

An Aprori Algorithm in Distributed Data Mining System

Dr C.Sunil Kumar¹

1

Received: 13 December 2012 Accepted: 31 December 2012 Published: 15 January 2013

Abstract

Many existing data mining (DM) tasks can be proficient effectively only in a distributed condition. The ground of distributed data mining (DDM) has therefore gained growing weightage in the preceding decades. The Apriori algorithm (AA) has appeared as one of the greatest Association Rule mining (ARM) algorithms. It also provides the foundation algorithm in majority of parallel algorithms (PAs). The size and elevated dimensionality of datasets characteristically existing as a key to difficulty of AR finding, makes it perfect difficulty for solving on numerous processors in parallel. The main causes are the computer memory and central processing unit pace constraints looked by single workstations. This paper is based on an Optimized Distributed AR mining algorithm for biologically distributed information is used in similar and distributed surroundings so that it decreases communication costs.

Index terms— association rules (ars), apriori algorithm (aa), distributed data mining (ddm), xml data, parallel.

1 Introduction

RM has turn out to be one of the hub DM tasks and has attracted marvelous interest among DM investigators. ARM is an unsupervised DM method which works on variable length data, and produces understandable results. There are two foremost approaches for utilizing numerous workstations that have appeared in distributed computer memory in which each CPU has a confidential memory; and public memory in which all CPUs access universal memory [4]. Collective memory planning has many gorgeous assets. Each CPU has straight and identical right to use memory in the computer system. Equivalent applications are easy to execute on such a distributed system. In allocated memory design each CPU has its own restricted memory that can simply right to use directly by that CPU [10]. For a CPU to have contact with facts in the restricted memory of another CPU a replica of the preferred data ingredient must be sent from one CPU to the other throughout message passing. XML information is used with the optimized distributed association rule mining (ODAM) algorithm. A similar application could be divided into numeral of jobs and implemented in parallel on different CPUs in the system [2,9]. Though the performance of a similar function on a distributed system is mainly depe-ndent on the distribution of the jobs contains the applic ation onto the obtainable CPUs in the system.

Modern associations are biologically dispersed. Classically, each location locally stores its ever growing amount of every day data. Using centralized DM to find out useful patterns in such institutions' data isn't always practicable because integration of data sets from dissimilar locations into a centralized location earns enormous communication system costs. Information from these institutions' is not only spread to a variety of sites but also vertically incoherent, making it complex if not unfeasible to merge them in an essential site.

Distributed DM therefore emerged as vigorous subarea of DM investigation. In this paper an ODAM Algorithm is used for executing the mining procedure.

2 II. Association Rule Mining Algorithms

An AR is a rule which implies certain association relationships among a set of objects in a database.

Given a set of transactions, where each transaction is a set of items, an AR is an expression of the form $X \rightarrow Y$, where X and Y are sets of items. The intuitive meaning of such a rule is that transactions of the database which contain X tend to contain Y [1].

3 a) Apriori Algorithm

Apriori has been urbanized for rule mining in large business databases by Quest project team of IBMs. An item set (IS) is a non-empty set of things.

4 They have decayed the difficulty of mining ARs into two (2) pieces

? Find all groupings of items that have business support above smallest support. Call those groupings frequent ISs [5]. ? Use the recurrent IS to produce the preferred rules.

The universal idea is that if, say, EFGH and EF are recurrent ISs, then we can determine if the rule $EF \rightarrow GH$ holds by computing the ratio $R = \text{support}(EFGH) / \text{support}(EF)$. The rule holds only if $R \geq \text{min assurance}$. The algorithm is extremely scalable [7]. The AA used in Quest to find all recurrent ISs is specified below.

Procedure AprioriAlg () begin $M1 := \{\text{frequent 1-itemsets}\}$; for ($k := 2$; $Mk-1 \neq \emptyset$; $k++$) do { $Nk = \text{Apriori-gen}(Mk-1)$; // new candidates for all transactions t in the dataset do {for all candidates $c \in Nk$ contained in t do c : count++} In the first pass, the algorithm simply counts item occurrences to determine the frequent 1-itemsets. A succeeding pass, say pass k , consists of two phases. First, the recurrent ISs $Mk-1$ found in the $(k-1)$ th pass are used to make the candidate ISs Nk , using the Apriori-gen () function. This task first joins $Mk-1$ with $Mk-1$, the joining condition being that the lexicographically ordered first $k-2$ items are the same. Next, it deletes all those ISs from the join result that have some $(k-1)$ subset that is not in $Mk-1$ yielding Nk . The algorithm now examines the database. For each deal, it concludes which of the candidates in Nk are limited in the deal using a hash-tree data structure and increases the count of those candidates [8], at the end of the pass, Ck is observed to decide which of the candidates' frequent, yielding Mk . The algorithm quits when Mk becomes vacant.

5 III. Optimized Distributed Mining Algorithm

The presentation of Apriori agent rule mining algorithms humiliates for diverse causes. It requires 'N' number of database examines to produce a common n -IS. In addition, it doesn't differentiate operations in the data set with identical ISs if that data set is not burdened into the memory.

Consequently, it needlessly occupies resources for frequently generating ISs from such matching transactions. For e.g., if a data set has 10 matching transactions, the AA not only indicates the same candidate ISs 10 times but also informs.

The support counts for those candidate ISs 10 times for every iteration. Furthermore, openly loading a unprocessed data set into the memory won't find an significant number of equal transactions because each business of a raw data set contains both recurrent and non-recurrent items. To overcome these difficulties, candidate support counts can't be supported from the raw data set following the first pass. This method not only decreases the average business length but also decreases the data set size considerably. The number of items in the data set might be huge, but only a few will gratify the support threshold (TH).

Consider ODA eliminates all globally infrequent litemsets from every transaction and inserts them into the main memory; it reduces the transaction size and finds more alike transactions. This is because the data set initially contains both frequent and rare items. However, total transactions could surpass the main memory limit.

ODAM removes rare items and inserts each transaction into the main memory. While inserting the transactions, it checks whether they are already in memory. If yes, it increases that transaction's counter by one. Otherwise, it inserts that transaction into the main memory with a count equal to one. Finally, it writes all main-memory entries for this separation into a temporary file. This process continues for all other divisions.

IV.

6 Pada Rule with Xml Data

Parallelism is predictable to relieve current ARM methods from sequential blockages, providing the ability to scale to enormous datasets and improving the response time. The parallel and Distributed Association rule (PADA) design space spans three main components including the hardware platform, the kind of parallelism broken and the load balancing strategy used [8]. Shared memory architecture has all the processors access common memory. Each processor has direct and equal access to all the memory in the system.

Parallel programs are easy to execute on such a system. The data warehouse (DW) is partitioned among 'P' processors logically. Each processor works on its local partition of the database but performs the same computation of counting support. Dynamic load balancing seeks to address this issue by balancing the load and reassigning the loads to the lighter ones. The development of distributed rule mining is a challenging and vital task, since it requires knowledge of all the data stored at different locations and the ability to combine partial results from individual databases into a single result.

The AR from XML data with a sample XML document is considered. For example, the set of transactions are identified by the tag <transactions> and each transaction in the transactions set is identified by the tag <transaction>. The set of items in each transaction is: Transaction document (transactions.xml) is identified by the tag <items> and an item is identified by the tag <item>. Consider the problem of mining all ARs among items that emerge in the transactions document. With the understanding of traditional AR mining is expected to obtain the large item sets document and ARs document from the source document.

Let the minimum support (minsupp) = 35% and minimum confidence (minconfi) = 99%.

V.

7 Performance Assessment

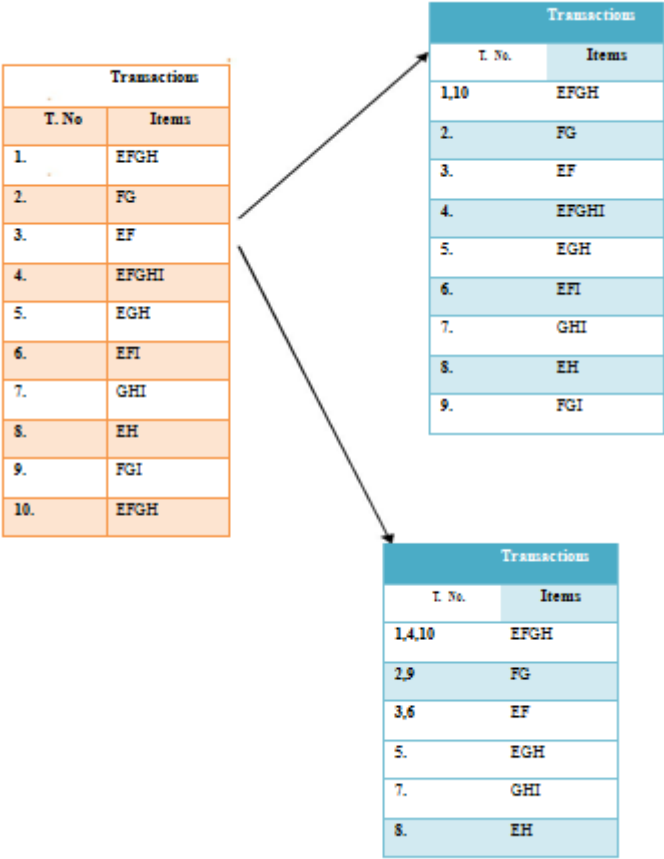
The number of messages that ODAM exchanges among various locations to generate the globally frequent item sets in a distributed environment, the original data set is partitioned into five partitions. To decrease the dependency among dissimilar partitions, each one contains only 25 percent of the original data set's transactions. So, the number of identical transactions among different partitions is very low.

ODAM provides a proficient method for generating ARs from different datasets, distributed among various locations.

The datasets are generated arbitrarily depending on the number of different items, the maximum number of items in each transaction and the number of transactions. The performance of the XQuery implementation is dependent on the number of large item sets found and the size of the dataset as shown in the The running time for dataset-1 with minimum support 20% is much higher than the running time of dataset-2 and dataset-3, since the number of large item sets found for dataset-1 is about 2 times more than the other datasets. The Response time of the parallel and distributed data mining task on XML data is carried out by the time taken for communication, computation cost involved [6]. Communication time is largely dependent on the DDM operational model and the architecture of the DDM systems. The computation time is the time to perform the mining process on the distributed data sets. AR mining is a vital problem of DM. It's a new and challenging area to perform AR mining on XML data due to the difficulty of XML data. In our approach, numerous problems in XML data is handled suitably to assure the correctness of the result. The ODAM Algorithm is used for the mining process in a parallel and distributed setting. The response time with the communication and computation factors are measured to achieve an improved response time. The performance examination is done by increasing the number of processors in a distributed environment. As the mining process is done in parallel an optimal solution is obtained.



Figure 1: A © 2013



1

Figure 2: CFigure 1 :

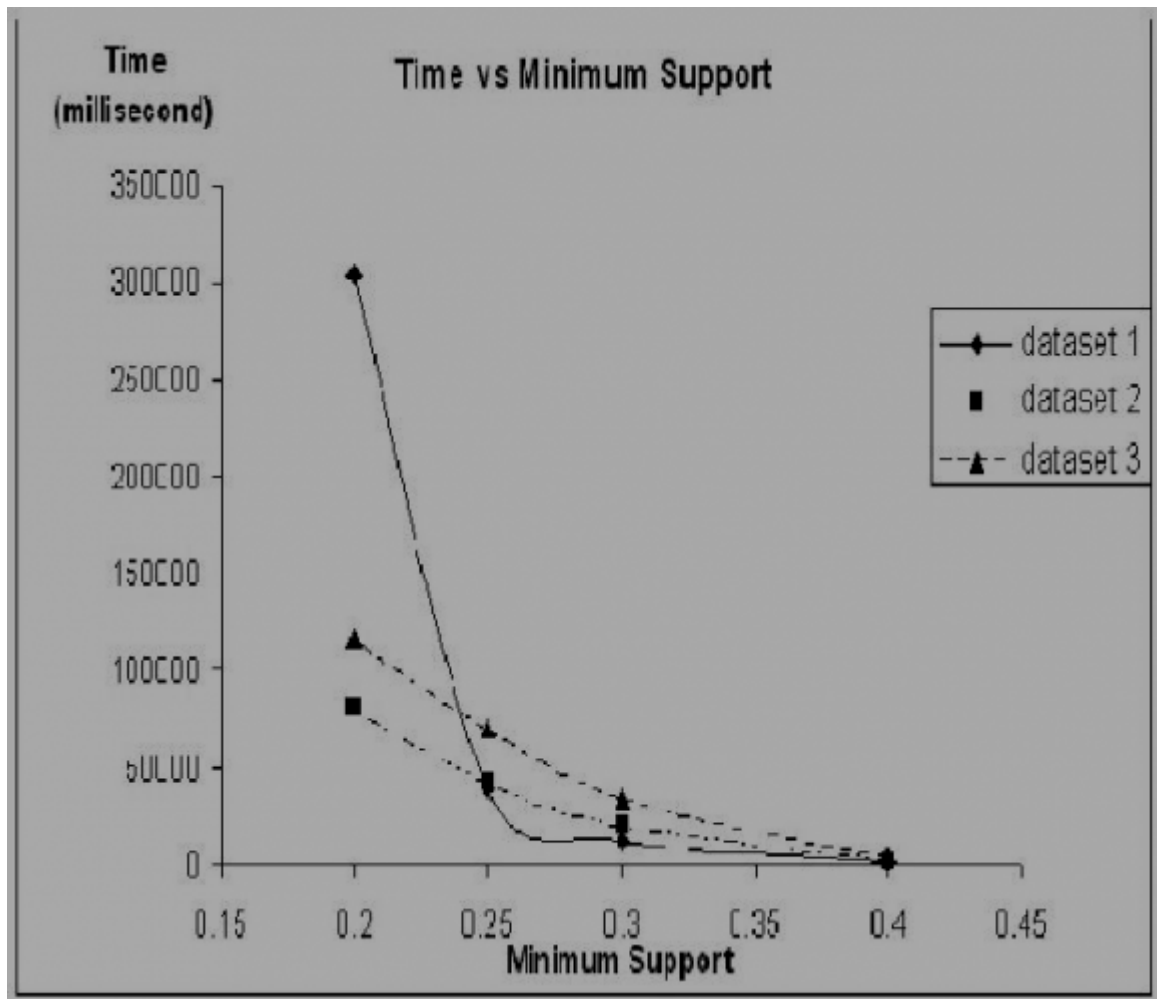


Figure 3:

```

NF= {Non-frequent global 1-itemset} for all transaction
t ∈ D {for all 2-subsets s of t
if (s ∈ C2) s.sup++;
t/=delete_nonfrequent_items(t);
Table.add (t/);}
Send_to_receiver (C2);
/* Global Frequent support counts from receiver*/
F2=receive_from_receiver (F?);
C3= (Candidate item set);
T=Table.getTransactions (); k=3;
While (Ck ?{ }) {For all transaction t ∈ T
For all k-subsets s of t
If (s ∈ Ck) s.sup++;
k++; Send_to_receiver (Ck);
/* Generating candidate item set of k+1 pass*/
Ck+1={Candidate item set); }

```

Volume XIII Issue XII Version I
Global Journal of Computer Science and Technology
© 2013 Global Journals Inc. (US)

Figure 4:

-
- [Cheung ()] ‘A Fast Distributed Algorithm for Mining Association Rules’. D W Cheung . *Proc. Parallel and Distributed Information Systems*, (Parallel and Distributed Information Systems) 1996. IEEE CS Press. p. .
- [Park et al. ()] ‘An Effective Hash Based Algorithm for Mining Association Rules’. J S Park , M Chen , P S Yu . *Proc. 1995 ACM SIGMOD Int’l Conf. Management of Data*, (1995 ACM SIGMOD Int’l Conf. Management of Data) 1995. ACM Press. p. .
- [Savasere et al. ()] ‘An Efficient Algorithm for Mining Association Rules in Large Database’. A Savasere , E Omiecinski , S B Navathe . *Proc. 21st Int’l Conf. Very Large Databases (VLDB 94)*, (21st Int’l Conf. Very Large Databases (VLDB 94)) 1995. Morgan Kaufmann. p. .
- [Cheung ()] ‘Efficient Mining of Association Rules in Distributed Databases’. D W Cheung . *IEEE Trans. Knowledge and Data Eng* 1996. 8 (6) p. .
- [Agrawal and Srikant ()] ‘Fast Algorithms for Mining Association Rules in Large Database’. R Agrawal , R Srikant . *Proc. 20th Int’l Conf. Very Large Databases*, (20th Int’l Conf. Very Large Databases) 1994. Morgan Kaufmann. p. .
- [Zaki and Pin ()] ‘Introduction: Recent Developments in Parallel and Distributed Data Mining’. M J Zaki , Y Pin . *J. Distributed and Parallel Databases* 2002. 11 (2) p. .
- [Han et al. ()] ‘Mining Frequent Patterns without Candidate Generation’. J Han , J Pei , Y Yin . *Proc. ACM SIGMOD Int’l. Conf. Management of Data*, (ACM SIGMOD Int’l. Conf. Management of Data) 2000. ACM Press. p. .
- [Zaki ()] *Parallel Data Mining for Association Rules on Shared-Memory Multiprocessors*, M J Zaki . TR 618. 1996. Computer Science Dept., Univ. of Rochester (tech. report)
- [Shafer ()] ‘Parallel Mining of Association Rules’. R , J C Shafer . *IEEE Tran* March 2004. 1996. 8 (6) p. . (Data Eng.)
- [Zaki ()] ‘Scalable Algorithms for Association Mining’. M J Zaki . *IEEE Trans. Knowledge and Data Eng* 2000. 12 (2) p. .