

An Efficient Market Basket Analysis based on Adaptive Association Rule Mining with Faster Rule Generation Algorithm

Dr. M. Dhanabhakya* & Dr. M. Punithavalli**

*Assistant Professor, Department of Commerce, Bharathiar University, Coimbatore, Tamilnadu, INDIA. E-Mail: dhana_giri@rediffmail.com

**Director and Head of the Department, Department of Computer Science, Ramakrishna College of Engineering, Vattamalaipalayam, Coimbatore, Tamilnadu, INDIA.

Abstract—Data mining is the process of extracting relatively useful information from a large data base. Majority of the recognized business organizations, super markets, etc., have accumulated huge amount of information from their customers. A vital sub problem of data mining is to identify frequent sets to assist mine association rules for Market Basket Analysis (MBA). Market Basket Analysis is an effective data mining tool utilized to discover the co-occurrence or co-existence of nominal or categorical observations. MBA is extensively used to identify purchasing pattern of customers in a supermarkets using transaction level data. However, it is very tough to find the valuable information hidden in large databases. Many researches were done by the database community based on association rule mining and classification technique to find the related information in large databases. The most widely used technique to conduct Market Basket Analysis is association rules technique. In this paper, an effective MBA based on Adaptive Association Rule Mining with Faster Rule Generation Algorithm (FRG-AARM) is proposed based on Adaptive Association Rule Mining. This algorithm speeds up the rule mining process with better accuracy and effectiveness.

Keywords—Association Rule Mining; Adaptive Association Rule Mining; Faster Algorithm; Market Basket Analysis; Rule Generation Algorithm.

Abbreviations—Adaptive-Support Association Rules (ASAR); Adaptive Association Rule Mining with Faster Rule Generation Algorithm (FRG-AARM); Market Basket Analysis (MBA).

I. INTRODUCTION

THE extensive application of distributed information systems provides structure of large data collections in departmental stores, shopping mall etc. These data collections enclose healthy information that required to be discovered. Businesses can learn from their transaction data more about the activities of their customers and thus can enhance their business by using this knowledge.

Data mining offers approaches that facilitate extracting from huge data collections unidentified relationships among the data items that are helpful for making decision. Thus data mining generates novel, unsuspected interpretations of data.

Information is collected roughly everywhere in our day-to-day lives. For instance, at supermarket checkouts information about customer purchases is recorded. When discount cards are used information about purchasing activities of the customer and individual details can be

obtained. Estimation of this information can assist retailers plan more efficient and customized marketing strategies.

Physical analysis of these huge amount of information stored in modern databases is very difficult. Data mining provides tools to reveal unknown information in large databases which are stored already. A well-known data mining technique is association rule mining [Kazienko & Kuzminska, 2005]. It is able to discover all interesting relationships which are called as associations in a database. Association rules are very efficient in revealing all the interesting relationships in a relatively large database with huge amount of data. The large quantity of information collected through the set of association rules can be used not only for illustrating the relationships in the database, but also used for differentiating between different kinds of classes in a database.

Association rule mining [Xie Wen-xiu et al., 2010] identifies the remarkable association or relationship between a large set of data items. With enormous quantity of data

regularly being acquired and stored in databases, numerous industries are becoming concerned in mining association rules from their databases. For instance, the finding of interesting association relationships between huge quantities of business operation data can help in catalog design, cross-marketing and different business decision making processes. A distinctive example of association rule mining is market basket analysis [Cunningham & Frank, 1999; Chiu et al., 2002; Trnka, 2010]. This technique inspects customer buying behaviors by recognizing associations among various items that customers place in their shopping baskets. The recognition of such associations can help retailers develop marketing strategies by gaining insight into which items are frequently purchased jointly by customers. It is useful to observe the customer purchasing manners and helps in increasing the sales and preserve inventory by focusing on the point of sale transaction data. This work is a wide area for the researchers to build up a better data mining algorithm for market basket analysis.

Mining Association Rules [Dexing Wang et al., 2010] is one of the most important application fields of Data Mining. Given a set of consumer transactions on items, the key objective is to determine correlations among the sales of items. Mining association rules, also known as market basket analysis, is one of the application fields of Data Mining. Think a market with a gathering of large amount of customer transactions. An association rule is $X \Rightarrow Y$, where X is referred as the antecedent and Y is referred as the consequent. X and Y are sets of items and the rule represents that customers who purchase X are probable to purchase Y with probability %c where c is known as the confidence. Such a rule may be: "Eighty percent of people who purchase cigarettes also purchase matches". Such rules assists to respond questions of the variety "What is Coca Cola sold with?" or if the customers are intended in checking the dependency among two items A and B it is required to determine rules that have A in the consequent and B in the antecedent.

In this paper, a new fast algorithm for mining association rule to provide efficient Market Basket Analysis has been proposed. The proposed approach is capable of making the customers more comfortable in making effective purchase. By combining similarity between rules and active user and confidence of the weighted rules, the shop keepers will place the most relevant items in the appropriate places. Therefore, it will increase the sales and moreover it will be more comfortable for the customers. The proposed algorithm also satisfies the customers by providing the most related items that is needed.

II. RELATED WORKS

Association rule mining is a data mining task that identifies relationships among items in a transactional database. Association rules have been widely investigated in the literature for their role in several application domains such as Market Basket Analysis (MBA), recommender systems,

diagnosis decisions support, telecommunication, intrusion detection, etc. The competent discovery of such rules has been a key focus in the data mining research community. The standard apriori algorithm has been modified for the improvements of association rule mining algorithms [Pei-ji Wang et al., 2009].

In Chris Anderson's latest book, "The Long Tail" [Anderson, 2006], the author illustrates why the future of retail business involves selling smaller quantities of more products. Anderson concludes with his concept as the "98% rule" contrasting with the well known "80/20 rule". The "98% rule" means that in a statistical distribution of the products, only 2% of the items are very frequent and 98% of the items have very low frequencies, creating a long tail distribution.

Weiyang Lin et al., (2004) described an efficient Adaptive-Support Association Rule Mining for Recommender Systems. The author examined the usage of association rule mining as a fundamental technique for collaborative recommender systems. Association rules have been used with sensation in other domains. Nevertheless, most currently existing association rule mining algorithms were designed with market basket analysis in mind. They described a collaborative recommendation technique based on a novel algorithm distinctively designed to excavate association rules for this rationale. The main advantage of their proposed approach is that their algorithm does not require the minimum support to be specified in advance. To a certain extent, a target range is specified for the number of rules, and the algorithm alters the minimum support for all customers with the intention of acquiring a rule set whose size is in the desired range. Moreover they employed associations between customers as well as associations between items in making recommendations. The experimental evaluation of a system based on their algorithm revealed that its performance is significantly better than that of traditional correlation-based approaches.

Shyue-Liang Wang et al., (2003) proposed an effective data-mining approach for discovering Adaptive-Support Association Rules (ASAR) from databases. Adaptive-support association rules are constrained association rules with application to collaborative recommendation systems. To find out association rules for recommendation systems, a particular value of target item in association rules is normally assumed and no minimum support is specified in advance. Depending on the size monotonicity of association rules, specifically the number of association rules reduces when the minimum support increases, an effective algorithm using variable step size for determining minimum support and as a result adaptive-support association rules is generated.

An incremental updating method is proposed by Jin Qian & Xiang-Ping Meng (2003) for maintaining the association rules discovered by database mining. There have been several investigations on effective finding of association rules in huge databases. On the other hand, it is nontrivial to preserve the association rules when the two thresholds transform. In this approach, an adaptive algorithm (IMARA) for

incremental mining of association rules in light of threshold changes.

The primary task in any associative classification approach is mining of the association rules. Many investigations have revealed that the minimum support measure plays an important part in constructing a perfect classifier. With no information regarding the items and their frequency, user provided support measures are unsuitable, not often may they coincide. Kanimozhi Selvi & Tamilarasi (2007) developed a technique called Dynamic Adaptive Support Apriori (DASApriori) to compute the minimum support for obtaining class association rules and to construct an uncomplicated and perfect classifier.

The association rules characterizes a significant class of knowledge that can be revealed from data warehouses. Present research attempts are concentrated on discovering well-organized ways of determining these rules from huge databases. At the same time as these databases develop, the discovered rules are required to be confirmed and it is necessary to discover new rules to the knowledge base. As mining afresh each time the database develops is incompetent, approaches for incremental mining are being studied. Their main objective is to reduce scans of the older database by exploiting the intermediary data built during the previous mining activities. Sarda & Srinivas (1998) used large and candidate itemsets and their counts in the older database and examined the growth to discover which rules maintain to overcome and which one is not succeeded in the complex database. It is also found that new rules for the incremental and updated database. The algorithm is adaptive in nature, as it concludes the nature of the increment and avoids altogether if possible, multiple scans of the incremental database. Another salient feature is that it does not need multiple scans of the older database.

A method for mining association rules that reflect the behaviors of past customers is proposed by Takama & Hattori (2007) for an adaptive search engine. The logs of the customers' retrieving behaviors are described with the resource description framework model, from which association rules that reflect successful retrieving behaviors are extracted. The extracted rules are used to improve the performance of a metadata-based search engine. The document repository with adaptive hybrid search engine is also developed based on the proposed method. The repository consists of a document registration module, hybrid search engine, and reasoning base. The document registration module is designed to reduce the cost of adding metadata to documents, and the hybrid search engine combines full-text search with metadata-based search engine to improve the recall of retrieval result. The reasoning base is implemented based on the association rule mining method, which contributes to improve both precision and recall of the hybrid search engine. Experiments are performed with a virtual user model, of which results show that appropriate rules can be extracted with the proposed method. The proposed technologies will contribute to realize the concept of humatronics in terms of establishing symmetric relation

between humans and systems, as well as sharing information, knowledge, and experiences via computer networks.

Min Wang et al., (2011) developed a technique for mining multi-dimension association rule depending on the Adaptive Genetic Algorithm (AGA) with crossover matrix and mutation matrix. In this association rule mining system, selection, mutation, and crossover are all parameter-free in evolution process. Results show that: combined with the adaptive genetic algorithm, the precision and efficiency of mining association rules is improved.

Taboada et al., (2006) proposed a method of association rule mining using genetic network programming (GNP) with a self-adaptation mechanism in order to improve the performance of association rule extraction systems. GNP is a kind of evolutionary methods, whose directed graphs are evolved to find a solution as individuals. Self-adaptation behavior in GNP is related to adjust the setting of control parameters such as crossover and mutation rates. It is called self-adaptive because the algorithm controls the setting of these parameters itself - embedding them into an individual's genome and evolving them. The aim is not only to find suitable adjustments but to do this efficiently. Our method can measure the significance of the association via the chi-squared test and obtain a sufficient number of important association rules. Extracted association rules are stored in a pool all together through generations and reflected in three genetic operators as acquired information. Further, our method can contain negation of attributes in association rules and suit association rule mining from dense databases.

III. PROPOSED ASSOCIATION RULE MINING METHODOLOGY

It is very complicated to select a proper minimum confidence and support for each item before the mining process because customers' interest and popularities of the items vary widely.

3.1. Adaptive Association Rule Mining

Sufficient rules for accurate recommendation will not be acquired if the minimum confidence and support for mining are set too high. If they are set extremely low, the runtime may be inappropriately long. In addition, an extreme number of rules may guide to decreased performance. This suggests the following new goal for association rule mining for recommender systems: "Given a transaction dataset, a target item, a specified minimum confidence and a desired range [minNumRules, maxNumRules] for the number of rules, find a set S of association rules with the target item in the heads of the rules such that the total number of rules in S is in the given range, the rules in S satisfy the minimum confidence constraint, and the rules in S have higher support than all rules not in S that are of the given form and satisfy the minimum confidence constraint" [Weiyang Lin et al., 2004].

The algorithm to mine association rules is as follows. This technique regulates the minimum support of the rules during mining to attain a suitable number of important rules

for the target item. This mining algorithm contains an outer loop and an inner loop. The overall procedure contains three parts. Initially, the minimum support count is started by the outer loop (product of the minimum support and the total number of transactions) based on the frequency of the target item and calls the inner loop to mine rules. If the inner loop ends, then the outer loop checks if the number of rules returned goes beyond maxNumRules. If so, the minimum support count is increased by the outer loop and calls the inner loop until the number of rules is maxNumRules or less. Ultimately, the outer loop checks if the number of rules is less than minNumRules; if it is, it decreases the minimum support count until the rule number is greater than or equal to minNumRules. For a given support, rules with shorter bodies are mined first. If the number n of rules is out of range for successive values of the minimum support count, with $n > \text{maxNumRules}$ at support count s and $n < \text{minNumRules}$ at support count $s + 1$, then the shortest maxNumRules rules with the smaller support count are returned [Weiyang Lin et al., 2000].

The inner loop of this approach is a variant of CBA-RG and hence of the Apriori algorithm as well. It is a variant of CBA-RG in the sense that rather than mining rules for all target classes, it only mines rules for one target item. It differs from CBA-RG in that it will only mine a number of rules within a certain range.

3.2. Proposed Association Rule Mining Methodology

In this section, two new algorithms, Apriori and AprioriTid are given, that differ fundamentally from ordinary algorithms. Later discussion is on how the best features of Apriori and AprioriTid can be combined into a hybrid algorithm, called AprioriHybrid.

The problem of discovering all association rules can be decomposed into two sub problems:

1. Find all sets of items that have transaction support over minimum support. The support for an itemset is defined as the number of transactions that include the itemset. Itemsets that has minimum support are considered as large itemsets, and remaining are small itemsets. In this section, novel approaches Apriori and AprioriTid are presented for solving this difficulty.
2. Use the large itemsets to produce the preferred rules. The algorithms for this difficulty are specified in this section. The universal idea is that if, say, ABCD and AB are large itemsets, one can then decide if the rule $AB \Rightarrow CD$ holds by computing the ratio $\text{conf} = \text{support}(ABCD)/\text{support}(AB)$. If $\text{conf} \geq \text{minconf}$, the rule then holds. (The rule will surely have minimum support because ABCD is large).

Apriori Algorithm is already discussed in this section. Now the following discussion is about AprioriTid algorithm.

The AprioriTid algorithm, shown in figure 1, also uses the Apriori-gen function to find the candidate items earlier than the pass initiates. The interesting feature of this algorithm is that the database D is not used for counting

support after the first pass. Rather, the set $\tilde{C}_k$ is used for this purpose. Each member of the set $\tilde{C}_k$ is of the form $\langle \text{TID}, \{X_k\} \rangle$, where each X_k is a potentially large k -itemset present in the transaction with identifier TID. For $k = 1$, $\tilde{C}_1$ corresponds to the database D , although conceptually each item i is replaced by the itemset $\{i\}$. For $k > 1$, $\tilde{C}_k$ is generated by the algorithm (step 10). The member of $\tilde{C}_k$ corresponding to transaction t is $\langle t, \text{TID}, \{c \in C_k | c \text{ contained in } t\} \rangle$. If a transaction does not contain any candidate k -itemset, $\tilde{C}_k$ will then not have an entry for this transaction. Thus, the number of entries in $\tilde{C}_k$ may be smaller than the number of transactions in the database, especially for large values of k . In addition, for large values of k , each entry may be smaller than the corresponding transaction because very few candidates may be contained in the transaction. However, for small values for k , each entry may be larger than the corresponding transaction because an entry in C_k includes all candidate k -itemsets contained in the transaction.

```

1)  $L_1 = \{\text{large 1-itemsets}\};$ 
2)  $\tilde{C}_1 = \text{database } D;$ 
3) for( $k=2; L_{k-1} \neq \emptyset; k++$ ) do begin
4)    $C_k = \text{apriori-gen}(L_{k-1});$  // New
   candidates
5)    $\tilde{C}_k = \emptyset;$ 
6)   For all entries  $t \in \tilde{C}_{k-1}$  do begin
7)     //determine candidate itemsets in  $C_k$  contained
     in the transaction with identifier  $t$ . TID
8)      $C_t = \{c \in C_k | (c - c[k]) \in t.\text{set-of-itemsets}$ 
       $\wedge (c - c[k-1]) \in t.\text{set-of-itemsets}\};$ 
9)     For all candidates  $c \in C_t$  do
10)       $C.\text{count}++;$ 
11)      If( $C_t \neq \emptyset$ ) then  $\tilde{C}_k += \langle t, \text{TID}, C_t \rangle;$ 
12)   End
13)    $L_k = \{c \in C_k | c.\text{count} \geq \text{minsup}\}$ 
14) End
15) Answer =  $\bigcup_k L_k$ 

```

Figure 1: Algorithm AprioriTid

To generate rules, for a group of item l , find all the non-empty subsets of l . For every such subset a , output a rule of the form $a \Rightarrow (l - a)$ if the ratio of support (l) to support (a) is at least minconf. Then, consider all subsets of l to generate rules with multiple consequents. Since the large itemsets are stored in hash tables, the support counts for the subset itemsets can be found efficiently. The above procedure can be improved by generating the subsets of a large itemset in a recursive depth-first fashion. For example, given an itemset ABCD, first consider the subset ABC, then AB. If a subset a of a large itemset l does not generate a rule, the subsets of a need not be considered for generating rules using l . For example, if $ABC \Rightarrow D$ does not have enough confidence, there is no need to check whether $AB \Rightarrow CD$ holds. No rules are missed because the support of any subset $\bar{a}$ of a must be as great as the support of a . Therefore, the confidence of the rule $\bar{a} \Rightarrow (l - \bar{a})$ cannot be more than the confidence of $a \Rightarrow (l - a)$. Hence, if a did not yield a rule involving all the

items in l with a as the antecedent, neither will $\bar{a}$. The following algorithm embodies these ideas:

```

forall large itemsets  $l_k, k \geq 2$  do
  call genrules( $l_k, l_k$ );

//The genrules generates all valid rules  $\hat{a} \Rightarrow (l_k - a)$ ,
for all  $\hat{a} \subset a_m$ 

procedure genrules( $l_k$ : large k-itemset,  $a_m$ : large m-
itemset)

1)  $A = \{(m-1)\text{-itemsets } a_{m-1} | a_{m-1} \subset a_m\}$ 
2) forall  $a_{m-1} \in A$  do begin
3)    $\text{conf} = \text{support}(l_k) / \text{support}(a_{m-1})$ ;
4)   if ( $\text{conf} \geq \text{minconf}$ ) then
       begin
7)     output the rule  $a_{m-1} \Rightarrow (l_k - a_{m-1})$ ,
       with confidence= conf and support= suppo
8)   if ( $m-1 > 1$ ) then
9)     call genrules( $l_k, a_{m-1}$ ); // to generate rules
       with subsets of  $a_{m-1}$  as the antecedents
10)  end

```

Figure 2: Simple Algorithm

The faster algorithm for association rule mining is discussed below:

It is showed earlier that if $a \Rightarrow (l - a)$ does not hold, neither does $\bar{a} \Rightarrow (l - \bar{a})$ for any $\bar{a} \subset a$. By rewriting, it follows that for a rule $(l - c) \Rightarrow c$ to hold, all rules of the form $(l - \bar{c}) \Rightarrow \bar{c}$ must also hold, where $\bar{c}$ is a non-empty subset of c . For example, if the rule $AB \Rightarrow CD$ holds, then the rules $ABC \Rightarrow D$ and $ABD \Rightarrow C$ must also hold.

Consider the above property for a particular large itemset, if a rule with consequent c holds, then so do rules with consequents that are subsets of c . This is similar to the property that if an itemset is large, then so is all its subsets. As a result, from a large itemset l , initially all rules with one item in the consequent is produced.

Then, use the consequents of these rules and the function Apriori-gen to produce all feasible consequents with two items that can appear in a rule generated from l . An algorithm using this idea is given below. The rules having one-item consequents in step 2 of this algorithm can be found by using a modified version of the preceding genrules function in which steps 8 and 9 are deleted to avoid the recursive call.

Faster algorithm for association rule mining is given below:

```

1) forall large k-itemsets  $l_k, k \geq 2$  do begin
2)    $H_1 = \{\text{consequents of rules derived from } l_k \text{ with one item in the consequent}\}$ ;
3)   call ap-genrules( $l_k, H_1$ );
4)   end
procedure ap-genrules ( $l_k$ : large k-itemset,  $H_m$ : set
of m-item consequents)
  if ( $k > m+1$ ) then begin
     $H_{m+1} = \text{apriori-gen}(H_m)$ ;
    forall  $h_{m+1} \in H_{m+1}$  do begin
       $\text{conf} = \text{support}(l_k) / \text{support}(l_k - h_{m+1})$ ;
      if ( $\text{conf} \geq \text{minconf}$ ) then
        output the rule  $(l_k - h_{m+1}) \Rightarrow h_{m+1}$  with
        confidence = conf and support = support( $l_k$ )
      else
        delete  $h_{m+1}$  from  $H_{m+1}$ ;
    end
    call ap-genrules ( $l_k, H_{m+1}$ );
  end
End

```

Figure 3: Faster Algorithm

As an example of the advantage of this algorithm, consider a large itemset (items) $ABCDE$. Assume that $ACDE \Rightarrow B$ and $ABCE \Rightarrow D$ are the only one-item consequent rules derived from this itemset that have the minimum confidence. In the usage of the simple algorithm, the recursive call $\text{genrules}(ABCDE, ACDE)$ will test if the two-item consequent rules $ACD \Rightarrow BE$, $ADE \Rightarrow BC$, $CDE \Rightarrow BA$, and $ACE \Rightarrow BD$ hold. The first of these rules cannot hold, because $E \subset BE$, and $ABCD \Rightarrow E$ does not have the smallest amount of confidence. The second and third rules cannot hold for similar reasons. The call $\text{genrules}(ABCDE, ABCE)$ will examine if the rules $ABC \Rightarrow DE$, $ABE \Rightarrow DC$, $BCE \Rightarrow DA$ and $ACE \Rightarrow BD$ hold, and will discover that the first three of these rules do not hold. Actually, the only two-item consequent rule that can probably hold is $ACE \Rightarrow BD$, where B and D are the consequents in the suitable one-item consequent rules. This is the only rule that will be tested by the faster algorithm.

IV. EXPERIMENTAL RESULTS

A simulation study is carried out to empirically compare the proposed MBA approach which uses fast adaptive association rule mining and traditional association-rule mining methods. The main objective of the simulation study is to identify the conditions under which the proposed method significantly outperforms the traditional method in identifying important purchasing patterns in a multi-store environment.

Prediction accuracy and average run time is taken as the performance measure to evaluate the performance of the proposed approach.

Figure 4 shows the comparison of the prediction accuracy of the proposed FRG-AARM is very high when compared with the other two approaches such as association rule mining and Adaptive association rule mining.

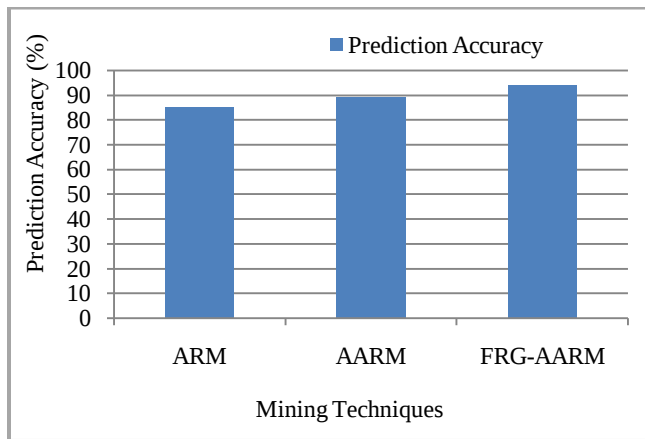


Figure 4: Prediction Accuracy Comparison

Figure 5 shows the run time comparison of the proposed FRG-AARM approach with other approaches. It is observed that the proposed approach takes very less time run time when compared with the other two approaches such as association rule mining and Adaptive association rule mining.

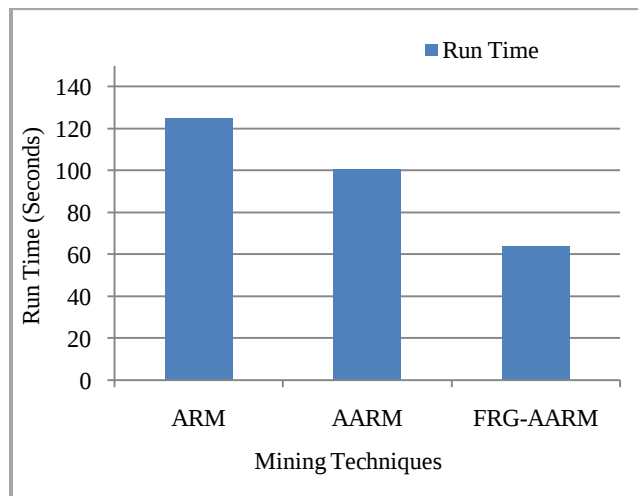


Figure 5: Run Time Comparison

V. CONCLUSION

There are various limitations in the existing association rule mining algorithms. The main drawback of the conventional Apriori algorithm is generating numerous candidate itemsets that must be repeatedly contrasted with the whole database.

This approach uses adaptive association rule mining algorithm with Faster Rule Generation Algorithm for effective Market Basket Analysis. This approach helps the customers in purchasing their products with more comfort which in turn increases the sales of the products. This is achieved with the help of using the confidence of the weighted rules in addition to the possible combination of similarity between rules and active user.

REFERENCES

- [1] N.L. Sarda & N.V. Srinivas (1998), "An Adaptive Algorithm for Incremental Mining of Association Rules", *Proceedings of Ninth International Workshop on Database and Expert Systems Applications*, Pp. 240–245.
- [2] S.J. Cunningham & E. Frank (1999), "Market Basket Analysis of Library Circulation Data", *International Conference on Neural Information Processing*, Vol. 2, Pp. 825–830.
- [3] Weiyang Lin, Sergio A. Alvarez & Carolina Ruiz (2000), "Collaborative Recommendation via Adaptive Association Rule Mining", *Data Mining and Knowledge Discovery*.
- [4] K.S.Y. Chiu, R.W.P. Luk, K.C.C. Chan & K.F.L. Chung (2002), "Market-basket Analysis with Principal Component Analysis: An Exploration", *IEEE International Conference on Systems, Man and Cybernetics*, Vol. 3.
- [5] Jin Qian & Xiang-Ping Meng (2003), "An Adaptive Algorithm for Incremental Mining Association Rules", *International Conference on Machine Learning and Cybernetics*, Vol. 1, Pp. 67–70.
- [6] Shyue-Liang Wang, Mei-Hwa Wang, Wen-Yang Lin & Tzung-Pei Hong (2003), "Efficient Generation of Adaptive-Support Association Rules", *IEEE International Conference on Systems, Man and Cybernetics*, Vol. 1, Pp. 894–899.
- [7] Weiyang Lin, Sergio A. Alvarez & Carolina Ruiz (2004), "Efficient Adaptive-Support Association Rule Mining for Recommender Systems", *Journal on Data Mining and Knowledge Discovery*, Vol. 6, No. 1, Pp. 83–105.
- [8] Xie Wen-xiu, Qi Heng-nian & Huang Mei-li (2010), "Market Basket Analysis based on Text Segmentation and Association Rule Mining", *First International Conference on Networking and Distributed Computing (ICNDC)*, Pp. 309–313.
- [9] P. Kazienko & P. Kuzminska (2005), "The Influence of Indirect Association Rules on ISDA Recommendation Ranking Lists", *Proceedings of 5th International Conference on Intelligent Systems Design and Applications*, Pp. 482–487.
- [10] C. Anderson (2006), "The Long Tail: Why the Future of Business is Selling Less of More", *Hyperion Books*.
- [11] K. Taboada, K. Shimada, S. Mabu, K. Hirasawa & J. Hu (2006), "Self-Adaptive Mechanism in Genetic Network Programming for Mining Association Rules", *International Joint Conference SICE-ICASE*, Pp. 6007–6012.
- [12] Y. Takama & S. Hattori (2007), "Mining Association Rules for Adaptive Search Engine based on RDF Technology", *IEEE Trans. on Industr. Electron.*, Vol. 54, No. 2, Pp. 790–796.
- [13] C.S. Kanimozhi Selvi & A. Tamilarasi (2007), "Association Rule Mining with Dynamic Adaptive Support Thresholds for Associative Classification", *International Conference on Conference on Computational Intelligence and Multimedia Applications*, Vol. 2, Pp. 76–80.
- [14] Pei-ji Wang, Lin Shi, Jin-niu Bai & Yu-lin Zhao (2009), "Mining Association Rules based on Apriori Algorithm and Application", *International Forum on Computer Science-Technology and Applications*, Vol. 1, Pp. 141–143.
- [15] Dexing Wang, Lixia Cui, Yanhua Wang, Hongchun Yuan & Jian Zhang (2010), "Association Rule Mining based on Concept Lattice in Bioinformatics Research", *International Conference on Biomedical Engineering and Computer Science (ICBECS)*, Pp. 1–4.
- [16] A. Trnka (2010), "Market Basket Analysis with Data Mining Methods", *International Conference on Networking and Information Technology (ICNIT)*, Pp. 446–450.
- [17] Min Wang, Qin Zou & Caihui Liu (2011), "Multi-dimension Association Rule Mining based on Adaptive Genetic Algorithm", *International Conference on Uncertainty Reasoning and Knowledge Engineering*, Vol. 2, Pp. 150–153.