

Web Mining

S. Murugan* and K. Kuppusamy

Computer Science Engineering Department, Alagappa University, Karaikudi, Tamil Nadu, India

Abstract

The advent of the World-Wide Web has overwhelmed home computer users with an enormous flood of information. To any point one can think about, one can discover bits of data that are made accessible by other web natives, going from individual clients that post a stock of their record accumulation, to significant organizations that work together over the Web. To be able to cope with the abundance of available information, users of the Web need assistance of intelligent software agents for finding, sorting, and filtering the available information. Beyond search engines, which are already generally used, research concentrates on the development of agents that are general, high-level interfaces to the Web. Many of these systems are based on machine learning and data mining techniques. Pretty much as information mining goes for finding important data that is covered up in ordinary databases, the rising field of web mining goes for discovering and removing applicable data that is covered up in Web-related information, specifically in hyper-content records distributed on the Web. Like information mining, web mining is a multi-disciplinary exertion that draws strategies from fields like data recovery, insights, machine learning, characteristic dialect preparing, and others. Web mining is generally comprises of the following three areas: Web content mining, Web structure mining and Web usage mining.

The World-Wide Web gives each web resident access to a plenitude of data; however it turns out to be progressively hard to recognize the pertinent bits of data. Research in web applying so as to mine tries to address these issue procedures from information mining and machine figuring out how to Web information and archives. This paper gives a brief diagram of web mining methods and examination zones, most outstandingly hypertext characterization, wrapper instigation, recommender frameworks and web use mining.

Keywords: WWW, web mining, web structure, URL

*Corresponding Author

E-mail: murugans@mail.com

INTRODUCTION

We have really touched base in the threadbare Information Age. There is an always extending measure of data "out there". In addition, the development of the Internet into the Global Information Infrastructure, combined with the colossal prominence of the Web, has likewise empowered the conventional native to end up a purchaser of data, as well as its disseminator.^[1] The Web, then, is turning into the fanciful Vox Populi. Given that there is this incomprehensible and steadily

developing measure of data, how does the normal client rapidly find what s/he is searching for—an errand in which the present day web search tools don't appear to help much! One conceivable methodology is to customize the web space make a framework which reacts to client questions by possibly collecting data from a few sources in a way which is reliant on who the client is. As an inconsequential case – a European questioning on gambling clubs is most likely better served by URLs indicating

Monaco, while somebody in North America ought to get URLs indicating Las Vegas. A scientist questioning on cricket no doubt needs an option that is other than a games aficionado would.

Existing business frameworks try to do some insignificant personalization in light of decisive data straightforwardly gave by the client, for example, their postal district, or catchphrases depicting their hobbies, or particular URLs, or even specific bits of data they are occupied with (e.g. cost for a specific stock). Our research aims at creating systems that (semi) automatically tailor the content delivered to the user from a web site. We do so by mining the web both the contents, as well as the users' interaction.

Web mining, when looked upon in information mining terms, can be said to have three operations of intrigues bunching (discovering common groupings of clients, pages and so forth.), affiliations (which URLs have a tendency to be asked for together), and consecutive examination (the request in which URLs have a tendency to be gotten to).

As in most genuine issues, the groups and relationship in Web mining don't have fresh limits and frequently cover impressively. What's more, awful models (exceptions) and inadequate information can undoubtedly happen in the information set, because of a wide mixture of reasons natural to web scanning and logging. Along these lines, Web Mining and Personalization obliges demonstrating of an obscure number of covering sets in the vicinity of noteworthy clamor and anomalies,^[4] (i.e., terrible models). Moreover, the data sets in Web Mining are extremely large.

The aim of our research is to develop scalable robust fuzzy techniques to model noisy data sets containing an unknown

number of overlapping categories. Specifically, in this work we are:^[2,3]

1. Developing new adaptable hearty fluffy grouping strategies for displaying information.
2. Exploring new strategies to handle semantic and printed elements creating so as to validate our strategies model web mining and personalization frameworks.
3. Our initial efforts have been to mine web access logs and to cluster search engine results on the fly.

WEB MINING CONCEPTS

The underlying concepts of web mining are the same as those for Data Mining only with a broader scope. Data Mining is primarily concerned with the discovery of knowledge and aims to provide answers to questions that we do not know how to ask. It is not a programmed process but rather one that comprehensively investigates vast volumes of information to focus generally concealed connections. The procedure removes top notch data that can be utilized to reach determinations in light of connections or examples inside of the information. Utilizing the procedures utilized as a part of Data Mining, Web Analyzing so as to mine applies the strategies to the Internet server logs and other customized information gathered from clients to give significant data and learning. Web Mining can be sub partitioned into the three fundamental zones of Web Content Mining, Web Usage Mining and Web Structure Mining.

WEB MINING TECHNIQUES

The common techniques for Web mining are: clustering/classification, association rules, path analysis, and sequential patterns.

Clustering/Classification

It intends to create profiles of thing with comparative qualities. This capacity upgrades the revelation of connections that are generally not self-evident. For instance, arrangement of Web access logs

permits an organization to find the normal period of clients who request a certain item. This data can be profitable when creating promoting procedures.

Association Rules

Decides that oversee "databases" of exchanges where every exchange comprises of an arrangement of things. This technique is used to predict the correlation of items "where the presence of one set of items in a transaction implies (with a certain degree of confidence) the presence of other items." For example, prediction of the percentage of clients assessing a particular URL who will place online orders for a certain product.

Path Analysis

A procedure that includes the era of some type of diagram that "speaks to" relation[s] characterized on Web pages." This can be the physical format of a Web website in which the Web pages are hubs and the hypertext connections between these pages are coordinated edges. Most diagrams are included in deciding incessant traversal examples or substantial reference successions from physical format, for example, the most habitually went by ways in a Web webpage. Case in point, what ways do clients set out before they go to a specific URL?

Sequential Patterns

Applied to web access server transaction logs. The purpose is to discover sequential patterns that indicate user visit patterns over a certain period. For example, "30% of clients, who visited /company/products/, had done a search in Yahoo within the past week on keyword W".

Web Mining Technology as a Tool

Web mining can be a promising instrument to address ineffectual web crawlers that deliver fragmented indexing, recovery of immaterial data or unconfirmed dependability of recovered

data. It is crucial to have a framework that helps the client find significant and solid data effortlessly and rapidly on the Web. Web mining finds data from hills of information on the www,^[5,6] however it likewise screens and predicts client visit propensities. This gives fashioners more solid data in organizing and planning a Web webpage. Web mining innovation can help custodians outline Web destinations with ways that can be voyage effectively by end clients, sparing time and exertion. Scenarios in this paper are used to envision what this technology could do for library reference services. These situations are intended to conceptualize new operational contemplations - if Web mining innovation is picked as an instrument to sort out data from the Internet.

Scenario 1

A gathering of custodians applies Web mining innovation to find data from the WWW that is apropos to distinctive examination intrigues as of now maintained in diverse controls in their scholastic establishment. This innovation finds data naturally and bunches it in gatherings, as indicated by branch of knowledge (e.g., subject-particular), permitting simple access. The librarians ensure that this information is from authoritative sources and of research quality by building sets of attributes (or profiles) for different subject areas. These qualities encourage the procedure of bunching and present data in consistent gatherings for precise retrieval. As an outcome, the bookkeepers invest next to no energy keeping up this database in light of the fact that Web mining innovation can find and sort out data naturally without human mediation. The end clients are fulfilled on the grounds that they can discover anything they need in a brief while without needing to screen a huge number of records. Above all, they can get

to data at whatever time and from anyplace on the planet. It is really a virtual operation in which curators exchange their reference aptitudes from the work area to the Internet.

Reality

Perfect circumstances looking like the situation specified above is once in a while the case. In the event that the above situation were genuine, no human-machine association would be required. As indicated by this situation, the machine is doing all the work and settling on all the essential choices (e.g., what sources to utilize). Sadly, regardless of how refined the framework is, the unforeseen dependably introduces it. Intricacies can be complex. The second situation recommends two contemplations that curators need to deliver to manage unanticipated issues later on.

These two considerations are:

1. What part ought to custodians play to encourage smoother operations so that library benefactors can get to important and solid data rapidly?
2. What conceivable outcomes do curators see that this innovation can be connected to improve their expert tries?

Scenario 2

A scholarly library chooses to consolidate Web mining innovation into its reference administrations (e.g., mining information from the WWW and utilizing the information to construct subject-particular databases) to help benefactors access library reference benefits whenever and anyplace.

A usage board, made up of reference and inventorying curators and additionally individuals from the PC staff, is accused of finishing the task.

This panel understands that all endeavors must be focused on encouraging machine knowledge, consequently accomplishing

the maximum capacity of the innovation, particularly in preparing reference exchanges. The experience of the librarians and their commitment to public service is essential in applying Web mining technology to enhance reference services.

To satisfy this dedication, on the other hand, curators should be ready to take an interest from the begin in the execution process. This does not suggest that custodians need to end up PC masters; however that powerful correspondence and coordination all through the procedure will bring about full support of all concerned, including commitments from individuals as indicated by individual expertise. The council chooses to handle this errand by considering two focuses:

Point A

What ought to curators do to encourage smooth operations so that library supporters can get to pertinent and dependable data quickly? To tackle this thought, the advisory group sets up meetings to generate new ideas went for distinguishing vital steps important to set up an operation in which data recovered is significant and solid. Thus, six stages are drafted. They are:

Step 1: Build up rules determining the extension for distinctive subject databases with the assistance of reference bookkeepers in their own particular branches of knowledge.

This rule would not be not at all like the gathering approach for every branch of knowledge. Sites would be underlined that are in accordance with the exploration hobbies of the foundation.

Step 2: Explore end-clients' qualities, for example, demographics, geographic areas, research targets and inquiry propensities

Step 3: Explore the system ability, plausibility for scaling, and similarity of

equipment stages between the organization and its end clients.

Step 4: Research information mining programming in the business sector including qualities, capacity, value, proceeding with change (new forms) and common sense. Most essential, the board of trustees ought to arrange a stock of programming that exists in the organization, investigating the adequacy and the coordination conceivable outcomes of focused on information mining programming with existing programming, so at last all would work gainfully.

Step 5: Set up a preparation arrangement for supporters and also library staff.

Step 6: Set up a gadget for client input and correspondence between substance makers and end clients.

Point B

What potential outcomes do custodians see that this innovation can improve their expert attempts? Numerous thoughts to increase the value of Internet assets are all around recorded. Proposals, for example, giving "bundled" answers, seek help, and basic assessment of applicable assets – are normal. Other than these thoughts, the advisory group needs to add one more advantage to this database, which is to bolster and support a scholarly digital trade environment.

For years, groups of prestigious scholars have shared preprints and findings in conferences to advance their research interests. Many professionals have benefitted from that sharing. Traditionally, however, libraries do not participate in such circles except for subscribing to publications generated by professional societies and conference proceedings.

As innovation develops and changes, the circle of imperceptible universities expands. Right now, numerous learning circles are being created in the internet and are expanding their participation past built up insightful gatherings.

The board of trustees sees that scholarly trade exercises are not that unique in relation to those occurring inside of the physical environment of the library. It sees the open door for beginning new undetectable universities in the internet. Individuals think it might be a decent time for the library to take an interest in an imperceptible school and to wind up more included in the procedures included in creating data. Along these lines, it consents to give an instrument to encourage correspondence between substance makers (analysts) and end clients (who may be scientists themselves). The thought is to let researchers (or invested individuals) build up contacts they could call their own, so that individuals with comparable premiums can share encounters, thoughts, disappointments and triumphs. To encourage correspondence, Web-based meeting programming (e.g. Discussion) is coordinated into every subject door. Content makers and clients can convey straightforwardly, either in a gathering setting or exclusively (guaranteeing security). Data shared can be presented on the database for open utilization or be ensured for access by individuals' just.

Formal or casual productions created can be incorporated in the database, which will reinforce the passage itself by giving looked for after data. The board of trustees additionally suspects that keeping insights on subject watchwords for every subject entryway may oblige the survey of exploration necessities. Case in point, the subject of wolf correspondence may indicate numerous solicitations are made

to the creature science door. This may recommend that interest for data on this theme exists. Specialists can then contrast this data with their own particular exploration intrigues. Some may discover this welcome new research heading and begin to take a gander at this territory.

Result

The advisory group uses Web mining innovation to manufacture a database that encourages element research exercises. More vital, the library is at long last a piece of the undetectable school, permitting it to tackle another part in the changing data age on grounds.

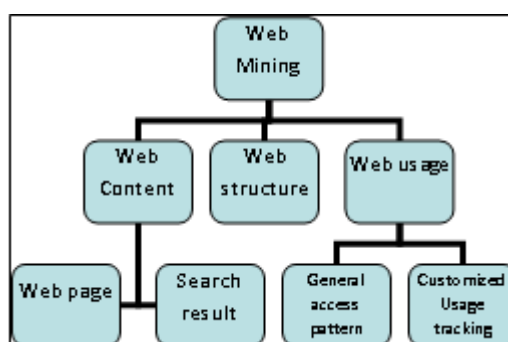
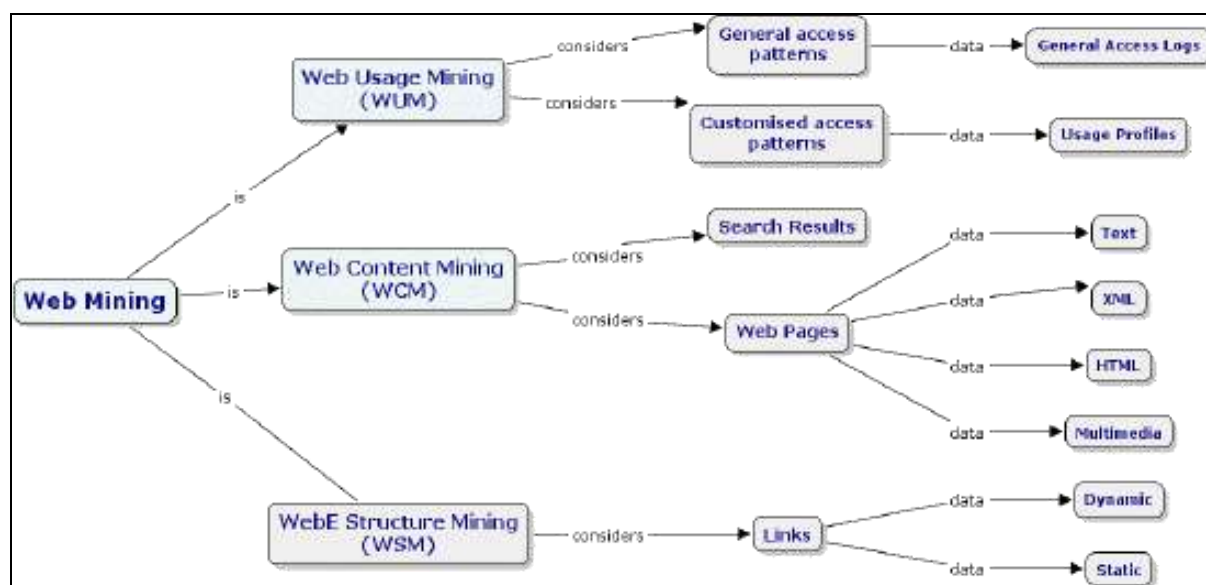


Fig. 1: Web Mining Process.

Web Structure Mining (WSM)

This strength means to uncover the genuine structure of sites through the social event of structure related information, and for the most part about its network. Ordinarily it considers two sorts of connections: static and element.

Web Content Mining (WCM)

Web content mining is the "process of information or resource discovery from millions of sources across the World Wide

Web". There are two methodologies in Web substance mining: the specialists based and database approaches.

Web Content Mining is the processing of information, or resource discovery, from multiple sources across the Internet.

1. An agent based approach may be used where a software agent acts autonomously in performing the analysis.
2. A database approach may also be used where an enterprises operational data

stores and Internet based information are processed into more structured high level collection of resources.

The agent-based approach involves artificial intelligence systems that can "act autonomously or semi-autonomously on behalf of a particular user, to discover and organize Web-based information ". Some intelligent Web agents can use a user profile to search for relevant information, then organize and interpret the discovered information (e.g., Harvest). Some use various information retrieval techniques and the characteristics of open hypertext documents to organize and filter retrieved information .Another kind of agent is programmed to learn user preferences and use those preferences to discover information sources for those particular users.The database methodology concentrates on "incorporating and sorting out the heterogeneous and semi-organized information on the Web into more organized and high-level collections of resources." These organized resources can then be accessed and analyzed. These "metadata, or generalizations, are then organized into structured collections (e.g., relational or object-oriented databases) and can be analyzed".

Its goal is gathering data and identifying patterns related to the contents of the web and the searches performed on them. There are two main strategies: Web Content Mining

1. **Web page** mining, extracting patterns directly from the contents those exist in web pages. In this case the data in use can be.
 - i. Free text
 - ii. HTML pages
 - iii. XML pages
 - iv. Multimedia elements
 - v. Any other type of contents existing in the web site.

2. **Search results mining**, intending to identify patterns in the results and the search engines.

Web Usage Mining (WUM)

The other measurement of information mining is Web utilization mining. This is the procedure of finding client access examples (or client propensities), as information are naturally gathered in every day access logs. As of late, referrer logs, which gather data about alluding pages for every reference and client enrollment, likewise have been incorporated. Web use mining is critical in building up client profiles for a superior organized Web webpage. "As the way in which the Web is utilized keeps on growing, there is a constant need to make sense of new sorts of information about client conduct that should be dug for". Here the goal is to dive into the records of the servers (log files) that store the information transactions that are performed in the web in order to find patterns revealing the usage the customers make of it.

For example the most visited pages, usual visiting paths, etc. We can also distinguish here:

1. **General access pattern** tracking. Here the hobby doesn't depend on the entrance examples of a specific guest yet on the joining of them into patterns permitting us to re-structure the web keeping in mind the end goal to encourage our client's entrance and use of our site.
2. **Customized** access pattern tracking. Here what we look for is gathering data about the individual visitor's behavior and their interaction with the website.

This way we can establish access/purchase profiles so that we can offer a customized experience to every customer. An archetypical case of this is amazon.com and its purchase advice and suggestions

Web Data Applications

1. Marketing 18%
2. Customer Service 16%
3. Don't Use Data 72%

In their frenzy to be the next Amazon, eBay, or Yahoo, companies are scrambling to set up their E-Commerce sites. Regularly focusing on the mechanics of value-based preparing, setting up their stock and shopping baskets however neglect to anticipate the immeasurable measure of client information their site will produce. Most companies fail to see that in E-Commerce, long term success depends on how this web data is leveraged to convert visitors into customers and customers into loyal clients.

Customer Relationship Modeling Via Web Metrics

The web data that is generated with a single sale is of more value than the sale itself since it can lead to a long and profitable relationship with that customer. The goal of marketers today is not to capture market share, but instead capture a share of a customer over a long period of time on a one-to-one basis: the Web provides an ideal marketplace for doing this. Each visit to retailing site produces critical purchaser behavioral information, paying little respect to whether a deal is made. Each guest activity is an advanced motion displaying propensities, inclinations and inclinations. This cooperation uncovers imperative patterns and examples that can help an organization plan a site that adequately imparts and markets its items and administrations. The beauty of E-Commerce is that you are able to create your own customer data. For example, your home page should quickly establish a dialog with your visitors in order to find out what their needs are. Focus on interacting with your customers in order to learn what their needs are so that you can service them better over time, and subsequently retain them.

Data Collection

One key to arranging and catching this customer data is an one of a kind identifier: a guest ID number. A demonstrated procedure is having your new client enroll at first at your site by tempting them with an uncommon administration or motivator. Offer access to an extraordinary segment of your site. Have challenges or entryway prizes. The fact of the matter is that you require them to enlist with a specific end goal to set a treat, which can be utilized as the novel ID number. A treat is a header that servers go to programs. Starting there, the one of a kind key can empower you to track each communication with that guest. This interesting key will permit your site to connection log documents and structures database with your organization's information distribution center and other outsider demographic and family unit data, advertisement server systems or community separating motors.

Mining Web Data

So far, most analyzes of web data have involved log traffic reports; most of which provide cumulative accounts of server activity but do not provide any true business insight about customer demographics and on-line behavior. Most of the current log analyzers provide pre-defined reports about server traffic activity. This limits the scope of these tools to statistics about domain names, IP addresses, browsers and other TCP/IP specific machine-to-machine activity. On the other hand, the mining of web data for an E-Commerce site yields visitor behavior analyses and profiles, rather than simple server statistics. An E-Commerce site needs to know about the preferences and life-styles of its visitors. Data mining in this context is about addressing such business questions as, "Who is buying what items and at what rates." You have to know how to offer and what motivating forces, offers and advertisements work, and how you ought to outline your site to

upgrade your benefits. The utilization of information mining calculations can independently separate connections among web information to figure out whether examples exist which can yield noteworthy business and promoting knowledge.

The Time Is Now

The Web is a fast, competitive marketplace that doubles every 100 days. A marketplace where on-line shoppers browse retailing sites with their fingers poised over their mouse: ready to buy if they find what they are looking for or move on should the content, wording, incentive, promotion, product or service of that site not meet their preferences.

The Web is a commercial center where programs are pulled in and held in view of how well the retailer recalls the clients' requirements and impulses. The objective is to know and serve each client, each one in turn, and to construct long haul, commonly helpful connections. Information mining is the way to client learning and closeness in this kind of aggressive and swarmed commercial center.

The data that a vendor accumulates from its site can uncover what items have cross-offering open doors, or what data and motivations the trader ought to give to its guests in view of their sexual orientation, age, demographics and way of life hobbies. The process involves capturing important visitor attributes from server logs, cookies and CGI forms, then appending to them household and demographic information. Using powerful pattern-recognition technologies such as neural networks, machine-learning and genetic algorithms, customers may be profiled by their propensity to buy.

BENEFITS OF WEB MINING

Understand Customer Behavior

1. Companies can optimize e-business sites for maximum commercial impact by understanding the dynamic behavior of visitors to their Web sites.
2. E-tailers can now gain knowledge on the individual tastes and preferences of the visitors to their sites.
3. Determine the conversion rate of visitors to buyers on your site.
4. Determine the repeat frequency of existing buyers (i.e., the likelihood of customers repurchasing your brand).
5. Calculate the rate of new customer acquisition.
6. Discover actionable browsing and buying patterns of customers.
7. Learn who is buying what from your site.
8. Discover cross connections between customers in your e-business locales.

Focus Web Site Effectiveness:

1. Discover high and low effect zones of your e-business site.
2. Web executives no more need to depend on instinct when outlining the format of a Web website.
3. E-tailers can now build up the look and feel of the Web webpage and customize online substance.

Measure the Success of Marketing Efforts:

1. In the physical world it is hard to get solid input on showcasing battles. In any case, on the Internet you can get genuine estimations of the accomplishment of an advertising crusade.
2. Companies can group clients with comparable examples, and the Web webpage can adjust to perceived clients. Portions can then be focused with crusades and uncommon offers.
3. Effectively gauge the arrival on venture (ROI) of flag publicizing.

THE FUTURE

With the change from 'bricks and mortar' to 'clicks and mortar' it is essential that in order to remain competitive ebusiness enterprises know as much as possible about their customers including their tastes, behavior and preferences. Enterprises must realize that there is more to ebusiness than just building a web site, sitting back and waiting for customers to turn up and spend their money. By using Web Mining techniques enterprises are able to understand their customers, their buying patterns and behaviors which is essential if they are to retain them as customers. Further more, Web Mining provides enterprises with the opportunity to optimize their sites for maximum commercial impact and the ability to provide personalized online content of their web site.

There are certain critical factors that must be considered before an enterprise defines its Web Mining strategy. Whilst personalization enhances the customer browsing experience, it does rely heavily on the collection of personal data from the customer in the first place. This may be done using browser based forms which the user completes or alternatively using cookie based technologies where the customer may or may not be aware what data, if any, is being collected. Recently there has been a great deal of publicity about the exposure of personal details from several high profile web sites due to security loopholes. There has also been a similar level of publicity regarding the privacy policies of web sites and what exactly can be done with personal data once it has been entered. These issues must be carefully considered by enterprises as the desire to obtain data and therefore market presence could easily lead to enterprises down fall.

It is becoming increasingly important for enterprises to develop and retain their customer base. Web Mining techniques

provide the opportunity for developing customer insight at an unprecedented level of detail. The enterprises that make the best use of traditional data sources and Internet sourced data through Web Mining technique will be placed to deliver and succeed with their ebusiness strategy.

CONCLUSION

Web mining is a very dynamic research area. A survey like this can only scratch on the surface. We tried to include references to the most important works in these areas, but we necessarily had to be selective. Nevertheless, we hope to have provided the reader with a good starting point for her own explorations into this rapidly expanding and exciting research area.

REFERENCES

1. Maimon O., Rokach L. Data Mining and Knowledge Discovery Handbook. *Springer Science & Business Media*. 2nd Edn.
2. Kaur C., Aggarwal R.R. Web Mining Tasks and Types: A Survey. *International Journal of Research in IT & Management*. 2012 February; 2(2): 547–58p.
3. <http://www.csee.umbc.edu/~joshi/web-mine/>
4. Claudia Elena Dinucă, Dumitru Ciobanu. Web Content Mining. *Annals of the University of Petroșani, Economics*. 2012; 12(1): 85–92p
5. Jicheng W., Yuan H., Gangshan W. Web mining: knowledge discovery on the Web. Systems, Man, and Cybernetics. *International Conference. IEEE*. 1999 Oct 12-15. 2; 137 – 141p.
6. Srivastava J. Web Mining: Accomplishments & Future Directions. *National Science Foundation Workshop on Next Generation Data Mining NGDM02*. 2002; 1–148p.