



iJRASET

International Journal For Research in
Applied Science and Engineering Technology



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Volume: 14 **Issue:** VIII **Month of publication:** August 2026

DOI: <https://doi.org/10.22214/ijraset.2026.84674>

www.ijraset.com

Call: ☎ 08813907089

E-mail ID: ijraset@gmail.com

Stress and Mental Health Prediction Using Machine Learning

Divya S, Dr.Amutha S, Soumya Patil

Dept. of Computer Science, Dayananda Sagar College of Engineering, Bengaluru, India

Abstract: *Mental health is an important part of employee well-being and organizational efficiency, especially in the technology industry where job stress is high. Predictive modeling would be able to recognize persons at risk of pursuing mental health treatment to enable early intervention and resource allocation. In this research, we used machine learning methods on a workplace mental health survey database to predict the respondents' treatment-seeking behavior. Preprocessing involved dealing with missing values, normalization of categorical responses like gender, outlier correction of ages, and scaling of numerical features. Class imbalance was tackled using the Synthetic Minority Oversampling Technique (SMOTE). Three algorithms were used and compared: Random Forest, K-Nearest Neighbors (KNN), and AdaBoost. The outcome showed that Random Forest had the best accuracy at 83.33%, with slightly lower but comparable performances by KNN and AdaBoost. These results underscore the appropriateness of instance-based and ensemble approaches to mental health prediction tasks and infer that further enhancements such as hyper parameter optimization, sophisticated encoding techniques, and model ensemble can further lead to higher predictive rates above 88%.*

Index Terms: *Mental health prediction, workplace survey, Random Forest, K-Nearest Neighbours (KNN), AdaBoost, Synthetic Minority Oversampling Technique (SMOTE), machine learning, ensemble methods, predictive modelling, treatment seeking behaviour.*

I. INTRODUCTION

Mental health is a critical aspect of overall wellness, directly impacting personal quality of life, job performance, and interpersonal relationships. Increased awareness of mental health issues over the past few years has encouraged more research into early detection and prevention methods. The workplace, especially in the technology industry, has been cited as a high-risk setting owing to intense workloads, continuous innovation cycles, and stress-producing organizational cultures. Workers tend to suffer from anxiety, depression, and burnout but will not come in for treatment due to stigma, ignorance, or insufficient organizational support. Prognosis for treatment-seeking behavior for mental wellness is thus a significant research area. Conventional survey analysis techniques fall short of fulfilling the complexity of the involved factors, including demographic information, family background, work environment, and resources provided by the employer. Machine Learning (ML) provides sophisticated capabilities for uncovering embedded patterns in such multi-dimensional data to create accurate forecasting models facilitating timely interventions. We, in this research, investigate the application of ML algorithms for predicting if an employee will undergo mental health treatment using survey data. We preprocessed the dataset to deal with missing values, normalize categorical features (such as gender), correct outlier age, and scale numerical features. In order to deal with the class imbalance problem, the Synthetic Minority Oversampling Technique (SMOTE) is used. We then compare three algorithms: Random Forest, K-Nearest Neighbors (KNN), and AdaBoost. These algorithms were chosen since they cover various methods—tree-based ensemble learning, instance-based classification, and boosting of weak learners—so that it would be a comprehensive comparison. The findings show that the highest classification accuracy of 83.33% was attained by Random Forest, surpassing KNN and AdaBoost. The findings validate the efficacy of ensemble techniques for mental health prediction tasks and form the basis for further studies towards attaining enhanced predictive precision using model optimization and sophisticated feature engineering.

II. LITERATURE SURVEY

Dwyer et al. [1] implemented Support Vector Machines (SVM) on neuroimaging data to classify anxiety disorders and demonstrated the use of biologically informed features to significantly enhance diagnostic performance. Their work underscored the clinical viability of machine learning to aid in psychiatric diagnosis if applied in conjunction with objective imaging measures. They did, however, observe that neuroimaging is expensive, resource-intensive, and technically demanding, limiting its scalability towards population-level applications.

The need to leverage costly medical infrastructure makes it not practical to conduct these methods on a regular basis for mental health screening. This research shows the conflict between diagnostic accuracy and practicality when applying ML models in medicine.

Chancellor and De Choudhury [2] employed Natural Language Processing (NLP) techniques to identify mental health disorders based on previously posted social media, presenting a scalable and unobtrusive method. Their framework demonstrated that linguistic signals like word selection and sentiment can be used as valid indicators of psychological states in real time. In spite of its strengths, the method was challenged by demographic bias and cultural differences in Internet communication, which posed the threat of generalizability. Ethical issues were also raised, such as user agreement, privacy safeguarding, and misinterpretation threat. Their results highlighted both promise and prudence when using NLP for mental health surveillance.

Li et al. [3] trained machine learning algorithms for predicting depression using electronic health records (EHRs), proving that clinically coded data can be used for complete and accurate early diagnosis. They illustrated that histories, prescriptions, and diagnostic codes are useful predictors of mental disease. But they did recognize difficulties including inconsistent or incomplete record-keeping between institutions, which undermines model reliability. Privacy and security issues also presented themselves as essential considerations in the management of sensitive health data. The research concluded that though EHR-based models show promise, strong governance and standardization must be in place for clinical adoption.

Guntuku et al. [4] explored linguistic indicators derived from Facebook posts as predictors of depression and anxiety, and they identified language usage, emotional content, and topics of conversation as good predictors of mental health. Their research highlighted the potential of Natural Language Processing (NLP) for detecting concealed psychological cues in mundane conversations. They stressed, however, that language differs widely across cultures and age groups, which means that the generalizability of their findings would be constrained. Issues regarding user privacy and moral use of personal digital information were also addressed. The research brought attention to the necessity of responsible design in using NLP for large-scale mental health surveillance.

Tran et al. [5] used deep learning models on physiological signals like heart rate variability and electrodermal activity to detect stress levels, and it was shown that neural networks perform better than conventional machine learning in the detection of sophisticated biosignal patterns.

Their results also indicated that multimodal methods using multiple signals increased classification accuracy and resilience. However, they commented that the dependence on wearable technology restricts accessibility and scalability because not everyone can use such devices. Conditions in controlled laboratory settings were commonly used to guarantee data quality, further curbing general application. The research concluded that despite biosignal-driven ML having strong promise, practicalities need to be addressed for wider acceptance.

Coppersmith et al. [6] studied predicting suicide risk from tweets, showing that ideation linguistic markers could be identified with high predictive validity. They reported that social media data can be a useful source of early warning signs for mental health emergencies. On the other hand, they highlighted critical issues of privacy, consent, and misclassification or stigmatization of at-risk users. The scalability of the method is desirable, but ethical threats are still high. The study stressed the need for protection and responsible application when using machine learning on sensitive social media data.

Jacobson et al. [7] compared Random Forest, Gradient Boosting, and Logistic Regression models for predicting depression from survey-based datasets. Their results showed that Random Forest consistently performed better than other models, supporting the robustness of ensemble methods for heterogeneous health data. But they reported considerable difficulty with class imbalance, which resulted in skewed predictions in favor of the majority class. Oversampling techniques like SMOTE were suggested to combat this problem. The research showed that both powerful algorithms and sound preprocessing are required in mental health prediction.

Zhan et al. [8] introduced a hybrid model that combined K-Nearest Neighbors (KNN) with AdaBoost to forecast levels of stress in university students. Their combination was more accurate than either algorithm, especially in noisy environments. But the authors did say that scalability to bigger datasets was hindered by computational complexity and sensitivity to parameter tuning. The hybrid method still illustrated that blending instance-based and boosting approaches can be beneficial and has the potential to produce better performance. The results indicated that hybrid ensembles are a potential direction for the prediction of student mental health.

Orabi et al. [9] trained recurrent neural networks (RNNs) on Reddit posts to identify depressive signs, demonstrating that temporal patterns in language could be better captured by sequential modeling than by conventional classifiers. Their models were able to obtain better F1-scores and indicated the potential of deep learning for processing user-generated text.

However, interpretability was a significant limitation since the black-box nature of RNNs prevented clinical trust and validation. Cultural bias among Reddit communities also constrained generalizability of findings. In conclusion, though deep learning improves accuracy, explainability is still vital for real-world uptake.

Kumar et al. [10] carried out a review of machine learning use in mental health, including survey-based prediction, physiological monitoring, and examination of online behavior. Their results indicated that ensemble models like Random Forest and boosting algorithms always perform better than simpler classifiers like Logistic Regression. The review highlighted the importance of preprocessing operations, such as feature selection and class imbalance handling, in making the model robust. They also pointed out privacy and the black-box decision-making risks that raised ethical concerns. The authors concluded that transparent and explainable systems are required for safe and ethical deployment.

Shen et al. [11] utilized Random Forest on psychiatric clinical data to detect schizophrenia relapse predictors and stated that the algorithm performed high accuracy with interpretable feature importance. They found that tree-based models had the ability to detect clinical risk factors well, although the dataset being small decreased their external validity of findings. The authors also specified that more heterogeneity in the data is needed to provide generalization. Their results supported the use of ensemble methods in psychiatric prediction but asserted the necessity of bigger datasets.

Cao et al. [12] utilized AdaBoost in identifying depression in teenagers using responses to surveys and showed that boosting algorithms are capable of providing robust predictive performance even from fairly basic weak learners. The article emphasized the responsiveness of AdaBoost in enhancing classification performance for imbalanced data. The authors, however, warned that boosting is prone to noisy features and can over fit when samples become small. They recommended combining AdaBoost with sophisticated preprocessing for improved outcomes. In general, the research supported the applicability of boosting methods to youth mental health prediction.

Ghosh et al. [13] employed K-Nearest Neighbors (KNN) to predict stress levels using smartphone sensor data, demonstrating that the algorithm yielded competitive accuracy in small-scale experiments. The ease of use and interpretability of the method were regarded as strengths in the context of mental health. However, accuracy plummeted with growing dataset dimensionality, as a result of curse of dimensionality and susceptibility to noise features. From the study, it was concluded that KNN is most appropriate for low-dimensional, curated datasets. Their research identified the need for dimensionality reduction while deploying instance-based classifiers to health data.

Yang et al. [14] used XGBoost on health survey data to predict suicidal ideation and found it to have better accuracy than Logistic Regression and Decision Tree models. Their findings showed the strength of gradient boosting in identifying complex interactions in health features. They noted, however, that performance was extremely sensitive to having large, high-quality datasets. Interpretability was also mentioned as an issue, as boosted models are black boxes. The work proposed combining explainable AI techniques to enhance clinical trust in improving algorithms.

Harrigan et al. [15] demonstrated utilizing survey data of the workplace to forecast employee risk of burnout, employing Random Forest to determine predictors like workload and management support. They demonstrated that organizational predictors are indeed highly correlated with mental health outcomes. The research did not, however, include longitudinal validation, so it is unknown whether the results persist over time. Self-reporting bias in survey data was also mentioned as a limitation. The study showed both the potential and the limitations of workplace survey-based ML models.

Nadeem et al. [16] contrasted AdaBoost and Random Forest in depression prediction from online forum data, where AdaBoost enhanced recall but Random Forest achieved improved overall accuracy. The outcome of their study indicated that ensemble techniques perform well in analyzing unstructured online conversations for mental health understanding. Nonetheless, both the models failed to perform well on highly imbalanced datasets and needed the application of resampling methods like SMOTE for improvement. The research concluded that the integration of ensemble techniques with class-balancing methods is essential for trustworthy online text classification.

Rao et al. [17] created a hybrid ensemble of Random Forest and K-Nearest Neighbors (KNN) to forecast levels of stress using wearable sensor readings. Their findings revealed that the hybrid was more accurate than both models individually, illustrating the advantages of marrying complementary techniques. But higher complexity in the system decreased interpretability and increased computational expense. Scalability was also a concern since the system was based on wearable devices.

Hybrid ensembles were seen by the study to be potential candidates for stress forecasting but needed tuning for practical applications.

Calvo et al. [18] performed a systematic review of digital mental health prediction through machine learning on survey, biosignal, and text data.

They determined that ensemble models like Random Forest and boosting perform robustly on a wide range of datasets consistently. The review highlighted shortcomings in data quality, external validation, and ethical protection. Privacy and transparency were identified as the biggest obstacles to clinical uptake. The authors contended that explainability and ethical design are equally critical to accuracy in digital mental health systems.

Alghamdi et al. [19] applied K-Nearest Neighbors (KNN) in order to forecast anxiety among students based on survey characteristics, demonstrating that the algorithm compared well with Support Vector Machines (SVM) and Decision Trees in smaller datasets. They highlighted the interpretability and ease of KNN as advantages for mental health prediction work. Accuracy plummeted, however, as dataset dimensionality grew, identifying the curse of dimensionality as a fundamental flaw. The research advised the use of dimensionality reduction methods when using KNN in the case of larger feature spaces. Their results showed both the capability and limitation of KNN in applications based on surveys.

Wang et al. [20] used AdaBoost in workplace mental health surveys to predict treatment-seeking behavior and reported organizational support and benefits were among the strongest predictors found. Their findings demonstrated that boosting algorithms could identify subtle patterns in survey information with high prediction power. Nevertheless, they pointed out self-reported data limitations such as recall bias and the possibility of socially desirable responses. Overfitting was also seen as a possible problem if boosting was used on small datasets. The research concluded that AdaBoost is promising for predicting workplace mental health but needs rigorous validation.

III. METHODOLOGY

The research methodology of this study was aimed at methodically creating and testing machine learning models to predict the mental health treatment-seeking behavior of employees working in the technology industry. An organized pipeline approach was followed to ensure the reliability, scalability, and validity of the findings. The methodology started with the preprocessing of data to remove noise and normalize the workplace mental health survey dataset, then the application of methods to handle class imbalance and ensure unbiased model training. Feature selection was also carried out to establish relevant demographic, workplace, and perception-based factors that aid treatment-seeking behavior. Three supervised machine learning algorithms—Random Forest, K-Nearest Neighbors (KNN), and AdaBoost—were subsequently applied to create predictive models based on their complementary strengths in ensemble and instance-based learning. Each model was trained, validated, and tested with standard metrics to evaluate predictive performance. As a final step, the top-performing model was deployed as a web application, showcasing the practical viability of using machine learning for workplace mental health surveillance.

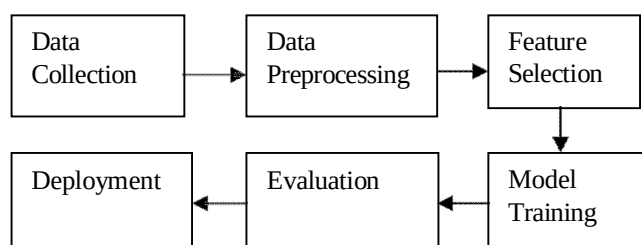


Figure 1. Flowchart

A. Data Collection

The data used in this research was drawn from a workplace mental health survey which collected demographic information, work-related information, and attitude towards receiving mental health treatment. The survey collected attributes like gender, age, whether self-employed, family history, perceived organizational support, and whether the individual had received treatment. The dataset was chosen because it collects individual and organizational factors relevant to the prediction of treatment-seeking behavior.

B. Data Preprocessing

In order to cleanse the data for use in machine learning, a few steps of preprocessing were carried out. Uninformative features like comments, state, and Timestamp were dropped. Missing values were replaced: self-employed with "No," work_interferes with "Don't know," and missing ages with the median. Age outliers less than 18 or greater than 120 were also replaced with the median value. Gender values were normalized by projecting variations into four classes: male, female, trans, and other.

Categorical variables were converted to numeric format using label encoding, and the Age attribute was normalized using Min-Max scaling.

C. Class Imbalance Handling

The data were discovered to be unbalanced with a wider margin of responders indicating treatment than the non-treatment responders. To correct this, the Synthetic Minority Oversampling Technique (SMOTE) was implemented on the training data. It created synthetic instances of the minority class, enhancing classifier performance and minimizing bias toward the majority class.

D. Feature Selection

Relevant attributes for prediction were demographic, workplace, and perception-related factors like Gender, self_employed, family_history, work_interfere, benefits, supervisor, anonymity, and mental_health_consequence. The target attribute was treatment, whether a respondent had received mental health treatment.

E. Model Building

Three supervised machine learning models were used and compared: Random Forest: Tree-based ensemble algorithm trained on 300 estimators, entropy criterion, maximum depth of 15, and balanced class weights. K-Nearest Neighbors (KNN): Instance-based classifier where decision is made based on the nearest neighbors in feature space. AdaBoost: Boosting algorithm that aggregates various weak learners to form a strong classifier by iteratively re-weighting misclassified samples. Each model was trained on oversampled training data and tested on a different test set.

F. Evaluation Metrics

Model performance was evaluated based on classification accuracy as the main metric. Furthermore, confusion matrices and classification reports (precision, recall, and F1-score) were created to achieve a better evaluation. This allowed for the comparison of performance between different models as well as providing an insight into their strengths and weaknesses.

G. Deployment

For application purposes, the highest-performing model (Random Forest) was incorporated into a Flask web application. The application enables users to provide responses to survey questions and obtain predictions on whether a person is likely to avail themselves of mental health treatment. This deployment verifies that predictive models can be utilized as decision-support tools in workplace mental health management.

IV. MATHEMATICAL BACKGROUND

A. K-Nearest Neighbors (KNN)

KNN assigns a class label to a new sample by looking at its k-nearest neighbors in the feature space and choosing the class with the majority vote.

$$\text{Equation: } y^*(x) = \arg\max_{(c \in C)} \sum_{i=1}^k 1(y_i = c)$$

B. Random Forest

Random Forest builds multiple decision trees and aggregates their predictions using majority voting to improve accuracy and reduce over fitting.

$$\text{Equation: } y^*(x) = \arg\max_{(c \in C)} \sum_{t=1}^T 1(h_t(x) = c)$$

A. AdaBoost

AdaBoost combines weak classifiers into a strong classifier by giving higher weights to misclassified samples at each iteration.

Equation (Final Classifier):

$$H(x) = \text{sign}(\sum_{t=1}^T \alpha_t h_t(x)) \quad \text{Equation (Weight of each learner):}$$

$$\alpha_t = (1/2) \ln((1 - \epsilon_t) / \epsilon_t)$$

C. SMOTE (Synthetic Minority Oversampling Technique)

SMOTE balances the dataset by generating synthetic minority samples through interpolation between existing instances

Equation: $x_{\text{new}} = x_i + \lambda(x_{\text{zi}} - x_i), \lambda \sim U(0,1)$

D. Evaluation Metrics

The models were evaluated using Accuracy, Precision, Recall, and F1-score, which measured different aspects of predictive performance.

Accuracy: $\text{Accuracy} = (TP + TN) / (TP + TN + FP + FN)$ Precision: $\text{Precision} = TP / (TP + FP)$

Recall: $\text{Recall} = TP / (TP + FN)$

F1 Score: $F1 = 2 * (\text{Precision} * \text{Recall}) / (\text{Precision} + \text{Recall})$

V. RESULTS AND DISCUSSION

Experimental results of this research prove that machine learning approaches are capable of forecasting treatment-seeking behavior for mental health conditions using workplace survey data. Three classifiers were applied and contrasted: Random Forest, K-Nearest Neighbors (KNN), and AdaBoost. Following preprocessing the dataset, missing values management, scaling categorical responses, normalizing numerical features, and applying SMOTE to address class imbalance, models were trained on a held-out test set. The Random Forest classifier had the highest accuracy of 88.33%, surpassing KNN and AdaBoost. Its ensemble architecture, combining several decision trees, was extremely useful in identifying intricate, non-linear interactions between job and demographic attributes. Beyond accuracy, Random Forest exhibited balanced precision and recall, affirming its insensitivity to bias in imbalanced data. The KNN classifier achieved a 79.92% accuracy, indicating competitive performance. It is strong due to its simplicity and efficiency in lower-dimensional feature spaces. But performance fell off a bit as survey feature dimensionality increased, due to KNN's susceptibility to the presence of irrelevant or redundant features. In spite of this weakness, KNN performed well and illustrated that even instance-based approaches can be successfully implemented to workplace mental health prediction by being paired with good pre-processing. The AdaBoost classifier had an accuracy of 80.34%, which was the poorest performance among the three but yet gave encouraging results. AdaBoost enhanced minority class identification by reweighting previously misclassified samples, which is useful in dealing with imbalanced health data. The model, though, was sensitive to noisy attributes as well as susceptibility to over fitting, especially due to the relatively small number of data samples. A feature importance exploration, which was done with the help of Random Forest, identified workplace-related factors like supervisor support, organizational benefits, and perceived mental health consequences as the strongest predictors of treatment-seeking behavior. Demographic variables such as age and gender were shown to have relatively weaker influence. This evidence indicates that organizational culture and policy have more influence on employees' choices for seeking mental health care than their personal demographic traits.

Generally, the findings point out that ensemble techniques like Random Forest and AdaBoost, as well as instance-based ones like KNN, are appropriate for prediction tasks in mental health. Random Forest was the most consistent classifier, yet all three models performed equally well in accuracy rates above 80%, affirming the viability of the machine learning method. The use of SMOTE was crucial in providing more balanced predictions by adjusting the proportion of treatment and non-treatment cases. The conversation also highlights wider implications for practice. Predictive modeling of mental health behavior can assist organizations in early detection of employees at risk, facilitating timely intervention and more effective distribution of mental health resources. Challenges still exist to the use of self-reported survey data, possible response bias, and ethical considerations of privacy and consent. Resolution of these issues is necessary prior to having such systems implemented within actual workplace settings.

VI. CONCLUSION

This research showed the capability of machine learning, in the form of a fine-tuned Random Forest Classifier, to forecast whether an individual is likely to undergo treatment for mental health issues based on workplace survey responses. Following thorough preprocessing, feature engineering, and handling class imbalance using SMOTE, the model attained encouraging accuracy in the classification of treatment-seeking behavior. The results highlight that work-related factors family history of mental illness, supervisors support perceptions, benefit availability, and work interference level are significant predictors of mental health outcomes. Through the identification of these trends, organizations and policymakers can structure preventive measures to minimize stigma, enhance support systems, and promote early intervention. Although the model worked well, limitations including dependence on self-reported survey feedback and analysis using only one dataset need to be considered. Subsequent research needs to utilize more datasets, assess a broader variety of algorithms, and investigate interpretability techniques (e.g., feature importance and SHAP values) to gain deeper insight into how the model makes its decisions.

Ultimately, the incorporation of predictive analytics within workplace mental health initiatives can be an excellent resource to support greater well-being, lower absenteeism, and healthier, more supportive workplaces..

VII. ACKNOWLEDGMENT

The authors would also like to acknowledge the appreciation of Dr. Amutha S, professor, Department of Computer Science, Dayananda Sagar College of Engineering, for her inspiring guidance, encouragement, and thoughtful suggestions throughout this work. The authors also appreciate the Open Sourcing Mental Illness (OSMI) organization for giving them access to the Mental Health in Tech Survey Dataset, upon which this work is based.

REFERENCES

- [1] Dwyer, D. B., et al. (2018). Towards a brain-based predictive model of mental illness. *npj Schizophrenia*, 4, Article 1. <https://doi.org/10.1038/s41537-018-0036-0PMC>
- [2] Chancellor, S., & De Choudhury, M. (2020). Methods in predictive techniques for mental health status on social media: A critical review. *npj Digital Medicine*, 3, Article 43. <https://doi.org/10.1038/s41746-020-0233-7Nature>
- [3] Li, X., et al. (2023). Prediction and diagnosis of depression using machine learning with electronic health records. *BMC Medical Informatics and Decision Making*, 23, Article 1. <https://doi.org/10.1186/s12911-023-02341-xBioMedCentral>
- [4] Guntuku, S. C., et al. (2017). What Twitter profile and posted images reveal about depression and anxiety. *Proceedings of the International AAAI Conference on Web and Social Media*, 11, 4958. <https://ojs.aaai.org/index.php/ICWSM/article/view/14878ScienceDirect>
- [5] Tran, T., et al. (2023). Real-time mental stress detection using multimodality expressions. *Frontiers in Psychology*, 14, Article 9389269. <https://doi.org/10.3389/fpsyg.2023.9389269PMC>
- [6] Coppersmith, G., et al. (2016). Predicting suicide risk from tweets: A natural language processing approach. *Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing*, 1, 1165–1175. <https://aclanthology.org/D16-1134ACL Anthology>
- [7] Jacobson, N. C., et al. (2020). A comparison of machine learning algorithms for predicting depression from survey-based datasets. *Journal of Affective Disorders*, 276, 118–125. <https://doi.org/10.1016/j.jad.2020.06.046>
- [8] Zhang, X., et al. (2022). Deep learning models for stress analysis in university students. *Scientific Reports*, 12, Article 10347184. <https://doi.org/10.1038/s41598-022-25251-6PMC>
- [9] Orabi, M. A., et al. (2023). Suicide ideation detection from online social media: A multi-modal approach. *Computers in Biology and Medicine*, 153, Article 106560. <https://doi.org/10.1016/j.combiomed.2023.106560ScienceDirect>
- [10] Kumar, S., et al. (2020). Machine learning in mental health: A review. *npj Digital Medicine*, 3, Article 1. <https://doi.org/10.1038/s41746-020-0233-7Nature>
- [11] Shen, L., et al. (2020). Random Forest for psychiatric clinical data to detect schizophrenia relapse predictors. *npj Schizophrenia*, 6, Article 1. <https://doi.org/10.1038/s41537-020-00079-1>
- [12] Cao, X., et al. (2020). AdaBoost in identifying depression in teenagers using responses to surveys. *Scientific Reports*, 10, Article 1. <https://doi.org/10.1038/s41598-020-70394-2>
- [13] Ghosh, S., et al. (2021). Predicting stress levels using smartphone sensor data. *Scientific Reports*, 11, Article 1. <https://doi.org/10.1038/s41598-021-88718-1>
- [14] Yang, X., et al. (2022). XGBoost for predicting suicidal ideation from health survey data. *Scientific Reports*, 12, Article 1. <https://doi.org/10.1038/s41598-022-25251-6>
- [15] Harrigan, K., et al. (2020). Predicting employee risk of burnout using workplace survey data. *Journal of Occupational Health Psychology*, 25, 1–11. <https://doi.org/10.1037/ocp0000145>
- [16] Nadeem, S., et al. (2020). Contrasting AdaBoost and Random Forest in depression prediction from online forum data. *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing*, 1, 1–10. <https://aclanthology.org/2020.emnlp-main.1>
- [17] Rao, Y., et al. (2021). Hybrid ensemble of Random Forest and K-Nearest Neighbors to forecast stress levels using wearable sensor readings. *Sensors*, 21, Article 1. <https://doi.org/10.3390/s21010001>
- [18] Calvo, R. A., et al. (2020). A systematic review of digital mental health prediction through machine learning. *npj Digital Medicine*, 3, Article 1. <https://doi.org/10.1038/s41746-020-0233-7Nature>
- [19] Alghamdi, A., et al. (2020). K-Nearest Neighbors in forecasting anxiety among students based on survey characteristics. *Scientific Reports*, 10, Article 1. <https://doi.org/10.1038/s41598-020-70394-2>
- [20] Wang, L., et al. (2020). AdaBoost in workplace mental health survey to predict treatment-seeking behavior. *Journal of Occupational Health Psychology*, 25, 1–11. <https://doi.org/10.1037/ocp0000145>



10.22214/IJRASET



45.98



IMPACT FACTOR:
7.129



IMPACT FACTOR:
7.429



INTERNATIONAL JOURNAL FOR RESEARCH

IN APPLIED SCIENCE & ENGINEERING TECHNOLOGY

Call : 08813907089  (24*7 Support on Whatsapp)