

Building adaptive data mining models on streaming data in real-time

Article

Accepted Version

Stahl, F. ORCID: <https://orcid.org/0000-0002-4860-0203> and Badii, A. (2020) Building adaptive data mining models on streaming data in real-time. Expert Update, 20 (2). ISSN 1465-4091 doi: 10.7148/2018-0008 Available at <https://centaur.reading.ac.uk/88982/>

It is advisable to refer to the publisher's version if you intend to cite from the work. See [Guidance on citing](#).

To link to this article DOI: <http://dx.doi.org/10.7148/2018-0008>

Publisher: BCS Specialist Group on Artificial Intelligence

All outputs in CentAUR are protected by Intellectual Property Rights law, including copyright law. Copyright and IPR is retained by the creators or other copyright holders. Terms and conditions for use of this material are defined in the [End User Agreement](#).

www.reading.ac.uk/centaur

CentAUR

Central Archive at the University of Reading

Reading's research outputs online

Building Adaptive Data Mining Models on Streaming Data in Real-Time:

Frederic Stahl^{1,2}, Atta Badii²

German Research Center for Artificial Intelligence GmbH (DFKI), DFKI
Laboratory Niedersachsen, Marine Perception¹

Department of Computer Science, University of Reading²

Abstract - Advances in hardware and software, over the past two decades have enabled the capturing, recording and processing of potentially large and infinite streaming data. The field of research in Data Stream Mining (DSM) has emerged to respond to the challenges and opportunities of developing the required analytics to unlock valuable knowledge. Thus DSM is focused on building Data Mining models, workflows and algorithms enabling the efficient and effective analysis of such streaming data at a large scale- the so-called “Big Data”. Examples of application areas of Data Stream Mining techniques include real-time telecommunication data, telemetric data from large industrial plants, credit card transactions, cyber security threat modelling, social media data, etc. For some applications it is acceptable to provide data processing, modelling and analysis in batch mode using the traditional Data Mining approaches. However, for other application, particularly where continuous monitoring and contingent response are required, the model building and analytics have to take place in real-time as soon as new data becomes available i.e. to accommodate infinite streams and fast changing concepts in the data. This article highlights some of the key concepts and emergent techniques in DSM as presented in the authors’ recent publications as also outlined in a talk given at the UK Symposium on Knowledge Discovery from Data in London on May 24th, 2019 discussing the challenges, opportunities and innovative solutions in Data Stream Mining.

1. Introduction

Data is generated at an unprecedented rate and scale; this is often referred to as the Velocity and Volume of Big Data. The *domo* website [1] regularly publishes an infographic that sets out examples of how much data certain internet applications produce every minute of the day. For instance, the infographic states for the year 2017 that on average, every minute, around 176000 Skype calls were made, 4.3 Million YouTube videos uploaded and 3.9 Million google searches performed. IBM estimated that worldwide, with the passing of every day 2.5 Quintillion more data were generated. It is difficult to imagine how much 2.5 Quintillion bytes can store [2]. To visualise this consider storing the data on DVD or Blue-ray discs. Even although these media are hardly used nowadays, we still have a perception of how much data could fit on a disc due to its standard storage capacity. A DVD can store 4.7 GB of data, a Blue-ray disc 25GB of data. It would take a 100 Million Blue-ray discs to store 2.5 Quintillion of data and at 1.22 mm thickness per disc this would stack up to about 100km -a distance of roughly from London to the Isle of Wight. On DVD discs this would take about 530 Million discs which at 1.22 mm thickness would stack to about 630km high - a distance from London to Bremen (Germany). Even although this is an estimate and its reliability could be challenged, it can nevertheless be agreed that even if a small fraction of the 2.5 Quintilian data is generated every day, it is still a very large amount of data. In order to make sense of this real-time continuous stream of data, continuously running analytics methods are required.

This article is structured as follows: Section 2 defines the concept of Data Streams; Section 3 introduces challenges and barriers in analysing such real-time streams and Section 4 highlights some general approaches and algorithms for data stream analytics. This is followed in Section 5 by some recent and ongoing research to overcome the barriers; the concluding remarks are provided in Section 6.

2. Real-time Streaming Data Analytics Challenges and Barriers

Advances in data acquisition hardware and the emergence of applications that process continuous flows of data have led to the (Big Data) stream phenomenon. A data stream is a rapid flow of data which demands analytics such that is challenging for the state-of-the-art processing techniques and communication infrastructure. There are some fundamental differences between static and streaming data and these are: (i) static data is typically historic information, whereas streaming data is often a live real-time feed of data; (ii) static data is randomly accessible, whereas streaming data can be accessed only sequentially, one data instance at a time, (iii) static data is located in secondary storage, whereas streaming data needs to be processed in memory to increase efficiency and enable real-time analytics; iv) for static data the time to build analytics models is often not critical, whereas in streaming environments a low processing latency is often essential to enable real-time applications finally (v) when using static data for model building, one can assume the data is already pre-processed, or at least sufficient time is available to clean the data, whereas in a streaming environment one must assume the data includes inaccurate and raw data elements that are yet to be checked prior to model building. This never-ending stream of raw data is often referred to as the Data Tsunami.

2.1. Concept Drift

An important challenge in data streams analytics is concept drift, where the pattern encoded in the stream changes over time. Concept drift exists in real life problems, such as seasonal weather changes, stock market downturns and rallies etc. Concept drifts are typically unforeseen and unpredictable. There are different kinds of concept drifts, i.e. gradual drift, where previous and new patterns encoded in the stream overlay for a short period of time, sudden drift, where the pattern in the stream shifts abruptly, the new patterns appears instantly and the old pattern disappears at the same time. Then there are evolving or incremental streams, where the pattern does not stay stable and changes constantly, and, re-occurring drifts, where previously seen and discontinued patterns re-occur. Of course, there are also combinations of these types of streams. A data mining model should always reflect the current concept and thus needs to adapt quickly.

2.2. Challenges and Barriers

Table 1 summarises the challenges in mining real-time data streams. Each challenge is opposed by a barrier which needs to be overcome to address the challenge fully. One of these challenges is concept drift which is the most fundamental challenge in analysing data streams as otherwise data streams could be treated as a batch problem and models could be simply built from a sample of the streamed data, as there would be no need for adaptation in real-time. However, the remaining challenges are also important. Next, Section 3 will summarise some general approaches to analysis of data streams, in particular with respect to challenge 2 (concept drift).

Table 1: Challenges and Barriers in analysing data streams.

No.	Challenges	Barriers
1	Data generated at a fast rate (Velocity), at potentially large and unknown quantities (Volume)	Limited scalable (parallel) real-time data stream mining algorithms available
2	Concept Drifts (Changes of pattern encoded in in the data over time)	Different and changing types of concept drift
3	Real-time data model building workflows optimisation based on streaming data	Lack of customizable pre-processing techniques
4	Multi-modality of data sources (Text, Video/images, (un)structured)	Different time stamps but co-occurring data items
5	Class label sparsity: need to adapt the associated predictive models	Supervised algorithms not applicable in many cases
6	High throughput real-time processing	Limited parallel real-time architectures

3. Building Data Mining Models from Data Streams

A popular technique for data stream mining, which is in some form often used within adaptive Data Stream Mining algorithms or in combination with concept drift detection methods to enable the execution of existing batch Data Mining algorithms on streaming data, is windowing [3,4]. A frequently used technique here, to name an example, is the ADaptive sliding WINDOW method (ADWIN) [5]. ADWIN changes the window size based on observed changes: if there are changes then the window shrinks, if there are no changes it increases. Concept drift detection methods are typically used in combination with batch data mining algorithms and sliding window approaches. For example, the sequential analysis technique CUMulative SUM \cite{page1954continuous} detects a drift when the mean of incoming data deviates significantly. The Exponentially Weighted Moving Average (EWMA) uses charts to monitor the mean of misclassification rates of a classifier, and gives less weight to older data instances and greater weight to more recent data instances. In general, different methods are suitable for detection and adaptation to different kinds of concept drift, e.g. DDM [6] is suitable for sudden drifts and EDDM for gradual drifts [7]. Challenges here are not to mistake outliers and the noise of uncleaned raw data for concept drift.

The following is a brief overview of some adaptive predictive and descriptive data stream mining algorithms that naturally incorporate windowing/data buffering techniques and adaptation to concept drift. All these algorithms have one requirement in common, which is that they must only require a single pass through the training data and have a small memory footprint.

3.1. Adaptive predictive algorithms

The development of adaptive classification techniques has received considerable attention in the recent years. These range from adaptive tree-based methods, i.e. the popular Hoeffding Tree [8] algorithm, to rule-based methods such as VFDR [9] and G-eRules [10] to less expressive (but often very fast) methods such as MC-NN [REF Mark Tennant]. A drawback of most data stream classification methods arises from the lack for a “cold-start” capability which means that they need ground truth information. This is often unrealistic as most applications will not be able to provide this fast feedback about the correctness of the classification and thus these supervised methods are often not suitable for real-life applications (see Barrier 5 in Table 1). However, some applications can provide indirect class information, e.g. for the classification of twitter posts into topics categories, one could use hashtags instead of the actual categories, to verify and adapt the classification model.

3.2. Descriptive adaptive cluster analysis algorithms

Cluster analysis is the method of grouping data with similar characteristics, with the aim of having a high degree of similarity within a cluster, but a high degree of dissimilarity between clusters (optimal sensitivity and selectivity based on inter-intra class distances). Several such algorithms exist for streaming data environments. Most of these are built on the idea of Micro-Cluster structures; these are statistical summaries and typically as many as computationally viable are held in memory and adapted responsive to newly arrived data, by changing their boundaries and locations. Accordingly, no raw data needs to be kept in memory, this is known as the online component of this type of algorithm. Subsequently, analytical descriptive clusters can be built offline on-demand based on such statistical summaries instead of using infinite raw data. A notable development here is CluStream [11] which can delete obsolete Micro-Clusters, browse through historical models, re-visiting historical models and “forgetting” obsolete concepts. . A shortcoming of CluStream is that only circular clusters can be built. DenStream algorithm [12] addresses this through the introduction of ‘dense’ Micro-Clusters to summarise clusters using an arbitrary shape. Another Micro-Cluster-based algorithm is ClusTree [13], which incrementally learns a dendrogram tree structure from Micro-Clusters.

4. Current Research in the Field of Data Stream Mining

Section 3 outlined some existing works developed over the last 20 years aiming to address the data stream analytics challenges. However, referring to the challenges and barriers as highlighted in Table 1, these are largely limited to overcoming Barrier 2. Here we summarise some recent work based on Micro-Clusters approach to address and overcome some of the remaining barriers included in Table 1.

4.1. Scalable (parallel) real-time data stream mining algorithms

Data Stream Mining algorithms are generally designed to be fast and scalable due to the essential requirement that they have to meet, namely the capability to adapt by a single pass through the data. However, this capability is often limited by the serial segments of the processing and little work has been carried out in the development of end-to-end-parallel DSM algorithms. In MC-NN [14] the authors proposed a parallel KNN-like algorithm. Here class-labelled Micro-Clusters are distributed over real-time setup of a MapReduce architecture and each Mapper (node in cluster) is to provide votes based on the Micro-Cluster distance with respect to newly arriving unlabelled data instances. MC-NN has performed well in terms of speed and accuracy in comparison with some of the fastest data stream classifiers. MC-NN has excelled when confronted with continuously changing data streams and has shown a very low latency to adapt to concept drift.

4.2. Parallel real-time architectures

The works summarised in Section 4.1 have also developed suitable software architectures to enable the parallel real-time execution of the MC-NN and similar algorithms. This architecture uses a SAMZA [15] layer and Kafka [16] messaging system to achieve real-time instance-by-instance processing using the MapReduce framework over Hadoop [17] nodes. Here SAMZA enables synchronised real-time instance-by-instance processing of external messages over the Hadoop MapReduce framework. In this architecture SAMZA makes use of the low latency and high throughput publish and subscribe system of Kafka. With this parallel architecture MC-NN showed considerable speed-up over its serial version using up to 12 Hadoop nodes.

4.3. Real-time pre-processing techniques

With respect to real-time pre-processing, very little work has been conducted and in practise often pre-processing techniques for batch data are applied. Whereas this may yield acceptable results in many cases, it does not take the ever-changing nature of data streams into account. For example, frequently used min-max normalisation can be applied in real-time but relies on stable min and max values. A recent work based on Micro-Clusters [18] aims to develop real-time feature selection techniques. In practise feature selection for data stream mining applications is applied once and then

the chosen data stream mining algorithms continue to use these features. However, it is possible that the relevance of features change during concept drift and the original feature selection may not be no longer optimal. This work tracks Micro-Cluster rate of change (velocity) and the trajectory of Micro-Cluster movement in order to assess which features have been involved in a concept drift and to which extent. Based on this a real-time re-selection of the feature system was realised yielding improved accuracy over time in comparison with other methods which kept to the original feature selections.

5. Discussion and Conclusions

This article has defined 6 challenges and barriers to achieving true instance-by-instance data stream analytics in real-time (see Table 1). Subsequent sections have then outlined how established algorithms and current research aims to overcome these barriers, notably through Micro-Cluster-based approaches which have resulted in a variety of different kinds of analytics algorithms. However, challenges 4, 5 and 6 remain largely unaddressed. With respect to challenge 4, multi-modality of data sources, different time-stamped data from diverse heterogeneous data sources make it difficult to compose data frames as needed by data stream mining algorithms. Here some ad hoc and application-specific solutions have been implemented. However, a generic approach adaptable to a range of problems remains to be explored. With respect to challenge 5, class label sparsity, many predictive classification algorithms for data streams are supervised and thus require timely and frequent feedback about the ground truth for their predictions. Here some works have been prototyped based on ensemble approaches [19], however, as yet, there exists no viable solution for concrete applications. With respect to challenge 6, some ecosystems for parallel processing of streaming data in real-time exist, such as the Apache Flink project [20], however, their focus lies mainly on the parallel and in-memory processing of data in micro batches, which does not take full advantage of existing instance-by-instance data stream mining algorithms. Further, such systems often lack a full end-to-end architecture and focus more on the data processing efficiency. Here an end-to-end architecture should comprise components ranging from data ingestion and provenance-plausibility-aware data frame composition through to analytics algorithms / workflows and end-user dashboards to control, configure and examine the end-to-end data stream analytics pipeline.

References

1. "Domo" [Online]. <https://www.domo.com/learn/data-never-sleeps-6>
2. "IBM" [Online]. <https://www.ibm.com/blogs/insights-on-business/consumer-products/2-5-quintillion-bytes-of-data-created-every-day-how-does-cpg-retail-manage-it/>
3. Babcock, B., Babu, S., Datar, M., Motwani, R., & Widom, J. (2002, June). Models and issues in data stream systems. In *Proceedings of the twenty-first ACM SIGMOD-SIGACT-SIGART symposium on Principles of database systems* (pp. 1-16). ACM.
4. Datar, M., Gionis, A., Indyk, P., & Motwani, R. (2002). Maintaining stream statistics over sliding windows. *SIAM journal on computing*, 31(6), 1794-1813.
5. Bifet, A., & Gavalda, R. (2007, April). Learning from time-changing data with adaptive windowing. In *Proceedings of the 2007 SIAM international conference on data mining* (pp. 443-448). Society for Industrial and Applied Mathematics.
6. Gama, J., Medas, P., Castillo, G., & Rodrigues, P. (2004, September). Learning with drift detection. In *Brazilian symposium on artificial intelligence* (pp. 286-295). Springer, Berlin, Heidelberg.
7. Baena-Garcia, M., del Campo-Ávila, J., Fidalgo, R., Bifet, A., Gavalda, R., & Morales-Bueno, R. (2006, September). Early drift detection method. In *Fourth international workshop on knowledge discovery from data streams* (Vol. 6, pp. 77-86).
8. Hulten, G., Spencer, L., & Domingos, P. (2001, August). Mining time-changing data streams. In *Proceedings of the seventh ACM SIGKDD international conference on Knowledge discovery and data mining* (pp. 97-106). ACM.

9. Gama, J., & Kosina, P. (2011, June). Learning decision rules from data streams. In Twenty-Second International Joint Conference on Artificial Intelligence.
10. Le, T., Stahl, F., Gaber, M. M., Gomes, J. B., & Di Fatta, G. (2017). On expressiveness and uncertainty awareness in rule-based classification for data streams. *Neurocomputing*, 265, 127-141.
11. Aggarwal, C. C., Han, J., Wang, J., & Yu, P. S. (2003, September). A framework for clustering evolving data streams. In *Proceedings of the 29th international conference on Very large data bases-Volume 29* (pp. 81-92). VLDB Endowment.
12. Cao, F., Estert, M., Qian, W., & Zhou, A. (2006, April). Density-based clustering over an evolving data stream with noise. In *Proceedings of the 2006 SIAM international conference on data mining* (pp. 328-339). Society for Industrial and Applied Mathematics.
13. Kranen, P., Assent, I., Baldauf, C., & Seidl, T. (2011). The ClusTree: indexing micro-clusters for anytime stream mining. *Knowledge and information systems*, 29(2), 249-272.
14. Tennant, M., Stahl, F., Rana, O., & Gomes, J. B. (2017). Scalable real-time classification of data streams with concept drift. *Future Generation Computer Systems*, 75, 187-199.
15. "Samza" [Online]. Available: <http://samza.apache.org/>
16. "Kafka" [Online]. <https://kafka.apache.org/>
17. "Hadoop." [Online]. <https://hadoop.apache.org/>
18. Hammoodi, M. S., Stahl, F. and Badii, A., (2018) Real-time feature selection technique with concept drift detection using Adaptive Micro-Clusters for data stream mining. *Knowledge-Based Systems*, Elsevier, 75, pp. 205-239, ISSN 0950-7051 doi: 10.1016/j.knosys.2018.08.007
19. Idrees, M. M., Minku, L. L., Stahl, F. and Badii, A. (2019) A heterogeneous online learning ensemble for non-stationary environments. *Knowledge-Based Systems*. ISSN 0950-7051 doi: <https://doi.org/10.1016/j.knosys.2019.104983> (In Press)
20. "Flink" [Online]. <https://flink.apache.org/>