

Career Path Insights through Advanced Random Forest Techniques: A Comprehensive Analysis

Aditya Jadhav¹, MohitKumar Pandey², Sahil Zine³, Muzfaar Baksh⁴ and Deepti Pawar⁵

¹⁻⁵Department of Computer Engineering Shah and Anchor Kutchhi Engineering College Mumbai India

Email: aditya.jadhav16005@sakec.ac.in, mohitkumar.pandey15716@sakec.ac.in, sahil.zine15620@sakec.ac.in, muzfaar.bakhsh15588@sakec.ac.in, deepti.pawar@sakec.ac.in

Abstract—“Career Path Insights” is a digital platform that helps students navigate their path to success through a blend of technology, user engagement, and educational tools. It provides a user-friendly website with various features to support students in making well-informed career choices. Some of its key features include user registration, data input forms for academic and skill details, prediction of eligibility using machine learning, personalized career suggestions, user dashboards, interactive data display, a resume creator, and additional features to improve the overall user experience.

Index Terms— Digital, Cross Platform Logistic Regression, Decision Tree.

I. INTRODUCTION

In a time where information-driven arrangements reclassify professional scenes, our noteworthy venture begins to lead the way, offering a fundamental toolkit for early career searches. Destined to enhance the profession options of users, our drive not only attempts to predict the potential of individuals in various jobs but also grants some form of professional experience on choosing prospects. With state-of-the-art calculations and industry information, we mean to increase occupation openings by identifying the students with accuracy and exactness. Engaged with computerized reasoning, our venture tends to the emerging difficulties of the contemporary work market, situating itself as a central member in steering the path towards career paths of the future.

The main aim for this job is to innovate for the development of an AI model that performs correctly in predicting prospects for successful performance at the chosen job at an early stage, making it possible to practice better mediation processes towards vocation enhancement and career upgrading. We focus on model development that yields the most accurate and precise results possible, on levels that can effectively cut down jumbles among competitors and open positions. This involves advanced AI processes to help create a powerful prescriptive system that recognizes and predicts the probability of an understudy's outcome in their future professional ways. Beyond this primary model building, our job imagines expanding its abilities by consolidating more boundaries and other sources of information. This essential approach means to improve the predictions of capability expectations to ensure an accurate device for finding the best career paths in a dynamic job market. Find out your true potential, explore your vocation process with confidence, and choose the career way that is uniquely made up of your special qualities and desires with our leading ability understanding venture.

II. LITERATURE REVIEW

This writing survey coordinates an enormous number of studies fully intent on using AI procedures for

foreseeing understudy situation and scholastic achievement, along with related results. The examinations are diary articles and meeting papers, which have been embraced to increment interest in the utilization of information driven ways to deal with make instructive cycles more productive. Here is an integrated outline of the writing: This writing covers the utilization of AI calculations for foreseeing understudy positions and scholastic execution. The examinations center around anticipating understudy arrangements in organizations or enterprises in light of scholastic records, individual ascribes, and other applicable elements.

Kandi and Sharan (2022) [6] present a review, "Placement Prediction and Analysis Using Machine Learning," which is possible among the many zeroed in on the use of prescient models to foresee the result of understudy positions. Additionally, Thangavel et al. [7] (2017) and Rao et al. [2](2018) talk about suggestion frameworks and instructive information mining procedures to foresee understudy arrangement.

Harinath et al. [3](2019) and Hussain and Singh [11] (2021) investigate AI application for understudy arrangement expectation, most likely recommending experiences into calculation decision and assessment of model execution.

Besides, Beaulac and Rosenthal [12] (2019) and Çakıt and Dağdeviren [13] (2022) dig into foreseeing scholastic achievement and majors utilizing AI, offering a more extensive viewpoint past forecast of understudy situation.

Different examinations, for example, those by Aravind et al. [14] (2019) and Spandana and Pallavi [16] (2023), think about various AI calculations for anticipating position data, giving helpful benchmarks to display determination and execution appraisal.

Besides, Gupta et al. (2021) and Patel et al. [19] (2017) examine enlistment frameworks and information digging strategies for grounds arrangement expectation, separately, featuring functional applications and difficulties in genuine situations.

Furthermore, different investigations, for example, those by Chandrahase et al. (2023) and Sreenivasa Rao et al. [2] (2015), can give a premise to highlight extraction procedures and information preprocessing strategies applied in the models foreseeing understudy position. All in all, the writing gives sign of developing interest in anticipating understudy position and scholastic achievement utilizing AI models as well as related instructive results. These examinations add to the headway of prescient demonstrating procedures and proposition important experiences for teachers, scientists, and policymakers trying to improve understudy results through information driven approaches.

III. METHODOLOGY

A. Problem Statement

Therefore, for understudies to get properly customized direction in pursuing informed profession decisions or figure out their arrangement qualification, more need to be done. To meet this, the "Profession Way Experiences" stage is going to use computerized reasoning and AI in offering undergraduates customized vocation direction and sharp position projections.

B. Approach

In the briskly changing position market, understudies commonly require custom direction to follow through informed professional decisions and to understand their situation-specific qualifications. For this purpose, we create the "Profession Way Bits of Knowledge" stage, which incorporates human brainpower and AI to offer customized career direction and precise position predictions. The first step is data preprocessing for the purposes of going on with the progresses:

Data Cleaning: This stage includes the discarding of irrelevant information, handling missing characteristics, and overseeing inconsistencies. It ensures that the data is consistent and accurate.

Data Mix: Data mix is a very important system that integrates information from different sources into a single dataset. This way, with additional information, the datasets become more complete, thereby bringing out more accurately what the data scene is. The procedure mixes data from different sources to provide a more complete view and examination of the available information.

Data Change: Information preprocessing is readied for man-made intelligence models by changing the information into a convenient configuration. Operating with the following elements: scaling for uniform element ranges, normalization for normalizing information, and dimensionality decrease for effective model utilization.

Data Decline: Information decline involves reducing the size of the dataset without losing fundamental information. This includes subtractive partial integrate confirmation and joining extraction, which form new parts. These methods improve the datasets for on-site execution of worked-on models and capacity.

Data Information Discretization: The most striking strategy used to change consistent information into discrete information for showing enormous datasets or situations when models require overall information.. It changes the data view to make it more realistic for examination and exhibitions.

Data Testing: The most widely recognized approach to filtering a subset of the data to set up the model is the data reviewing method. This is useful while watching over huge datasets that the model cannot be managed by. These preprocessing stages are really basic for building convincing computer-based intelligence models. Preprocessing ensures that the data is good, consistent, and good enough for simulated intelligence models, provoking better model execution and precision.

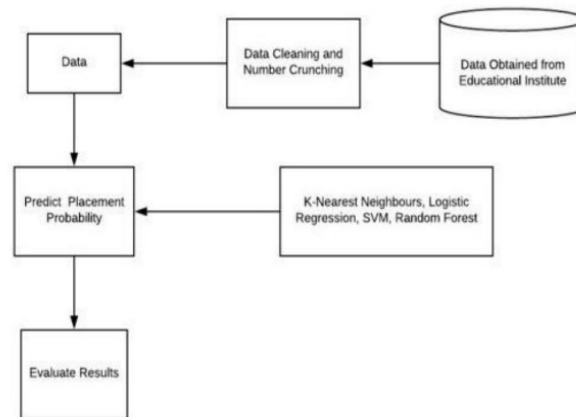


Fig.1: Proposed Model

C. Dataset

- The "Rating" column is probably the overall rating or reputation of the companies, probably based on employee reviews, industry standings, or other performance metrics. Information of this kind can be rather instrumental in figuring out which organizations are highly rated among job seekers who want to join reputed workplaces.
- The "Company Name" column gives us the names of the companies associated with the job listings. This helps in knowing the organizations that provide these vacancies, understanding the profiles, values, and culture. The 'Job Title' column is full of several job positions for which salaries have been reported. It helps the user in understanding the variety of job chances while scrutinizing the trendiness of job titles, also highlighting certain high-demand roles in the list.
- 'Salary' consists of the compensation associated with each job title. Based on this column, it can be analyzed for wage ranges, trends, and possible correlations with other things such as location, company rating, or educational background.
- 'Salaries Reported' may specify the number of reported salaries for the particular job title. This information is necessary in order to assess the reliability and representativeness of the salary data for that particular job position.
- The 'Location' column specifies where the jobs are located. It is important for people considering relocation or for assessing regional differences in salaries and employment prospects.
- 'Employment Status' classifies the employment conditions, distinguishing between full-time, part-time, or contractual positions. This column is important for job seekers looking for particular types of work. SSC Marks, HSC Marks, Sem-1 to Sem-8 CGPA columns give a complete picture of the academic background of people in the dataset. This information can be useful in examining possible correlations between academic achievements and job positions or salaries. The 'Skills' column lists the various skills present with these individuals. This column can be analyzed to find out which skills are in high demand in the job market, helping job seekers in acquiring those skills required for any of the jobs.

Programming Languages' speaks about the specific programming languages relevant to the listed positions. This is very important for an individual who has special knowledge in programming, for example, as it makes them able to relate their skills to the expectations of the industry. This dataset provides a general picture of different parts of the gig market, such as company notorieties, work titles, pay rates, scholastic foundations, and required abilities.

$$MSE = \frac{1}{N} \sum_{i=1}^N (f_i - y_i)^2$$

Where N is the number of data points,
 f_i is the value returned by the model and
 y_i is the actual value for data point i .

Fig 5: Numerical Equation of Random Forest

Stage 2: Preprocessing the Information Since we have the information preprocessed, we can move to our second step of setting up the information for AI calculations. The following are the particular assignments that are done:

- Taking care of missing qualities: We decide to drop the lines that contain missing information since this is the straightforward yet successful answer for handle them.
- Envision the dispersion of one specific mathematical component, 'Sem-1 CGPA', utilizing a histogram.
- One-hot encoding clear cut factors: This interaction changes unmitigated factors over completely to twofold vectors so they might be utilized by AI models.
- Scaling mathematical highlights: We utilize the StandardScaler to standardize mathematical elements with the goal that they will have zero mean and unit change.

Stage 3: Parting the Dataset Before we can prepare our AI model, we should part our dataset into preparing and testing sets. That is achieved by utilizing scikit-learn's 'train_test_split' capability. We likewise separate our highlights (X) from our objective variable (y) at this stage. Stage 4: Preparing an Random Forest The following stage is to utilize our dataset to prepare an Random Forest. Random Forest calculation is a troupe learning technique, which addresses choice trees through haphazardly choosing tests from preparing information. Here are the boundaries of the Arbitrary Woodland Classifier:

- Number of assessors: 100
- Arbitrary state: 42

Stage 5: Assessing the Model At last, in this last step, we assess the exhibition of our model on the testing dataset. We do the accompanying:

- Guarantee consistency in encoding clear cut factors for the testing dataset.
- Get forecasts on the testing information utilizing our prepared model.

Evaluate the precision and a characterization report for model execution. The exactness score lets us know how well our model predicts the right class names; a characterization report gives more nitty gritty measurements like accuracy, review, and F1-score for each class. These measurements empower us to assess the adequacy of the model on the testing information.

1. Logistic Regression

Strategic relapse is a standard simulation intelligence technique of the worldview Regulated Learning. It is based on expecting a hidden variable using several independent factors. The results of this dependent variable are usually binary, like Yes or No, 0 or 1, Valid or Bogus, and so on.

Calculated relapse, on the other hand, instead of giving accurate values like 0 and 1, gives likelihood values between 0 and 1. This is what both strategic and direct relapse share: calculated relapse fundamentally used for classification tasks, while direct relapse is used for classification issues. In calculated relapse, a sigmoid-molded bend is fitted to the information to foresee the likelihood of the occurrence.

The bend addresses the probability of different results, for example, whether cells are destructive, or in the event of a mouse's weight on view. Calculated relapse is important as it can create probabilities and arrangement with new information using both persisted and discrete datasets. It's powerful in characterizing perceptions using various sorts of information and in recognizing the most compelling factors for grouping tasks.

i. Approach:

Step 1: Splitting the Dataset

To proceed, we need to split the dataset into two parts - preparation and test sets - for making predictions. We do this using a function called `train_test_split()`. This function takes the entire dataset and separates it into two sets: one for training and one for testing.

In the training set, we have two important components: elements (X) and the target variable (y). Elements are the features or variables that we will use to make predictions, while the target variable is the variable we want to predict. We assign these components to two variables: X_train and y_train.

Similarly, in the testing set, we have the same components: elements (X) and the target variable (y). We assign these components to two different variables: X_test and y_test.

To ensure that our results are reproducible, we set the random_state parameter to 0. This means that every time we run the code, we will get the same split of the dataset into the training and testing sets. This is helpful for consistency in our analysis.

The next step involves applying one-hot encoding to our data, which has taken place in step two.

Stage 2: One-Hot Encoding

In this stage, we take factors like "Organization Name" and "Occupation Title" and turn them into a form that is easier to work with using one-hot encoding.

One-hot encoding takes each unique name or title and creates separate sections or columns for it. Then, for each entry in your data, it assigns a value of 1 if that entry corresponds to that particular name or title, and 0 otherwise.

For example, let's assume our dataset has organization names like "Apple," "Microsoft," and "Google." One-hot encoding will result in three separate columns, one for each organization name. If an entry in the dataset corresponds to "Apple," the value in the "Apple" column would be 1, and the values in the "Microsoft" and "Google" columns would be 0.

The same process applies to occupation titles. We create separate columns for each unique title and assign values of 1 or 0 based on whether an entry matches that title. We use a parameter called "drop_first=True" so as not to suffer from multicollinearity, where the columns become too correlated. This helps keep our data accurate and reliable.

Stage 3: Finding Important Features

Imagine you have a magic machine that automatically picks out the most convincing elements from a bunch of options. Well, that's what the RandomForestClassifier model does! It helps us identify the most persuasive factors. But we don't stop there - we also use the SelectFromModel technique to select the top highlights based on their importance. This helps make our model even more effective, like adding the perfect ingredients to a recipe to make it taste amazing.

Stage 4: Model Preparation

In this stage, we prepare the Calculated Relapse model using a special set of features (X_train_selected) and the corresponding target variable (y_train). The model helps us understand how these features are connected to the desired outcome, making it a valuable tool in various AI applications. It's like having a map that shows us the important connections between different factors and the end result.

Stage 5: Testing and Evaluation of the Model

In this stage, we will put our Calculated Relapse model to the test. We will use a set of selected testing highlights called X_test_selected.

Once we have the test case ready, we will evaluate it using the model. We will calculate the accuracy score, which tells us how well the model predicts the outcome. Additionally, we will generate a characterization report that provides us with specific measurements such as accuracy, recall, and F1-score.

Just like trying out a recipe, before we decide if it really tastes good or not, after baking it up according to the recipe. Maybe it didn't turn out as sweet as you wanted, so you tested it out, tried adding a pinch of salt or sugar to adjust the taste. In this way, you'd probably end up with a recipe that tasted good and looked great.

Similarly, in this stage, we will put our model to the test. We want to see how accurate it is in predicting outcomes. Just like evaluating a recipe, we will measure the model's performance using various measurements to get a clear picture of its effectiveness.

If we are satisfied with the model's accuracy, we can consider deploying it or using it for further predictions. It's like giving the model an ultimate test for it to pass with flying colors.

The specific approach followed here organizes new information in a clear and efficient way. It takes into consideration both numerical data and personal opinions or feelings.

This precise methodology works with exact and productive order of new pieces of information, taking into account both quantitative and subjective elements.

Instead, strategic relapse can be seen as a carefully thought-out decision to determine outcomes within a clear range, like a scale from 0 to 1. This decision is not dependent on whether it is supported by a complex mathematical concept like a hyperplane or a simple straight line.

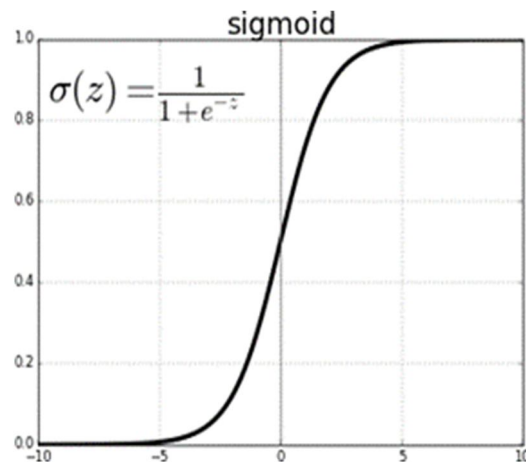


Fig 6: Numerical Representation Of Logistic Regression

Advantages:

- Logistic regression is a robust measure that focuses on analyzing data that is targeted at one variable in view. It is easy to apply, easy to understand, and easy to use. Consider it an efficient and fast way to uncover hidden patterns and relationships in your data.
- Logistic regression is one of the huge advantages in that it does allow us to gauge how strong the connections between a variety of factors are and what elements are more important to us. It's like shining a light on the key drivers and insights of our model, making it possible to better understand all that interconnection.

2. Evaluation Metrics

Accuracy: 1.00					
Classification Report:					
	precision	recall	f1-score	support	
Android	1.00	1.00	1.00	564	
Backend	1.00	0.99	0.99	271	
Database	1.00	1.00	1.00	181	
Frontend	1.00	1.00	1.00	445	
iOS	1.00	1.00	1.00	342	
Java	1.00	1.00	1.00	335	
MobIle	0.98	0.98	0.98	48	
Python	0.97	0.99	0.98	186	
SDE	1.00	0.99	1.00	1636	
Testing	0.98	1.00	0.99	348	
Web	1.00	0.99	1.00	198	
accuracy			1.00	4554	
macro avg	0.99	0.99	0.99	4554	
weighted avg	1.00	1.00	1.00	4554	

Fig 7: Evaluation Metrics

E. FrontEnd

Executing structure styling in JavaScript includes controlling the Report Article Model (DOM) to apply styles to HTML components powerfully. To begin, make use of methods like `getElementById`, `getElementsByClassName`, or `querySelector` access the targeted HTML elements. Utilize JavaScript to create and modify style attributes after selecting the elements. For example, you can use the `style` property to set an explicit style: you can set CSS properties like `backgroundColor`, `fontSize`, or `edge`. You can add or remove CSS classes powerfully using the `classList` property to apply predefined styles. Occasionally, audience members can be used to cause style changes based on client communications, increasing the intelligence of the page.. This dynamic styling approach empowers designers to make responsive and versatile points of interaction, modifying the visual show of components in view of client activities or changing information states. In general, by utilizing JavaScript to control styles, designers can make more powerful and drawing in client encounters on the web.

F. BackEnd

To integrate a Django backend with a Node.js application for job placement, there has to be communication between the two technologies. Django is the backend framework and so acts to handle the processing of data, authentication, and business logic, while Node.js is a frontend that tends to be more concerned with the interaction of users and shows information. The Flask backend of Django exposes endpoints to the frontend Node.js application, which is more often done with HTTP requests, to get data from the backend. Node.js can use popular libraries such as Axios or the built-in HTTP module to make HTTP requests to the Flask endpoints of Django to fetch or submit data as per the need.

This combination empowers a dynamic and responsive work position stage where Node.js handles client-side rationale, UI communications, and ongoing updates, and Django deals with the server-side tasks, data set collaborations, and confirmation. This cooperative execution guarantees a durable and effective framework that places the qualities of the two innovations into great use for a proficient work position application. The huge measure of information related with work postings, competitor profiles, and other significant data is overseen and coordinated inside a MySQL data set while setting up position. Tables for work listings, candidate subtleties, organizations, and client records can be remembered for the information base outline. For example, the work postings table might store data, for example, work title, organization subtleties, and required abilities. Up-and-comer profiles can incorporate instructive foundation, abilities, and work insight.

Organizations tables can hold insights concerning selecting associations. Client accounts tables can store data for validation and approval purposes.

The execution includes making and keeping up with these tables, laying out connections among them, and advancing inquiries for effective information recovery. SQL inquiries are used to embed, update, recover, and erase information in view of client associations inside the arrangement entry. Ordering can be applied to upgrade inquiry execution, guaranteeing speedy and exact information recovery. Besides, MySQL's help for exchanges guarantees information trustworthiness during complex tasks, giving a dependable groundwork to basic cycles in the situation entry, for example, requests for employment and up-and-comer shortlisting.

By using MySQL as the information base backend, the position entrance benefits from a vigorous, versatile, & social data set administration framework that guarantees consistency of the information trustworthiness, and effective recovery, eventually adding to a consistent and dependable experience for both work searchers and bosses utilizing the gateway.

IV. RESULT

KNN, Logistic Regression, Random Forest, and SVM are some of the evaluated machine learning algorithms for job placement and prediction. The placement prediction system is supposed to forecast the likelihood of successful job placement for individuals. In the Python environment, these algorithms are used for data analysis and prediction.

The results reveal that Random Forest has reached a highest accuracy; thus, our study concludes. It is quite obvious that Random Forest has shown up much better performance than the other algorithms, keeping an accuracy rate of up to 97% in predicting job placement. The high accuracy of Random Forest is a good enough reason for such applications.

A. Features

- a. Visualization: While creating a data portrayal that will be used in a task-position entryway, it is important to take care of the needs of key partners such as work seekers, bosses, and executives. Key execution markers (KPIs) will include work postings, skill needs, and industry pattern should be developed with imagery such as bar diagrams and line charts used to present connections and patterns. Ease of interpretation has been highlighted with the segments of work postings, up-and-comer profiles, and interview plans to be included on the dashboard. Incorporate experiences, such as market patterns, topographical development, and ability planning. Carry out intelligent channels for client evaluation, promise portable responsiveness, and conduct standard updates and examination to scratch areas of interest all through the expedition.
- b. Resume Builder: A resume creator is a contraption helping people with making capable resumes through guiding them through the collaboration and making composed records. It generally includes an

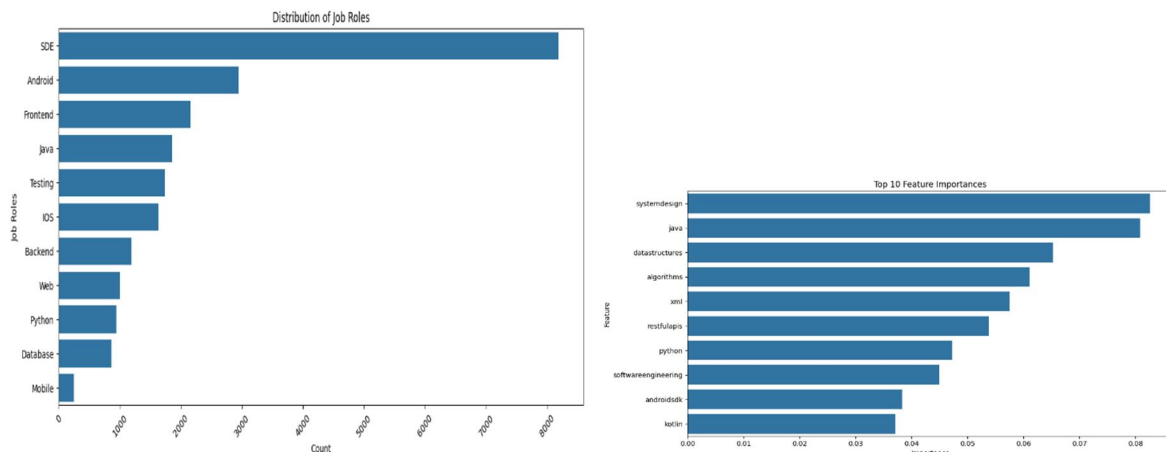


Fig 8: Graph of Each Job Distribution

easy to-comprehend interface where clients enter their own and skilled data. The construction then, uses predefined plans and arrangements to synchronize information into an ideal outcome.

- c. **ChatBot:** Integrating the Chatbot interface into Gig Way Encounters upgrades client commitment and backing. Use the OpenAI library in Python to get a Programming Connection point key that would allow communication between the Chatbot. Define features according to client prompts; decide the Chatbot's responses; and integrate it into the application for input/yield handling. Store history of communication for logically important reactions; ensure secure Programming interface key administration; and constantly test for accuracy, with reference to OpenAI documentation in regard to model determination.
- d. **Google Calendar:** Using the Google Schedule Programming interface can accommodate booking of new employee screenings inside your association. Coordinate the Programming interface to make booking of occasion possible, such as creation, refreshing, and recovery of meeting dates. Confirm application access through OAuth 2.0. Clients can easily make HTTP solicitations to be able to effectively handle the interview booking activities. The Google Schedule Programming interface smooths out the coordination of booking functionalities that enhance proficiency in planning arrangements or meetings inside the association.
- e. **Community Forum:** Setting up a local area includes choosing the stage, preparing the facade, and programming the local area. For example, using Talk, which is a well-known local area stage, establishes a commonality of correspondences with the team member.
- f. **Progress Tracking:** Modifying the Vocation Way Experiences stage involves handling its appearance and design, for instance, through the administrator board like getting sorted out a meeting on Talk. Enable client enlistment, ensure security conventions, and construct the local area for an open and helpful climate. Enhance commitment by starting appropriate conversations and cultivating client collaboration. For continued maintenance, guarantee standard updates, dissect client input, and direct effect appraisals to fine-tune and streamline the Vocation Way Bits of Knowledge experience. This structure, guarded by Docker for Talk, gives a platform to making a stage where people can communicate and find out their ways of doing their jobs efficiently.

B. Future Scope

- a) Collaborative efforts constitute gathering around a common goal, with the use of joint liabilities and beneficial interactions. They stimulate imagination by leveraging shared capacities and viewpoints within a group. A portfolio is a coordinated combination showing individual or group achievements and capacities. Planning useful projects into a portfolio works on the master profile, showing both individual responsibility and collaboration capacities with respect to possible managers or partners.
- b) Modified learning personalizes the learning process and frames such things to suit an individual's prerequisites, tendencies, and learning pace. It embeds a flexible methodology of altered content, versatile plans, and different educational strategies that provide continuous observation of express learning styles and goals. It licenses students to attract with material at their own speed, overhauling appreciation and support. This approach allows for the incorporation of development and data

assessment, which regularly changes learning toward maintaining smooth opportunities for growth for each part.

- c) A LinkedIn Flask Streamline smoothes out the joining of LinkedIn highlights into applications, interfacing experts and growing organizations. The Programming Interface empowers admission to client profiles, computerized network development, and consistent cooperation among applications and LinkedIn. For instance, it might be combined with a task stage to enable clients to apply via LinkedIn profiles or get work proposals from their organization. Moreover, it upholds highlights such as sharing updates, accessing organization data, and upgrading joint effort and systems administration inside the LinkedIn system. Essentially, it serves as a scaffold, enhancing client experiences and organization development inside the LinkedIn system.
- d) Carrying out a mastery-based contest contemplates middle space, spreading out clear principles and measures, and making a streamlined enrollment and evaluation process through development. Posse the opposition to draw individuals, developing skill improvement and strong contest. Skill-based challenges fill various requirements from capability confirmation to proficient achievement, providing a coordinated stage to showing limits and adding to capacity improvement within unambiguous spaces.

V. CONCLUSION

This makes exact characterization results using the numerous calculations, hence contrasting the results of different calculations. The Irregular Backwoods Calculation is used to make the most extreme precision with the preparation of information in the given model. Afterward, we will establish major areas of strength in the framework and use instruments that will enable clients to track their achievements, competency development, and overall growth. For systems administration, going to occasions, and utilizing online stages. A situation occasion schedule directs a client's efforts. It gives helpful criticism on mock meetings. Criticism and cooperation between clients, together with continuous self-reflection, will streamline your arrangement.

ACKNOWLEDGMENT

The Shah and Anchor Kutchhi Engineering College authorities are appreciated by the writers for their assistance throughout the research project. Without their persistent support, conducting research in this subject would not have been possible.

REFERENCES

- [1] Machine Learning” 2017 International Conference on advanced computing and communication systems(ICACCS-2017), Jan06-07,2017,Coimbatore, INDIA.
- [2] K. Sreenivasa Rao, N. Swapna, P. Praveen Kumar Educational data mining for student placement prediction using machine Learning algorithms Research Paper, International Research Journal of Engineering and Technology (IRJIET) 20182015, 153,286–299. [CrossRef]Kohavi, R and F, Provost (1998). Machine Learning 30:271-274.
- [3] Shreyas Harinath, Aksha Prasad, Suma H and Suraksha A. Student Placement Prediction using Machine Learning, International Research Journal of Engineering and Technology (IRJIET) Volume : 06 Issue: 04 April 2019
- [4] Pavani, N. M. Pujitha, P. V. Vaishnavi, K. Neha and D. S. Sahithi, "Feature Extraction based Online Job Portal," 2022 International Conference on Electronics and Renewable Systems (ICEARS), Tuticorin, India, 2022, pp.1676-1683,doi:10.1109/ICEARS53579.2022.9752295.
- [5] V. Pavani, N. M. Pujitha, P. V. Vaishnavi, K. Neha and D. S. Sahithi, "Feature Extraction based Online Job Portal," 2022 International Conference on Electronics and Renewable Systems (ICEARS), Tuticorin, India, 2022, pp.1676-1683,doi:10.1109/ICEARS53579.2022.9752295.
- [6] Naresh Patel K M, Goutham N M, Inzamam K A, Suraksha V Kandi, Vineet Sharan V R, 2022, Placement Prediction and Analysis using Machine Learning, INTERNATIONAL JOURNAL OF ENGINEERING RESEARCH & TECHNOLOGY (IJERT) ICEI – 2022 (Volume 10 – Issue 11).
- [7] Student Placement Analyzer: A Recommendation System Using Machine Learning”, Senthil Kumar Thangavel , Divya Bharathi P, Abijith Sankar, International Conference on Advanced Computing and Communication Systems (ICACCS -2017), Jan. 06 - 07, 2017, Coimbatore, INDIA
- [8] "Class Result Prediction using Machine Learning", Pushpa S K, Associate Professor, Manjunath T N, Professor and Head, Mrunal T V, Amartya Singh, C Suhas, International Conference on Smart Technology for Smart Nation, 2017
- [9] S. Chandrasana, S. M. S. Ganeshan, R. C. Maddineni, M. S. Divya, P. Tumuluru and B. Suneetha, "Machine Learning Algorithms based Student Performance Prediction based on Previous Records," 2023 7th International Conference on Computing Methodologies and Communication (ICCMC), Erode, India, 2023, pp. 181-186, doi: 10.1109/ICCMC56507.2023.10084099.

- [10] "Campus placement Prediction Using Supervised Machine Learning Techniques" International Journal of Applied Engineering Research ISSN 0973-4562 Volume 14, Number 9 (2019) pp. 2188-2191
- [11] Md Shadab Hussain , Sarita Singh "Developing classifiers through Machine learning algorithms for student placement prediction based on academic performance", Applied Artificial Intelligence, Vol 35(2021), Issue 6
- [12] Beaulac, C., & Rosenthal, J. S. (2019). Predicting university students' academic success and major using random forests. *Research in Higher Education*, 60(7), 1048–1064.
- [13] Çakıt, E., Dağdeviren, M." Predicting the percentage of student placement: A comparative study of machine learning algorithms". *Educ Inf Technol* 27, 997–1022(2022).
- [14] T. Aravind, B. S. Reddy, S. Avinash and J. G., "A Comparative Study on Machine Learning Algorithms for Predicting the Placement Information of UnderGraduate Students," 2019 Third International conference on I-SMAC
- [15] C. S. K and K. S. Kumar, "Data Preprocessing and Visualizations Using Machine Learning for Student Placement Prediction," 2022 2nd International Conference on Technological Advancements in Computational Sciences (ICTACS), Tashkent, Uzbekistan, 2022, pp. 386-391, doi: 10.1109/ICTACS56270.2022.9988247. A
- [16] G. M. Spandana and L. Pallavi, "Placement Prediction System using Machine Learning," 2023 2nd International Conference on Edge Computing and Applications (ICECAA), Namakkal, India, 2023, pp. 903-907, doi: 10.1109/ICECAA58104.2023.10212409.
- [17] S. Gupta, A. Hingwala, Y. Haryan and S. Gharat, "Recruitment System with Placement Prediction," 2021 International Conference on Artificial Intelligence and Smart Systems (ICAIS), Coimbatore, India, 2021, pp. 669-673, C
- [18] N. K. Sharma, A. K. Singh, S. Salvi and S. S. More, "College Kart and k-NN Algorithm based Placement Prediction," 2021 Second International Conference on Electronics and Sustainable Communication Systems (ICESC), Coimbatore, India, 2021, pp. 1271-1276, doi: 10.1109/ICESC51422.2021.9532987.
- [19] Patel T., et al "Data Mining Techniques for Campus Placement Prediction in Higher Education." *India J.Sci.Res.* 14 (2) 2017.
- [20] "Class Result Prediction using Machine Learning", Pushpa S K, Associate Professor, Manjunath T N, Professor and Head, Mrunal T V, Amartya Singh, C Suhas, International Conference on Smart Technology for Smart Nation, 2017
- [21] Aharva Kulkarni , Tanuj Shankarwar , Siddharth Thorat, 2021, Personality Prediction Via CV Analysis using Machine Learning, INTERNATIONAL JOURNAL OF ENGINEERING RESEARCH & TECHNOLOGY (IJERT) Volume 10, Issue 09 (September 2021)
- [22] G. Eason, B. Noble, and I. N. Sneddon, "On certain integrals of Lipschitz-Hankel type involving products of Bessel functions," *Phil. Trans. Roy. Soc. London*, vol. A247, pp. 529–551, April 1955.
- [23] J. Clerk Maxwell, *A Treatise on Electricity and Magnetism*, 3rd ed., vol. 2. Oxford: Clarendon, 1892, pp.68–73.
- [24] I. S. Jacobs and C. P. Bean, "Fine particles, thin films and exchange anisotropy," in *Magnetism*, vol. III, G. T. Rado and H. Suhl, Eds. New York: Academic, 1963, pp. 271–350.
- [25] K. Elissa, "Title of paper if known," unpublished.
- [26] R. Nicole, "Title of paper with only first word capitalized", *J. Name Stand. Abbrev.*, in press.
- [27] Y. Yorozu, M. Hirano, K. Oka, and Y. Tagawa, "Electron spectroscopy studies on magneto-optical media and plastic substrate interface," *IEEE Transl. J. Magn. Japan*, vol. 2, pp. 740–741, August 1987 [Digests 9th Annual Conf. Magnetism Japan, p. 301, 1982].
- [28] [7] M. Young, *The Technical Writer's Handbook*. Mill Valley, CA: University Science, 1989.