

A Novel Approach for Scalability -Two Way Sequential Pattern Mining using UDDAG

Dr.P.Raguraman¹

¹ SRI SANKARA ARTS AND SCIENCE COLLEGE

Received: 9 December 2012 Accepted: 3 January 2013 Published: 15 January 2013

Abstract

Traditional pattern growth-based approaches for sequential pattern mining derive length- $(k + 1)$ patterns based on the projected databases of length- k patterns recursively. At each level of recursion, they unidirectionally grow the length of detected patterns by one along the suffix of detected patterns, which needs k levels of recursion to find a length- k pattern. In this paper, a novel data structure, UpDown Directed Acyclic Graph (UDDAG), is invented for efficient sequential pattern mining. UDDAG allows bidirectional pattern growth along both ends of detected patterns. Thus, a length- k pattern can be detected in $\lceil \log_2 k + 1 \rceil$ levels of recursion at best, which results in fewer levels of recursion and faster pattern growth. When minSup is large such that the average pattern length is close to 1, UDDAG and PrefixSpan have similar performance because the problem degrades into frequent item counting problem. However, UDDAG scales up much better. It often outperforms PrefixSpan by almost one order of magnitude in scalability tests. UDDAG is also considerably faster than Spade and LapinSpam. Except for extreme cases, UDDAG uses comparable memory to that of PrefixSpan and less memory than Spade and LapinSpam. Additionally, the special feature of UDDAG enables its extension toward applications involving searching in large spaces.

Index terms— data mining algorithm, directed acyclic graph, performance analysis, sequential pattern, transaction database.

1 Introduction

SEQUENTIAL pattern mining is an important data mining problem, which detects frequent subsequences in a sequence database. A major technique for sequential pattern mining is pattern growth. Traditional pattern growth-based approaches (e.g., PrefixSpan) derive length- $(k + 1)$ patterns based on the projected databases of a length- k pattern recursively. At each level of recursion, the length of detected patterns is grown by 1, and patterns are grown unidirectionally along the suffix direction. Consequently, we need k levels of recursion to mine a length- k pattern, which is expensive due to the large number of recursive database projections. In this paper, a new approach based on UpDown Directed Acyclic Graph (UDDAG) is proposed for fast pattern growth. UDDAG is a novel data structure, which supports bidirectional pattern growth from both ends of detected patterns. With UDDAG, at level i recursion, we may grow the length of patterns by $2i - 1$ at most. Thus, a length- k pattern can be detected in $\lceil \log_2 k + 1 \rceil$ levels of recursion at minimum, which results in better scale-up property for UDDAG compared to PrefixSpan. Our extensive experiments clearly demonstrated the strength of UDDAG with its bidirectional pattern growth strategy. When minSup is very large such that the average length of patterns is very small (close to 1), UDDAG and PrefixSpan have similar performance because in this case, the problem degrades into a basic frequent item counting problem. However, UDDAG scales up much better compared to PrefixSpan. It often outperforms PrefixSpan by one order of magnitude in our scalability tests. UDDAG is also considerably faster than two other representative algorithms, Spade and LapinSpam. Except for some extreme

cases, the memory usage of UDDAG is comparable to that of PrefixSpan. UDDAG generally uses less memory than Spade and LapinSpam. UDDAG may be extended to other areas where efficient searching in large searching spaces is necessary.

2 II.

3 Related Work

The problem of sequential pattern mining was introduced by Agrawal and Srikant [1]. Among the many algorithms proposed to solve the problem, GSP [17] and PrefixSpan[13], [14] represent two major types of approaches: a prioribased and pattern growth-based. A priori principle states that any supersequence of a nonfrequent sequence must not be frequent. A priori based approaches can be considered as breadth-first traversal algorithms because they construct all length-k patterns before constructing length-(k+1) patterns. The AprioriAll algorithm [1] is one of the earliest a prioribased approaches. It first finds all frequent item sets, transforms the database so that each transaction is replaced by all frequent item sets it contains, and then finds patterns. The GSP algorithm [16] is an improvement over AprioriAll. To reduce candidates, GSP only creates a new length-k candidate when there are two frequent length-(k-1) sequences with the prefix of one equal to the suffix of the other. To test whether a candidate is a frequent length-k pattern, the support of each length-k candidate is counted by examining all the sequences. The PSP algorithm [12] is similar to GSP except that the placement of candidates is improved through a prefix tree arrangement to speed up pattern discovery. The SPIRIT algorithm [9] uses regular expressions as constraints and developed a family of algorithms for pattern mining under constraints based on a priori rule. The SPaRSe algorithm [3] improves GSP by using both candidate generation and projected databases to achieve higher efficiency for high pattern density conditions.

4 III.

Problem Definition a) Updown Directed Acyclic Graph-Based Sequential Pattern Mining UDDAG-based pattern mining approach, which first transforms a database based on frequent item sets, then partitions the problem, and finally, detects each subset using UDDAG. The absolute support for an item set in a sequence database is the number of tuples whose sequences contain the item set. An item set with a support larger than minSup is called a frequent item (FI) set. Based on frequent item sets, we transform each sequence in a database D into an alternative representation.

5 i. Transformed Database Definition

Let D be a database and P be the complete set of sequential patterns in D, D' be its transformed database, substituting the ids of each item pattern contained in D' with the corresponding item sets, and denoting the resulted pattern set by P', we have $P = P'$.

Based on frequent item sets, we transform each sequence in a database D into an alternative representation. Steps involved in Database Transformation: 1. Find the set of frequent items in D. 2. Assign a unique id to each FI in D and then replace each item set in each sequence with the ids of all the FIs contained in the item set. For the database in Table 1, the FIs are: (1),(??), (3), (??), (??), (??), (1,2), (2,3).

By assigning a unique id to each FI, e.g., (1)-1, (1,2)-2, (2)-3, (2,3)-4, (3)-5, (4)-6, (5)-7, (6)-8, we can transform the database as shown in Table 2 (infrequent items are eliminated).

ii. Problem Partitioning Definition Let $\{x_1, x_2, \dots, x_t\}$ be the frequent item sets in a database D, $x_1 < x_2 < \dots < x_t$, the complete set of patterns (P) in D can be divided into t disjoint subsets. The ith subset (denoted by P_{x_i} , $1 \leq i \leq t$) is the set of patterns that contains x_i and FIs smaller than x_i .

iii. (Projected database) collection of all the tuples whose sequences contain an item set in a database D is called x-projected denoted by x . The total number of different Frequent Items FIs in the Sequential Database are found by Database Transformation module. For a database with n different frequent items, its patterns can be divided into n disjoint subsets. The subset (1 < i < n) is the set of patterns that contain (the root item of the subset) and items smaller than i. Each subset i of the problem is mapped into a projected database denoted by (x_i D).

P is partitioned into eight subsets as there are 8 FIs in table-2, the one contains 1 (P1), the one contains 2 and smaller ids (P2), ..., and the one contains 8 and smaller ids (P8).

Given the following database (P8) alone is found as given by 8

6 iii. UDDAG based Pattern Mining

In the ith subset, each pattern in the projected database (x_i D) can be divided into two parts, prefix and suffix of i.

The collection of all the prefix/suffix tuples of a frequent item set X in x_i D is called the prefix/suffixprojected database of x_i , denoted by $Pre(x_i D) / Suf(x_i D)$.

To detect the sequential pattern in projected database ($x D$) $P x$, Detect patterns in $Pre(x D)$ called pattern prefix (PP). Detect pattern in $Suf(x D)$ called pattern suffix (PS) The above steps are repeated recursively until no frequent items are found in the $pre(x D) / suf(x D)$. Combine the patterns of all the to derive $P x$.

The complete set of patterns the union of patterns of the all subsets or projected database ($x D$) detected above.

The Apriori property is used to reduce the number of candidate sets to be considered

7 Example for Pattern Mining

8 Performance Evaluation

We conducted an extensive set of experiments to compare our approach with other representative algorithms. All the experiments were performed on a windows Server 2003 with 3.0 GHz Quad Core Intel Xeon Server and 16 GB memory. The algorithms we compared are PrefixSpan, Spade, and LapinSpam, which were all implemented in C++ by their authors (Minor changes have been made to adapt Spade to Windows). Two versions of UDDAG were tested. UDDAG-bv uses bit vector to verify candidates and UDDAG-co uses co-occurrences to verify candidates whenever possible. We perform two studies using the same data generator as in [14]: 1) Comparative study, which uses similar data sets as that in [14]; 2) Scalability study. The data sets were generated by maintaining all except one of the parameters shown in Table ?? fixed, and exploring different values for the remaining ones.

V.

9 Conclusion

In this paper, a novel data structure UDDAG is invented for efficient pattern mining. The new approach grows patterns from both ends (prefixes and suffixes) of detected patterns, which results in faster pattern growth because of less levels of database projection compared to traditional approaches. Extensive experiments on both comparative and scalability studies have been performed to evaluate the proposed algorithm.

One major feature of UDDAG is that it supports efficient pruning of invalid candidates. This represents a promising approach for applications involving searching in large spaces. Thus, it has great potential to related areas of data mining and artificial intelligence. In the future, we expect to further improve UDDAG-based pattern mining algorithm as follows: 1) Currently, FI detection is independent from pattern mining. Practically, the knowledge gained from FI detection may be useful for pattern mining. In the future, we will integrate the solutions of the two so that they can benefit from each other. 2) Different candidate verification strategies may have different impacts to the efficiency of the algorithm. In the future, we will study more efficient verification strategy. 3) UDDAG has big impact to the memory usage when the number of patterns in a subset is extremely large. In the future, we will find an efficient way to store UDDAG. ¹

¹© 2013 Global Journals Inc. (US) Global Journal of Computer Science and Technology



2

Figure 1: C 2 .

Seq. Id	Sequence
1	<1 (1,2,3) (1,3) 4 (3,6)>
2	<(1,4) 3 (2,3) (1,5)>
3	<(5,6) (1,2) (4,6) 3 2>
4	<5 7 (1,6) 3 2 3>

8

Figure 2: Assuming 8

Seq. Id	Sequence
1	<1 (1,2,3,4,5) (1,5) 6 (5,8)>
2	<(1,6) 5 (3,4,5) (1,7)>
3	<(7,8) (1,2,3) (6,8) 5 3>
4	<7 (1,8) 5 3 5>

Figure 3: C

1

Figure 4: Table 1 :

2

Figure 5: Table 2 :

-
- [Berkovich et al. ()] ‘A Bit-Counting Algorithm Using the Frequency Division Principle’. S Berkovich , G Lapir , M Mack . *Software: Practice and Experience* 2000. 30 (14) p. .
- [Chen and Xiao] ‘BISC: A Binary Itemset Support Counting Approach Towards Efficient Frequent Itemset Mining’. J Chen , K Xiao . *ACM Trans. Knowledge Discovery in Data*
- [Grahne and Zhu ()] ‘Efficiently Using Prefix-Trees in Mining Frequent Itemsets’. G Grahne , J Zhu . *Proc. Workshop Frequent Itemset Mining Implementations (FIMI '03)*, (Workshop Frequent Itemset Mining Implementations (FIMI '03)) 2003.
- [Agrawal and Srikant ()] ‘Fast Algorithms for Mining Association Rules’. R Agrawal , R Srikant . *Proc. 20th Int'l Conf. Very Large Data Bases (VLDB)*, (20th Int'l Conf. Very Large Data Bases (VLDB)) 1994. p. .
- [Antunes and Oliveira ()] ‘Generalization of Pattern-Growth Methods for Sequential Pattern Mining with Gap Constraints’. C Antunes , A L Oliveira . *Proc. Int'l Conf. Machine Learning and Data Mining*, (Int'l Conf. Machine Learning and Data Mining) 2003. 2003. p. .
- [Chen and Cook (2007)] ‘Mining Contiguous Sequential Patterns from Web Logs’. J Chen , T Cook . *Proc. World Wide Web Conf. (WWW '07) Poster Session*, (World Wide Web Conf. (WWW '07) Poster Session) May 2007.
- [Agrawal and Srikant ()] ‘Mining Sequential Patterns’. R Agrawal , R Srikant . **CHEN: ANUPDOWNDIRECTEDACYCLICGRAPHAPPROACHFOREQUENTIALPATTERNMINING927** *Proc. Int'l Conf. Data Eng. (ICDE '95)*, (Int'l Conf. Data Eng. (ICDE '95)) 1995. p. .
- [Ayres et al. ()] ‘Sequential Pattern Mining Using a Bitmap Representation’. J Ayres , J Gehrke , T Yu , J Flannick . *Proc. Int'l Conf. Knowledge Discovery and Data Mining*, (Int'l Conf. Knowledge Discovery and Data Mining) 2002. 2002. p. .
- [Garofalakis et al. ()] ‘SPIRIT: Sequential Pattern Mining with Regular Expression Constraints’. M Garofalakis , R Rastogi , K Shim . *Proc. Int'l Conf. Very Large Data Bases (VLDB '99)*, (Int'l Conf. Very Large Data Bases (VLDB '99)) 1999. p. .