

Association Rules Approaches in Heterogeneous Databases to Protect M-Players

Banupriya R, Sowmiya R, Sowmya S, Suganya T

Department of Computer Science and Engineering,
Rathinam Technical Campus,
Coimbatore, Tamilnadu, India

Abstract: - One of the technological processes that is progressing at the rate that is now considered to be the most rapid is known as data mining. This technique is applied in order to sift through massive datasets in search of the information that is regarded as being of the utmost significance and importance. On the other hand, databases may be sectioned off using either a horizontal or a vertical partitioning technique, allowing for the databases to be broken up into subsections. In addition to this, the ownership of these database collections is held by a wide number of various organizations. [Citation needed] In the field of data mining, mining for association rules is rapidly being one of the most important approaches, and it is on its way to become the most important strategy. This research presents a protocol for safe association rule mining approaches that may be utilized in horizontally distributed databases. The method is implemented on the Fast Distributed Mining (FDM) algorithm, which is an unsecure distributed system application of both the Apriori algorithm. The basic components of this protocol are made up of two different multiparty algorithms, each of which is assured to have a high level of privacy and protection. Mining association rules is what these algorithms do. The first of these algorithms calculates the union of the private subsets that are owned according to each participant inside the interaction, and the second of these algorithms determines whether or not a component that is held by one person is contained in a selection that is held by another participant. In comparison to techniques involving several parties, the level of privacy that end users enjoyed as a result of our solution was significantly improved. In addition, our protocol is substantially simpler to comprehend and is significantly more effective than the approaches that came before it in proportion to the number of link utilization, the cost, and the computational cost.

Keywords: Rule base, Apriori Algorithm, inter Algorithm, FDM, and Consolidated Database

1. INTRODUCTION

Data mining is a process that involves extracting meaningful information from massive data collections through the use of a variety of different methods. On the other hand, these collections may be dispersed among a lot of distinct parties on occasion thanks to the utilization of distributed databases. The term "data mining" refers to the practice of extracting previously unknown information that may be reasonably predicted from a vast number of small datasets. New technology has emerged as a direct result of this as a method of recognizing patterns and trends within large amounts of data. This technology has been developed as a means of recognizing patterns and trends within enormous volumes of data. As a result of this, the term "big data" was developed in order to adequately explain the circumstance. The culmination of this process is the accumulation of knowledge, which can be thought of as the significant information obtained from previously unknown characteristics. The information that one has acquired as a result of the process of learning is another definition of knowledge. One of the topics that is explored in this paper is the challenge of mining datasets that have been horizontally partitioned in order to find association rules. There are a significant number of websites that host homogeneous

databases, which are also usually referred to as databases with the same schema that contain information about a variety of various sorts of things. Homogenous databases store information about a wide range of these entities. In light of the fact that the unified database contains just the most fundamental degrees of support and confidence, the objective is to unearth all of the association regulations while revealing as little as possible about the players' personal databases. Our objective is to define the challenge as a protected approach to multiparty computing involving a number of different parties. Because of this, we will be able to move forward with our research. If it is in any way possible, we would also like to secure data that is of a more general type. For example, the locally stored association rules that are contained within each of these databases would be a good example of this. This specific component is one of the components of the undertaking that has been proposed that piques our interest and makes us curious. In situations such as these, M players each have their own private inputs, which are indicated by $x_1, x_2, x_3, \dots, x_M$, and they attempt to compute $y = f(x_1, x_2, x_3, \dots, x_M)$ for any function that is accessible to the general public. These inputs are marked by $x_1, x_2, x_3, \dots, x_M$. The participants would be able to hand over their contributions to a trustworthy third party in the event that they were present. The third party would then

evaluate the function and provide the participants with the results of his assessment. In order to obtain the desired result y[1, it is essential to develop a protocol that players may independently carry out even in the absence of a reliable third party. This is because doing so is essential in order to get the results that are intended. The participants are expected to be able to independently carry out the operation. This need is a prerequisite for participation. Then, such a protocol is taken into consideration if it can be demonstrated that no participant may learn more from his perspective on the protocol than he would in an idealized environment in which the computation is carried out by a trustworthy third party. This is the condition that must be met before a protocol can be taken into consideration.

The proposed system accepts as inputs a variety of datasets that are only partially homogenous, and its ultimate objective is to provide a list of association rules that have a particular amount of support and confidence. Because generic solutions are dependent on a representation of the function f as a Boolean circuit, the scope of problems that can be solved by using these solutions is restricted. These methods can also only be used to inputs and functions that are relatively easy to design using fundamental circuits, and they are limited to those applications. In light of these difficult circumstances, extra strategies are required in order to complete this mathematical procedure. When confronted with such conditions, it may be impossible to resist the concept of perfect security when seeking for practical answers, thinking that the additional knowledge may be seen as being risk-free. Kantarcioglu and Clifton came up with a protocol in order to address the issue that more acceptable security definitions that allow multiple parties to choose their preferred level of security are required for cost-effective solutions that maintain the intended security [2]. Kantarcioglu and Clifton came up with the protocol in order to address the issue that more acceptable security definitions that enable multiple parties to choose their preferred level of security are required for The two researchers came up with the protocol so that they might find a solution to this problem. The most important feature of the protocol is a sub-protocol that enables a secure computation of the union of the private subsets of each of the parties that are engaged. Cryptographic primitives such commutative encryption, oblivious transmission, and hash functions are some examples of those that are necessary for the successful implementation of this component of the protocol. Players only have the ability to acquire knowledge about other databases from their view of the protocol during this phase of the protocol, which is also the only phase of the protocol in which this option exists. In addition to what is

inferred from the end product and from their own input, this information will also be shown. It has been asserted that the disclosure of such information does not pose any threat and is, hence, appropriate from a pragmatist's point of view. This is in spite of the fact that the dissemination of such information does, in fact, compromise the confidentiality of the protocol. Clearly, there are restrictions placed on the additional information.

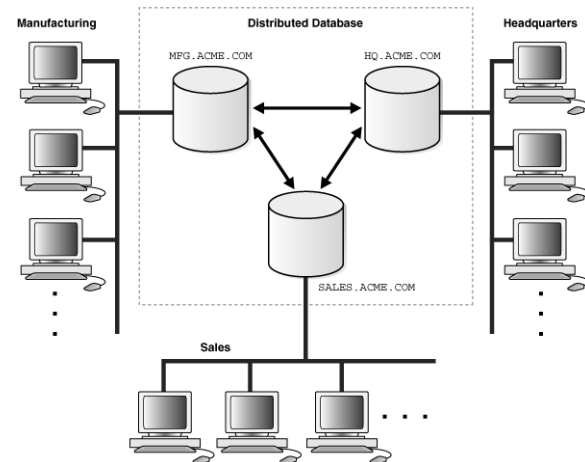


Figure 1: Architecture of Distributed Database

We offer an alternate protocol for the result that would be obtained from the union of limited subsets if this approach were used, which allows us to satisfy the constraints of the methodology. Therefore, the recently suggested protocol is preferable to the protocol that is being used at the moment in terms of both its possibility to be useful and its capacity to preserve the users' responsibility to preserve their anonymity. This is because the recently suggested protocol has the capacity to preserve the users' right to maintain their anonymity. The protocol does not implement all varieties of cryptographic primitives since it is insecure. Commutative encryption and oblivious transfer are two examples of this type of primitive. The method does not provide a safety guarantee that is completely bulletproof, but it does so for a smaller variety of groups (three) than the protocol does. In contrast to this, the protocol discloses information not just to coalitions but also to specific single players. This information is shared with both of these types of players. You should also mention that the surplus information that may be exposed by our protocol contains a lesser level of sensitivity than the surplus information that was spilt by the protocol. This should be stated as an additional point in your argument. You need to put out an argument supporting this point. You should argue that the information that was disclosed through the protocol was more sensitive than the information that was

leaked through other ways, to put it another way. This method computes a parameterized collection of functions, which from this point forward will be referred to as threshold functions. The term "threshold functions" will be used throughout the rest of this discussion. The functions in this family that entail computing the union or intersection of private subsets are particularly troublesome examples of the family as a whole and stand out as the two most problematic instances of the family. In point of fact, such processes are relevant everywhere in the globe and can be utilized in a wide variety of settings.

II. DATA MINING

Data mining is the process of sifting through massive amounts of previously accumulated data in order to identify patterns that were not previously known about. In the process of turning raw data into knowledge, the technique known as "data mining" is turning into a stage that is becoming an increasingly crucial component. [Take a look at this great illustration:] [Take a look at this great illustration:] The database makes it possible to extract data sets of any size, which can then be mined for insights. This mining process can take place at any time. This can be done whenever it is convenient. It is feasible to unearth buried patterns, but it is unable to find patterns that are not already dependent on the data gathering. [Case in point] Due to this constraint, the tool cannot develop into one that is truly useful. Applications such as scientific research, marketing, and the identification of fraudulent behavior [5] are all examples of common uses of this technology. The practice of data mining, which comprises the extraction of current and relevant information from unstructured data, has developed as a useful instrument for the analytical processes and decision-making procedures that are carried out within businesses. It is essential that both types of data be protected before being made accessible to the general public. This is because sensitive information and confidential regulations are often maintained within huge data repositories. The research field of "privacy preserving data mining" has become a focal point in the domains of data mining and database security as a direct result of the conflicting priorities of "data sharing," "guarding privacy," and "knowledge discovery," all of which are entirely incompatible with one another. This is a direct result of the fact that all three of these priorities are completely incompatible with one another. Because of this, "privacy preserving data mining" has emerged as a central topic of discussion in the fields of database security and data mining respectively. An issue that is associated with PPDM is the security of personal information as well as the safety of

sensitive rules (knowledge) that are stored inside the data. This is an issue since PPDM is tied to the data. Both of these obstacles are quite tough to circumvent. The first method determines how to acquire regular mining results when private information can be accessed precisely; the second method determines how and when to prevent sensitive rules contained in the data from becoming discovered while still allowing non-sensitive rules to be mined properly [6]. The first method determines how to acquire regular mining results when private information can be accessed precisely; the second method determines how and when to prevent sensitive rules contained in the data from becoming discovered. In the first way, it is determined how to achieve frequent mining results when precise access to confidential data is not possible. The first method studies how to produce consistent mining results in circumstances in which private data cannot be accessed in a precise manner. This technique is intended to be used in conjunction with the second strategy. The first strategy is figuring out how to get normal mining outcomes under conditions in which private information can be accessed with pinpoint accuracy.

Data mining is a technique that has been developed in recent years and is known as "data mining." The objective of this strategy is to unearth hidden patterns and trends that are buried within massive volumes of data. To get started with the process of data mining, the first thing that needs to be done is to assemble all of the relevant data in a single spot. This is a prerequisite for continuing. After that, the data are processed by the appropriate algorithm, and the procedure is then carried out an unlimited number of times until the desired result is achieved. The terms "knowledge discovery in databases" and "data mining" are often used interchangeably by a sizeable portion of the general population. This is not an accurate representation of the two concepts. Another phrase that sees a fair amount of usage is "Knowledge Discovery in Databases" (KDD). On the other hand, a great number of people are of the belief that the data mining process is the aspect of KDD that has the most weight and should be given the most attention. There are three different procedures that need to be adhered to in the great majority of scenarios. Pre-processing is the name of both the first stage in the process, which is completed directly before applying data mining algorithms to the appropriate data, as well as the name of the phase in which this step is included. Pre-processing is also the name of both the initial stage in the process and the name of the phase in which this step is included. During the preprocessing step, a number of tasks, including transformations, selection, integration, and cleaning of the data, are carried out [7]. Other duties include

integrating the data. The mining of previously buried material and the discovery of new information both rely heavily on the data mining process, which is the most essential component of both processes. In the course of this operation, a range of different analytic methodologies are performed in order to discover data that had not been made available to the general public in the years prior to this one. After that follows the article, which consists of examining the findings of the data mining based on the requirements of the client and the expertise that is already accessible in the relevant subject. The prior step is immediately followed by this one in the sequence. If the inquiry turns up results that are favorable to one's position, then the material can be disclosed to the proper persons once it has been uncovered. If, on the other hand, the result is not to our liking, we will have to start over and do any one of these stages, or even all of them again, until we reach a point where we are pleased with the results. The following is a summary of the methods that are utilized in the process of data mining on the most regular basis:

- Clustering:

The following constitutes the overarching principle for associations:

- Sequential patterns

Arrangements of synthetic brain tissue in the form of networks (also known as ANNs)

"Structures of computations that are derived from genetic material."

Making use of decision trees throughout the process:

The process of determining who among your neighbors lives the nearest to you.

The establishment of guidelines:

- Data visualization

3. a database that may be accessed from a number of different locations

A database management system is something that does not share any of its physical pieces and is made up of websites that are only indirectly connected to one another. It is made up of websites that are only indirectly connected to one another. A decentralized database system is another name for this particular category of computer system. The database

management systems that are utilized on each website are entirely unique from one another and cannot in any way be compared to one another; this is the case regardless of the context. There is no way that these two systems can be compared to one another. Transactions have the capability of acquiring data from either a single location or simultaneously from a number of different locations. A distributed database management system, or DDBMS for short, is a piece of software that manages a distributed database, also known as a DDB, and provides users with an access method while concealing the fact that the database is distributed from the users [9]. A distributed database is also sometimes referred to as a DDB. The terms "distributed database" (DDB) and "distributed database management system" (DBMS) are often used interchangeably.

A distributed database management system is a form of data structure that is responsible for managing distributed databases. Distributed databases are a type of data structure. Traditional datasets can be differentiated from distributed databases by the fact that the storage devices in distributed databases are not all connected to a central processing unit. The administration of distributed databases is under the purview of distributed database management systems (CPU). When discussing distributed database systems, it is customary to refer to distributed database systems and distributed databases collectively [11]. This is because distributed database systems and distributed databases are essentially synonymous with one another. There is a large selection of distributed database types available for users to select from.

A centralized database that maintains a format that is consistent across multiple locations

III. DATABASE THAT FEATURES A STRUCTURE THAT IS NOT JUST DIVERSIFIED BUT ALSO DECENTRALIZED.

Each of the websites that make up a homogeneous distributed system makes use of the same piece of software, is aware of the existence of the other sites, and has come to an agreement to handle user requests in a cooperative manner. Additionally, each of these websites is aware of the presence of the other sites. In addition, each website is fully aware of the existence of the other sites in the network. When a website gives up the ability to modify the schema or the software, it gives up a piece of the control it originally had over its own domain. This can be seen of as a loss of autonomy. It is possible to convey the impression that there is only one system in place by adopting a database

management system that is consistent throughout the organization. This is something that can be accomplished. In comparison to the consistent system, the alternative is not only more complicated to develop but also more straightforward to administer.

A heterogeneous database system is a type of information (or semi-automatic) solution that integrates many types of heterogeneous data management in order to supply a user with a single, unified query interface. This type of solution is also known as a distributed database system. The term "distributed database management system" can also be used to refer to this category of system (DBMS). This group of computer systems can also be referred to as a "distributed database management system," which is another title for it (DBMS). Another name for this type of system is a "distributed database system," which is also a word that may be used to refer to this sort of system. The computational models and software implementations that make it possible to integrate heterogeneous databases are referred to collectively as heterogeneous database systems, or HDBs for short. The combination of several kinds of databases is what's meant to be referred to by the adjective "heterogeneous."

IV. ADDITIONAL EXERCISES THAT ARE ASSOCIATED WITH THIS

The mining of association rules identifies potentially noteworthy linkages and/or correlations between vast groupings of data components. [Case in point:] [Case in point:] [Case in point:] [Here's a good example:] [Here's a good example:] One method for accomplishing this objective is to compare and analyze the data in order to search for patterns in the information. The principles of the association lay forth the components of a collection, such as its values and the events that surrounded its acquisition, that are most frequently encountered in conjunction with one another and are hence referred to as "components that go together." In order to conduct out a market basket analysis, a decentralized system employed a technique known as association rule mining. This was done for the goal of gathering data. Mining for association rules is a technique that can be used to uncover rules that can predict the occurrence of an item based on the occurrences of other objects in the transaction. These rules can be used to make predictions about future transactions. This objective can be attained by investigating the relationships that exist between the various elements of the scene. [10], search patterns gave association rules where the support will be counted as the fraction of transactions that

contain both an item X and an item Y, and confidence in a transaction can be measured by determining not just whether an item I appears in a transaction that also contains an item X. [11], search patterns gave association rules where the support will be counted as the fraction of transactions that contain both an item X and an item Y. [12], search patterns gave association rules where the support will be counted as the fraction of transactions that contain both an item X and an item Y. [12], search patterns provided association rules, and the support will be counted as the fraction of transactions that contain both an item X and an item Y. [12], search patterns provided association rules, and [13], search patterns provided association rules. [12], search patterns led to the discovery of association rules where the support [11], search patterns offered association rules, and the support will be counted as the fraction of transactions that contain both an item X and an item Y. [12], search patterns supplied association rules, and [13] search patterns gave association rules. Both [12] and [13] indicated that search patterns offered association rules. [12] indicated that search patterns offered association rules. [12], search pattern analysis led to the finding of association rules where the support was located. [You are needed to provide references and citations] The Apriori Algorithm has the potential to explore a given data set and find items that appear frequently by making use of the ant monotone constraint. This capability is made possible by the Apriori Algorithm. These results can be obtained since the program possesses search capabilities. A well-known example of an algorithm that mines common item sets for Boolean association rules is known as the Apriori algorithm. One well-known example of the algorithm is the Apriori algorithm, which relates to the subject of market basket analysis and serves as an example of the algorithm. This method, the specifics of which can be found in [13], involves carrying out an extensive quantity of searches across the entirety of the database in an iterative fashion. This method searches for the minimal support criterion in each and every viable frequent item set of length k. It will keep searching until it finds one that fits the criterion, at which point it will stop exploring. The pass k process is the name that has been given to this particular kind of operation. Apriori was developed to work in tandem with databases that include information about many types of financial activities. The Apriori Algorithm's [11] primary objective is to discover connections between separate data sets by investigating the areas in which these data sets are similar to and diverge from one another. [Citation needed] This method is usually known as a "Market Basket Analysis," which is a word that is frequently used interchangeably with notions that are associated with it. A transaction is a collection of data that has its own unique name and is made up of a variety of

different parts and pieces. Transactions can be thought of as separate data collections in their own right. Following the execution of Apriori, we are presented with sets of rules that, when taken together, provide an estimate of the frequency with which particular objects are found in the data collection under consideration. The rule of association Mining is used to establish rules that would forecast the occurrence of an item, and based on the occurrences of other items in the transaction, search patterns revealed insights into the relationship between the occurrences of the items. Association rules that state that support will be counted as the proportion of transactions that contain both an item X and an item Y, and that confidence can be measured in a transaction if an item I appears in a transaction that also contains an item X. [Citation needed] Mining association rules is used to uncover rules that will forecast the existence of an object, and mining search patterns has given us a lot of information about how people search for things. Mining association rules is a technique that is used to find rules that will forecast the occurrence of an item, and mining association rules is a technique that is used to find rules that will predict the occurrence of an item. In addition, search patterns included association criteria as a method of gauging support. This criterion was defined as the percentage of transactions that included both item X and the criterion being measured. Those Who Are in Favor Of: - The percentage of total commercial activities that comprise X and Y activities as separate components. - The percentage of total commercial activities that include both X and Y activities. Those who Support - The percentage of total commercial activities that contain both X and Y activities as independent components. - The percentage of total commercial operations that include both X and Y activities. Helping someone out ($X \rightarrow Y$) can be broken down into the following formula: $P(XY)/T$ Determine the frequency with which products from Y are included in deals that also contain X, and determine the frequency with which this happens by first determining the frequency with which it does happen. Confidence ($X \rightarrow Y$) = $\text{Support}(XY) / \text{Support}(X)$

The data are combed through in search of recurring if-then patterns, and after that, support and confidence are used to hone down on the connections that are the most important to concentrate on. The establishment of association rules is going to be the end result of carrying out this approach to its logical conclusion. The quantity of support that something receives is directly proportional to the number of times that it is entered into the database. There is a one-to-one relationship between the two. A ratio that is exactly proportionate to the number of times that the if/then

statements have been demonstrated to be correct can be used to describe the amount of certainty that can be obtained from the data. This ratio can be used to describe the amount of certainty that can be obtained from the data.

V.THE FAST DISTRIBUTED MINING ALGORITHM COMES IN AT NUMBER FIVE IN THE OVERALL RANKINGS.

The protocols are organized in a fashion that is comparable to the method known as Fast Distributed Mining (FDM), which is outlined in [2.] The protocols are constructed with this algorithm serving as their basis and being developed upon it. The word "FDM" refers to a variation of the Apriori technique known for its distributed nature but lack of security features. Any itemset that is recognized as being s-frequent must, in addition, be considered to be locally s-frequent at least one of the sites. This criterion has to be satisfied in order to proceed. This is the primary idea that functions as the basis for it. In order to acquire information about all of the s-frequent item sets that might be found in other regions of the world, it is necessary for each player to reveal the s-frequent item sets that are situated inside their own zone. After this has been completed, the remaining players will examine each of these item sets to determine whether or not they are also s-frequent on a global scale. In the event that it is required to do so, the FDM process can be divided into the following stages:

1.Initialization: It is assumed that all of the participants have already coordinated their efforts in order to compute F_{k-1} s. This is the case because it is required for the task. This must be completed before moving on to the next stage.

At this juncture, the objective is to carry on with the computation of F_k s and bring it to a fruitful and satisfactory finish.

2.The Production of Possible Candidate Groups: Using the Apriori technique, each P_m generates a collection of potential k-item sets called B_k , m s. This process begins with F_{k-1} , m s F_{k-1} s. These possible k-item sets consist of the (k 1)-item sets that are shared by the global and the local contexts.

3.Editing at the Neighborhood Level If X is less than B_k , then the m and s values need to be deducted from the total. P_m is the component that is in charge of computing $\text{supp}_m(X)$, and it only keeps track of the item sets that are considered to be locally s-frequent. Moreover, it is the only component that computes $\text{supp}_m(X)$. The letters C_k , m, and

s stand for the collection of movable and immovable sets that can be obtained through various means.

4. Combining the multiple candidate item sets Once all of the players have disclosed their C_k ms, the players who are still participating in the game will next compute their respective C_k s. = $SM_{m=1} C_k, ms$.

5. Computing the local supports At this stage of the game, all of the players are accountable for calculating the local supports of each and every itemset that is a component of C_k s.

6) The next stage, which all players must complete, is to broadcast the local supports that they have computed for themselves in the previous phase, referred to as "Broadcast Mining Results." On the basis of this knowledge, any individual is capable of computing the global support for each and every itemset that is contained in C_k s. Last but not least, the subset of C_k s known as F_k s is the one that includes all of the globally s-frequent k-itemsets. This subset may be found here. This particular subgroup is referred to as F_k s.

It is essential to conduct research into efficient methods for the distributed mining of association rules given the availability of a large number of large transaction databases, the enormous volumes of data, the high scalability of distributed systems, and the ease with which a centralized database can be partitioned and distributed. In light of these factors, it is essential to conduct research. Conducting research into effective ways for the decentralized mining of association rules is another activity that is absolutely necessary. The findings of this research have put light on a number of intriguing relationships that exist between locally big item collections and globally large item collections. In addition to that, it provides an unusual method for mining association rules known as FDM, which is provided (Fast Distributed Mining of association rules). When it comes to mining for association rules, this method generates a manageable number of candidate sets and drastically reduces the amount of communication that needs to take place between processes. According to the findings of our investigation into performance, FDM is superior in terms of output quality to the straightforward implementation of a conventional sequential technique. This was discovered through our investigation into performance. This was determined by looking at the two strategies side by side and contrasting them. When further work is done to improve the performance of the algorithm, many distinct variants of the algorithm are developed as a result.

VI. CONCLUSION

The challenge that arises from attempting to compute association rules while working with a database that is completely consistent all the way through. Assume that each website use the same schema, but that the individual websites do not divulge any information regarding the various entities. The purpose of this endeavor is to establish association regulations that are applicable in all sites while at the same time putting restrictions on the amount of information that is shared for each location. Recent events have resulted in the implementation of a sizeable number of protocols. The well-known method of association rule mining is implemented in this case, with the primary emphasis being placed on data that is horizontally partitioned and spread across multiple locations. The underlying problem is only of interest when there are more than two actors participating in it; as a result, protocols take advantage of the fact that there are more than two actors involved in the situation.

References

- [1]. Tassa, T. (2013). Secure mining of association rules in horizontally distributed databases. *IEEE Transactions on Knowledge and Data Engineering*, 26(4), 970-983.
- [2]. Jin, Y., Su, C., Ruan, N., & Jia, W. (2016, November). Privacy-preserving mining of association rules for horizontally distributed databases based on FP-tree. In *International conference on information security practice and experience* (pp. 300-314). Springer, Cham.
- [3]. Marodkar, N., & Chaudhari, M. Mining of Association Rules in Distributed Databases.
- [4]. More, V. R., Patil, B. K., & Auti, R. A. (2017, March). Secure extraction of association rules in horizontally distributed database using improved UNIFI. In *2017 International Conference on Big Data Analytics and Computational Intelligence (ICBDAC)* (pp. 205-210). IEEE.
- [5]. JAGINI, N., & RAVI, C. (2014). Secure Association Rule Mining in Horizontally Distributed Databases using THRESHOLD-C and UNIFY Protocols.
- [6]. Shaikh, I. R. (2015). Secure Extraction of Association Rule from Distributed Database. *International Journal of Computer Applications*, 117(3).
- [7]. KRISHNAMACHARI, S. M., & RAO, A. N. (2014). Mining of Association Rules in Parallel Distributed Databases with Multiparty Algorithms.
- [8]. Thiripal, K., & Harshavardhan, V. A Novel Association Rules for Secure Mining in Horizontally Distributed Databases.
- [9]. Jyothi, B., & Rao, K. V. Secure Mining of Association Rules in Horizontally Distributed Databases (Protecting Sensitive Labels in Social Network Data Anonymization).
- [10]. Sumana, C. M., & Rajeswari, R. P. AN EFFICIENT FDM-KC ALGORITHM FOR SECURE MINING IN HORIZONTALLY DISTRIBUTED DATABASES USING ASSOCIATION RULES