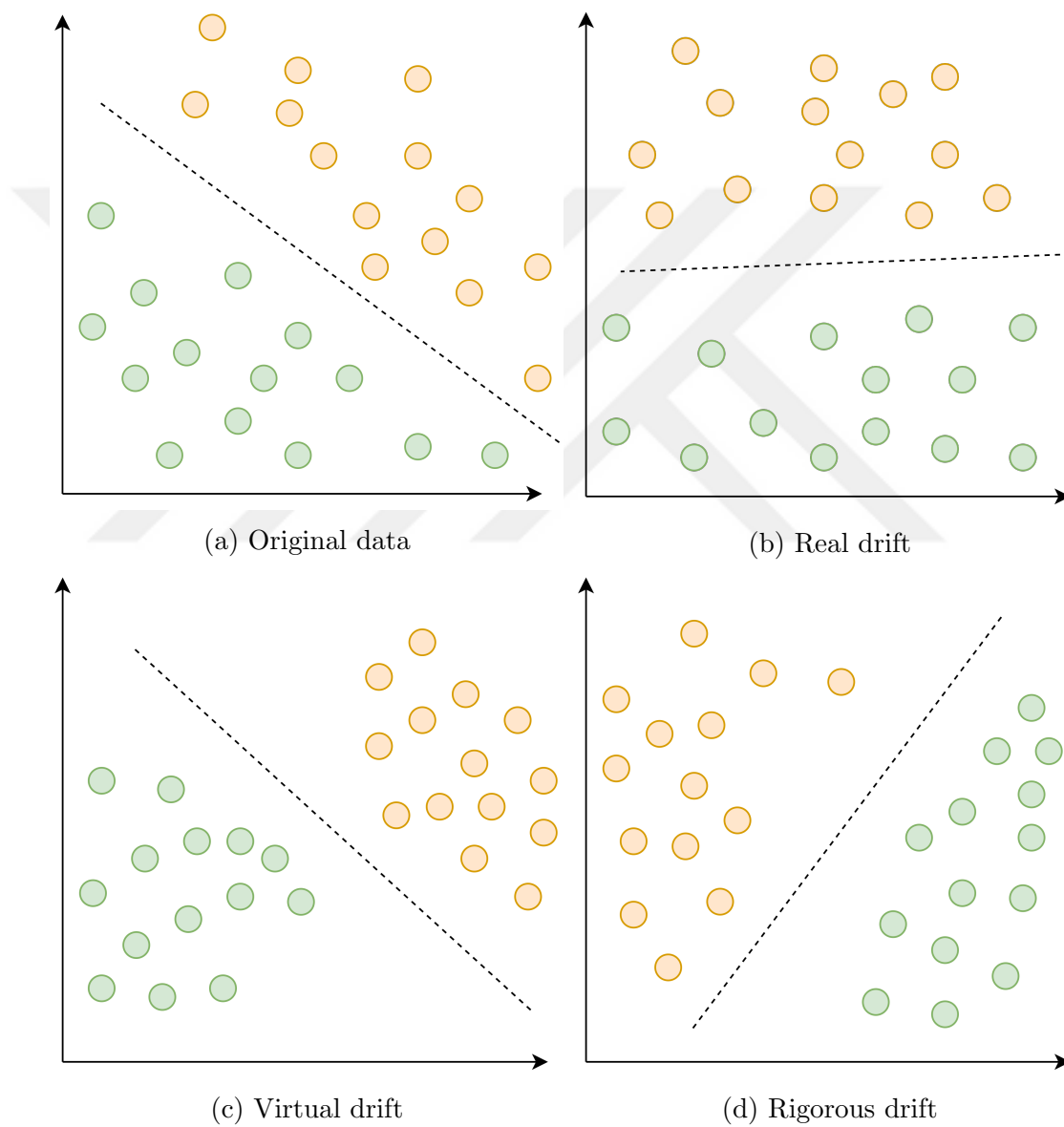


Figure 2.1: Sources of concept drift based on how they occur probabilistically. Different colors represent different classes and the dotted line is the decision boundary.

According to this definition (Eq. 2.1), Lu et al. [16] describe three potential sources of concept drift:

- The change may happen due to a change in posterior probabilities $P_t(y|X)$, meaning, $P_t(y|X) \neq P_{t+1}(y|X)$. This is called real or actual drift and it usually results in a shift in the decision boundary, causing a significant decrease in effectiveness.
- If the cause of the change is $P_t(X) \neq P_{t+1}(X)$, while $P_t(y|X) = P_{t+1}(y|X)$, it is called a virtual drift because it does not cause a shift in the decision boundary.
- The final source is when the change occurs as both $P_t(y|X)$ and $P_t(X)$ change over time which is called rigorous drift [17].

Figure 2.1 illustrates the three sources of concept drift along with the original distribution.

Lu et al. [16] further define four types of concept drift in terms of how they happen over time, which are illustrated in Figure 2.2. In the figure, the vertical axis represents data distribution and the horizontal axis represents time. Depending on the nature of data, this change may be sudden, incremental, gradual, or reoccurring. For example, with the changing of seasons, climate data may show reoccurring concepts or big events may cause sudden drift in news data.

Among the types presented in Figure 2.2, outlier is not actually a drift type. Outliers usually happen as a result of noise in the data. As they are false positives, a drift detector should also account for their presence to maintain accurate detection.

In this study, we aim to design an unsupervised concept drift detector and measure its boosting impact in classifier prediction accuracy. Our method LD3 exploits dependencies among predicted labels and grants drift detection capabilities to multi-label stream classifiers with no drift detection facility. In this work, we consider a supervised classification environment. The use of LD3 with an

unsupervised multi-label [13] or a self-tuning classifier [18] is beyond the scope of this thesis.

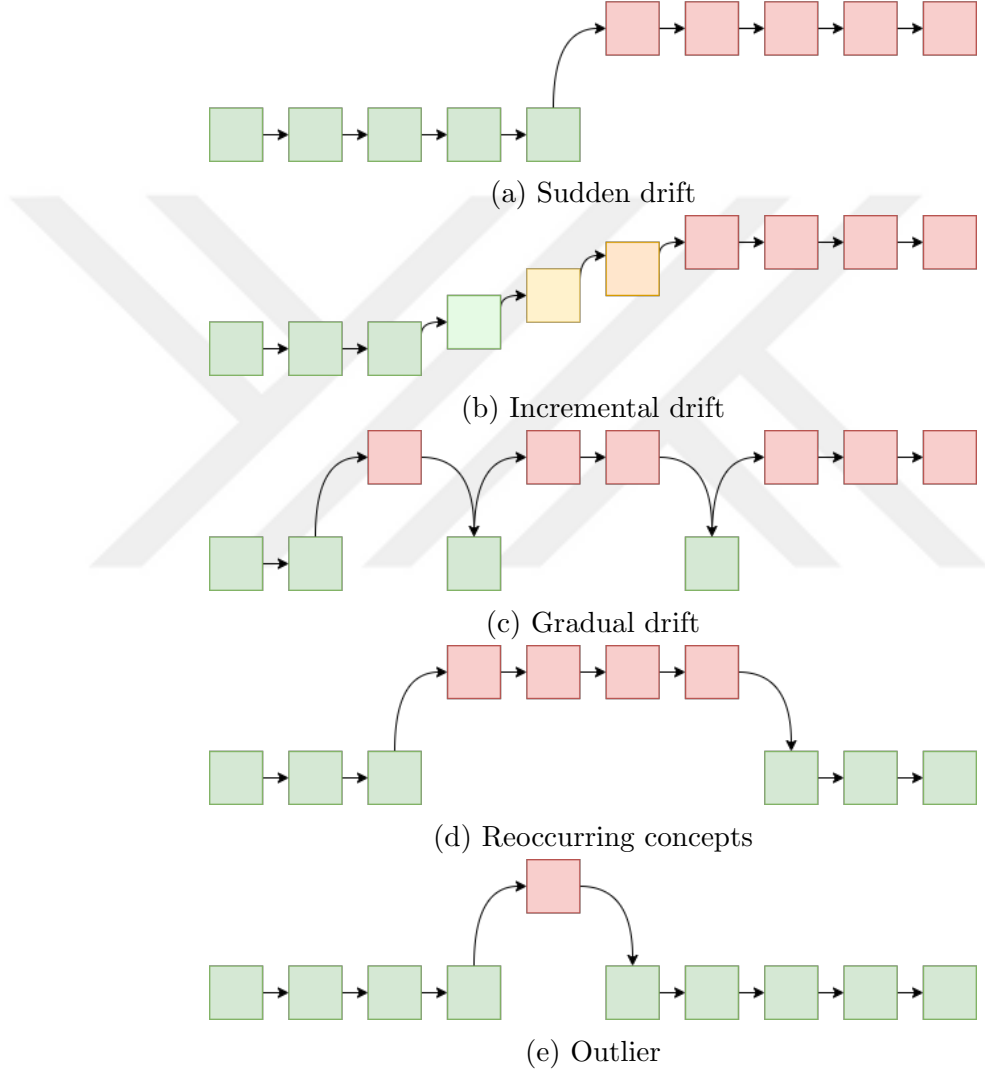


Figure 2.2: Types of concept drift based on what happens to the concept over time (outlier is not a drift). Different colors represent different concepts.

Chapter 3

Related Works

In the following sections, we first briefly introduce previously developed supervised detection algorithms based on online error rate monitoring and multi-label specific algorithms. Then, previous work on some of the unsupervised concept drift detection algorithms is presented.

3.1 Supervised Concept Drift Detection Algorithms

Concept drift is a prevalent problem in data stream mining, where many detectors are proposed to combat it within the literature. Based on previously developed algorithms, Lu et al. [16] propose an overall framework for concept drift detection which is comprised of four stages: Data retrieval, data modeling, test statistics calculation, and the hypothesis test. They further divide the detectors into three categories: Error rate-based, data distribution-based, and the multiple hypothesis test. Among these three, we focus on the first category as we were unable to find available source codes for multiple hypothesis test algorithms and data distribution-based algorithms. Table 3.1 displays the categories and methods used for each baseline algorithm we tested. The data modeling stage is not included in the table since all of the detectors except LD3 share the same data

modeling scheme which is learner-based modeling.

Table 3.1: Error rate-based drift detection algorithm characteristics (They are all our baselines; for those with more than one version, the number of variants is given in parentheses after the detector name). The symbols w_{hist} and w_{new} indicate the windows for old and new data in respective order. Landmark uses the entire set of data before drift and flushes this data when drift is detected to collect the new distribution.

Category	Detector	Data Retrieval	Test Statistics	Hypothesis Test
Error rate-based	ADWIN [19]	Auto cut w_{hist} , w_{new}	Error rate difference	Hoeffding's Bound
	DDM [20]	Landmark	Online Error Rate	Distribution Estimation
	EDDM [21]	Landmark	Online Error Rate	Distribution Estimation
	FHDDM [22]	Sliding w_{hist} , w_{new}	Correct prediction probability	Hoeffding's Bound
	FHDDMS (2) [23]	Stacked windows	Correct prediction probability	Hoeffding's Bound
	HDDM (2) [24]	Landmark	Online Error Rate	Hoeffding's Bound
	KSWIN [25]	Sliding w_{hist} , w_{new}	Distribution distance	KS-Test
	MDDM (3) [26]	Sliding w_{hist} , w_{new}	Weighted mean difference	McDiarmid's Bound
	RDDM [27]	Landmark	Online Error Rate	Distribution Estimation
	SeqDrift2 [28]	Sliding w_{hist} , w_{new}	Online error rate	Bernstein Bound
Data distribution-based	LD3 (Our method)	Sliding w_{hist} , w_{new}	Rank correlation	Distribution Estimation

3.1.1 Error Rate-based Algorithms

The detection algorithms in this category usually focus on the changes in the online error rate of the base classifiers, i.e, whether changes in the error rate of a classifier are statistically significant enough, the detector concludes that there is drift. Since these detectors monitor the online error rate and related metrics, they usually employ learners to model the data, which is the case for all of the algorithms presented in Section 3.1.1, and they tend to be efficient algorithms due to the simplicity of their input data (error-rate).

One of the most frequently used algorithms is the Drift Detection Method (DDM) proposed by Gama et al. [20]. It uses a separate classifier to check whether the change in the online error rate is significant enough, which is done through distribution estimation. If the change is statistically significant, drift is detected. There are also a large number of algorithms that build upon the DDM algorithm, which usually change the test statistics and hypothesis test stages. For instance, Early Drift Detection Method (EDDM) [21] improves DDM by also considering the sample-wise distance between erroneous classifications.

In order to improve the hypothesis test stage, Frías-Blanco et al. [24] propose HDDM, which utilizes Hoeffding’s inequality [29] to obtain probabilistic guarantees to further increase effectiveness. It has two variants, namely HDDM_A and HDDM_W, where, in the former, they use bounded moving averages (A-test) and in the latter, they use bounding weighted moving averages (W-test). Furthermore, Pesaranghader and Viktor [22] introduce Fast Hoeffding Drift Detection Method (FHDDM) which requires the base detectors to either stay at a steady level or improve in accuracy. FHDDM checks this by monitoring the most recent probability of correct predictions instead of the error rate. They further improve FHDDM as FHDDMS by adopting a stacking window scheme of various sizes [23]. The stacking windows are a short and long window that has overlapping content that is used to detect sudden and gradual drifts separately. It also has another variant, FHDDMS_add, where the authors employ additive summaries for better efficiency in memory and execution time.

Apart from Hoeffding’s inequality-based improvements to DDM, Pesaranghader et al. [26] developed McDiarmid Drift Detection Method (MDDM). MDDM is similar to HDDM as they also make improvements in hypothesis testing, by including the McDiarmid inequality to check the statistically significant differences. MDDM also implements a weighting scheme within its window where the most recent elements are given more importance. According to alternate weighting, it has three variants: MDDM_A (Arithmetic weighting), MDDM_E (Euler weighting), and MDDM_G (Geometric weighting).

Instead of modifying a stage, Barros et al. [27] propose Reactive Drift Detection Method (RDDM), in which they define a type of soft concept drift based on the number of samples accumulated in the window, called RDDM drift. RDDM re-calculates their DDM-based statistics to solve problems that arise from the high number of accumulated samples in the window.

Aside from DDM-based detectors, Bifet and Gavalda [19] introduce ADWIN, which monitors the difference in error rate between two adaptive windows. The difference is bounded by Hoeffding’s inequality and if the difference exceeds the Hoeffding bound, a drift is detected.

Furthermore, in [28], Pears et al. propose SeqDrift2. It employs an adapting sampling strategy to sample data from a window to get old and new concepts. Then, they detect the drift by comparing the differences according to the Bernstein bound.

Finally, Raab et al. [25] introduce KSWIN, which detects concept drift by applying “the Kolmogorov-Smirnov test” (KS-Test) which finds the distance between the estimated data distribution and the empirical distribution. These two distributions in KSWIN’s case would be the error rates stored in a sliced sliding window for new and old data. If the distance between the distributions exceeds the confidence interval, a drift is detected.

3.1.2 Algorithms for Multi-label Concept Drift Detection

In the multi-label classification problem domain, concept drift detection is a very thinly researched area. To the best of our knowledge, there are two previously developed concept drift detectors. However, in both cases the authors did not provide source codes thus we were unable to compare our results with them which is the reason why we did not include them in Table 3.1.

Shi et al. [30] propose a method in which they use label grouping and class entropy to detect concept drift. Initially they group the labels using clustering methods. Then for each label group, they calculate the multi-label entropy values within two sliding windows where they apply a threshold method to detect concept drift.

Furthermore, Wang et al.[31] introduce DDM-FP-M, a multi-label focused concept drift detector aimed at data streams on the Internet of Things problem domain. They propose modifications to DDM in which they add false positive classifications to make it more suitable for multi-labeled data stream environments.

3.2 Unsupervised Concept Drift Detection

Although many concept drift detectors exist, one area that is insufficiently researched is unsupervised concept drift detection. Iwashita and Papa [32] found that unsupervised drift detectors only make up 3% of developed concept drift detectors. However, as Gemaque et al. point out, many of the real-world problems are better suited for unsupervised drift detection, since the swift acquisition of labels is often not possible [33]. Since concept drift detection requires fast detection to enable rapid recovery after drift, the topic of unsupervised concept drift detection requires more attention.

Some of the previous works on unsupervised concept drift detection include detectors such as IKS-bdd [34], CD-TDS [35], Plover [36], DSDD [37], D3 [38], and OCDD [17].

IKS-bdd applies the Incremental Kolmogorov-Smirnov test, which is a modified Kolmogorov-Smirnov test that is better suited for online learning, to each of the data features in order to detect drift. CD-TDS detects two types of drift: Local drift and global drift. For local detection, the sample means of the new and old data are compared, bounded by the Hoeffding Bound. In the case of global drift, a pairwise statistical test is applied to find differences between two tree structures representing new and old data.

Plover monitors the input data behavior and detects changes based on observed instabilities. Moreover, DSDD detects drifts based on classification error simulation using an ensemble of classifiers.

Furthermore, Gözüaık et al. introduce D3, a drift detector that utilizes a discriminative classifier that is trained on auto labeled samples based on how recent a sample was seen within a window (“0” for old and “1” for new samples). Drift detection is made by measuring the AUC score of the classifier. Likewise, OCDD uses a one-class classifier to distinguish between changing concepts, in which drift is detected when the ratio of false predictions is higher than a threshold.

Our work proposes an unsupervised drift detection algorithm that utilizes the predicted labels of the paired classifier; however, we were unable to compare our detector with the described unsupervised algorithms, excluding D3 and OCDD, as we could not find readily available source code.

Furthermore, D3 and OCDD are not compared with LD3 since we only used original codes provided by the frameworks we used. The provided source codes are incompatible with multi-label classification and we chose not to modify the original code for this reason.

Chapter 4

Label Dependency Drift Detector

In this study, we propose LD3, a data distribution-based, unsupervised concept drift detector which exploits the dependencies among predicted labels. It works alongside any online multi-label classifier that does not have inherent drift detection properties to assist in handling the changes in data distribution. We evaluate the changes in the label dependencies through the predicted labels provided by the paired model, and if a significant change occurs in the dependencies, we detect concept drift. In the remainder of this chapter, Algorithm 1 is referenced for explanations.

Table 4.1: Symbols table for Algorithm 1.

Symbol	Description
l	Predicted label
L	Threshold for number of anomalies
t	Standard deviation multiplier
w	Number of samples in a window
W_{new}	Label window for new samples
W_{old}	Label window for old samples
W_{corr}	Past correlation window

Algorithm 1 LD3: Label Dependency Drift Detector

Initialize $\mathbf{W}_{new}$, $\mathbf{W}_{old}$ and $\mathbf{W}_{corr}$ as $\emptyset$

procedure LD3(l , t , w , L)

$drift \leftarrow \text{False}$

$l' \leftarrow \text{Insert_Element}(l, \mathbf{W}_{new})$ $\triangleright l'$ is the last element removed

if $|\mathbf{W}_{new}| = w$ **then**

$\text{Insert_Element}(l', \mathbf{W}_{old})$

end if

if $|\mathbf{W}_{new}| \neq w$ and $|\mathbf{W}_{old}| \neq w$ **then**

 return $drift$

else

$\mathbf{M}_{new} \leftarrow$ Co-occurrence matrix created from $\mathbf{W}_{new}$

$\mathbf{M}_{old} \leftarrow$ Co-occurrence matrix created from $\mathbf{W}_{old}$

$\mathbf{R}_{new} \leftarrow$ Reciprocal ranking of $\mathbf{M}_{new}$ $\triangleright$ Global ranking for $\mathbf{W}_{new}$

$\mathbf{R}_{old} \leftarrow$ Reciprocal ranking of $\mathbf{M}_{old}$ $\triangleright$ Global ranking for $\mathbf{W}_{old}$

$\mathbf{C} \leftarrow \text{WS}(\mathbf{R}_{new}, \mathbf{R}_{old})$ $\triangleright$ Rank Correlation of $\mathbf{R}_{new}$ and $\mathbf{R}_{old}$

$\text{Insert_Element}(\mathbf{C}, \mathbf{W}_{corr})$

if $|\mathbf{W}_{corr}| \neq w$ **then**

 return $drift$

end if

$len \leftarrow \text{Sigma_Rule}(\mathbf{W}_{corr}, t)$ $\triangleright$ Number of anomalies in $\mathbf{W}_{corr}$

if $len > L$ **then**

$drift \leftarrow \text{True}$

 Clear windows $\mathbf{W}_{new}, \mathbf{W}_{old}, \mathbf{W}_{corr}$

end if

end if

 return $drift$

end procedure

4.1 Evaluating the Changes in the Label Dependencies

In an ideal environment, labels can be chosen such that each label is independently distributed, i.e, dependencies do not exist. However, in real-world datasets, this is usually not the case, as some labels tend to occur more frequently together (e.g, in movie categories, comedy and drama tags occur more commonly together than comedy and thriller tags). In our study, we hypothesize that this correlation causes dependencies, and changes in these dependencies are a precursor to concept drift.

Given a multi-label classifier, we buffer the labels predicted by the classifier in two fixed-sized moving windows, one each for the new and old data. Apart from providing up-to-date statistics about label distribution, these windows allow the classifier to boot itself up, i.e, learn enough of the data distribution to generate consistent predictions, as the windows are filled (Alg. 1, lines 4-6), which functions as a warmup scheme.

After the windows are full, we first generate two co-occurrence matrices from the windows (Alg. 1, lines 10-11). The matrices are obtained by counting the number of times each class label occurs as “1” alongside other labels. The generated matrices are then ranked within each row, which we call local ranking, by creating a ranking for each label based on their co-occurrence frequencies.

Following the local ranking, the ranks are aggregated by utilizing a data fusion algorithm to obtain a representation of label dependencies for new and old samples (Alg. 1, lines 12-13). We call the resulting aggregated ranking as global ranking. Through this, we obtain the most influential labels among all of them and use these labels to monitor changes in the stream. We use reciprocal rank fusion, which is seen in Equation 4.1, where r_i is the global ranking of a class label within the predicted label l , which is denoted by l_i . Moreover, we use n to represent the number of classes and r_{ij} for the local ranking of l_i , where $\{j | 1 \leq j \leq n\}$. We justify our selection of reciprocal rank fusion, rather than some other commonly

used approaches, by experiments in Section 6.9.

$$r_i = \frac{1}{\sum_{j=1}^n \frac{1}{r_{ij}}} \quad \{i \mid 1 \leq i \leq n\} \quad (4.1)$$

4.2 Measuring Similarity Between Two Rankings

Subsequent to the global ranking, we calculate the rank correlation between the two global rankings we obtained (Alg. 1, line 14), which evaluates the similarity between rankings. We use the WS coefficient as our rank correlation measure which is a ranking similarity method developed by [39]. It calculates the weighted similarity between two rankings and returns a value bounded within $[-1, 1]$, where two identical rankings are scored as 1, whereas the score is -1 for the opposite case, from which can be deduced Equation 4.2. In this equation, R_{xi} and R_{yi} represent the rank position of a label l_i within R_{new} and R_{old} .

$$C = 1 - \sum_{i=1}^n \left(2^{-R_{xi}} \cdot \frac{|R_{xi} - R_{yi}|}{\max\{|1 - R_{xi}|, |n - R_{xi}|\}} \right) \quad (4.2)$$

Within the summation, $2^{-R_{xi}}$ is the weight of the label l_i which ensures that higher ranking labels have more influence. The numerator $|R_{xi} - R_{yi}|$ is the ranking distance of l_i within the two rankings and the denominator scales this distance.

We chose to use the WS coefficient instead of other popular rank correlation measures such as Kendall's τ [40] or Spearman's ρ [41] because it provides a way to weigh the labels. Since we measure the influences labels have on each other, a change in the higher ranked labels is more important than other labels. Furthermore, it measures the similarity based on the distance between the two rankings. In a possibly volatile environment like a data stream, rankings could change by a few places temporarily; however, if we measure the similarity based on distance, such irregularities are tolerated more easily, which is why measures

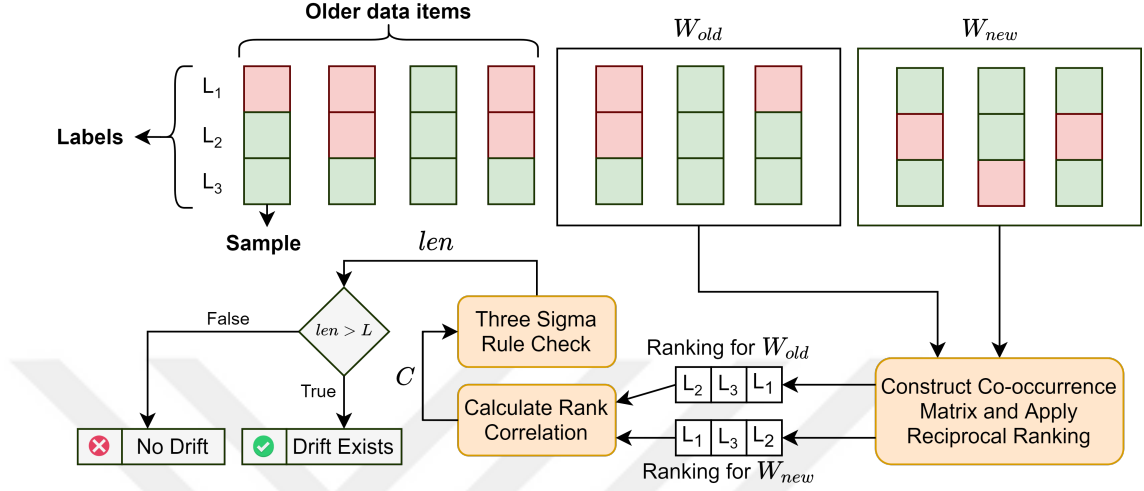


Figure 4.1: Workflow of LD3. Red boxes represent false labels (0) and green boxes represent true labels (1).

such as *weighted* τ [42] are not suitable for our case.

Although the generated rank correlation value implies whether the current ranking indicates a drift, it needs to be checked for statistical significance. This is done by utilizing the three sigma rule [43], which translates to a simple left tailed test for our case (Alg. 1, line 15-19). A separate moving window (W_{corr}) is utilized to accumulate the past rank correlations. If the current correlation is less than the mean (μ) within the window by t times the standard deviation (σ), it is considered an anomaly. However, such anomalies may be the result of noise in the data, i.e, outliers described in Chapter 2. To prevent false detections, we store these anomalies in a list, and if the length of this list is greater than a chosen length L , the detector signals that a drift has been encountered. The three sigma rule is used here to increase the computational efficiency of the algorithm as the desired result is obtained without having to estimate the distribution of the recent data in its entirety.

The general representation of all of the stages is illustrated in Figure 4.1.

4.3 Detailed Calculation Example

To illustrate and clarify the calculation of LD3, we present a detailed calculation of the algorithm in this section. Assume that a given data stream $S_{0,t} = d_0, \dots, d_t$, in a time window $[0, t]$, where $d_i = (X_i, y_i)$, the window size $w = 3$ and three classes, the most recent six predictions are:

$$S = ((0, 0, 1), (1, 1, 1), (0, 1, 1), (\mathbf{1}, \mathbf{0}, \mathbf{1}), (\mathbf{1}, \mathbf{1}, \mathbf{0}), (\mathbf{1}, \mathbf{0}, \mathbf{1})) \quad (4.3)$$

We first calculate two co-occurrence matrices M_{new} and M_{old} . For M_{new} we use the most recent w labels which are highlighted in bold. The next w labels after that are used for M_{old} . The matrices are calculated by counting the number of times each label occurs together. For instance, for the third sample $(0, 1, 1)$, the second and third labels are true together, so $M_{old}(2, 3)$ and $M_{old}(3, 2)$ are incremented by one. Each row of the matrices represents the co-occurrence for a different label. For the given stream, the resulting matrices are the following:

$$M_{old} = \begin{bmatrix} 0 & 1 & 1 \\ 1 & 0 & 2 \\ 1 & 2 & 0 \end{bmatrix} \quad (4.4) \quad M_{new} = \begin{bmatrix} 0 & 1 & 2 \\ 1 & 0 & 0 \\ 2 & 0 & 0 \end{bmatrix} \quad (4.5)$$

Next, we calculate the ranking of the co-occurrences for each label, i.e, which labels occur more frequently together with the currently ranked label. A.4 is for M_{old} and A.5 is for M_{new} , where L_i represents the rankings for each past predicted label and L'_i is used for the most recently predicted labels.

$$\begin{aligned} L_1 &\leftarrow l_2 = l_3 & L'_1 &\leftarrow l_3 > l_2 \\ L_2 &\leftarrow l_3 > l_1 & L'_2 &\leftarrow l_1 > l_3 \\ L_3 &\leftarrow l_2 > l_1 & L'_3 &\leftarrow l_1 > l_2 \end{aligned} \quad (4.6) \quad (4.7)$$

We perform reciprocal rank fusion on the obtained rankings to get a global, aggregated ranking for the old and new labels. For tied rankings, we assume they

each have the same ranking. r_i and r'_i are the reciprocal rank of each label for old and new predictions in their respective order. The calculated ranks are sorted in ascending order to get the global rankings R_{old} and R_{new} . If two reciprocal ranks are equal, they are sorted based on their label indexes.

$$\begin{aligned}
r_1 &= \frac{1}{\frac{1}{2} + \frac{1}{2}} = 1 & r'_1 &= \frac{1}{\frac{1}{1} + \frac{1}{1}} = \frac{1}{2} \\
r_2 &= \frac{1}{\frac{1}{1} + \frac{1}{1}} = \frac{1}{2} & r'_2 &= \frac{1}{\frac{1}{2} + \frac{1}{2}} = 1 \\
r_3 &= \frac{1}{\frac{1}{1} + \frac{1}{1}} = \frac{1}{2} & r'_3 &= \frac{1}{\frac{1}{1} + \frac{1}{2}} = \frac{2}{3}
\end{aligned} \tag{4.8}$$

$$\begin{aligned}
R_{old} &= [l_2, l_3, l_1] & R_{new} &= [l_1, l_3, l_2]
\end{aligned} \tag{4.9}$$

The obtained R_{old} and R_{new} rankings are then used to calculate the WS coefficient to measure the rank correlation between the two rankings. In A.8, the ranks represent the ranking position of the i -th label where $R^{old} = [2, 0, 1]$ and $R^{new} = [0, 2, 1]$.

$$C = 1 - \sum_{i=1}^3 \left(2^{-R_i^{new}} \cdot \frac{|R_i^{new} - R_i^{old}|}{\max\{|1 - R_i^{new}|, |3 - R_i^{new}|\}} \right) \tag{4.10}$$

$$\begin{aligned}
C = 1 - \left(2^0 \cdot \frac{|0 - 2|}{\max\{|1 - 0|, |3 - 0|\}} + 2^{-2} \cdot \frac{|2 - 0|}{\max\{|1 - 2|, |3 - 2|\}} + \right. \\
\left. 2^{-1} \cdot \frac{|1 - 1|}{\max\{|1 - 1|, |3 - 1|\}} \right)
\end{aligned} \tag{4.11}$$

$$C = 1 - \left(\frac{2}{3} + \frac{1}{2} + 0 \right) = -0.167 \tag{4.12}$$

For the remaining part, if the obtained correlation C is less than $\mu - t\sigma$, where μ is the mean, σ is the standard deviation and t is the σ multiplier used in three sigma rule, the obtained correlation is added to an anomaly list. The drift is

detected when the length of the anomaly list generated from W_{corr} is greater than the set L threshold.



Chapter 5

Evaluation and Experimental Setup

5.1 Datasets

For evaluation, we used nine well-known real-world multi-label datasets, which we obtained in MEKA [44] formats, and three synthetic data streams [45]. Table 5.1 represents the datasets and their properties. The datasets are chosen among the ones tested in the study of [18]. We used all datasets from that list shown to have drift by any of the tested detectors and had better effectiveness results than a baseline classifier without drift detection functionalities. Since we measure the detectors' benefit based on the effectiveness improvement on the classifier, if the baseline classifier shows higher predictive performance, we conclude that detectors only make false positive detections. In such a case, we assume the dataset does not have drift.

The three synthetic data streams are generated to measure the difference in effectiveness results among sudden, incremental, and reoccurring drifts, which are generated by changing between three different multi-label data streams that have different label cardinalities (for reoccurring concepts, the third data stream has the same properties and data distribution as the first data stream). The change between streams is smoothed by the sigmoid function over a specified amount

Table 5.1: Table of multi-label datasets used in the experiments¹. The upper group contains the real-world datasets. N represents the number of samples, D is the number of features, n is the number of labels and $LC(\mathcal{D})$ and $LD(\mathcal{D})$ represents label cardinality and label density which are the average number of true labels of the samples in dataset $\mathcal{D}$ and $LC(\mathcal{D})$ divided by the number of labels, indicating the average label cardinality per label [1].

Dataset Name	Domain	N	D	n	$LC(\mathcal{D})$	$LD(\mathcal{D})$
20NG	Text	19,300	1,006	20	1.029	0.051
Birds	Audio	645	260	19	1.014	0.053
Enron	Text	1,702	1,001	53	3.378	0.064
EukaryotePseAAC	Biology	7,766	440	22	1.146	0.052
Imdb	Text	120,900	1,001	28	2.000	0.071
Ohsumed	Text	13,930	1,002	23	1.663	0.072
PlantPseAAC	Biology	978	440	12	1.079	0.090
Tmc2007-500	Text	28,600	500	22	2.220	0.101
Yeast	Biology	2,417	103	14	4.237	0.303
Synthetic Sudden Drift	Generic	20,000	200	50	1.587	0.0397
Synthetic Incremental Drift	Generic	20,000	200	50	1.588	0.0397
Synthetic Re-occurring Drift	Generic	20,000	200	50	1.588	0.0397

^aThe real-world datasets can be accessed from: <http://www.uco.es/kdis/mlresources/>

of samples. To simulate sudden drift, we applied the change over one sample whereas in the reoccurring and incremental drifted streams the change is applied over 500 samples. The data streams have 20,000 samples and the drifts are at sample positions 4,000 and 10,000. The properties of the streams are shown in Table 5.1.

5.2 Experimental Setup and Evaluation Metrics

The experiments are performed using the Scikit-Multiflow [45] and Tornado [23] frameworks. Scikit-multiflow is employed for synthetic data stream generation and it contains six (ADWIN, DDM, EDDM, HDDM_A, HDDM_W, KSWIN) of the tested baseline detectors. Moreover, the Tornado framework is utilized as it contains the remaining eight baseline detectors.

To demonstrate the effectiveness benefits of the detectors, a Classifier Chain (CC) [46] using Gaussian Naive Bayes (NB) classifiers [47] is used as a base classifier. CCs build binary classifiers for each class label that are linked in a chain to preserve inter-label dependencies. The reason for using a CC is that it provides accurate results with reasonably fast execution. The same reasoning applies to the use of NBs as base classifiers and it is also a popular base classifier that is frequently used in the literature [48]. Furthermore, our concept drift adapting strategy resets the classifier if drift is detected.

[15] propose three possible metrics for change detection algorithms: (1) Probability of true change detection, (2) Probability of false alarms, and (3) Delay of detection. In our experiments and discussions, in order to incorporate the suggested measures, we first perform effectiveness tests on the 12 datasets we presented. Then, through plots, we make an in-depth analysis on the obtained effectiveness measures on the top four performing detectors, in which we discuss the effects of the false detections (alarms) and true detections. Lastly, to measure the impact of detection delay, we perform experiments using six synthetic data streams with varying drift speeds, i.e, how fast does the change happen between old and new concepts. This experiment allows us to assess the detectors’ response within different drift environments.

The tests are conducted using prequential evaluation [49]. We apply the following four metrics to evaluate the effectiveness of the algorithms which are used in previous literature to measure multi-label classification effectiveness [14, 50].

- Example-based metrics: Example-based accuracy (Eq. 5.1), Hamming score (Eq. 5.2), and example-based F1 score (Eq. 5.3). The formulas for these are given below in equations with $\hat{y}_i$ being the prediction, Y_i as the ground truth, n as the number of labels, and N being the number of samples.
- Label-based metrics: Micro-averaged F1 score (Eq. 5.4). TP_i , FP_i , and FN_i are true positive, false positive, and false positive counts for the i -th label [51].

$$Accuracy_{example} = \frac{1}{N} \sum_{i=1}^N \frac{|Y_i \cap \hat{y}_i|}{|Y_i \cup \hat{y}_i|} \quad (5.1)$$

$$HammingScore = 1 - \frac{1}{N} \sum_{i=1}^N \frac{1}{n} |\hat{y}_i \Delta Y_i| \quad (5.2)$$

$$F1_{example}^2 = \frac{2 \cdot Precision_{example} \cdot Recall_{example}}{Precision_{example} + Recall_{example}} \quad (5.3)$$

$$F1_{micro} = F1 \left(\sum_{i=1}^n TP_i, \sum_{i=1}^n FP_i, \sum_{i=1}^n FN_i \right)^3 \quad (5.4)$$

LD3 is compared with the detection algorithms displayed in Table 3.1 and their variations. Although these detection algorithms are not specifically designed for multi-label use-case, they are eligible baseline algorithms because they are error rate-based detectors. For this type of algorithms, the input passed into the detector is whether or not the prediction is correct, which means that their input is “1” for correct and “0” for false predictions (For “1”, all predicted labels must match exactly to the true labels). For this reason, they can also be used in the multi-label concept detection problem domain since this information is also generated in multi-label data streams. [31] also test their multi-label concept drift detection algorithm against DDM.

² $Precision_{example} = \frac{1}{N} \sum_{i=1}^N \frac{|Y_i \cap \hat{y}_i|}{|\hat{y}_i|}$
³ $F1(TP, FP, FN) = \frac{2 \cdot TP}{2 \cdot TP + FN + FP}$

$Recall_{example} = \frac{1}{N} \sum_{i=1}^N \frac{|Y_i \cap \hat{y}_i|}{|Y_i|}$

Chapter 6

Experimental Results and Discussion

In the following sections, we discuss and analyze our results based on effectiveness measures, provide simple guidelines for hyperparameter optimization, discuss the effects of different data fusion algorithms, and the influence of various drift types and speeds.

6.1 Effectiveness Analysis

We assess our effectiveness results by comparing our experimental results with the baseline detectors. Furthermore, we visually analyze the effectiveness results in the time scale. Lastly, we discuss the advantages of unsupervised detection with pointers on efficiency.

Table 6.1: Experiment results for the detectors¹. ND is a classifier without any detection methods for the baseline. Average result is calculated from the average result of each algorithm, except LD3, for a dataset. *LD3 Imp.* represents LD3’s improvement over the average. *Avg. Imp.* illustrates the average improvement for each metric. The best results are highlighted in bold.

(a) Example-based accuracy and Hamming score results

Detector	20NG	Birds	Enron	Enkatype PseAAC	Imdb	Ohsumed	Plant PseAAC	Tmc2007- 500	Yeast	Synthetic Incremental	Synthetic Reoccurring	Synthetic Sudden	Average Rank
Example-based Accuracy													
LD3	0.1616	0.0992	0.1403	0.1946	0.1140	0.0693	0.3198	0.2029	0.3780	0.0788	0.0808	0.0788	1.00
ADWIN	0.1386	0.0550	0.1245	0.0327	0.1093	0.0361	0.1431	0.1962	0.3622	0.0724	0.0730	0.0724	8.46
DDM	0.0525	0.0550	0.1255	0.0327	0.0799	0.0364	0.1126	0.1529	0.3557	0.0749	0.0759	0.0749	10.58
EDDM	0.0654	0.0550	0.1355	0.0327	0.0894	0.0368	0.1272	0.1593	0.3707	0.0787	0.0762	0.0744	6.25
FHDDM	0.1453	0.0550	0.1236	0.0327	0.0797	0.0381	0.1240	0.1555	0.3622	0.0659	0.0761	0.0659	9.83
FHDDMS	0.1444	0.0550	0.1236	0.0327	0.0797	0.0381	0.1136	0.1555	0.3622	0.0659	0.0761	0.0659	10.42
FHDDMS-Add	0.1483	0.0550	0.1239	0.0327	0.0797	0.0373	0.1165	0.1555	0.3622	0.0659	0.0761	0.0659	9.96
HDDM.A	0.1505	0.0550	0.1371	0.0327	0.0797	0.0367	0.1235	0.1555	0.3622	0.0698	0.0761	0.0722	8.08
HDDM.W	0.1469	0.0550	0.1371	0.0327	0.0797	0.0367	0.1070	0.1555	0.3622	0.0659	0.0761	0.0659	9.79
KSWIN	0.1438	0.0550	0.1245	0.0327	0.0797	0.0590	0.1793	0.1555	0.3622	0.0659	0.0761	0.0659	8.67
MDDM.A	0.1396	0.0550	0.1237	0.0327	0.0797	0.0381	0.1218	0.1555	0.3622	0.0659	0.0761	0.0659	10.25
MDDM.E	0.1419	0.0550	0.1242	0.0327	0.0797	0.0381	0.1180	0.1555	0.3622	0.0659	0.0761	0.0659	10.00
MDDM.G	0.1419	0.0550	0.1242	0.0327	0.0797	0.0381	0.1180	0.1555	0.3622	0.0659	0.0761	0.0659	10.00
RDDM	0.1562	0.0574	0.1237	0.1225	0.1130	0.0382	0.2225	0.2016	0.3622	0.0722	0.0746	0.0720	5.38
SEQDRIFT2	0.1190	0.0550	0.1283	0.0327	0.0797	0.0363	0.1709	0.1555	0.3622	0.0717	0.0739	0.0750	9.25
ND	0.1469	0.0550	0.1371	0.0327	0.0797	0.0608	0.1369	0.1555	0.3622	0.0659	0.0761	0.0659	8.08
Average Result	0.1321	0.0552	0.1278	0.0387	0.0846	0.0403	0.1357	0.1614	0.3623	0.0689	0.0756	0.0689	Avg. Imp.
LD3 Imp. (%)	22.3	79.7	9.8	402.8	34.8	72.0	135.7	25.7	4.3	14.4	6.9	14.4	68.6
Hamming Score													
LD3	0.8578	0.7302	0.8033	0.7979	0.7673	0.8023	0.8559	0.7340	0.7193	0.8098	0.7670	0.8306	1.50
ADWIN	0.8426	0.4996	0.7428	0.5773	0.7499	0.6520	0.7765	0.6867	0.6970	0.7396	0.7342	0.7408	6.67
DDM	0.2271	0.4996	0.7398	0.5773	0.7000	0.1627	0.5688	0.5019	0.7087	0.8283	0.8016	0.7888	10.17
EDDM	0.5486	0.4996	0.7594	0.5773	0.7039	0.4412	0.5984	0.6007	0.7221	0.7969	0.7743	0.7830	8.00
FHDDM	0.8550	0.4996	0.7423	0.5773	0.7088	0.7520	0.6887	0.5069	0.6970	0.6773	0.6971	0.6781	9.21
FHDDMS	0.8599	0.4996	0.7423	0.5773	0.7088	0.7485	0.6463	0.5069	0.6970	0.6773	0.6971	0.6781	9.67
FHDDMS-Add	0.8340	0.4996	0.7404	0.5773	0.7088	0.7490	0.6575	0.5069	0.6970	0.6773	0.6971	0.6781	10.54
HDDM.A	0.8069	0.4996	0.7632	0.5773	0.7088	0.4764	0.6022	0.5069	0.6970	0.7026	0.6971	0.7412	9.38
HDDM.W	0.7953	0.4996	0.7632	0.5773	0.7088	0.5013	0.5768	0.5069	0.6970	0.6773	0.6971	0.6781	10.58
KSWIN	0.8343	0.4996	0.7445	0.5773	0.7088	0.6916	0.7489	0.5069	0.6970	0.6773	0.6971	0.6781	9.46
MDDM.A	0.8518	0.4996	0.7415	0.5773	0.7088	0.7520	0.6835	0.5069	0.6970	0.6773	0.6971	0.6781	9.67
MDDM.E	0.8504	0.4996	0.7428	0.5773	0.7088	0.7523	0.6714	0.5069	0.6970	0.6773	0.6971	0.6781	9.29
MDDM.G	0.8504	0.4996	0.7428	0.5773	0.7088	0.7524	0.6714	0.5069	0.6970	0.6773	0.6971	0.6781	9.21
RDDM	0.8696	0.5140	0.7417	0.6560	0.7654	0.7669	0.7847	0.7227	0.6970	0.7272	0.7117	0.7286	4.58
SEQDRIFT2	0.7942	0.4996	0.7421	0.5773	0.7088	0.6304	0.7763	0.5069	0.6970	0.7266	0.7327	0.7778	8.83
ND	0.7953	0.4996	0.7632	0.5773	0.7088	0.6685	0.8168	0.5069	0.6970	0.6773	0.6971	0.6781	9.25
Average Result	0.7744	0.5006	0.7475	0.5825	0.7144	0.6332	0.6845	0.5392	0.6995	0.7078	0.7150	0.7109	Avg. Imp.
LD3 Imp. (%)	10.8	45.9	7.5	37.0	7.4	26.7	25.0	36.1	2.8	14.4	7.3	16.8	19.8

^dWe will make the implementation of LD3 available on Github after the review process.

Table 6.1: Continued.

(b) Example-based F1 score and Micro-averaged F1 score results

Detector	20NG	Birds	Enron	Eukatyote PseAAC	Imdb	Ohsumed	Plant PseAAC	Tmc2007- 500	Yeast	Synthetic Incremental	Synthetic Reoccurring	Synthetic Sudden	Average Rank
Example-based F1 Score													
LD3	0.3173	0.1242	0.3496	0.4593	0.2240	0.1031	0.5596	0.4048	0.5299	0.1387	0.1719	0.1325	2.75
ADWIN	0.2999	0.1527	0.3150	0.1074	0.2077	0.0424	0.2479	0.3858	0.5204	0.1371	0.1455	0.1364	7.46
DDM	0.1207	0.1527	0.3017	0.1074	0.1424	0.0646	0.2547	0.3010	0.5072	0.1124	0.1201	0.1191	12.08
EDDM	0.1932	0.1527	0.3155	0.1074	0.1606	0.0557	0.2737	0.3458	0.5159	0.1236	0.1340	0.1265	8.17
FHDDM	0.2965	0.1527	0.3050	0.1074	0.1409	0.0424	0.2408	0.3024	0.5204	0.1210	0.1556	0.1211	10.38
FHDDMS	0.2915	0.1527	0.3050	0.1074	0.1409	0.0435	0.2253	0.3024	0.5204	0.1210	0.1556	0.1211	10.25
FHDDMS_Add	0.3024	0.1527	0.3077	0.1074	0.1409	0.0439	0.2283	0.3024	0.5204	0.1210	0.1556	0.1211	9.29
HDDM_A	0.3145	0.1527	0.3176	0.1074	0.1409	0.0530	0.2549	0.3024	0.5204	0.1264	0.1556	0.1386	6.46
HDDM_W	0.3099	0.1527	0.3176	0.1074	0.1409	0.0526	0.2104	0.3024	0.5204	0.1210	0.1556	0.1211	8.75
KSWIN	0.2893	0.1527	0.3030	0.1074	0.1409	0.0974	0.3273	0.3024	0.5204	0.1210	0.1556	0.1211	8.96
MDDM_A	0.2992	0.1527	0.3066	0.1074	0.1409	0.0426	0.2422	0.3024	0.5204	0.1210	0.1556	0.1211	9.79
MDDM_E	0.3040	0.1527	0.3065	0.1074	0.1409	0.0428	0.2375	0.3024	0.5204	0.1210	0.1556	0.1211	9.62
MDDM_G	0.3040	0.1527	0.3065	0.1074	0.1409	0.0428	0.2375	0.3024	0.5204	0.1210	0.1556	0.1211	9.62
RDDM	0.3021	0.1461	0.3029	0.3905	0.2239	0.0429	0.4141	0.3986	0.5204	0.1432	0.1536	0.1436	6.54
SEQDRIFT2	0.2828	0.1527	0.3241	0.1074	0.1409	0.0429	0.3099	0.3024	0.5204	0.1299	0.1410	0.1339	8.21
ND	0.3099	0.1527	0.3176	0.1074	0.1409	0.1044	0.2430	0.3024	0.5204	0.1210	0.1556	0.1211	7.67
Average Result	0.2813	0.1523	0.3102	0.1263	0.1523	0.0543	0.2632	0.3108	0.5192	0.1241	0.1500	0.1259	Avg. Imp.
LD3 Imp. (%)	12.8	-18.5	12.7	263.7	47.1	89.9	112.6	30.2	2.1	11.8	14.6	5.2	48.7
Micro-averaged F1 Score													
LD3	0.2783	0.1804	0.2461	0.3258	0.2047	0.1297	0.4846	0.3374	0.5486	0.1462	0.1495	0.1461	1.00
ADWIN	0.2435	0.1043	0.2215	0.0633	0.1970	0.0697	0.2504	0.3280	0.5318	0.1351	0.1360	0.1351	8.42
DDM	0.0997	0.1043	0.2230	0.0633	0.1480	0.0702	0.2024	0.2652	0.5247	0.1393	0.1411	0.1393	10.58
EDDM	0.1227	0.1043	0.2387	0.0633	0.1641	0.0711	0.2257	0.2749	0.5409	0.1460	0.1417	0.1385	6.25
FHDDM	0.2537	0.1043	0.2200	0.0633	0.1476	0.0734	0.2206	0.2691	0.5318	0.1236	0.1414	0.1236	9.92
FHDDMS	0.2523	0.1043	0.2200	0.0633	0.1476	0.0735	0.2040	0.2691	0.5318	0.1236	0.1414	0.1236	10.29
FHDDMS_Add	0.2583	0.1043	0.2205	0.0633	0.1476	0.0719	0.2087	0.2691	0.5318	0.1236	0.1414	0.1236	9.96
HDDM_A	0.2616	0.1043	0.2411	0.0633	0.1476	0.0708	0.2199	0.2691	0.5318	0.1306	0.1414	0.1346	8.12
HDDM_W	0.2561	0.1043	0.2411	0.0633	0.1476	0.0709	0.1933	0.2691	0.5318	0.1236	0.1414	0.1236	9.75
KSWIN	0.2514	0.1043	0.2214	0.0633	0.1476	0.1115	0.3041	0.2691	0.5318	0.1236	0.1414	0.1236	8.71
MDDM_A	0.2450	0.1043	0.2201	0.0633	0.1476	0.0734	0.2171	0.2691	0.5318	0.1236	0.1414	0.1236	10.38
MDDM_E	0.2485	0.1043	0.2209	0.0633	0.1476	0.0734	0.2111	0.2691	0.5318	0.1236	0.1414	0.1236	10.08
MDDM_G	0.2485	0.1043	0.2209	0.0633	0.1476	0.0735	0.2111	0.2691	0.5318	0.1236	0.1414	0.1236	9.88
RDDM	0.2701	0.1087	0.2202	0.2183	0.2030	0.0736	0.3640	0.3356	0.5318	0.1348	0.1389	0.1344	5.33
SeqDrift2	0.2127	0.1043	0.2274	0.0633	0.1476	0.0700	0.2919	0.2691	0.5318	0.1338	0.1376	0.1395	9.25
ND	0.2561	0.1043	0.2411	0.0633	0.1476	0.1147	0.2408	0.2691	0.5318	0.1236	0.1414	0.1236	8.08
Average Result	0.2320	0.1046	0.2265	0.0736	0.1557	0.0774	0.2378	0.2776	0.5319	0.1288	0.1406	0.1289	Avg. Imp.
LD3 Imp. (%)	20.0	72.5	8.7	342.7	31.5	67.6	103.4	21.5	3.1	13.5	6.3	13.3	58.7

6.2 Analysis with Different Effectiveness Measures

The results of our experiments are presented in Table 6.1. Each row represents the results of a detector and columns represent the datasets. The table is divided into four sections for each effectiveness measure. The baseline classifier without a detector is labeled as “ND”, i.e., no detector. While calculating the average rankings, the ties are handled by averaging the ranks that would be assigned to all the tied values and assigning that value to each tied element, e.g., for the given list $[1, 2, 3, 3, 4]$, the average ranking is $[1, 2, 3.5, 3.5, 5]$.

The results show that LD3 achieves the best results overall on all metrics. For all metrics, LD3 achieves at least 19.8% improvement over the baseline averages.

However, average rank-wise, LD3 displays comparatively worse predictive performance on the example-based F1 score than other metrics. We found that the reason behind this is mostly due to the classifier resetting procedure. Example-based F1 score punishes miss-classifications more harshly than example-based accuracy since miss-classified samples are multiplied in its numerator and they are divided by additive components. A recently reset classifier is prone to incorrect predictions and it is usually the case that none of the predicted labels match with the true labels for the first samples after resetting which produces the outcome of worse example-based F1 scores.

6.3 Statistical Significance Evaluation of Effectiveness Results

To further investigate the effectiveness of LD3, we performed the “*Friedman test with Nemenyi post-hoc analysis*” in which we evaluate the statistical significance of our results, which is presented in Figure 6.1. We applied the two-tailed Nemenyi test to find our critical distance for Nemenyi Significance. Our critical distance is $CD = 6.659$, which is calculated using Equation 6.1, where $q_{\alpha,k}$ is the critical value acquired from the Critical Values Table from [52] with $\alpha = 0.05$ and k is the number of algorithms (16) and K is the number of datasets (12).

$$CD = q_{\alpha,k} \sqrt{\frac{k(k+1)}{6K}} \quad (6.1)$$

The Nemenyi Significance tests show that LD3 achieves the overall best rankings with a statistically significant difference from most of the detectors. Among the different effectiveness measures, only RDDM and EDDM consistently display statistically indistinguishable predictive performance.

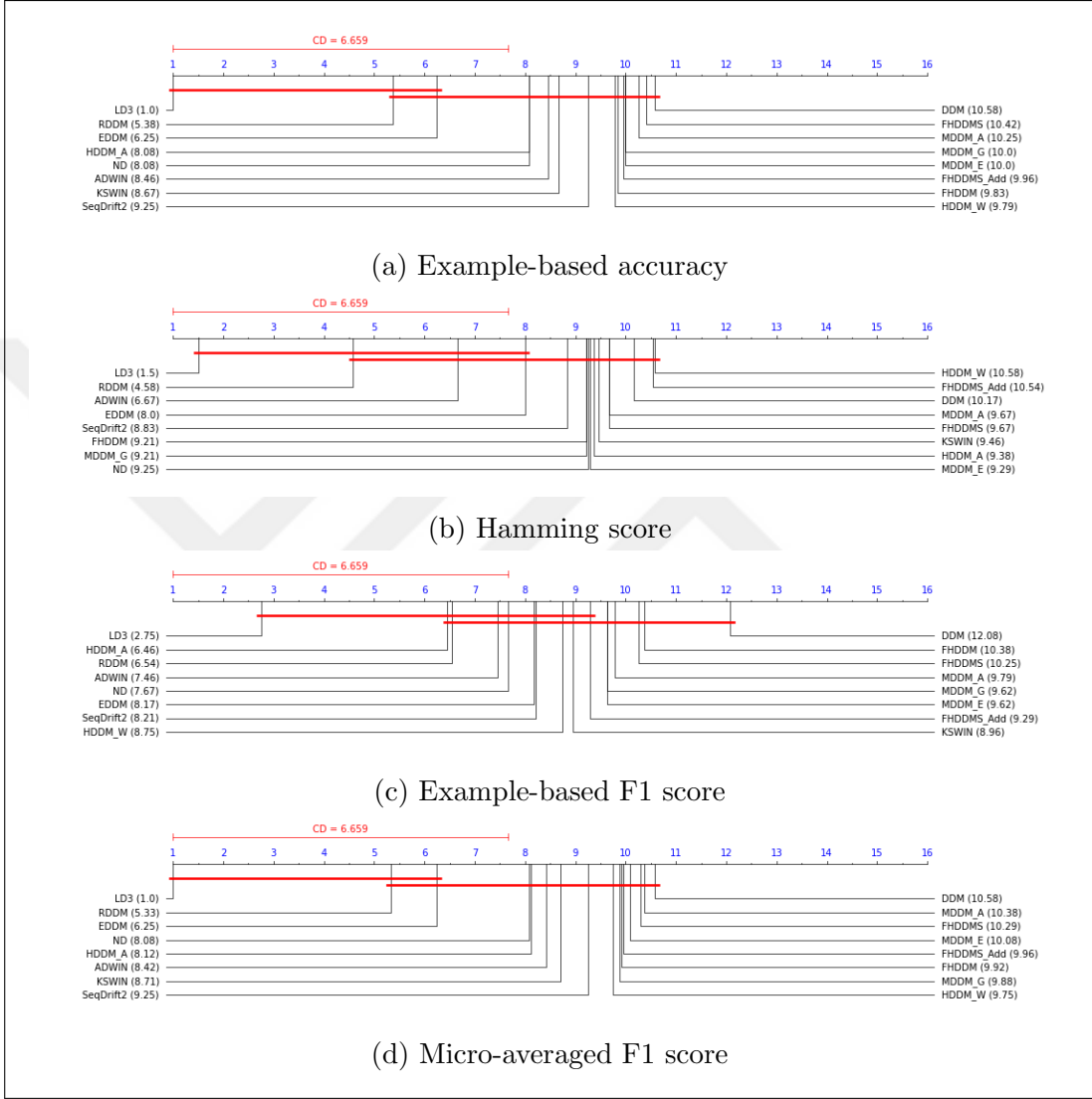


Figure 6.1: Nemenyi critical distance diagrams for all metrics. The number given within parentheses after the method name indicates rank position of the method.

It should be observed that contrary to prior expectation, ND does not show the worst results and it achieves better results compared to more than half of the tested detectors. Our experiments show that apart from ADWIN, EDDM, HDDM_A, and RDDM, the baseline detectors either do not detect drift or perform worse than ND in general. We identify the reason behind this as frequent false positive drift detections. An incorrectly detected drift leads the classifier to discard the progress it made which severely reduces the effectiveness results. In

addition, after a reset, the error rate of the classifier is increased which causes an error rate-based detector to miss-interpret the miss-classification statistics and may result in oscillating drift detections. The positive feedback loop generated from both of these factors causes the detectors to produce worse results than ND.

However, our experiments confirm that ADWIN, EDDM, HDDM_A, and RDDM are suitable for drift detection on multi-label data streams. For the following discussions, we use the top four detectors in terms of average rank, namely, LD3, ADWIN, EDDM, and RDDM.

6.4 Time Change Analysis of Effectiveness

One aspect we would like to highlight is LD3’s ability to detect drifts even when tested with a dataset that has low label cardinality. Our initial assumption was that LD3 would perform worse on datasets with lower label cardinality, since there would be less co-occurrence. However, in datasets with low label cardinality such as Birds, 20NG, and PlantPseAAC; LD3 still outperforms the baselines in general. We found that even if the label cardinality is low, as long as the data stream is a multi-label stream, some labels still have more influence over other labels and LD3 continues to be effective. One problem that arises from this is that how the algorithm should respond if the nature of the stream changes from multi-label to multi-class stream. Our suggestion in such a case would be to switch to a different detector suited to multi-class streams since LD3 would lose its ability to detect drifts. However, we were unable to find any detectors that could detect a change in the data stream’s nature and recommend further research for such cases.

Figure 6.2 plots the time scale example-based accuracy results of the top four detectors and ND. These datasets are chosen as most of the detectors co-detect drifts and their plots are more clear, allowing better inspection. The datasets are divided into 25 subsets and the average accuracy within those subsets is plotted.

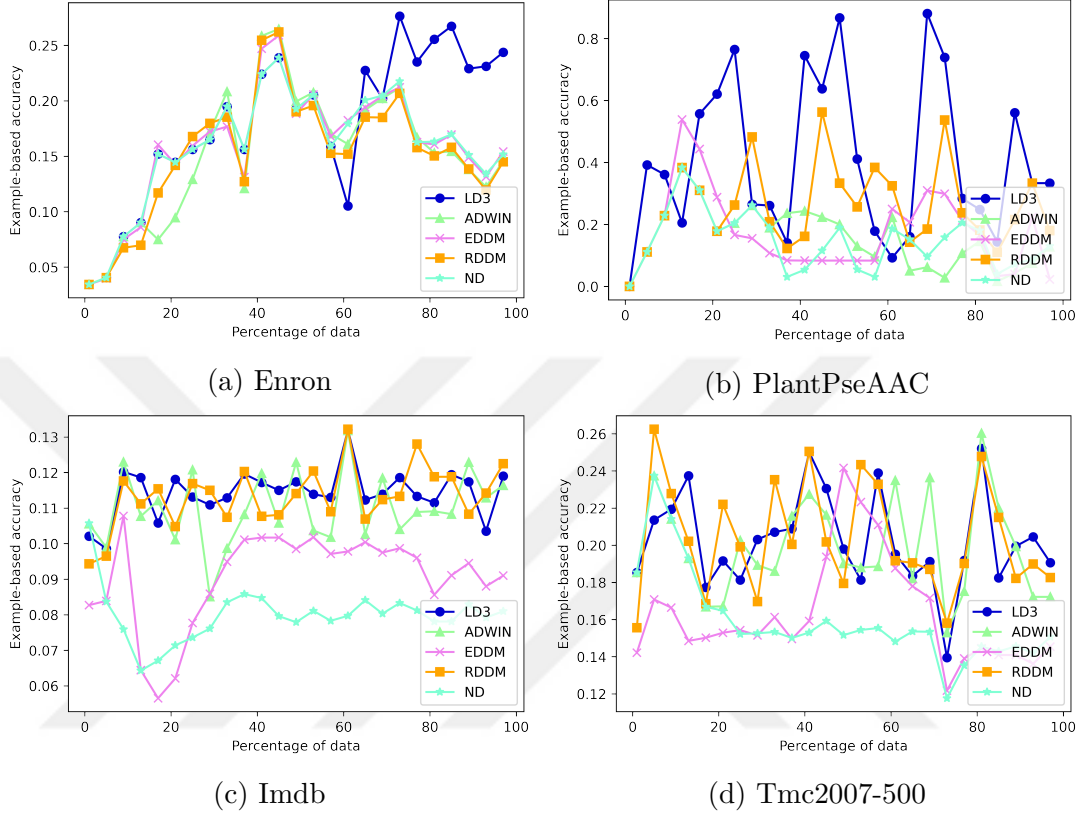


Figure 6.2: Example-based accuracy results of the top detectors and ND on four datasets. Y-axis of the plots is given with different scales.

Most of the sudden decreases in accuracy are caused by the classifier resets which means that at those points, the algorithms deduce that drift is present. Noticeably on the PlantPseAAC dataset, the sudden decreases in accuracy occur earlier for LD3, after which the classifier produces higher accuracy values than the baselines. LD3’s early detections allow the classifier to learn larger amounts of samples and adapt to the new concept faster. This is also observed in the second half of the Enron dataset where the other detectors stagnate in accuracy while LD3 resets to learn the new drift.

Furthermore, the plots illustrate that LD3 is more resistant to the noisy data originating from a recently reset classifier. Notably with ADWIN and RDDM on Tmc2007-500 and Imdb datasets, we observe that there are numerous consecutive drops in accuracy, which is the result of a new detection following a recent reset. While the classifier builds a probabilistic model for the new distribution,

some noise is usually expected. Yet for error rate-based algorithms like ADWIN and RDDM, this noise causes false detections that result in oscillating resets as previously discussed. LD3 combats this problem by waiting until sufficient statistics are obtained while the windows get filled, which allows faster recovery. For instance, in the Tmc2007-500 dataset, RDDM generally reaches higher accuracy values. However, RDDM resets the classifier frequently as opposed to LD3’s more stable execution, which results in marginally worse overall accuracy (LD3: 0.2029, RDDM: 0.2016) for RDDM, despite having multiple higher local maximum values. A similar case is also displayed with the Imdb dataset.

We define robustness, in the context of concept drift detection, as the ability to aid classifiers to provide the best predictive performance within a data stream with changing concepts. Our investigation on Figure 6.2 and results on Table 6.1 demonstrate that LD3 is a robust detection algorithm that assists the classifier to achieve the best effectiveness results overall, within varying streaming environments with drift.

6.5 Efficiency Concerns and Advantages of Unsupervised Detection

In our study, we chose not to include an efficiency analysis, because, in our experiments, we observe that a drift detector’s influence on the execution time is much lower than the paired classifier. Furthermore, supervised detection algorithms assume that the true labels are readily available [11, 12]; however, that is not always the case in real life. The acquisition of labeled data is usually difficult and it is much easier to collect unlabeled data [53]. As LD3 is an unsupervised detector, the lack of labeled data does not prevent its continuous concept drift detection whereas supervised algorithms may halt while waiting for labels.

6.6 Hyperparameter Analysis

In our analysis of the hyperparameters, we propose a method to set the standard deviation multiplier t . Furthermore, we conduct experiments to determine the effect of window size w and threshold for the number of anomaly L .

6.7 Setting the Standard Deviation Multiplier t

We developed a simple way to tune the t parameter of LD3. Figure 6.3 plots the distribution of past correlations within W_{corr} , when a drift is detected for four datasets as histogram plots.

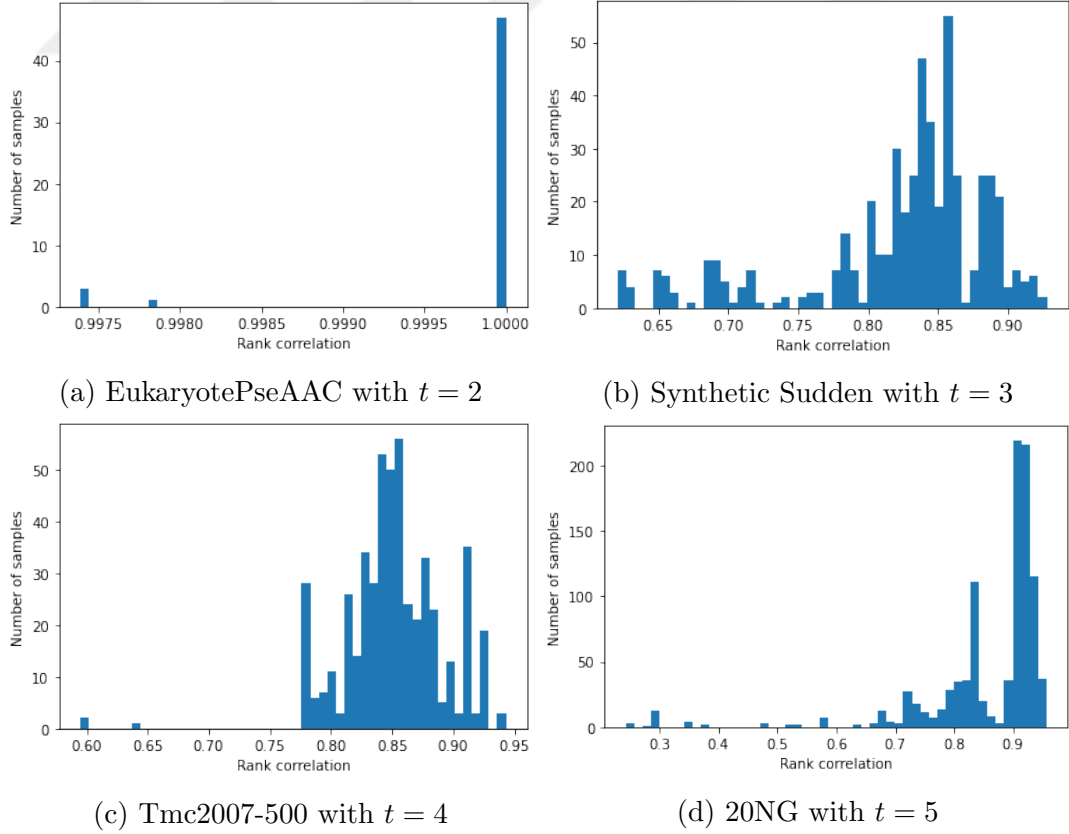


Figure 6.3: Distribution of the rank correlations within W_{corr} of four datasets after a drift is detected with different t values (σ multipliers). X-axis of the plots is given with different scales.

Since LD3 is an unsupervised detector, such plots could easily be obtained as a part of the test runs of a given data stream. In the figure, we see that with each subplot the range of correlation values expand; however the mean correlation does not decrease at the same rate, where they stay in $[0.80, 1]$. If the standard deviation of the correlations are very small, e.g, EukaryotePseAAC dataset in Figure 6.3.a, a low t is required such as $t = 2$. With this t value, if the current correlation value is smaller than $\mu - 2\sigma$, it lies outside of the 95% confidence interval and it should be considered as an anomaly [43]. For an accurate tuning of LD3, starting from $t = 2$, the confidence interval should be expanded with an increasing t , with guidance from the plots generated, similar to the examples in Figure 6.3. We used this method to determine the t values for our experiments on Table 6.1. Among the parameter values we used, $t = 4$ most frequently achieves the best effectiveness result. Therefore, we use this value as a default parameter.

6.8 Setting the Parameters Window Size w and Threshold for Number of Anomalies L

To investigate the effect of L and w parameters, we conducted experiments on the real-world datasets using $w = [50, 250, 500]$, $L = [0, 1, 2]$ values. Table 6.2 presents the results of our experiments. We found that $L = 0$ is a good default value since most of the best results are achieved with this value. This indicates that our concerns about outlier resiliency are not as significant as initially assumed. Though this is most likely not the case for every data stream as streams with extreme noise may be susceptible to outliers. We urge researchers who would utilize LD3 to better tune the L parameter if there are available resources such as labeled data and computational capacity.

Unlike the L parameter, our results show that setting the w parameter to a default value is not as straightforward. In Table 6.2 we observe that the datasets with low number of samples performs best with $w = 50$ and larger datasets with $w = 500$. This is related to the problem of sampling rate for data streams.

Table 6.2: Example-based accuracy of LD3 with varying parameter settings for each real-world dataset. The best scores are highlighted in bold. Datasets are listed in ascending order of the number of data items they contain.

Dataset	N	$w = 50$ $L = 0$	$w = 50$ $L = 1$	$w = 50$ $L = 2$	$w = 250$ $L = 0$	$w = 250$ $L = 1$	$w = 250$ $L = 2$	$w = 500$ $L = 0$	$w = 500$ $L = 1$	$w = 500$ $L = 2$
Birds	645	0.0618	0.0550	0.0550	0.0550	0.0550	0.0550	0.0550	0.0550	0.0550
PlantPseAAC	978	0.1383	0.1369	0.1369	0.1369	0.1369	0.1369	0.1369	0.1369	0.1369
Enron	1,702	0.1055	0.1089	0.1371	0.1403	0.1301	0.1371	0.1371	0.1371	0.1371
Yeast	2,417	0.3780	0.3644	0.3622	0.3724	0.3720	0.3718	0.3613	0.3613	0.3613
EukaryotePseAAC	7,766	0.1663	0.0727	0.0327	0.0627	0.0647	0.0658	0.0587	0.0578	0.0593
Ohsumed	13,930	0.0367	0.0370	0.0608	0.0444	0.0443	0.0443	0.0630	0.0630	0.0527
20NG	19,300	0.0667	0.0733	0.1469	0.1209	0.1219	0.1247	0.1516	0.1384	0.1497
Tmc2007-500	28,600	0.1485	0.1918	0.1555	0.1931	0.1934	0.1980	0.2029	0.2009	0.2007
Imdb	120,900	0.0717	0.0810	0.0797	0.0950	0.0954	0.1012	0.1129	0.1124	0.1140

Depending on the problem domain, the source of a stream can be sampled daily for stable streams such as music billboard charts and it can be sampled several times a minute for volatile data streams such as the price of a cryptocurrency. For this purpose it is possible to use a hyperparameter self-tuning scheme [54], to better adjust the classifier based on the current distribution. In our experiments, we observe that the value of the w parameter is highly dependant on the sampling rate of a stream and it should be set accordingly. For frequently sampled data streams we recommend $w = 500$ as a default value and $w = 50$ for infrequently sampled streams.

6.9 Analysis of Different Fusion Methods

In order to study the effects of different data fusion algorithms and to choose a suitable, default data fusion algorithm for LD3, we conducted experiments on the nine real-world datasets we used on effectiveness tests. We chose to exclude the synthetic data streams for these experiments as we aim to evaluate the general effectiveness of the tested data fusion algorithms in real-world datasets. In these experiments, in addition to reciprocal ranking, we used popular data fusion algorithms Borda Fuse, Condorcet’s Fuse [55] and Markov Chains Type 4 (MC4) [56]. The results are shown in Table 6.3. The tests are done by only changing the data fusion part of LD3.

The results show that reciprocal rank fusion provides better results (in six of

Table 6.3: Example-based accuracy results of LD3 using four different data fusion algorithms. Best results are highlighted in bold.

Dataset	Reciprocal	Borda	Condorcet	MC4
20NG	0.1616	0.1480	0.1497	0.1517
Birds	0.0992	0.0913	0.0984	0.0903
Enron	0.1403	0.1343	0.1316	0.1368
EukaryotePseAAC	0.1946	0.2096	0.0727	0.0685
Imdb	0.1140	0.1154	0.1110	0.1134
Ohsumed	0.0693	0.0568	0.0595	0.0594
PlantPseAAC	0.3198	0.3393	0.2397	0.2099
Tmc2007-500	0.2029	0.1954	0.1999	0.1997
Yeast	0.3765	0.3643	0.3577	0.3586

the nine cases). This is within our expectations as past studies such as [57] and [58] have shown similar observations in different problem domains.

6.10 Influence of Drift Types and Speed

Regarding drift types, since LD3 counts the number of co-occurrences to detect concept drift, incremental, sudden, and reoccurring drifts have the same structure as all of these types of drift result in changing rankings. As shown in Table 6.1, LD3 generally performs the same between different types of drifts and outperforms other detectors. Also, from the algorithm perspective, both incremental and gradual drifts show the same structure within the co-occurrence matrix. Therefore, we chose not to add an additional stream for gradual drift as the results were very similar.

In addition to drift types, the speed at which the drift occurs also influences a detector’s ability to find drifts. In Table 6.4, the accuracy results of the top four detectors are presented. The streams used are synthetic data streams with sudden and incremental drifts and the number after the underscore is the number of samples over which the drift occurs. We generate them using the Scikit-Multiflow framework [45] by the same method described in Section 5.1 where they have

Table 6.4: Example-based accuracy results of the top detectors with different drift types. Best results are highlighted in bold.

Data Stream	LD3	ADWIN	EDDM	RDDM	ND
Sudden_1	0.0788	0.0724	0.0744	0.0720	0.0659
Sudden_25	0.0767	0.0724	0.0741	0.0724	0.0659
Sudden_50	0.0776	0.0725	0.0737	0.0724	0.0659
Incremental_250	0.0769	0.0725	0.0733	0.0731	0.0659
Incremental_500	0.0788	0.0724	0.0787	0.0722	0.0659
Incremental_1000	0.0776	0.0717	0.0766	0.0713	0.0659

20,000 samples, 50 classes, and 200 features with drifts occurring at sample positions 4,000 and 10,000. Our findings indicate that regardless of the drift speed and type, LD3 still outperforms the baseline detectors, which demonstrate LD3’s robustness under different speed conditions.

Chapter 7

Conclusion

In this thesis, we present LD3, a novel unsupervised concept drift detection algorithm that exploits dynamic temporal dependencies between class labels in a multi-label classification environment. The algorithm is based on a new concept, label dependency ranking, which we introduce. We perform an extensive evaluation of LD3 against 14 prevalent drift detection algorithms using a Classifier Chain of Gaussian Naive Bayes classifiers. The experimental results show that LD3 provides the highest classifier accuracy: Without using true class labels, it statistically significantly outperforms several of the other drift detection algorithms that require such labels.

Our algorithm is able to accurately assist multi-label classifiers in the presence of a concept drift, which results in more robust classification models that are more adaptive to changing trends. In many streaming environments, true class labels required by supervised concept drift detection methods are likely to be unavailable due to several reasons. As LD3 is an unsupervised algorithm, this advantage illustrates the practical value of LD3.

In future work, we plan to examine the uses of LD3 in unsupervised classification and in environments that include concept evolution, i.e. the emergence of entirely new labels.

Bibliography

- [1] G. Tsoumakas and I. Katakis, “Multi-label classification: An overview,” *International Journal of Data Warehousing and Mining (IJDWM)*, vol. 3, no. 3, pp. 1–13, 2007.
- [2] X. Zheng, P. Li, Z. Chu, and X. Hu, “A survey on multi-label data stream classification,” *IEEE Access*, 2019.
- [3] M. Bahri, A. Bifet, J. Gama, H. M. Gomes, and S. Maniu, “Data stream analysis: Foundations, major tasks and tools,” *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, vol. 11, no. 3, p. e1405, 2021.
- [4] H. R. Bonab and F. Can, “GOOWE: Geometrically optimum and online-weighted ensemble classifier for evolving data streams,” *ACM Transactions on Knowledge Discovery from Data (TKDD)*, vol. 12, no. 2, pp. 1–33, 2018.
- [5] Z. Hammami, M. Sayed-Mouchaweh, W. Mouelhi, and L. Ben Said, “Neural networks for online learning of non-stationary data streams: a review and application for smart grids flexibility improvement,” *Artificial Intelligence Review*, vol. 53, pp. 6111–6154, 2020.
- [6] D. Xu, Y. Shi, I. W. Tsang, Y.-S. Ong, C. Gong, and X. Shen, “Survey on multi-output learning,” *IEEE Transactions on Neural Networks and Learning Systems*, 2019.
- [7] Y. Guo and S. Gu, “Multi-label classification using conditional dependency networks,” in *Twenty-Second International Joint Conference on Artificial Intelligence*, 2011.

- [8] J. Wang, Y. Yang, J. Mao, Z. Huang, C. Huang, and W. Xu, “Cnn-rnn: A unified framework for multi-label image classification,” in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, pp. 2285–2294, 2016.
- [9] M.-L. Zhang and K. Zhang, “Multi-label learning by exploiting label dependency,” in *Proceedings of the 16th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 999–1008, 2010.
- [10] X. Xue, W. Zhang, J. Zhang, B. Wu, J. Fan, and Y. Lu, “Correlative multi-label multi-instance image annotation,” in *2011 International Conference on Computer Vision*, pp. 651–658, IEEE, 2011.
- [11] T. S. Sethi and M. Kantardzic, “On the reliable detection of concept drift from streaming unlabeled data,” *Expert Systems with Applications*, vol. 82, pp. 77–99, 2017.
- [12] I. Žliobaite, “Change with delayed labeling: When is it detectable?,” in *2010 IEEE International Conference on Data Mining Workshops*, pp. 843–850, IEEE, 2010.
- [13] D. Wang and S. Zhang, “Unsupervised person re-identification via multi-label classification,” in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, June 2020.
- [14] A. Büyükçakır, H. Bonab, and F. Can, “A novel online stacked ensemble for multi-label stream classification,” in *Proceedings of the 27th ACM International Conference on Information and Knowledge Management*, pp. 1063–1072, 2018.
- [15] J. Gama, I. Žliobaitė, A. Bifet, M. Pechenizkiy, and A. Bouchachia, “A survey on concept drift adaptation,” *ACM Computing Surveys (CSUR)*, vol. 46, no. 4, pp. 1–37, 2014.
- [16] J. Lu, A. Liu, F. Dong, F. Gu, J. Gama, and G. Zhang, “Learning under concept drift: A review,” *IEEE Transactions on Knowledge and Data Engineering*, vol. 31, no. 12, pp. 2346–2363, 2018.

- [17] Ö. Gözüaık and F. Can, “Concept learning using one-class classifiers for implicit drift detection in evolving data streams,” *Artificial Intelligence Review*, vol. 54, no. 5, pp. 3725–3747, 2021.
- [18] M. Roseberry and A. Cano, “Multi-label knn classifier with self adjusting memory for drifting data streams,” in *Second International Workshop on Learning with Imbalanced Domains: Theory and Applications*, pp. 23–37, PMLR, 2018.
- [19] A. Bifet and R. Gavaldà, “Learning from time-changing data with adaptive windowing,” in *Proceedings of the 2007 SIAM International Conference on Data Mining*, pp. 443–448, SIAM, 2007.
- [20] J. Gama, P. Medas, G. Castillo, and P. Rodrigues, “Learning with drift detection,” in *Brazilian Symposium on Artificial Intelligence*, pp. 286–295, Springer, 2004.
- [21] M. Baena-García, J. del Campo-Ávila, R. Fidalgo, A. Bifet, R. Gavaldà, and R. Morales-Bueno, “Early drift detection method,” in *Fourth International Workshop on Knowledge Discovery from Data Streams*, vol. 6, pp. 77–86, 2006.
- [22] A. Pesaranghader and H. L. Viktor, “Fast hoeffding drift detection method for evolving data streams,” in *Joint European Conference on Machine Learning and Knowledge Discovery in Databases*, pp. 96–111, Springer, 2016.
- [23] A. Pesaranghader, H. Viktor, and E. Paquet, “Reservoir of diverse adaptive learners and stacking fast hoeffding drift detection methods for evolving data streams,” *Machine Learning*, vol. 107, no. 11, pp. 1711–1743, 2018.
- [24] I. Frías-Blanco, J. del Campo-Ávila, G. Ramos-Jimenez, R. Morales-Bueno, A. Ortiz-Díaz, and Y. Caballero-Mota, “Online and non-parametric drift detection methods based on hoeffding’s bounds,” *IEEE Transactions on Knowledge and Data Engineering*, vol. 27, no. 3, pp. 810–823, 2014.
- [25] C. Raab, M. Heusinger, and F.-M. Schleif, “Reactive soft prototype computing for concept drift streams,” *Neurocomputing*, 2020.

- [26] A. Pesaranghader, H. L. Viktor, and E. Paquet, “Mediarmid drift detection methods for evolving data streams,” in *2018 International Joint Conference on Neural Networks (IJCNN)*, pp. 1–9, IEEE, 2018.
- [27] R. S. Barros, D. R. Cabral, P. M. Gonçalves Jr, and S. G. Santos, “Rddm: Reactive drift detection method,” *Expert Systems with Applications*, vol. 90, pp. 344–355, 2017.
- [28] R. Pears, S. Sakthithasan, and Y. S. Koh, “Detecting concept change in dynamic data streams,” *Machine Learning*, vol. 97, no. 3, pp. 259–293, 2014.
- [29] W. Hoeffding, “Probability inequalities for sums of bounded random variables,” *Journal of the American Statistical Association*, vol. 58, no. 301, pp. 13–30, 1963.
- [30] Z. Shi, Y. Wen, C. Feng, and H. Zhao, “Drift detection for multi-label data streams based on label grouping and entropy,” in *2014 IEEE International Conference on Data Mining Workshop*, pp. 724–731, IEEE, 2014.
- [31] P. Wang, N. Jin, and G. Fehringer, “Concept drift detection with false positive rate for multi-label classification in iot data stream,” in *2020 International Conference on UK-China Emerging Technologies (UCET)*, pp. 1–4, IEEE, 2020.
- [32] A. S. Iwashita and J. P. Papa, “An overview on concept drift learning,” *Ieee Access*, vol. 7, pp. 1532–1547, 2018.
- [33] R. N. Gemaque, A. F. J. Costa, R. Giusti, and E. M. Dos Santos, “An overview of unsupervised drift detection methods,” *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, vol. 10, no. 6, p. e1381, 2020.
- [34] D. M. dos Reis, P. Flach, S. Matwin, and G. Batista, “Fast unsupervised online drift detection using incremental kolmogorov-smirnov test,” in *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 1545–1554, 2016.

- [35] Y. S. Koh, “Cd-tds: Change detection in transactional data streams for frequent pattern mining,” in *2016 International Joint Conference on Neural Networks (IJCNN)*, pp. 1554–1561, IEEE, 2016.
- [36] R. F. de Mello, Y. Vaz, C. H. Grossi, and A. Bifet, “On learning guarantees to unsupervised concept drift detection on data streams,” *Expert Systems with Applications*, vol. 117, pp. 90–102, 2019.
- [37] F. Pinagé, E. M. dos Santos, and J. Gama, “A drift detection method based on dynamic classifier selection,” *Data Mining and Knowledge Discovery*, vol. 34, no. 1, pp. 50–74, 2020.
- [38] Ö. Gözüaık, A. Büyükakır, H. Bonab, and F. Can, “Unsupervised concept drift detection with a discriminative classifier,” in *Proceedings of the 28th ACM International Conference on Information and Knowledge Management*, pp. 2365–2368, 2019.
- [39] W. Sałabun and K. Urbaniak, “A new coefficient of rankings similarity in decision-making problems,” in *International Conference on Computational Science*, pp. 632–645, Springer, 2020.
- [40] M. G. Kendall, “A new measure of rank correlation,” *Biometrika*, vol. 30, no. 1/2, pp. 81–93, 1938.
- [41] C. Spearman, “The proof and measurement of association between two things,” *The American Journal of Psychology*, vol. 100, no. 3/4, pp. 441–471, 1987.
- [42] S. Vigna, “A weighted correlation index for rankings with ties,” in *Proceedings of the 24th International Conference on World Wide Web*, pp. 1166–1176, 2015.
- [43] F. Pukelsheim, “The three sigma rule,” *The American Statistician*, vol. 48, no. 2, pp. 88–91, 1994.
- [44] J. Read, P. Reutemann, B. Pfahringer, and G. Holmes, “Meka: a multi-label/multi-target extension to weka,” *The Journal of Machine Learning Research*, vol. 17, no. 1, pp. 667–671, 2016.

- [45] J. Montiel, J. Read, A. Bifet, and T. Abdesslem, “Scikit-multiflow: A multi-output streaming framework,” *The Journal of Machine Learning Research*, vol. 19, no. 1, pp. 2915–2914, 2018.
- [46] J. Read, B. Pfahringer, G. Holmes, and E. Frank, “Classifier chains for multi-label classification,” *Machine Learning*, vol. 85, no. 3, p. 333, 2011.
- [47] G. John, “Estimating continuous distributions in bayesian classifiers,” in *Proc. 11th Conference on Uncertainty in Artificial Intelligence*, 1995.
- [48] J. T. Pintas, L. A. Fernandes, and A. C. B. Garcia, “Feature selection methods for text classification: a systematic literature review,” *Artificial Intelligence Review*, pp. 1–52, 2021.
- [49] J. Gama, R. Sebastião, and P. P. Rodrigues, “Issues in evaluation of stream learning algorithms,” in *Proceedings of the 15th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*, pp. 329–338, 2009.
- [50] J. Nam, E. L. Mencía, H. J. Kim, and J. Fürnkranz, “Maximizing subset accuracy with recurrent neural networks in multi-label classification,” in *Proceedings of the 31st International Conference on Neural Information Processing Systems*, pp. 5419–5429, 2017.
- [51] M.-L. Zhang and Z.-H. Zhou, “A review on multi-label learning algorithms,” *IEEE Transactions on Knowledge and Data Engineering*, vol. 26, no. 8, pp. 1819–1837, 2013.
- [52] J. Demšar, “Statistical comparisons of classifiers over multiple data sets,” *Journal of Machine Learning Research*, vol. 7, no. Jan, pp. 1–30, 2006.
- [53] L. Subhashini, Y. Li, J. Zhang, A. S. Atukorale, and Y. Wu, “Mining and classifying customer reviews: a survey,” *Artificial Intelligence Review*, pp. 1–47, 2021.
- [54] B. Veloso, J. Gama, B. Malheiro, and J. Vinagre, “Hyperparameter self-tuning for data streams,” *Information Fusion*, vol. 76, pp. 75–86, 2021.

- [55] R. Nuray and F. Can, “Automatic ranking of information retrieval systems using data fusion,” *Information Processing & Management*, vol. 42, no. 3, pp. 595–614, 2006.
- [56] C. Dwork, R. Kumar, M. Naor, and D. Sivakumar, “Rank aggregation methods for the web,” in *Proceedings of the 10th International Conference on World Wide Web*, pp. 613–622, 2001.
- [57] G. V. Cormack, C. L. Clarke, and S. Buettcher, “Reciprocal rank fusion outperforms condorcet and individual rank learning methods,” in *Proceedings of the 32nd International ACM SIGIR Conference on Research and Development in Information Retrieval*, pp. 758–759, 2009.
- [58] D. C. G. Pedronette and R. d. S. Torres, “Unsupervised effectiveness estimation for image retrieval using reciprocal rank information,” in *2015 28th SIB-GRAPI Conference on Graphics, Patterns and Images*, pp. 321–328, IEEE, 2015.