

Article

A Movie Recommender System: BIG SCREEN

Akash Gupta¹, Amit Gupta²

^{1,2}Research Scholar, MCA, Thakur Institute of Management Studies, Career Development & Research (TIMSCDR), Mumbai, India.

I N F O

Corresponding Author:

Akash Gupta, Thakur Institute of Management Studies, Career Development & Research (TIMSCDR), Mumbai, India.

E-mail Id:

akashgupta12375@gmail.com

How to cite this article:

Gupta A, Gupta A. A Movie Recommender System: BIG SCREEN. *J Adv Res Info Tech Sys Mgmt* 2020; 4(1): 9-13.

Date of Submission: 2020-05-02

Date of Acceptance: 2020-05-29

A B S T R A C T

Presently a day's suggestion framework has changed the style of looking through the things of our advantage. This is data sifting approach that is utilized to anticipate the inclination of that user. The most well-known zones where recommender framework is applied are books, news, articles, music, recordings, motion pictures and so on. In this paper we have proposed a film suggestion framework named BIG SCREEN. It depends on community oriented sifting approach that utilizes the data gave by clients, breaks down them and afterward suggests the films that is most appropriate to the client around then. The prescribed film list is arranged by the appraisals given to these motion pictures by past clients and it utilizes K-implies calculation for this reason. BIG SCREEN likewise help clients to discover the films of their decisions dependent on the film understanding of different clients in productive and powerful way without burning through much time in pointless perusing. This framework has been created in PHP utilizing Dreamweaver 6.0 and Apache Server 2.0. The introduced recommender framework produces proposals utilizing different sorts of information and information about clients, the accessible things, and past exchanges put away in modified databases. The client would then be able to peruse the proposals effectively and discover a film of their decision.

Keywords: K-means, Recommendation System, Recommender System, Data Mining, Clustering, Movies, Collaborative Filtering, Content-Based Filtering

Introduction

In today's world where internet has become an important part of human life, users often face the problem of too much choice. Right from looking for a motel to looking for good investment options, there is too much information available. To help the users cope with this information explosion, companies have deployed recommendation systems to guide their users. The research in the area of recommendation systems has been going on for several decades now, but the interest still remains high because of the abundance of practical applications and the problem rich domain. A number of such online recommendation systems implemented and used are the recommendation system for books at Amazon.com, for movies at MovieLens.org, CDs at CD Now.com (from Amazon.com), etc. Recommender

Systems have added to the economy of the some of the e-commerce websites (like Amazon.com) and Netflix which have made these systems a salient parts of their websites. A glimpse of the profit of some websites is shown in table below:

Table I. Companies Benefit Through Recommendation System

Netflix	2/3 rd of the movies watched are recommended
Google News	recommendations generate 38% more click-throughs
Amazon	35% sales from recommendations
Choicestream	28% of the people would buy more music if they found what they liked

Recommender Systems generate recommendations; the user may accept them according to their choice and may also provide, immediately or at a next stage, an implicit or explicit feedback. The actions of the users and their feedbacks can be stored in the recommender database and may be used for generating new recommendations in the next user-system interactions. The economic potential of these recommender systems have led some of the biggest e-commerce websites (like Amazon.com, snapdeal.com) and the online movie rental company Netflix to make these systems a salient part of their websites. High quality personalized recommendations add another dimension to user experience. The web personalized recommendation systems are recently applied to provide different types of customized information to their respective users. These systems can be applied in various types of applications and are very common now a day.

We can classify the recommender systems in two broad categories:

- Collaborative filtering approach
- Content-based filtering approach

Collaborative Filtering

Collaborative filtering system recommends items based on similarity measures between users and/or items. The system recommends those items that are preferred by similar kind of users. Collaborative filtering has many advantages

- It is content-independent i.e. it relies on connections only
- Since in CF people makes explicit ratings so real quality assessment of items are done.
- It provides serendipitous Recommendations because recommendations are base on user's similarity.

Content-based filtering

Content-based filtering is based on the profile of the user's preference and the item's description. In CBF to describe Items we use keywords apart from user's profile to designate (1) user's favored liked or dislikes. In other words CBF algorithms indorse those items or similar to those items that were liked in the past. It examines previously rated items and recommends best matching item. There are various approaches proposed in various research papers listed below. These approaches are often combined in Hybrid Recommender Systems. An earlier study by Eyjolfsson et al for the recommendation of movies through MOVIEGEN had certain drawbacks such as ,it asks a series of questions to users which was time taking . On the other hand it was not user friendly for the fact that it proved to be stressful to a certain extent. Keeping in mind these shortcomings, we have developed BIG SCREEN, a movie commendation system that recommends movies

to users based on the information provided by the users themselves. In the present study, a user is given the option to select his choices from a set of attributes which include actor, director, genre, year and rating etc. We predict the users choices based on the choices of the previous visited history of users. The system has been developed in PHP and currently uses a simple console based interface.

Relatedwork

Many recommendation systems have been developed over the past decades. These systems use different approaches like collaborative approach, content based approach, a utility base approach, hybrid approach etc. Looking at the purchase behavior and history of the shoppers, Lawrence et al. 2001 presented a recommender system which suggests the new product in the market. To refine the recommendation collaborative and content based filtering approach were used. To find the potential customers most of the recommendation systems today use ratings given by previous users. These ratings are further used to predict and recommend the item of one's choice. In 2007 Weng, Lin and Chen performed an evaluation study which says using multidimensional analysis and additional customer's profile increases the recommendation quality. Weng used MD recommendation model (multidimensional recommendation model) for this purpose. Multi-dimensional recommendation model was proposed by Tuzhilin and Adomavicius (2001).

Research Methodology

The Basic K-means Algorithm

The original K-means algorithm was proposed by MacQueen [20]. The ISODATA algorithm by Ball and Hall²² was an early but sophisticated version of k-means. Clustering divides the objects into meaningful groups. Clustering is unsupervised learning. Document clustering is automatic document organization.

In K-means clustering technique we choose K initial centroids, where K is the desired number of clusters. Each point is then assigned to the cluster with nearest mean i.e. the centroid of the cluster. Then we update the centroid of each cluster based on the points that are assigned to the cluster. We repeat the process until there is no change in the cluster center (centroid). Finally, this algorithm aims at minimizing an objective function, in this case a squared error function. The objective function:

Where, k is the number of clusters, n is the number of cases $\|x_i^{(k)} - c_k\|^2$ is a chosen distance measure between a data point and the cluster centre $x_i^{(k)}, c_k$, is an indicator of the distance of the n data points from their respective cluster centers.

The algorithm is composed of the following steps:

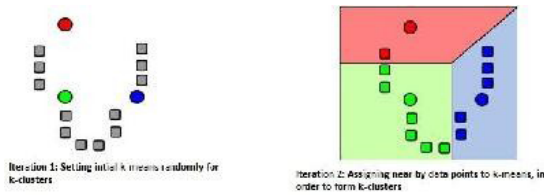
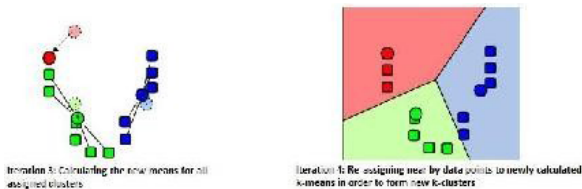


Figure 1.K-means Algorithm

Figure 2.The 5 steps of the K-Means algorithm³

Data Description

In proposed model we use a pre filter before applying K-means algorithm. The attributes used to calculate distance of each point from centroid are

- Genre
- Actor
- Director
- Year
- Rating

Different attributes have different weights. In our research we have found that the most appropriate recommendations that can be generated should be based on the ratings given to the movies by previous users, therefore we have given more importance to the rating attribute than other attributes. These ratings have been taken from www.imdb.com because perhaps it has the largest collection of movies along with the rating given to these movies by a large number of different users from different parts of the world. Another important parameter in our proposed model is total number of votes received by a particular movie. We have divided number of votes in to three categories that is less than or equal to 1000, more than 1000 but less than or equal to 10,000 and greater than 10,000.

In our research we have found that as the number of vote's increases the weight of rating should also increase respectively. Therefore we have used ratios of 1:1, 1:2, and 1:3 depending on total number of votes received by a movie. we have also found that the movies which have rating less than 5 are the ones which are least suitable for recommendation, and are least desirable by users. Users generally want to see a good movie and higher rating ensures that our predicted movie set are of those movies which are liked by a large number of users. Weights assigned to other attributes are generally based on the average of total movies associated with that particular attribute to the total number of movies in our dataset.

Simulation of BIG SCREEN

When any user enters our system MOVREC he has a couple of options. He /she can search a particular movie or see upcoming movies list or can go to our recommendation page. On recommendation page he is given the choice to select/input values for different attributes. On the basis of these input values, we search our search our database and prepare an array of suitable movies. Movies included in the array are those whose even one attribute value matches with the input value of the user. We then calculate the number of movies in our array with the help of a counter. If the counter value is less than or equal to twenty we display the movie list sorted according to ratings associated with the movies. If number of movies is greater than twenty then we apply a pre filter and select top twenty movies according to rating. If two movies have same rating then priority is given to the movie having a large number of votes. After filtering the movie list we match the attributes value to their respective weights and compute the total weight of each movie. Once we have calculated the total weight of each movie we apply K-means clustering algorithm on these group of movies. In our research we have also found that generally a user prefer a list with five movies so we assume K equal to be 4 so that an average every K has five movies, where K is the number of cluster to be formed.

For each cluster k1, k2 , k3, k4 we assume initial centroid c1, c2, c3, c4 which corresponds to the first, sixth, eleventh, and sixteenth movie in the movie array. After defining the initial centroid we compute the distance of all the other data points from each centroid and assign the remaining data points (movies) to closest centroid and form clusters. The distance measure we have used to calculate the distance between data points and centroid is the Euclidean Distance.

After forming initial clusters we take one cluster at a time. We again calculate centroids but this time each centroid corresponds to mean of the points in that cluster. After recalculating centroids we compute the distance of all data points with respect to these newly formed centroids and reassign them to form clusters. We repeat this process till there is no change in centroids. This ensures that the clusters finally formed are optimized and no further grouping is possible. Once final cluster are formed we compute the average rating of all points belonging to that cluster i.e. cluster rating, then according to the input user query we display the cluster having highest clusterrating.

Weightage and Matching of Attributes

Actor(Wa)

$$W_a = \frac{\text{No. of movies of Actor(a) in data set}}{\text{Total no. of movies in data set}}$$

Director(Wd)

$Wd = \frac{\text{No. of movies of Director(d) in data set}}{\text{Total no. of movies in data set}}$

Rating

Table I

Rating	Weight		
	If number of votes ≤ 1000	If number of votes $1000 < \text{votes} \leq 10000$	If number of votes > 10000
10	10	20	30
9	9	18	27
8	8	16	24
7	7	14	21
6	6	12	18
5	5	10	15
1-4.9	1	2	3

Genre (Wg)

$Wg = \frac{\text{No. of movies of Genre(g) in data set}}{\text{Total no. of movies in data set}}$

Year (Wy)

$Wy = \frac{\text{No. of movies in Year (Y) in data set}}{\text{Total no. of movies in data set}}$

Total weight of a particular movie m is given by:

Proposed Algorithm

Input: a number of movies: m

Output: a number of clusters: K

Step 1 Select n movies from m movies $n < m$

Step 2 If $n > 20$ then select top 20 movies from n movies based on ratings. Else display the output movies sorted by rating.

Step 3 If rating of movies x, y are equal i.e. If $R_x = R_y$ Then select those movies which have greater number of user votes.

Step 4 Assume $K=4$.

Step 5 REPEAT (6, 7)

Step 6 Chose initial centroid C1, C2, C3, C4. Step 7 Calculate Euclidean distance of all data points w.r.t. C1, C2, C3, C4 and re- compute the centroid of each cluster.

Step 8 UNTILL centroid does not change.

Where, m: Total number of movies in database n: Number of movies after user query x, y: Two random movies

R_x, R_y: Rating of movies x, y

K: Number of cluster

C1, C2, C3, C4: Initial Centroid.

Challenges Faced

In developing any system the biggest challenge is to satisfy the end users for which the system is being developed. We also faced certain challenges while developing our system. Some of the are:

Table 2.Excerpt of the Survey Conducted

Kritika Baranwal	1	3	2	5	4
Ranjana Rai	3	2	1	4	5
Dr. Siddhart h Singh	2	3	1	4	5
Alok Kumar	1	4	2	5	3
Shashank Pandey	3	2	1	5	4
Er.Pankaj Agarwal	3	2	1	4	5
Mr.Girjesh Mishra	3	1	2	4	5

In our survey we have found that almost every user has given maximum priority to the rating concerned with the movie so we have given the maximum priority to our rating attribute. Also we have deduced from our survey that movie rating has more weight if that movie has received a large number of:

Overcome the Problems

- The proposed system has been tested over a small group of people, and we have received a positive response from them. We have kept our system simple and interactive for this we have choose php and javascript.
- For collecting information we have intensively search free online movie data bases and extract the information which was useful for our proposed system.
- To accurately recommend movie to user we have applied K-means clustering algorithm along with a prefilter.
- We have included movies in our database irrespective of their language or location so that users from all across the globe can use our system.
- For assigning weights to attributes and for giving priority to them we have conducted a survey on a group of people and on the basis of the result obtained we have prioritize our attributes.

An excerpt of the survey is given below votes. So we have divided rating and votes in to three categories i.e. minimum, medium and maximum votes.

Table 3

Name	Actor/Cast	Genre	Rating	Director	Year
Mrs Malti Singh	2	5	1	4	3
K.K Singh	5	3	1	2	4

Conclusion

In this paper we have introduced BIG SCREEN, a recommender system for movie recommendation. It allows a user to select his choices from a given set of attributes and then recommend him a movie list based on the cumulative weight of different attributes and using K-means algorithm. By the nature of our system, it is not an easy task to evaluate the performance since there is no right or wrong recommendation; it is just a matter of opinions. Based on informal evaluations that we carried out over a small set of users we got a positive response from them. We would like to have a larger data set that will enable more meaningful results using our system. Additionally we would like to incorporate different machine learning and clustering algorithms and study the comparative results. Eventually we would like to implement a web based user interface that has a user database and has the learning model tailored to each user.

References

1. Han J, Kamber M. Data Mining: Concepts and Techniques. Morgan Kaufmann (Elsevier), 2006.
2. Ricci, Missier FD. Supporting Travel Decision making Through Personalized Recommendation. *Design Personalized User Experience for e-commerce* 2004; 221-251.
3. Steinbach M, Tan P, Kumar V. Introduction to Data Mining. Pearson, 2007.
4. Jha NK, Kumar M, Kumar A et al. Customer classification in retail marketing by data mining. *International Journal of Scientific & Engineering Research* 1998; 5(4), April-2014 ISSN 2229- 5518, Giles CL, Bollacker KD, Lawrence S. CiteSeer: An automatic citation indexing system. in Proceedings of the third ACM conference on Digital libraries, 1998: 89-98.
5. Beel J, Langer S, Genzmehr M et al. Introducing Docear's Research Paper Recommender System. in Proceedings of the 13th ACM/IEEE-CS Joint Conference on Digital Libraries (JCDL'13), 2013: 459-460.
6. Bethard S, Jurafsky D. Who should I cite: learning literature search models from citation behavior. in Proceedings of the 19th ACM international conference on Information and knowledge management, 2010: 609-618.
7. Bollacker KD, Lawrence S, Giles CL. CiteSeer: An autonomous web agent for automatic retrieval and identification of interesting publications. in Proceedings of the 2nd international conference on Autonomous agents, 1998; 116-123.
8. Erosheva E, Fienberg S, Lafferty J. Mixed- membership models of scientific publications. in Proceedings of the National Academy of Sciences of the United States of America, 2004; 101(1): 5220-5227.
9. Ferrara F, Pudota N, Tasso C. A Keyphrase- Based Paper Recommender System. in Proceedings of the IRCDL'11, 2011; 14-25.
10. Jiang Y, Jia A, Feng Y et al. Recommending academic papers via users' reading purposes. in Proceedings of the sixth ACM conference on Recommender systems, 2012; 241-244.
11. McNee SM, Kapoor N, Konstan JA. Don't look stupid: avoiding pitfalls when recommending research papers. in Proceedings of the 20th anniversary conference on Computer supported cooperative work, 2006; 171-180.
12. Middleton SE, De Roure DC, Shadbolt NR. Capturing knowledge of user preferences: ontologies in recommender systems. in Proceedings of the 1st international conference on Knowledge capture, 2001; 100-107.
13. Zarrinkalam F, Kahani M. SemCiR - A citation recommendation system based on a novel semantic distance measure. *Program: electronic library and information systems* 2013; 47(1): 92-112.
14. Schafer JB, Frankowski D, Herlocker J et al. Collaborative filtering recommender systems. *Lecture Notes In Computer Science* 2007; 4321: 291.
15. Seroussi Y. Utilising user texts to improve recommendations. *User Modeling, Adaptation and Personalization* 2010; 403-406.
16. Buttler D. A short survey of document structure similarity algorithms. in Proceedings of the 5th International Conference on Internet Computing, 2004.
17. Goldberg D, Nichols D, Oki BM et al. [Using collaborative filtering to weave an information Tapestry]. *Communications of the ACM* 1992; 35(12): 61-70.
18. Beel J, Langer S, Genzmehr M. Mind-Map based User Modelling and Research Paper Recommendations," in work in progress, 2014.
19. MacQueen J. Some methods for classification and analysis of multivariate observations. In *Proc. Of the 5th Berkeley Symp. On Mathematical Statistics and Probability*, University of California Press, 1967; 281-297.
20. Ball G, Hall D. A Clustering Technique for Summarizing Multivariate Data. *Behavior Science* 1967; 12: 153-155. Bowman M, Debray SK, Peterson LL. Reasoning about naming systems 1993.