

Understanding Rule Behavior through Apriori Algorithm over Social Network Data

Dr. S.S.Phulari¹ and Dr.S.D.Khamitkar²

¹ S.R.T.M.University , Nanded, MS, India.

Received: 7 December 2011 Accepted: 1 January 2012 Published: 15 January 2012

Abstract

APRIORI algorithm is a popular data mining technique used for extracting hidden patterns from data. This paper highlights practical demonstration of this algorithm for association rule mining over a survey data set of students related to social network usage. We concluded with discussions on the number of research observations including new rules generated during the process.

Index terms—

1 Introduction

Data mining is a technique that helps to extract important data from a large database. It is the process of sorting through large amounts of data and picking out relevant information through the use of certain sophisticated algorithms. As more data is gathered, with the amount of data doubling every three years, data mining is becoming an increasingly important tool to transform this data into information. Data mining techniques are the result of a long process of research and product development. This evolution began when business data was first stored on computers, continued with improvements in data access, and more recently, generated technologies that allow users to navigate through their data in real time. Data mining takes this evolutionary process beyond retrospective data access and navigation to prospective and proactive information delivery. The myth possible with data mining includes automated prediction of trends and behaviors and automated discovery of previously unknown patterns.

2 Analysis of Apriori Algorithm

Apriori was proposed by Agrawal and Srikant in 1994. The algorithm finds the frequent set L in the database D. It makes use of the downward closure property. The algorithm is a bottom search, moving upward level; it prunes many of the sets which are unlikely to be frequent sets, thus saving any extra efforts. Apriori algorithm is an algorithm of association rule mining. It is an important data mining model studied extensively by the database and data mining community. It Assume all data are categorical. It is initially used for Market Basket Analysis to find how items purchased by customers are related.

The problem of finding association rules can be stated as : Given a database of sales transactions, it is desirable to discover the important associations among different items such the presence of some items in a transaction will imply the presence of other items in the same transaction. As example of an association rule is: Contains (T, "baby food") ? Contains (T, "diapers") [Support= 4%, Confidence=40%]

The interpretation of such rule is as follows: ? 40% of transactions that contains baby food also contains diapers; ? 4% of all transactions contain both of these items. The above association rule is called singledimension because it involves a single attribute or predicate (Contains). The main problem is to find all association rules that satisfy minimum support and minimum confidence thresholds, which are provided by user and/or domain experts. A rule is frequent if its support is greater than the minimum support threshold and strong if its confidence is more than the minimum confidence threshold.

Discovering all association rules is considered as two phase process where we find all frequent item sets having minimum support. The search space to enumeration all frequent item sets is on the magnitude of 2^n . In

second step, we generate strong rules. Any association that satisfies the threshold will be used to generate an association rule. The first phase in discovering all association rules is considered to be the most important one because it is time consuming due to the huge search space (the power set of the set of all items) and the second phase can be accomplished in a straightforward manner.

III.

4 Algorithm for Apriori

The pseudo code for the algorithm is given below. For a transaction database T , and a support threshold of σ . Usual set theoretic notation is employed; though note that C_k is a multi set. C_k is the candidate set for level k . Generate C_k algorithm is assumed to generate the candidate sets from the large item sets of the preceding level, heeding the downward closure lemma.

$count(C_k)$ accesses a field of the data structure that represents candidate set C_k , which is initially assumed to be zero. Many details are omitted below, usually the most important part of the implementation is the data structure used for storing the candidate sets, and counting their frequencies.

IV.

Steps in finding the association rules using Apriori

A large supermarket tracks sales data by stockkeeping unit (SKU) for each item, and thus is able to know what items are typically purchased together. Apriori is a moderately efficient way to build a list of frequent purchased item pairs from this data. Let the database of transactions consist of the sets $\{1,2,3,4\}$, $\{1,2\}$, $\{2,3,4\}$, $\{2,3\}$, $\{1,2,4\}$, $\{3,4\}$, and $\{2,4\}$. Each number corresponds to a product such as "butter" or "bread". The first step of Apriori is to count up the frequencies, called the supports, of each member item separately: Table ?? explains the working of Apriori algorithm. We can define a minimum support level to qualify as "frequent," which depends on the context. For this case, let $\min\text{ support} = 3$. Therefore, all are frequent. The next step is to generate a list of all 2-pairs of the frequent items. Had any of the above items not been frequent, they wouldn't have been included as a possible member of possible 2-item pairs. In this way, Apriori prunes the tree of all possible sets. In next step we again select only these items (now 2-pairs are items) which are frequent and generate a list of all 3-triples of the frequent items (by connecting frequent pairs with frequent single items). In the example, there are no frequent 3-triples. Most common 3-triples are $\{1,2,4\}$ and $\{2,3,4\}$, but their support is equal to 2 which is smaller than our $\min\text{ support}$.

6 Global Journal of Computer Science and Technology

Volume XII Issue X Version I ??table 1: Table 2 : V.

7 Implementing Apriori algorithm in WEKA

WEKA is a collection of machine learning algorithms for data mining tasks. The algorithms can either be applied directly to a dataset or called from your own Java code. WEKA contains tools for data preprocessing, classification, regression, clustering, association rules, and visualization. It is also well-suited for developing new machine learning schemes. WEKA is open source software issued under the GNU General Public License.

VI.

9 Data set features

A closed questionnaire of 56 questions, labeled A,B,C?? BB was prepared and circulated among 56 students. Maximum questions were having four options to answer. answers caught as a1, a2, a3, a4 (a means answer) . These questionnaire were circulated randomly to avoid mass copying of the answers and then collected after one hour. Out of 56 , only 43 questionnaire were correct in all respects. Remaining 13 needed interactions with the corresponding students as few questions on them were not answered by them. Since 13 students refused to re answer these, we have rejected them out. Microsoft Excel was used to tabulate the data in the questionnaire and 43 rows were created . A CSV(Comma Separated Values) sheet was made from it which has been fed as input to the WEKA Algorithm.

VII.

11 Rules generated

12 Conclusions

In this paper, we have studied association rule mining over survey dataset. Our study shows that mining multiple-level association rules from databases has wide applications and efficient algorithms can be developed for discovery



Figure 1:

T

Figure 2:

T

Figure 3:

Figure 4:

12 CONCLUSIONS

95 of interesting and strong such rules in the database The larger the set of frequent item sets the more the number
96 of rules presented to the user, many of which are redundant. ¹

¹© 2012 Global Journals Inc. (US)

-
- 97 [Klementtinen ()] ‘Finding interesting rules from large sets of discovered association rules’. M Klementtinen .
98 *Proceedings of the CIKM*, (the CIKM) 1994.
- 99 [Bocca Jarke Zaniolo (ed.) (1994)] *Proc. of the 20th International Conference on Very Large Data Bases, VLDB*,
100 Jorge B Bocca, Matthias Jarke, Carlo Zaniolo (ed.) (of the 20th International Conference on Very Large Data
101 Bases, VLDBSantiago, Chile) September 1994. p. .