



Web User Analysis Using Hierarchical and Optimize K-meanAlgorithm for Online Market Analysis.

Dr. V. Venugopal

Associate Professor, Department of MECH

Sri Sai Institute of Technology and Science, Rayachoti

Email: venugopalphd@gmail.com

Abstract

Data mining plays really very important part in the analysis of business. Computational statistics and information retrieval has attracted attention of many researchers. As there is enormous data there are various challenges for the researchers to extract particular data as per users requirement. Considering large volumes of data which is produced due to various users visiting online web portal clustering of that data needs to be done and categorizing of that data on basis of some similar attributes of the users. When standard K-means partitional algorithm is not that efficient because when it is subjected to large data set it consumes considerable time. So use of Optimized K-Means and Hierarchical algorithm is to be done which will yield better results for on-line shopping web portal. Reports are generated from output of clustering algorithm. Results can be helpful for business owners or website owner in market research for deciding marketing strategies, customer retention strategies, production and operations taking place in the business. .

Key terms

Data mining, partitional algorithm, Market research

algorithm, hierarchical algorithm, K-means, optimized K-Means algo-

1. Introduction

1.1 Existing System

Data Mining Can be Used in Various business application for different purposes such as decision support system, customer retention strategies, selective marketing, business management, user profile analysis to name a few [1],[2],[3],[10]. Data mining is the process of discovering the knowledge. In today's electronic information era it becomes highly challenged to digital firms to manage customer data to retrieve useful information as per their requirement from that data, so market segmentation can be used. Market segmentation also includes customer retention strategies, allocation of resources for advertising, to check profit margins. So outcome of segmentation plays a big role in deciding price of the products, attracting new customers and identifying potential customers. Clustering analysis is able to find out data distribution and proper inter relationship between data items. Clustering is defined as "grouping of similar data"

. Clustering divides records in the database or data objects in the dataset into series of meaningful subclasses or groups. Data mining is basically a useful process in which the enterprise retrieves useful knowledge from a load of in-

formation which is incomplete and random that has been generated from various business tasks such as production, marketing, customer services of the enterprise.

A good analytical tool should be able to compare between different characteristics or attributes of different groups and identify different important characteristics of each segment to decide different business strategies. Clustering analysis can find out the distribution of different data

entities as well can find out proper inter relationships between the data objects so that it can divide the data set into series of meaningful subclasses.

The motivational state of the art:

Limitations of K-means algorithm:

The constraint on K-means algorithm is that if a data point is included in one cluster then it cannot be included in another cluster. K-means algorithm has significant disadvantages [9][4]. Time taken by K-means algorithm when large amount of data is applied to it is considerable. The result or the output i.e. clusters formed using K-means algorithm are not always same though same data is applied to the K-means algorithm. It means output is different for each run because result depends upon random initial assignment of initial k clusters. The application helps to determine users visiting details, which products he/she has viewed and which products he/she has purchased and keeps track of customer population visiting the online shopping portal. To avoid limitations of traditional K-means methods such as output dependency, so there is increasing need of improved clustering algorithm which will yield same result but will take less amount of time to process large amount of data in online shopping portal. So the optimization of K-means algorithm is to be done.

It becomes highly challenged to digital firms due to the increasing growth of customers data in today's electronic business world to analyse data. A good an-

alytical tool should be able to analyse the data by comparing the characteristics of different data entries so as to take very important business decisions.

This paper describes the existing system in section 1 and describes motivation of the new one (why there is need of new improved algorithm). Section 2 describes work related to the proposed system which is nothing but literature survey, associated challenges, drawbacks, efficiency of the system and experimental setup and software requirement specification. Section 3 describes Programmers design input to the system and outcome of the system. Section 4 is about results and validation of outcomes and section 5 describes conclusion.

2. Related Work

2.1 Literature survey:

Traditional K-means algorithm is traditional partitioning algorithm and can be used as clustering algorithm. and it is most popular algorithm because of its simplicity. Data Mining Can be Used in Various business application for different purposes such as decision support system, customer retention strategies, selective marketing, business management, user profile analysis to name a few [1],[2],[3],[10].

The premise of data mining is that it is necessary to establish a customer-information data warehouse which consists of customer data and will guide the business managers for designing better customer strategies as well as better customer retention strategies.

Data in digital business firm is increasing exponentially in today's electronic fast world. Hence, to manage such data, market segmentation is solution. Customer retention strategies [5], allocation of advertising resources and increase in profit margins are included in market segmentation [2],[5].

For that a comparisons of different characteristics of different segments(group) is to be done so as to identify important attributes of each group. This can be done using a good analytical tool. This tool will give opportunity for targeting customers from different locations, purchasing different products, having numerous different attributes and so on. Cluster analysis can yield the same result as mentioned above. clustering is nothing but "grouping similar kind of data together". Using clustering and the analytical tool proper interrelationships between data points can be found out.

Existing system uses traditional standard clustering algorithm such as K-means algorithm[6]. It is most commonly used traditional method for categorizing or clustering data and it is exclusive clustering algorithm. K-means uses concept of unsupervised learning. In K-means algorithm each centroid of the cluster is assigned by each data point in the data set to which ever centroid is nearest. The centre or centroid is computed by averaging all points in the cluster. Centroid has co-ordinates in the space and that co-ordinates are computed by taking arithmetic mean for each dimensions separately for all data points in the specific cluster[6].

K-means algorithm is simple and having high processing speed even though it is applied to large data sets but has significant disadvantages as discussed in [8][4].

K-means algorithm is simple and having high processing speed when applied to large amount of data. k-means calculates centroid of the clusters by averaging the data points in the data set.

Limitations of traditional K-means methods such as output of k-means algorithm depends upon depends upon order in which input are given and K-means tends to result in local minimum and limited applicability to only the data set consisting of isotropic clusters

.When K-means algorithm is used there is need to adjust the sample space continuously .so, we should calculate the centres of the clusters again and again. If data set applied to the K-means algorithm is large then more time is consumed by algorithm to produce the clusters. So more fast algorithm is required and there is need to improve efficiency of K-means algorithm. In K-means method while assigning each data object in the data set D to the cluster centres which are initially assigned some calculation regarding distance of the data object from the centre of cluster which was randomly assigned initially is to be made but if the distance of the data object is far from the centre then there is no need to measure the exact distance between that data point and centre of the cluster in order to know that that point

should not belong to the respective cluster. This was case for one data point as data set is large then such distance calculations are unnecessary in K-means method

.and these redundant calculations should be avoided in order to increase processing time of the K-means method.

2.2 Related algorithm

We are using optimized version of K-means algorithm. Hierarchical algorithm, OKH algorithm .all the algorithm are as follows: Optimized K-means Algorithm:

We use triangle inequality theorem to reduce redundant calculation in K-means algorithm. The triangle inequality theorem says that the sum of two sides of the triangle is always greater than the remaining side of the triangle. this theorem is applied to Euclidean space[9] and can be extended to multidimensional Euclidean space.

Take three random vectors in Euclidean space :p,q,r Then

Then

$d(p,q) + d(q,r) \geq d(p,r)$
 $d(q,r) - d(p,q) \leq d(p,r)$
 $d(A_i, A_j)$ is the distance between the two cluster centres namely i, j .
 if
 $2d(p,q) \leq d(A_i, A_j)$ then
 $2d(p, A_j) - d(p, A_j) \leq d(A_j, A_k) - d(p, A_j)$(1)
 Then according to above equation (1) $d(p, A_j) \leq d(p, A_k)$

The optimized K-means algorithm is as follows: 1) select centres for clusters and set lower bound $y(p,f)=0$ for each data object as well as for centre of the cluster. 2) Now start assigning each data object to centre of cluster by computing distance between cluster centre and data object [3]. we can use previously obtained data to avoid redundant calculations. each time $d(p,f)$ is calculated and set its value to $y(p,f)$.

Hierarchical Algorithm:

As the name of the algorithm suggests the output of hierarchical algorithm [14] is nothing but hierarchical structure formed using data entities. Single link or complete link hierarchical algorithm can be used. Consider the example below. With this we can get different hierarchical structure with different similarity requirement. From fig. 2.1 if the similarity requirement is set at

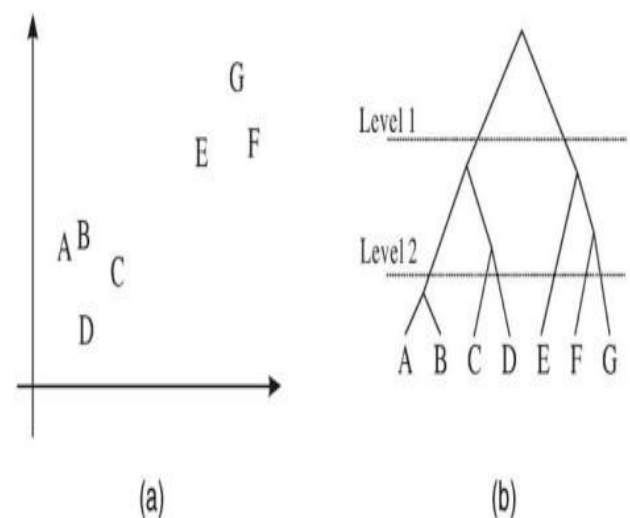


Figure 1: Level 1 Illustrative hierarchical clustering results for a data set of seven points. (a) The input (b) A possible Hierarchical tree

level 1, the input data set is partitioned into different clusters. As given in the above figure A, B, C, D, E, F, G are data objects given. when similarity requirements are given at level 1, the input data objects are divided into 2 clusters i.e. (A, B, C, D) and (E, F, G). when similarity requirement is set at level 2 then the data objects are divided into six clusters i.e. (A), (B), (C), (D), (E), (F) and (G). Most hierarchical clustering algorithms are variations of the single-link and complete linkage hierarchical algorithm. By good quality of clustering output.

The outline of a general hierarchical algorithm steps: Single link or complete link hierarchical algorithm can be used.

1) each data point forms a cluster.

2) and algorithm merges to closest cluster repetitively. 3) Hierarchical structure is outputted.

OKH Algorithm:

This algorithm is nothing but combination of optimized k-means algorithm and hierarchical clustering algorithm.

Algorithm OKH:

Input: data set D containing data objects, size of data set s , number of sub clusters m , desired number of clusters n

Output: hierarchical structure of n clusters.

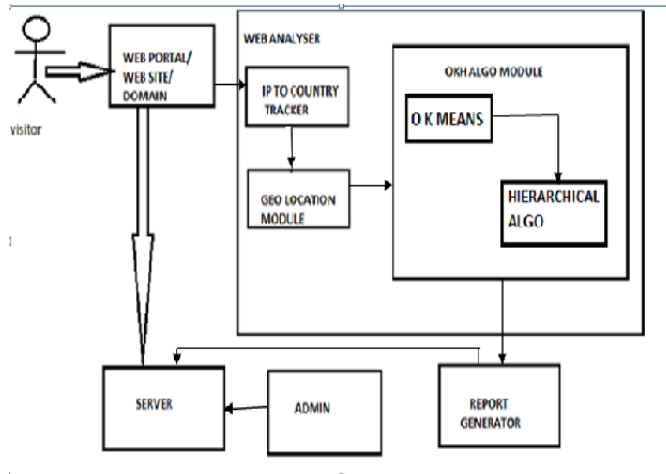


Figure 2: Architecture of web user analyser

- 1) apply O K-means on input data set to get m subcluster.
- 2) Apply single link algorithm on m subcluster to get n cluster based on some similarity.

Experimental setup including SRS :-

The proposed system is an Analytical tool developed for online shopping web portal. we the system builders are third party who are going to develop a system that will help a particular domain to use its users data in order to expand their business. developers are located at the server site of the domain. The users Of the domain can be anywhere in the world. We are going to collect data of users of particular domain ,then we are going to apply clustering algorithm on it and are generating reports according to purchase, enquiries, browsing details of users and then are sending reports to particular domain so that domain owner can use these reports for deciding business expansion and customer retention strategies. Instead of manual surveying the system is providing domain with ready reports for deciding business strategies, as the existing system uses K means algorithm for intended purpose

The experimental setup of overall system is as shown in figure 1.

Instead of manual surveying the system is providing domain with ready reports for deciding business strategies, as the existing system uses K means algorithm

for intended purpose and its performance degrades when subjected to large datasets. proposed system is efficient, fast and performs well when subjected to large datasets as it is using new improvised optimized k means hierarchical algorithms.

Normal Requirements:

They are also called as basic requirements. they are explicitly expressed by customers and may be simply placed into the SRS. Following are some of the normal/basic requirements.

NREQ1:

Description : System can be used by more than one user at the same time.

Other factors affecting the requirement: access control/privileges given by admin to different users.

NREQ2:

Description : System Is Supposed To Be Connected To Internet Through Wi-Fi Or Dial Up Or Broad Band Connection. Other factors affecting the requirement: Processing Speed

NREQ3:

Description: System Is Supposed To Work For Users Who Have Moderate Amount Of Computer Experience (E.G. Excel, Analytical Tools Etc.)

Other factors affecting the requirement: Versions Of Softwares (Excel 2003/2007)

NREQ4:

Description : Backup Is To Be Maintained At Regular Time Intervals So As To Avoid Loss Of Data During Crashes.

Expected Requirements

These are implicit but expected requirements. they are regarded by customers as obvious and rarely are placed in the first version of the SRS, however if they are not fulfilled, the customers are dissatisfied. Following are Expected Requirements: EREQ5:

Description: Owner of the domain as well as all authorized persons are allowed to use the output (Reports) of the system.

Other factors affecting the requirement: access control/privileges given by admin to different users.

EREQ6:

Description : System can be run on any platform

EREQ7 :
Description : Should Be Able To Handle Defined Amount Of Data Approx. Number Of Records Required

Initially Plus Anticipated Growth Eg 1,000 Customers Increasing At 100 Per Year; 10,000 Jobs Increasing At 1,000/Year; Etc.

EREQ8:

Description : System Should Be Maintainable. Maintainability/Future Expansion. Who Will Maintain The System In Future (Eg In-House), Any Plans For Future Developments, Etc.

Other factors affecting the requirement: Backup Facility During System Crash/Failure.

EREQ9:

Description : Source Code Should Be Provided So As To Maintain System In Future Or For Future Expansion.

EREQ10:

Description : System Should Be Reliable. It Should Have Downtime As Minimum As Possible..

Other factors affecting the requirement: Input Data Subjected To The System For Processing.

2.1.3 Excited Requirements

These are implicit and unexpected requirements. the customers may be not aware of them, but if they are fulfilled, the customers even may be excited.

XREQ11:

Description : System can be used by more than one domain concurrently.

Other factors affecting the requirement: Contents Of Domain And Server Location.

XREQ12:

Description : This System Need High Level Of Security (Eg Different Group With Different Access Rights)Is Required.

XREQ13:

Description : External Communications Should Be Involved. Eg Automated Faxing, Import Csv Data , Export Data To Particular System, Etc.

XREQ14:

Description : The Deliverables Should Contain Help Files,Manuals,Support Files,Source Code, Training To Use The System Efficiently Etc. Other factors affecting the requirement: Extra Cost Required.

XREQ15:

Description : Existing data should be converted to theformat of the current reports from proposed system.

Other factors affecting the requirement: Extra AmountOf Work Needed And Extra Time Needed.

Software Requirements: -

Operating System: windows 2xDatabase: sql server 2005

Language of Implementation: ASP .net,C sharpH/w Requirements:

Hard disk:80gb Ram:512 Mb Processor:2.1GHz

3. Programmer's design

3.1 Input to the proposed system:

Input to the the proposed system is nothing but the the various user visiting the website.so inputs are nothing but user navigation on web portal or web logs maintained by the server of the website . 3.2 Outcome of the system:

Ouput of the system is nothing but report generated by the report generator from clusters formed by applying OKH algorithm.as well as consisting details of the customer along with items he has viewed as well as items he has purchased .The proposed system consists of following components.each model does its work mentioned below.

A)Module for Admin :Whenever user visits a particular web his details (eg.ip,login detais or account details)are sent to server of the website .Admin module is used to give access control privileges to the different users of the system.. Admin performs functions like creation, deletion. upadation of account, restricted access control to other intruders or users.

B)Server Of Domain:Server of domain is nothing but the server of the web portal/wesite or domain for whichthe system is to be implemented.The server serves the user request.

C) IP to Country tracker: : Whenever user visits the domain .all the details of the user are tracked by the server(e.g. IP address,zip code,country code etc.) When a visitor requests a web page that contains tracking code, the visitor's browser will read this tracking code and call a JavaScript file installed on the web server. Once called, the JavaScript file will open a data collection session and set (or reset) a one-year, first-party cookie in the visitor's browser. This cookie will be stored in the visitor's browser until it expires after one year, or until the user deletes it.

D)Module for Geo Location : Apply this characteristics

on Google maps API for getting its geographical location (Exact longitude and latitude) And the output of that API is geographical representation of the user location Geo-Location Module helps to locate visitor's geographical location using inputs like longitude and latitude of the user from user details for locating him. E) Module for OKH algo: On that data values OKH is applied i.e. at first a hierarchical algo is applied then optimized K means algorithm is applied to produce clusters of these data values.

3.1. Mathematical Model

Web User Analysis Using Hierarchical And Optimized K-Means Algorithm For Online Market Analysis is of P Class because:

1. Problem can be solved in polynomial time.
2. GA always produces strong feature set.
3. Optimized K means hierarchical algorithm in this case gives us optimum solution.

Let S be the set of visitors, ip addresses, clusters, clustering reports $S = I, R, T, M$ where I represents the set of visitors who visits domain, which are input to web analyser, R represents the details of visitors which are input to the clustering algorithm and T are clusters which are generated as output of analyser as clustering output and M be the set of reports for respective clusters.

$I = (I_1, I_2, I_3, \dots)$

$R = (R_1, R_2, R_3, \dots)$

$T = (T_1, T_2, \dots)$

$M = (M_1, M_2, \dots)$

Input is mapped to output which is shown in the following vein diagram:

3.2. Dynamic Programming and Serialization

In proposed approach various dynamic programming aspects are used. detail is as follows:

The whole system divided into four parts according to its functionality. first part retrieves user data from navigation of user on web portal, all the data are applied to the OKH algorithm to get the desired result.

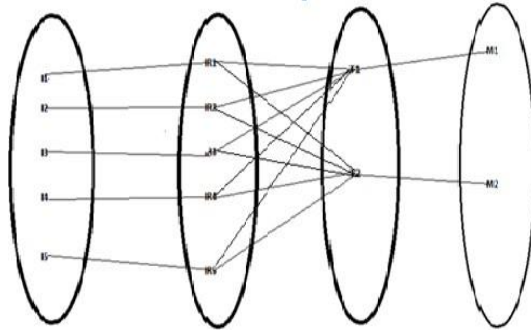


Figure 3: Venn Diagram

The first phase is again divided into subphases .To find location of the user longitude and latitude of the user is retrieved by the system then according to values of longitude and latitude city of the user will be decided. By applying those details to geo-location API.hence it makes use of optimal substructure approach of dynamic programming. With this technique problem is reduced and memorization is achieved.

In OKH phase Branch and bound method can be used for construction of clusters and for outlier detection.

3.3. Data independence and Data Flow architecture

A DFD allows you to identify the transformations that take place on data as it moves from input to output in the system.

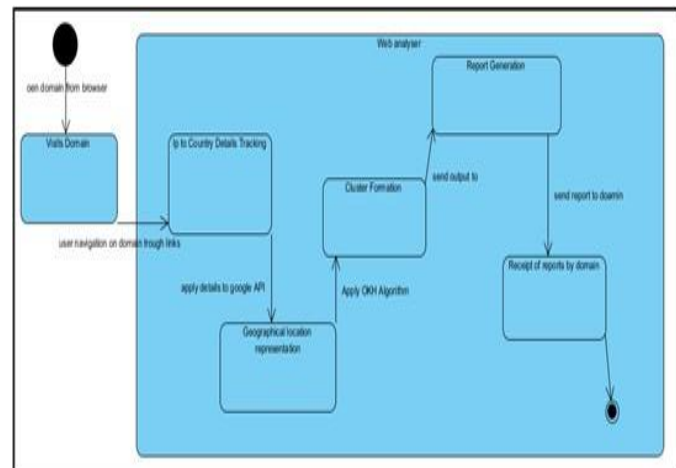
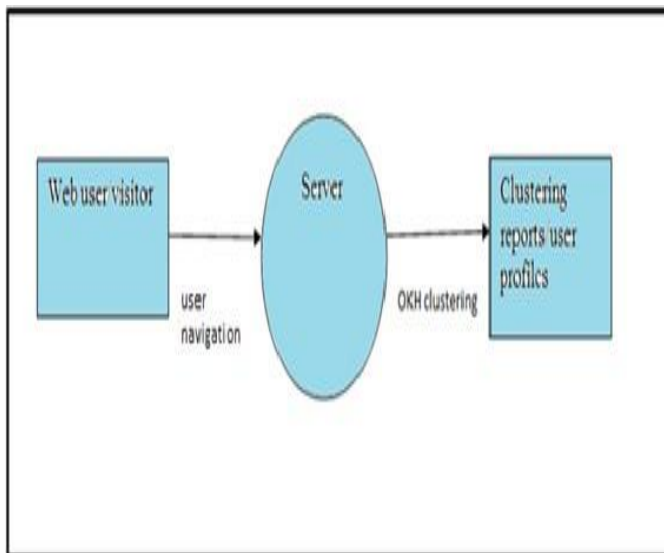


Figure 4: Data Flow Diagram

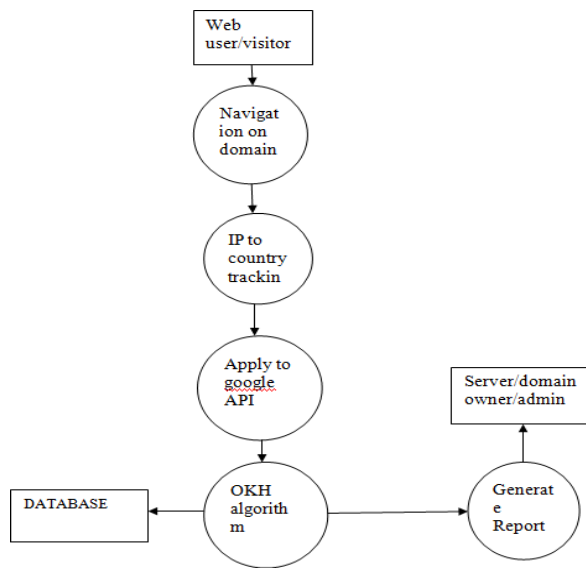


Figure 5: Level 1 Data Flow Diagram

Figure 6: State Diagram of web user analyser

3.4. Turing Machine

Current State	Description	Output	Next state
S1	Visiting domain	Web logs are tracked by domain	S2
S2	IP to country tracking	Details of user from ip addresss to country are tracked	S3
S3	Geo-location	Geographical Details of user are represented on maps	S4
S4	Cluster forming by applying OKH algorithm	Cluster formed	S5
S5	Report generator	Report generated as per user require-ment	S6
S6	Reception of report by owner of web	Admin can view details reports	[End State]

4. Results and Discussion

The performance of the proposed approach can be analyzed using five fold or Tenfold cross validation technique. For measuring performance of the proposed approach details of input to system is as follows

ID	List of Input
1	Browsing from the server located in Maharashtra
2	Browsing From Machine with IP 117.228.80.124
3	Browsing from machine with IP 10.147.1.0, 203.88.23.33

And the corresponding output is as follows:

IID	Output
I1	117.228.80.124 INDIA PUNE MAHARASHTRA IN 18.5196 73.8553 +05:30
I2	14.97.245.36, INDIA, MUMBAI, MAHARASHTRA, IN , 19.0144, 72.8471, +.5:30
I3	14.140.40.14 INDIA ,PUNE, MAHARASHTRA ,IN, 18.519 ,73.8553,+05:30

The final output can be validated by comparing quality of the outputs of the K-means algorithm as well

This application is designed by making use of improved (optimized) standard K-means algorithm. This optimized algorithm makes use of triangle inequality theorem and incorporating hierarchical clustering algorithm for robust data. This approach applies optimized algorithm on visitors data and applies hierarchical algorithm onto the products they requested. The proposed system is surely efficient than that of standard K-means Method. This system helps to identify target markets as well as customer density and potential customers of website or particular product, hence helps in target marketing.

After implementing the said algorithm time is saved for reformulation of the clusters.

References

- [1] A.G. Buchner and M. Mulvenna, "Discovery Internet Marketing Intelligence through Online Analytical Web Usage Mining," Proc. ACM SIGMOD'98, vol.27,no.4, PP.54-61, dec 1988
comparing time taken by algorithms to process same amount of data.

5. Conclusion

In this electronic information era, digital enterprises should make more or full use of information available from various information resources and instead of using product centric management model should use customer centric management model. So, business owner should focus on building

customer information data warehouses which will support business managers in deciding marketing strategies, customer retention strategies, operations as well as production strategies.

If the proposed system traces which application of the website requested by visitors then business intelligence takes place. With this real time clusters can be formed for every click of customer and it will generate reports in terms of different charts regarding the real time data. These different charts will reflect that the owner of the website should concentrate on which application (product) of the website.

M.S.Chen, J. Han and P.S. Yu, "Data mining – An Overview from database perspective", IEEE trans Knowledge and data engineering vol 5, no. 1, PP. 866-883, Dec. 1996

- [2] U.M.Fayyad, G. Piattetsky-shapiro, P. Smyth and R. Uthurasamy "Advances in Knowledge Discovery and data mining" Cambridge, Mass: MIT, 1996.
- [3] Mrs. G.P. Dharne, Mrs. S.A. Kinariwala, Mrs. A.S. Vaidya, MS.P.V. Pandit "A web user analyzer by hierarchical and optimized K-means algorithm", vol. 1, issue 7, dec. 2011
- [4] Google Latitude Website
<http://www.google.com/latitude/intro.html>
- [5] Wanghualin, "Data Mining and Its Applications in CRM", 978-0-7695-4043-6/10 2010 IEEE DOI 10.1109/ICCRD.2010.184
- [6] J. Han and M. Kamber, Data Mining: Concepts and Techniques., 2006.
- [7] Xiaoping Qin, Shijue Zheng, Tingting He, Ming Zou, Ying Huang, "Optimized K-means algorithm and application in CRM system" 2010 International Symposium on Computer, Communication, Control and Automation
- [8] G.P. Mohole, S.A. Kinariwala "Market user analyzer Using OKH algorithm" International journal of Computer application