

Walmart Sales Prediction Using Machine Learning Algorithms

¹Dr. C. Shyamala, ²P. Sabarish, ³T. Vignesh, ⁴S. Yogeendran

¹Associate Professor, K.Ramakrishnan College of Technology

^{2,3,4}UG Student, K.Ramakrishnan College of Technology

ABSTRACT: Machine Learning is changing each walk of life and has gotten to be a major giver in honest to goodness world scenarios. The dynamic applications of Machine Learning can be seen in each field checking instruction, healthcare, building, bargains, energy, transport and a couple of more; the list is never wrapping up. The conventional approach of bargains and showcasing destinations no longer offer assistance the companies, to oversee up with the pace of competitive grandstand, as they are carried out with no bits of knowledge to customers' securing plans. Major changes can be seen inside the space of bargains and showcasing as a result of Machine Learning headways. Owing to such headways, diverse fundamental points such as consumers' purchase plans, target gather of spectators, and anticipating deals for the afterward a long time to come can be easily decided, hence making a distinction the bargains gather in characterizing plans for a boost in their exchange. The point of this paper is to propose a measurement for foreseeing long run bargains of Tremendous Bazaar Companies keeping in see the bargains of past a long time. A comprehensive consider of deals estimate is done utilizing Machine Learning models such as Random Forest, K-Neighbors Regressor, XGBoostRegressor and LinearRegressor. The expectation incorporates data parameters such as thing weight, thing fat substance, thing deceivability, thing sort, thing MRP, outlet foundation year, outlet appraise and outlet region sort.

Keyword: Machine learning, bazaar companies, walmart, Random Forest, K Nearest Neighbour, XG Boost, Linear Regressor, outlet appraise.

I. INTRODUCTION

Deals estimating have continuously been an awfully critical region to concentrate upon. An effective and ideal way of estimating has gotten to be fundamental for all the merchants in arrange to support the adequacy of the promoting organizations. Manual invasion of this assignment may lead to extreme blunders driving to destitute administration of the organization, and most imperatively would be time devouring, which is something not alluring in this assisted world. A major portion of the worldwide economy depends upon the commerce segments, which are actually anticipated to deliver fitting amounts of items to meet the generally needs.

Focusing on the advertise gathering of people is the major center of trade divisions. It is hence critical that the company has been able to realize this objective by employing a framework of determining. The method of estimating includes analyzing information from different sources such as advertise patterns, buyer behavior and other variables. This examination would moreover offer assistance the companies to oversee the budgetary assets viably. The determining process can be utilized for numerous purposes, counting: anticipating long haul request of the items or benefit, anticipating how much of the item will be sold in a given period.

Typically where machine learning can be misused in a incredible way. Machine learning is the space where the machines pick up the capacity to outflank people in particular errands. They are utilized to do a few specialized errands in a consistent way and pick up superior comes about for the advance of the current society. The base of machine learning is the craftsmanship of arithmetic, with the assistance of which different standards can be defined to approach the ideal yield. In case of deals determining too machine learning has proved to be a boon. It is supportive in anticipating the long run deals more precisely.

The objective for this extend is to gauge accurately for the item unit deals estimating within the USA-based Walmart. In arrange to perform expectations on different items that are sold in Walmart, machine learning methods have been actualized at the side the conventional strategies to extend the exactness. Three diverse machine learning models are utilized to figure every day deals for taking after a 28-day period. The metric of assessing the models is Root Cruel Square Error(RMSE). The result from RMSE seem back the expressed theory and help the trade examiner to progress arranging on diverse perspectives of the trade level, for occurrence stock dispersion, dispersion administration, stock capacity arrangements, item fulfillment, etc.

II. LITERATURE SURVEY

Social media have changed the world in which we live is proposed by Tesco and Walmart [1]. In spite of the fact that a few considers have revealed shapes of client engagement on social media, there's a shortage of scholarly inquire about on client engagement inside the basic need segment. This consider hence points to address this hole within the writing and shed light on the different ways clients lock in with basic supply stores on Facebook. Netnography is utilized to pick up an understanding of the conduct of clients on the Facebook page of Tesco and Walmart. The discoveries of this ponder uncover that cognitive, emotional and behavioral client engagement are showed which clients can both make and annihilate esteem for the firm.

Data innovation in this 21st century is coming to the skies with large-scale of information to be prepared and examined to form sense of information where the conventional approach is no more compelling. Presently, retailers require a 360-degree see of their shoppers, without which, they can miss competitive edge of the advertise. Retailers need to make successful advancements and offers to meet its deals and promoting objectives, something else they will swear off the major openings that the current showcase offers. Numerous times it is difficult for the retailers to comprehend the advertise condition since their retail stores are at different topographical areas. Huge Information application empowers these retail organizations to utilize earlier year's information to superior estimate and anticipate the coming year's deals. It moreover empowers retailers with profitable and expository experiences, particularly deciding clients with wanted items at wanted time in a specific store at distinctive topographical areas.

The retail division [3] has broadly adjusted diverse stock administration applications and a few retail chains indeed utilize expectation computer program to analyze future deals. Be that as it may, a parcel of day-to-day shopping in India happens through nearby shops. The proprietors of such mom-and-pop shops don't essentially have the capital to contribute in exclusive applications for setting up an stock administration framework. Unnecessary to say that

same is the case for any deals forecast program. As a result, numerous of the businesspeople conclusion up accumulating a part of insignificant and non profitable items that lead to financial misfortunes. A really cost effective and available arrangement for this issue may be a portable application that provides all the highlights of a point-of-sale framework as well as gives future deals bits of knowledge. It'll empower businesspeople to oversee their current item buys and invoicing. The prescient deals examination will offer assistance them to alter their ventures on items and supplies subsequently guaranteeing greatest benefits.

This paper [7] presents a modern strategy to discover driving pointers of residential book deals from routine Web journal data. In spite of the fact that Web journal data is diverse from real buys, it'll impact client behaviors. They consider it would be valuable for choice making of businesses and organizations. Primary commitments of the paper are three overlay: 1) As well as think about on the US of a Web bookstore, there's a relationship between book deals information and Web journal data in Japan by separating judge of book with reference rate; 2) A few of the Web journal data are viable to foresee the lead-time between distributors and bookstores; and 3) Handle is proposed to judge whether the data is substantial as driving markers.

The most point of this paper [8] is to anticipate the deals of a vehicle utilizing opinion examination from different places on the web. The online nearness of a vehicle , as well as its brand, plays a key part within the deals of the vehicle. Be that as it may, numerous other parameters are required and will be talked about in this paper. Deals expectation in today'sadvertise isn't as it were useful for the producer but too for the different other companies that fabricate parts or embellishments for vehicles. It can moreover be considered as a boon for retailers, showroom proprietors, and benefit mechanics. In this application, we have made utilize of direct relapse for opinion investigation and polynomial relapse for deals expectation.

Brilliantly Choice Explanatory Framework [9] requires integration of choice examination and forecasts. Most of the trade organizations intensely depend on a information base and request expectation of deals patterns. The precision in deals estimate gives a huge affect in commerce. Information mining methods are exceptionally compelling apparatuses in extricating covered up information from an colossal dataset to upgrade precision and effectiveness of estimating. The nitty gritty consider and investigation of comprehensible prescient models to make strides future deals expectations are carried out in this inquire about. Conventional figure frameworks are troublesome to bargain with the enormous information and exactness of deals estimating. These issues may well be overcome by utilizing different information mining methods. In this paper, we briefly analyzed the concept of deals information and deals estimate. The different techniques and measures for deals forecasts are depicted within the afterward portion of the investigate work.

The point of this paper [10] is to analyze the deals of a enormous superstore, and foresee their future deals for making a difference them to extend their benefits and make their brand indeed way better and competitive as per the advertise patterns by creating client fulfillment as well. The method utilized for forecast of deals is the Linear Regression Algorithm, which may be a popular calculation within the field of Machine Learning. The deals information is from the year 2011-13 and forecast of data for the year 2014 is done. At that point, real-time information

of the year 2014 is additionally taken and the actual data of the year 2014 has been compared to the anticipated information to calculate the precision of expectation. This can be done so as to approve our comes about with the genuine ones. This in turn would offer assistance them take fundamental activities.

III. PROPOSED METHODOLOGY

In proposed Framework walmart deals forecast is anticipated with the machine learning calculations towards the information collected from the past deals of a basic supply store. The objective here is to imagine the design of deals and the amounts of the items to be sold based on a few key highlights accumulated from the crude information we have. In case of deals estimating totally foreseeing the deals level, stock and offers data are get known. It is supportive in anticipating long term deals more precisely. The capacity to anticipate information precisely is amazingly profitable in a endless cluster of spaces such as stocks, deals, climate or indeed sports. Displayed here is the consider and usage of a few gathering classification calculations utilized on deals information, comprising of week after week retail deals numbers from diverse divisions in Walmart retail outlets all over the Joined together States of America. The hyperparameters of each show were changed to get the finest Cruel Supreme Mistake (MAE) esteem and R2 score. The number of estimators hyperparameter, which specifies the number of choice trees utilized within the demonstrate, plays a especially critical part within the assessment of the MAE esteem and R2 score and is managed with in an attentive way. A comparative examination of the three calculations is performed to demonstrate the most excellent calculation and the hyper parameter values at which the leading comes about are gotten.

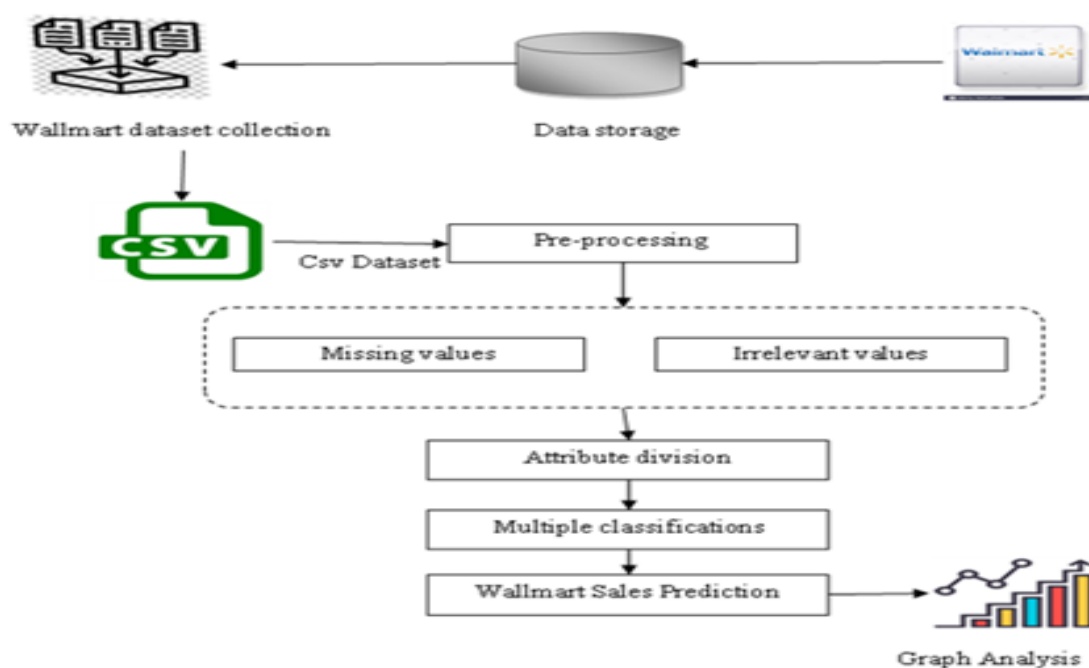


Figure 1 System architecture of the walmart sales prediction

(i) Dataset Collection

The dataset comes from the Kaggle stage and comprises of information from an American retail organization, Walmart Inc. The dataset was utilized for a machine learning competition in 2014. It comprises information from 45 Walmart division stores primarily centered around their deals on a week after week premise. The dataset has 282,452 passages that will be utilized for preparing the models. Each section has properties as takes after: the related store (recorded as a number), the comparing division (81 offices, each entered as a number), the date of the beginning day in that week, departmental week after week deals, the store measure, and a Boolean esteem indicating on the off chance that there's a major occasion within the week. The major occasions being one of Thanksgiving, Labor Day, Christmas or Easter. Together with the previously mentioned qualities may be a parallel set of highlights for each section counting Customer Cost List, unemployment rate, temperature, fuel cost, and special markdowns.

(ii) Dataset Acquisition

Information Securing comprises of two words: Information: Information alludes to the crude actualities, figures, or piece of realities, or measurements collected for reference or examination. Securing: Securing alludes to securing information for the venture. There are four strategies of securing information: collecting modern information; converting/transforming bequest information; sharing/exchanging information; and obtaining information. Retailer's to begin with need is ordinarily to get it their clients to be able to fulfill their needs so that these clients will return to the store for future needs, in this way expanding the item requests and including to the trade esteem. These businesses need this data to arrange where and when to contribute beneficially.

(iii) Data Pre-Processing

Information pre-processing could be a portion of information mining, which includes changing crude information into a more coherent organize. Crude information is more often than not, conflicting or fragmented and ordinarily contains numerous mistakes. The information pre-processing involves checking out for lost values, seeking out for categorical values, part the data-st into preparing and test set and finally do a highlight scaling to constrain the range o factors.

- Data preprocessing may be a information mining method which is utilized to convert the crude information in a valuable and effective format.

- Steps Included in Information Preprocessing: Data cleaning

(iv) Attribute Division

It can be seen as a information field that speaks to the characteristics or highlights of a information question. For a client, protest traits can be client Id, address, etc. Able to say that a set of traits utilized to portray a given question are known as property vector or include vector. One viewpoint is the handling of diminished collections of enormous information with less computing assets. In this manner, the think about analyzed 40 GB information to test different methodologies to decrease information handling. In this way, the objective is to decrease this

information, but not to compromise on the location and demonstrate learning in machine learning. A few choices were analyzed, and it is found that in numerous cases and sorts of settings, information can be decreased to a few degree without compromising location effectiveness

Subordinate factors are the targets or the yield factors which ought to be at last assessed and after that compared against each other. Autonomous factors are the highlights or the input factors which can't be changed by any means and appropriately the targets are anticipated. Feature scaling could be a strategy in which we scale the information into an exact and versatile estimate for the reason of increasing accuracy and decreasing blunder. It basically prevents the expansive fluctuation of data points to be utilized within the calculation and permits us to realize superior comes about. Standard Scaler could be a lesson imported from sklearn library.

(v) Multiple Classification

Four determining models were built in this inquire about on the taking after calculations: Irregular Woodland, Angle Boosting, Direct Relapse and Greatly Randomized Trees (Additional Trees). Other calculations such as Naïve Bayes and Versatile Boosting were scrutinized, but their exhibitions were not up to the stamp and experiences were trifling, so they will not be considered in this. All models were executed in Python 3.7. on the Boa constrictor dispersion utilizing pycharm Scratch pad. The code utilized for the usage has been transferred to GitHubTo maintain a strategic distance from over fitting, two isolated datasets are not imported for prepare and test. So, part is drained a single dataset. The preparing dataset are the information we got to prepare the demonstrate on.

Linear Regression:

Direct relapse may be a straight show, e.g. a show that accept a straight relationship between the input factors (x) and the single yield variable (y). More particularly, that y can be calculated from a direct combination of the input variables (x). In measurements, straight relapse may be a straight approach to modeling the relationship between a scalar reaction and one or more informative variables. Linear Relapse is the foremost commonly and broadly utilized calculation Machine Learning calculation. It is utilized for setting up a direct connection between the target or subordinate variable and the reaction or autonomous factors. The direct relapse demonstrate is based upon the taking after condition:

$$y = \theta_0 + \theta_1x_1 + \theta_2x_2 + \theta_3x_3 + \dots + \theta_nx_n \quad (1)$$

where, y is the target variable, θ_0 is the caught, $x_1, x_2, x_3, \dots, x_n$ are autonomous factors and $\theta_1, \theta_2, \theta_3, \dots, \theta_n$ are their individual coefficients. The most point of this algorithm is to discover the finest fit line to the target variable and the independent factors of the information. It is accomplished by finding the foremost ideal values for all θ . With best fit it is implied that the anticipated esteem ought to be exceptionally near to the real values and have least mistake. Mistake is the separate between the information focuses to the fitted relapse line and for the most part can be calculated by utilizing the taking after condition: Blunder = $y - \hat{y}$, where, y is the genuine esteem and $\hat{y}$ is the anticipated value

K Nearest Neighbor:

KNN calculation for Relapse may be a administered learning approach. It predicts the target based on the closeness with other accessible cases. The likeness is calculated utilizing the distance degree, with Euclidian separate being the foremost common approach. Forecasts are made by finding the K most similar occasions i.e., the neighbors, of the testing point, from the whole dataset. KNN calculation calculates the separate between scientific values of these focuses utilizing the Euclidean separate equation.

KNN works by finding the separations between a inquiry and all the illustrations within the information, selecting the desired number examples (K) closest to the query, at that point votes for the foremost visit name (within the case of classification) or midpoints the names (within the case of relapse). One of the only choice methods that can be utilized for classification is the closest neighbor (NN) run the show. It classifies a test based on the category of its closest neighbor

The esteem of K to be chosen shouldn't be exceptionally little because it might result into commotion within the information and in turn over fitting. The common arrangement is to save a portion of information for testing the accuracy of the show. At that point select $K=1$, and after that utilize the training part of modeling and calculate the precision of expectation utilizing all tests within the test set. Rehash this handle expanding the K and select K such that it is best for the model.

XG Boost:

Agreeing to Friedman et al, Slope Boosting successively fits a basic parameterized work to current pseudo-residuals by least-squares at each emphasis by building added substance relapse models. The angle of the misfortune work being minimized with regard to the show values at each preparing point are alluded to as the pseudo residuals within the demonstrate. Figure 4 outlines the operation of the Angle Boosting calculation at different emphasess. The highlights utilized for preparing the demonstrate were the same as the ones utilized within the Arbitrary Timberland classification. The calculation was actualized utilizing Python's Slope Boosting Relapse work from the scikit-learn course, and the cruel outright mistake, cruel squared mistake and R2 score were calculated for the anticipated values.

XGBoost moreover known as Extraordinary Slope Boosting has been utilized in arrange to urge an effective demonstrate with tall computational speed and adequacy. The equation makes expectations utilizing the gathering strategy that models the expected mistakes of a few choice trees to optimize final forecasts. Generation of this demonstrate moreover reports the esteem of each feature's impacts in deciding the final building execution score forecast. This highlight esteem shows that result in outright measures – each characteristic has on anticipating school execution. XGBoost underpins parallelization by making choice trees in a parallel design.

Dispersed computing is another major property held by this calculation because it can assess any expansive and complex demonstrate. It is an out-core-computation because it investigations gigantic and shifted datasets. Taking care of of utilization of assets is done very well by this

calculative show. An additional show ought to be executed at each step in arrange to diminish the blunder. XGBoost objective work at cycle t is:

$$L(t) = \sum_{i=1}^n L(y_{outi}, y_{out1i}(t-1) + f_t(x_i) + g(f_t)) \quad (2)$$

where, y_{out} = genuine esteem knowm from the preparing dataset, and the summation portion may well be said as $f(x + dx)$ where $x = y_{out1i}(t-1)$ We got to take the Taylor guess. Let's take the best straight estimation of

$$f(x) \text{ as: } f(x) = f(b) + f'(b)(x-b) \quad dx = f_t(x_i) \quad (3)$$

where, $f(x)$ is the misfortune work L , whereas b is the past step $(t-1)$ anticipated esteem and dx is the modern learner we ought to include in step t . Moment arrange Taylor guess is:

$$f(x) = f(b) + f'(b)(x-b) + 0.5f''(b)(x-b)^2 \quad (4)$$

These execution estimations were the most excellent fulfilled with the $n_estimators$ hyper parameter set at 150, whereas the $min_samples_split$ parameter, which decides the base number of tests required to portion an insides center, and $min_samples_leaf$ parameter which demonstrates the base number of tests required to be at a leaf center, are set at 2 and 1 independently. The $n_estimators$ parameter was set to 150, whereas the $min_samples_split$ and $min_samples_leaf$ parameters were set at 2 and 1 independently, to induce the finest results wherein the MAE was 1965.5 and R^2 score was 0.94.

Random Forest:

Random Forest is characterized as the collection of choice trees which makes a difference to donate rectify yield by making utilize of stowing instrument. Sacking along side boosting are two of the foremost common outfit procedures which proposed to tackle higher inconstancy and higher partiality. In stowing, we have numerous base learners, or able to say base models, which in turn takes different irregular tests of records from the preparing dataset.

In case of Random Forest Regressor choice trees are the base learners, and they are prepared on the information collected by them. Choice trees are itself not precise learners as, when it is executed up to its full profundity, for the most part there are chances of overfitting with tall preparing exactness, but moo genuine exactness. So, we grant out the tests of the most information record by utilizing push examining and include examining with substitution strategy to each of the choice trees and this strategy is alluded to as bootstrap. The result is that each demonstrate has been prepared on all of these data records and after that at whatever point we nourish a test information to each of the prepared one out there, the forecasts assessed by each of them are combined in a way such that the ultimate yield is the cruel of all of the comes about produced. The method of combining the person comes about here is known as conglomeration.

The hyper parameter that we have to be direct in this calculation is the no of decision trees to be considered to make a random forest. Let's calculate the Gini significance of a single hub of a decision tree:

$$M_{ij} = w_j C_j - w_l(j) C_l(j) - w_r(j) C_r(j) \quad (5)$$

where M_{ij} - significance of hub j , w_j - weighted no of tests coming to hub j , C_j - the entropy esteem of hub j , $l(j)$ - child hub from cleared out part, $r(j)$ - child hub from right part on hub j The significance of each include on a base learner is at that point found out as: $N_{ij} \sum_j$:node j parts on highlight i $M_{ij} / \sum_{k \in \text{all nodes}} M_{ik}$ where N_{ij} - significance of highlight i Normalized esteem will be: $\text{norm}N_{ii} = N_{ii} / \sum_{j \in \text{all features}} N_{ij}$ The last include significance, at the Irregular Timberland level, is it's normal over all the trees. $\text{RFON}_{ii} = \sum_{j \in \text{all trees}} \text{norm}N_{ij} / T$ where RFON_{ii} - significance of include i calculated from all trees within the irregular woodland demonstrate, $\text{norm}N_{ij}$ - the normalized highlight significance for i in tree j and T - add up to no of trees .

The highlights utilized are just like the ones utilized within the past calculations. Python's Additional Trees Regressor work from the scikitlearn lesson was utilized to execute the calculation, and the different execution measurements calculated for the past strategies are assessed and reported

(vi) Walmart Sales Prediction

The capacity to foresee information precisely is greatly profitable in a endless cluster of spaces such as stocks, deals, climate or indeed sports. Displayed here is the think about and execution of a few gathering classification calculations utilized on deals information, comprising of week after week retail deals numbers from diverse offices in Walmart retail outlets all over the Joined together States of America The total deals examination are recognized with the created usage.

(vii) Graph Analysis

A total exactness expectation is made with the comparison of this four calculation. The precision forecast of the 4 calculations appears distant better;a much better;ahigher;astronger;an improved">a higher chart examination. Here the deals examination are done with the linear regressor 72.9% , XG boost 81.34%, Random Forest 89.45% and KNN gives 92.72%. In this way the KNN gives the finest exactness ever.

(IV) RESULT AND DISCUSSION

Machine Learning calculations such as Straight Relapse, KNearest Neighbors calculation, XG Boost calculation and Arbitrary Woodland calculation have been utilized to anticipate the deals of different outlets of the Enormous Bazaar. Different parameters such as Root Cruel Squared Blunder (RMSE), Fluctuation Score, Preparing and Testing Exactnesses which decided.

The winning accommodation for the Kaggle competition had a Cruel Supreme Blunder (MAE) of around 2130 [11]. As a reference, a accommodation where all the anticipated values of week by week deals are 0's, the MAE is found to be roughly 21000. In this ponder, the final 20% of the preparing dataset was utilized as the neighborhood test-set. The Angle Boosting calculation was taken as a standard and the MAE was found to be 5771.5, with a R^2 score of 0.80 that suggests that 80% of the anticipated values were exact. These were the most excellent comes about gotten with the $n_{\text{estimatorshyperparameter}}$, which alludes to the number of choice trees that are utilized for relapse, set at 200. The other hyperparameters were set to their default values. Table 1 alludes to diverse values given to the parameters and the comes about that taken after

Algorithms	Accuracy	MAE
Random Forest	89.45%	0.45
Linear Regression	72.9%	1.24
K Nearest <u>Neighbour</u>	92.72%	0.19
XG Boost	81.34%	1.75

Table 1 Comparison chart of four algorithms

The Amazingly Randomized Trees calculation works marginally superior than the randomforest. This increment in execution may be ascribed to higher randomization within the preparing handle. The `n_estimators` parameter was set to 150, whereas the `min_samples_split` and `min_samples_leaf` parameters were put at 2 and 1 respectively, to get the finest comes about wherein the MAE was 1965.5 and R2 score was 0.94. Table 3 alludes to distinctive values given to the parameters and the comes about that taken after.

It was famous that as the esteem of the `n_estimatorshyperparameter` is expanded past the values given within the tables over, it was found that the MAE was expanding rather than diminishing, conceivably inferring that the models were overfitted. Also, increasing the number of relapse trees unpredictably isn't exhorted because it leads to expanded computational escalated coming about in a bigger sum of time went through in preparing the show without benefitting its precision.

This venture appears that there are exceedingly effective calculations to estimate deals in huge, medium or little organizations, and their utilize would be advantageous in giving important knowledge, in this way driving to superior decision-making.

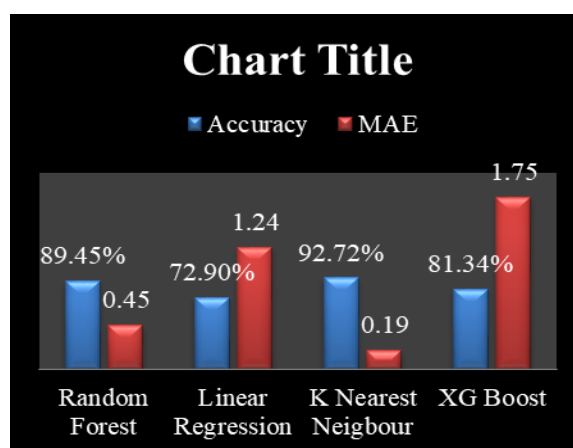


Figure 1 comparison chart of the machine learning algorithms

The system shows the random forest at the highest accuracy system.

(V) CONCLUSION

In conclusion, Wal-Mart is the number one retailer within the USA and it too works in numerous other nations all around the world and is moving into unused nations as a long time pass by. There, are other companies who are continually rising as well and would donate Walmarta extreme competition within the future in case Walmart does not remain to the best of their amusement. In order to do so, the individuals will have to be get it their commerce patterns, the client needs and oversee the assets shrewdly. In this time when the innovations are coming to out to unused levels, Enormous Information is taking over the conventional strategy of overseeing and analyzing information. These advances are always utilized to get it complex datasets in a matter of time with lovely visual representations. Through watching the history of the company's datasets, clearer thoughts on the deals for the past a long time was realized which is able be exceptionally accommodating to the company on its possess.

(VI) FUTURE ENHANCEMENT

Furthermore, regularity slant and arbitrariness and future estimates will offer assistance to analyze deal drops which the companies can maintain a strategic distance from by employing a more focused and efficient strategies to play down the deal drop and maximize the benefit and stay in competition. The exactness of the expectation can be improved in future so that the recognizable proof of deals can be found on the leading.

(VII) REFERENCES

- [1]Baba, Norio, and HidetsuguSuto. "Utilization of artificial neural networks and GAs for constructing an intelligent sales prediction system."In Proceedings of the IEEE-INNS-ENNS International Joint Conference on Neural Networks.IJCNN 2000. Neural Computing: New Challenges and Perspectives for the New Millennium, vol. 6, pp. 565-570. IEEE, 2000.
- [2]Cheriyana, Sunitha, Shaniba Ibrahim, SajuMohanana, and Susan Treesa. "Intelligent Sales Prediction Using Machine Learning Techniques."In 2018 International Conference on Computing, Electronics & Communications Engineering (iCCECE), pp. 53-58.IEEE, 2018.
- [3]Fawcett, Tom, and Foster J. Provost. "Combining Data Mining and Machine Learning for Effective User Profiling."In KDD, pp. 8-13. 1996.
- [4]Friedman, Jerome H. "Stochastic gradient boosting." Computational Statistics & Data Analysis 38.4 (2002): 367-378.
- [5]Giering, Michael. "Retail sales prediction and item recommendations using customer demographics at store level." ACM SIGKDD Explorations Newsletter 10, no. 2 (2008): 84-89.
- [6]Kulkarni, Vrushali Y., and Pradeep K. Sinha. "Random forest classifiers: a survey and future research directions." Int J AdvComput 36.1 (2013): 1144-53.
- [7]Panjwani, Mansi, Rahul Ramrakhiani, Hitesh Jumnani, Krishna Zanwar, and RupaliHande. Sales Prediction System Using Machine Learning.No. 3243.EasyChair, 2020.
- [8]Ragg, Thomas, Wolfram Menzel, Walter Baum, and Michael Wigbers. "Bayesian learning for sales rate prediction for thousands of retailers."Neurocomputing 43, no. 1-4 (2002): 127-144.
- [9]Sekban, Judi. "Applying machine learning algorithms in sales prediction." (2019).
- [10]Singh Manpreet, BhawickGhutla, Reuben Lilo Jnr, Aesaan FS Mohammed, and Mahmood A. Rashid."Walmart's Sales Data Analysis-A Big Data Analytics Perspective." In 2017 4th Asia-Pacific World Congress on Computer Science and Engineering (APWC on CSE), pp. 114-119. IEEE, 2017.