
PREDICTIVE MODELING OF WORK-LIFE BALANCE AMONG WOMEN TEACHERS IN TAMIL NADU USING MACHINE LEARNING ALGORITHMS

BHUVANESWARI J

DEPARTMENT OF MANAGEMENT STUDIES PERIYAR MANIAMMAI INSTITUTE OF SCIENCE & TECHNOLOGY
(DEEMED TO BE UNIVERSITY), THANJAVUR, INDIA

THAMILVANAN G

DEPARTMENT OF EDUCATION PERIYAR MANIAMMAI INSTITUTE OF SCIENCE & TECHNOLOGY (DEEMED TO
BE UNIVERSITY), THANJAVUR, INDIA

GURU P

DEPARTMENT OF MANAGEMENT STUDIES PERIYAR MANIAMMAI INSTITUTE OF SCIENCE & TECHNOLOGY
(DEEMED TO BE UNIVERSITY), THANJAVUR, INDIA

CHRISTY A N

DEPARTMENT OF COMMERCE PERIYAR MANIAMMAI INSTITUTE OF SCIENCE & TECHNOLOGY (DEEMED TO
BE UNIVERSITY), THANJAVUR, INDIA

LAKSHMI BALA M

DEPARTMENT OF BUSINESS ADMINISTRATION, KUNTHAVAI NAACHIYAAR GOVT. ARTS COLLEGE
(AUTONOMOUS), THANJAVUR, INDIA

SWATHI S

DEPARTMENT OF MANAGEMENT STUDIES PERIYAR MANIAMMAI INSTITUTE OF SCIENCE & TECHNOLOGY
(DEEMED TO BE UNIVERSITY), THANJAVUR, INDIA

KEERTHANA K

DEPARTMENT OF MANAGEMENT STUDIES PERIYAR MANIAMMAI INSTITUTE OF SCIENCE & TECHNOLOGY
(DEEMED TO BE UNIVERSITY), THANJAVUR, INDIA

Abstract

The work life balance (WLB) is one of the most important determinants of professional performance and personal wellbeing of women teachers in Tamil Nadu. As the needs of teaching, administration, and family requirements continue to rise, the forecasting of the levels of WLB may offer useful information to the policy-makers and the institutions, to develop an effective support system. The proposed research hypothesizes that it is possible to predict and model WLB among the women teachers with the help of real-time data based on the structured questionnaire that will be administered with the use of Google Forms. The numerical responses in the form of a dataset were pre-treated, applying the methods of working with missing values, normalizing features, and equal distribution of classes. They used two machine learning algorithms, namely, the k-Nearest Neighbors (K-NN) and the Logistic Regression (LR) to categorize teachers into three types of WLB: poor, moderate and good. Model analysis had been conducted through k-fold cross-validation and the measures of accuracy, precision, recall, F1-score and ROC-AUC. Comparative findings revealed the strengths of K-NN in the local data structure capture, whereas Logistic Regression could offer interpretable information on important predictors of WLB. The paper emphasizes that quantitative predictors of WLB outcome include workload, teaching hours, family commitments, and institutional support, which have a considerable effect on the outcomes. The findings do not only add to the usage of machine learning in educational and social studies, but also offer evidence-based suggestions to better the work-life balance of women educators in Tamil Nadu.

Keywords: Work-Life Balance, Women Teachers, Machine Learning, K-Nearest Neighbors (K-NN), Logistic Regression.

1. INTRODUCTION

One important field of study has become work life balance (WLB) especially when it comes to education especially when it comes to women in the profession as they tend to have two responsibilities in the form of both work and family. Women teachers in the Indian setting and Tamil Nadu in particular have a central role to play in not just life in learning institutions but also in influencing the social and cultural life of communities. Teaching and administrative demands, student mentoring and research demands, combined with household and caregiving tasks, tend to create problems of maintaining a healthy WLB [1]. Unaddressed these challenges may lead to stress, burnout, low job satisfaction and poor overall well-being.

The dynamic character of the work settings in the recent years, particularly due to the effects of digitization, the online methods of teaching, and the shift in institutional policies, has exacerbated the necessity to assess and forecast WLB among educators through systematic and data-driven methodologies. Descriptive statistics and qualitative analysis have been the major methods of studying WLB, which, although informative, do not have the predictive power of policy intervention and supporting mechanisms [2]. Incorporation of machine learning algorithms in social science studies offers the researcher with the potent instrument to model intricate connections among several predictors and WLB findings that gives a predictive power and analytical knowledge.

Real-time information was gathered in this research among women teachers in the state of Tamil Nadu by use of structured questionnaires that were distributed using Google Forms. The answers, which were mostly an outcome of the numerical Likert-scale questions, revealed some of the major facets of professional workload, family demands, the level of stress, and the support systems offered by the institution. Two machine learning algorithms, including k-Nearest Neighbors (K-NN) with its capacity to identify patterns in local data and Logistic Regression (LR) with its interpretability and the potential to determine meaningful predictors of categorical data, were used to analyze such a dataset. These models have been used to categorize the respondents into three types namely poor, moderate and good work-life balance thus offering a diagnostic and predictive insight into work-life balance among women educators [3]. Through these methods, the study does not only improve the methodological implementation of machine learning in educational studies, but it is also an empirical study in the sense that it provides indispensable information that may be used to inform the institutional policies, work practices, and personal coping mechanisms. Finally, the results are supposed to contribute to improving the welfare and professional performance of women educators in the educational organizations in Tamil Nadu.

Problem Statement

Work-life balance (WLB) is a decisive factor of both professional productivity and well-being especially in women teachers who often juggle between work and home among other obligations in their workplace. Female teachers in Tamil Nadu are making significant contributions to the education system, but they generally have difficulties, which are excessive workload, long working hours, administration, and family life. The combination of these two overlapping demands may result in stress, burnout, decreased job satisfaction, and work-life integration.

Whereas earlier research on WLB in the teaching field has offered useful descriptive data, the majority of studies have used traditional methods of statistics, which is limited in the sense that they fail to respond to diverse, non-linear connections among various contributing variables. Also, predictive modeling strategies that could be used to classify and predict the degree of work-life balance depending on real-time information are lacking. In the absence of this foresight, the policymakers and institutions cannot be able to come up with effective interventions based on data to meet the needs of women teachers.

With the growing access to digital means of data collection, e.g. online surveys, and the development of machine learning methods, such a gap can be bridged. Using the algorithm that includes k-Nearest Neighbors (K-NN) and the Logistic Regression (LR), one can construct the models of predictive validity to classify the outcomes of WLB (poor, moderate, good) and, additionally, determine the strongest contributors to these results. Nevertheless, the lack of this kind of research in the case of Tamil Nadu indicates a substantial gap in research that should be filled.

2. LITERATURE SURVEY

In [4], Holgado-Apaza et al., (2024) uses machine learning on ENDO-2020, which is the national survey of teachers of public basic-education in Peru that was conducted in 2020, to determine predictors of life satisfaction. Filtering (mutual information, ANOVA, chi-square, Spearman) used feature selection, embedded (CART, Random Forest, Gradient Boosting, XGBoost, LightGBM, CatBoost) and then modeling with Random Forest, XGBoost, Gradient Boosting, CART, CatBoost, LightGBM, SVM and MLP. The major predictors were satisfaction with health; working in a learning institution; family living conditions; the ability to carry out teaching responsibilities; age; trusting the Ministry of Education and the Local Management Unit (UGEL); engaging in the process of continuing training; reflective practice; work-life balance; and time spent preparing lessons and managing the administrative department. LightGBM and Random Forest have achieved the best results (LightGBM: accuracy 0.68, precision 0.55, F1 0.55,

0.42, Jaccard 0.41, 0.41). The results highlight the importance of ML in educational management by identifying dissatisfied teachers and notifying specific evidence-based policy responses.

In [5], Holgado-Apaza et al., (2025) perform a survey of the National Survey of Directors, in this research, ensemble feature selection was performed, and five ML models used (RF, CART, HistGB, XGBoost, LightGBM) were used to forecast job satisfaction of principals. Among the important drivers were the satisfaction with salaries, school location, student/teacher relationships, climate in the working place, student performance, benefits, as well as economic (gross/net income, minimum needed income) and time (training hours, working off-hours, UGEL travel/stay time, commute) aspects. The optimal model, which is the Histogram-Based Gradient Boosting, but tuned using Bayesian optimization and trained using the Random Oversampling achieved balanced accuracy of 0.63 on a test set with natural class balance, and with GANs balancing the training set only, the recall was 0.74, the precision was 0.72 and the F1 was 0.70. SHAP has shown financial reasons as the primary cause of dissatisfaction, whereas interpersonal reasons prevail in highly satisfied principals, indicating the role of hierarchy needs and providing practical indications to support policy-making based on the data.

In [6], the survey in the Expectancy Confirmation Model and TAM, a SEM study of 282 university users of generative AI discovered that application of the knowledge and perceived intelligence enhance perceived usefulness and confirmation; confirmation, in its turn, enhances perceived usefulness and satisfaction. Continuance intention is substantially influenced by perceived usefulness and satisfaction, as well as social influence, but not AI configuration. The model describes 64.1% of continuance intention variance, which can be interpreted as recommendations when developing and marketing effective AI-based language tools in education.

In [7], Diana, B., and Kvr, R reveal how it is possible to utilize synthetic data to understand the actual difficulty of work-life balance (WLB) frustrations in the sample of female higher-education teachers in the Thanjavur district of Tamil Nadu. Increasing data with artificially created data and implementing various classifiers of the ML, the authors forecast the stability of the WLB and demonstrate the methodological usefulness of synthetic data in the research of social science. The results are used both in educational administration and policy and provide practical suggestions on legislative and institutional changes that can enhance WLB among women.

In [8], Yoo, J.E., and Rho, M (2020) OECD TALIS 2013 group-level analysis using machine-learning predictor scores, this paper employed group Mnet (penalized regression) to cull 558 teacher job satisfaction predictors. Out of 100 random splits that used cross-validation, variables chosen more than half of the time produced 18 strong predictors. In line with the previous research findings, collaborative school climates and teacher self-efficacy were found to be important predictors; new predictors were teacher feedback, participatory school climates, and perceived barriers to professional development. The outcomes narrow down the set of predictors of job satisfaction and indicate school climate, feedback systems, and access to PD as the factors of action, and future research implications are mentioned.

3. PROPOSED METHODOLOGY

The research design adopted by this study is a quantitative, predictive research design where machine learning algorithms are used to predict the level of work-life balance (WLB) in a cohort of women teachers in Tamil Nadu. Survey-based data collection is combined with computational modeling to offer predictive insights as well as interpretability in the design. Below figure 1 elaborates the step by step procedure of proposed work.

3.1 Data Collection

The design of a properly designed Google Form, which was prepared with the help of the local anchoring committee, helped to gather the 750 valid responses of the women teachers working in the Thanjavur District. The first dataset consisted of 25 variables of both professional and personal nature regarding work life balance, including the average hours of work, the intensity in workload, work flexibility, family and caring roles, the support systems provided by the institution (e.g. child care facilities, leave policies), and the degree of job satisfaction. The survey was a mixture of quantitative Likert-scale items with qualitative inputs which made sure that there was a comprehensive evaluation on the conditioning factors that determine WLB. After an intensive data cleaning exercise that involved elimination of data inconsistencies and missing values, the dataset was narrowed down to 14 significant attributes which were chosen according to their relevance and use in the study. The qualitative answers were transformed into the numerical ones through label encoding and one-hot encoding methods where binary responses were encoded to Yes = 1, No = 0, Maybe = 0.5. The presence of missing entries (NaN) was filled with a neutral value (0) in order to ensure data integrity. The data was divided into training and testing (70 and 30 respectively) to develop a predictive model that was well-assessed. Even though the original data had an imbalance in classes, thereby decreasing the predictive accuracy, a synthetic data generation method (SMOTE) [9] was used to obtain balance in WLB categories. This systematic method was a guarantee to achieving quality and representative final dataset, which would be a solid basis to analyze machine learning.

3.2 Pre-Processing

A very important step in converting raw survey data into machine learning data is preprocessing. It makes sure that the data is clean, consistent and it is converted to form that can be handled by algorithms. Primary procedures involve cleaning of the data, coding, missing data management, normalization data and data partitioning.

i. Data Cleaning

A very important step in converting raw survey data into machine learning data is preprocessing. It makes sure that the data is clean, consistent and it is converted to form that can be handled by algorithms. Primary procedures involve cleaning of the data, coding, missing data management, normalization data and data partitioning.

ii. Handling Missing Values

There are common gaps in the responses of the survey (NaN = "Not a Number"). The missing values on the numerical values are dealt with with the mean / median method, and the categorical values are dealt with with the mode most frequent value or a default value 0 which is called as neutral value [10].

$$x'_i = \begin{cases} x_i & \text{if } x_i \neq \text{NaN} \\ \text{Median}(X) & \text{if } x_i = \text{NaN} \end{cases} \quad (1)$$

iii. Encoding Categorical Values

Label Encoding and One-Hot Encoding were the two methods to transform qualitative responses into numerical values.

- Label Encoding is employed in order to Turn categories into integers. The Yes=1, No=0 and Maybe=0.5 are used

instead of the entries in dataset.

$$f(x) = \begin{cases} 1 & \text{if } x = \text{yes} \\ 0 & \text{if } x = \text{No} \\ 0.5 & \text{if } x = \text{Maybe} \end{cases}$$

(2)

Representation of the categorical variables is the use of One-Hot Encoding to encode them as binary vectors. Assuming that there are three categories under the Institutional Support: Low, Medium and High, the low would be represented as: (1, 0, 0), Medium would be represented as (0,1, 0) and high would be represented as (0,0,1).

OHE(x) = e_i where i is the category index, and $e_i \in \{0,1\}^k$

iv. Normalization

The models of machine learning (in particular K-NN and Logistic Regression) work best when the numerical variables are scaled identically. Z-score standardization (also known as standard scaling or normalization to standard deviation units) is a method that is applied to re-scale numerical features to have an average of 0 and a standard deviation of 1. It also makes sure that all the variables are scaled, which is particularly significant in such algorithms as the K-NN, the Logistic Regression, and the SVM when the distance and weights magnitude is a key factor.

Repeat the length of feature X on data point x_i :

$$z_i = \frac{x_i - \mu}{\sigma} \quad (3)$$

Where,

x_i original data point, z_i standardized value and μ mean of the feature and σ of the feature.

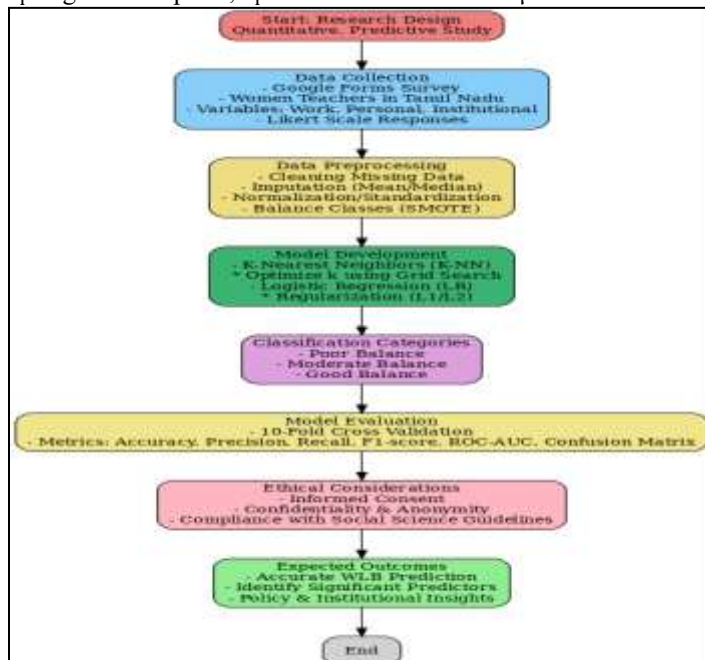


Figure 1: Flow Chart of Proposed work

3.3 Classification

Classification is a form of supervised machine learning that aims at identifying input data (feature vectors) that belong to one of a set of predefined categories (classes).

For an input training dataset

$$D = \{(x_1, y_1), (x_2, y_2), \dots, (x_n, y_n)\}$$

where x_i is the feature vector (e.g., workload, teaching hours, job satisfaction) and $y_i \in \{C_1, C_2, \dots, C_k\}$ (e.g., poor, moderate, good WLB) is the class label.

$$f: X \rightarrow C \quad (4)$$

that assigns new, unobserved instances x test to either of the class labels.

3.3.1 Classification using K-NN

k-Nearest Neighbors (K-NN) is a non-parametric, supervised and instance-based machine learning algorithm, which employs k-nearest neighbors as an algorithm to perform classification and regression.

K-NN in classification, assigns a new point in the feature space to a class which is most prevalent within its k nearest neighbors. Regression, it is used to predict the value of a new point by the average (or weighted average) of the k nearest neighbors of that point.

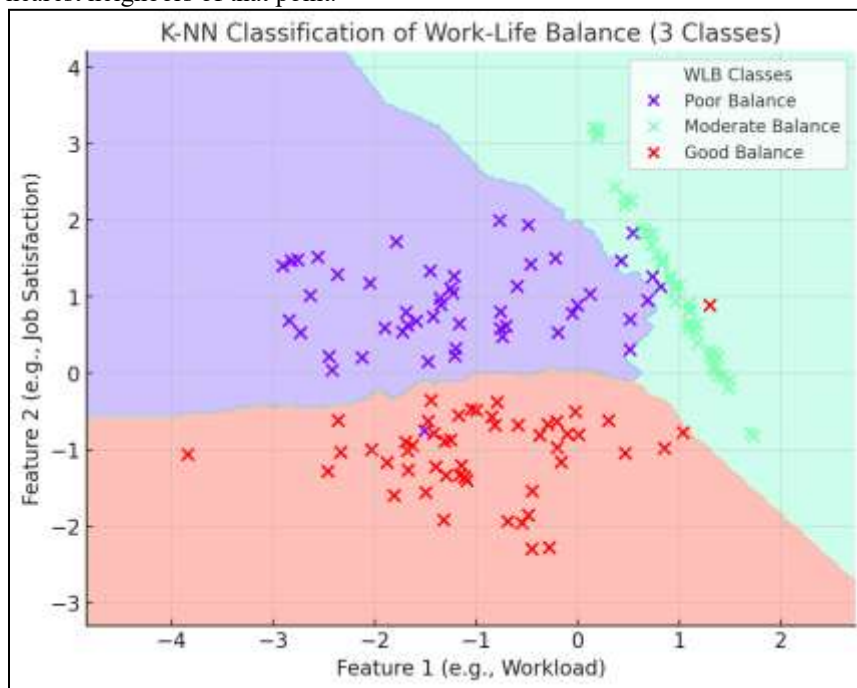


Figure 2: K-NN classification

Figure 2 shows a sample visualization of K-NN classification for three classes:

The K-NN model decision boundary is depicted by each color region. The dots are the points that represent the samples (teachers in your case). New points within an area will be considered poor, good, or moderate WLB based on the majority of the classes of that neighborhood.

Algorithm 1: Classification using K-NN

Step 1: Select the input data

$D = \{(x_i, y_i)\}_{i=1}^N$ and $X_i \in R^m$, $Y_i \in C$ where X_i is the teaching feature (or workload, flexibility, family, institution support, job satisfaction, ...).

Step 2: Choose the number of neighbors k.

Select k stratified K-fold cross-validation Maximize macro-F1 or accuracy (e.g. K=10). Typical search: $k \in \{3, 5, 7, 9, 11, 15\}$. Prefer odd k to reduce ties.

Calculate the Euclidean distance between new point of data and every training point.

$$d(x, y) = \sqrt{\sum_{i=1}^n (x_i - y_i)^2} \quad (5)$$

Step 3: Calculate the Neighbor Weighting

In weighted K-NN, the vote of all neighbors is weighted by some weight parameter which is typically determined by the distance between the test point and the neighbor. The prediction is affected more by the neighboring ones as compared to the distant ones.

Assume X_* is the new node and $N_k(X_*)$ are the k nearest neighbors of X_* .

Let the score of each class c be computed.

$$S_c = \sum_{i \in N_k(x_*)} w_i \cdot 1(y_i = c) \quad (6)$$

The neighboring weights are calculated with the help of the uniform weights method. In equalized weights all neighbors are equal contributors.

$$w_i = 1$$

(7)

Step 4: Select the k closest training points

When the data point of a new teacher x_* (with such features as workload, family responsibilities, job satisfaction, etc.) is supplied, the algorithm should determine which of the previous respondents in the training set are closest.

This is done by the following steps:

- i. Determining the distance between x_* and all training points x_i .

$$d_i = \{(x_1, y_1), (x_2, y_2), \dots, (x_N, y_N)\}$$

For each training point x_i

$$d_i = d(x_*, x_i) \quad (8)$$

- ii. Collect a majority vote among these neighbors based on the Euclidean distance measure in equation (5).

- iii. Sort all the distances to give each of the classes with the most number of votes the new point.

$$d_1 \leq d_2 \leq \dots \leq d_N \quad (9)$$

The training points that are related to then are the k nearest neighbors.

3.3.2 Classification by using Logistic Regression

Logistic regression is a parametric supervised classifier that computes the likelihood of binary outcome using a logistic (sigmoid) link on a linear sum of inputs.

For features $X \in R^m$

$$P(y = 1|x) = \sigma(\beta_0 + \beta_x) = \frac{1}{1 + \exp -(\beta_0 + \beta_x)} \quad (10)$$

The log-odds are linear in X :

$$\log \frac{P(y = 1|x)}{1 - P(y = 1|x)} = \beta_0 + \beta_x \quad (11)$$

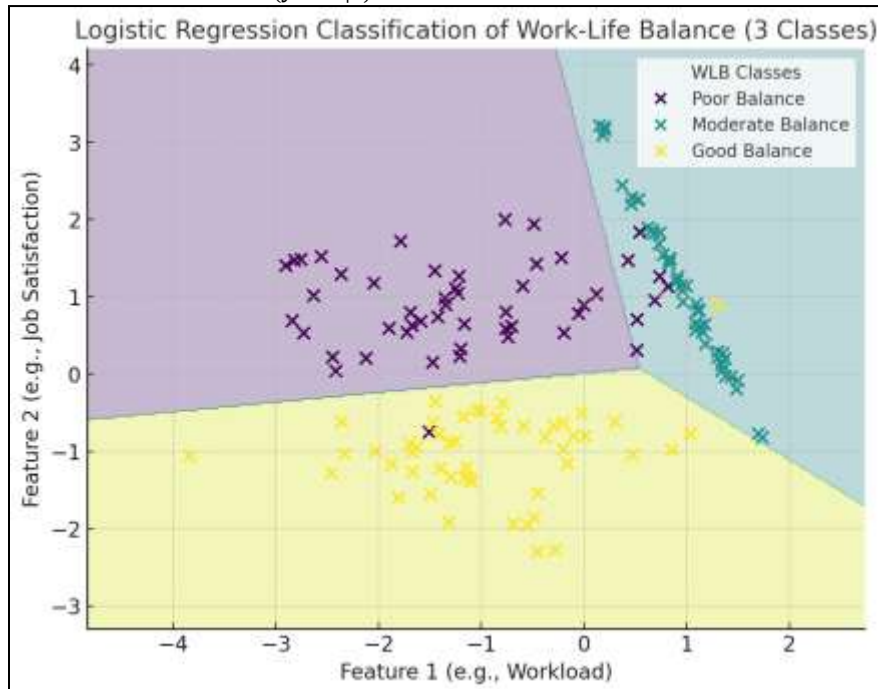


Figure 3: Classification using LR

Algorithm 2: Classification of WLB using LR

Step 1: impute the input data and feature vector $X_i = \{x_{i1}, x_{i2}, \dots, x_{im}\}$ to three classes classification of $C = \{C_1 = \text{Poor}, C_2 = \text{Moderate and } C_3 = \text{Good}\}$.

Step 2: One hot encoding in equation (2) define to encode Categorical features.

Step 3: Z-score normalize features with equation (3)

Step 4: Train Logistic Regression model by defining the following

- Set weights β_c and β_{c0} equals zero.
- Calculate softmax probabilities of every sample using equation (10).

- Calculations of Cross-entropy loss using the following equation (12)

Cross-entropy loss of a single sample (X_i, Y_i) is:

$$l(X_i, Y_i) = - \sum_{c=1}^K Y_{i,c} \log P(Y = c|X_i) \quad (13)$$

$Y_{i,c} = 1$ when the sample i is part of class c or 0 (one-hot encoding).

$P(Y = c|X_i)$ Predicted probability by softmax function.

On the entire data of N samples

$$L = \frac{1}{N} \sum_{i=1}^N l(X_i, Y_i) = - \frac{1}{N} \sum_{i=1}^N \sum_{c=1}^K Y_{i,c} \log P(Y = c|X_i) \quad (14)$$

- Update parameters using gradient descent.

In case of class $c \in \{1,2,3\}$ (Poor, Moderate, Good) and input $X \in R^m$.

$$p_c(X) = \frac{\exp(\beta_{c0} + \beta_c^T X)}{\sum_{k=1}^c \exp(\beta_{k0} + \beta_k^T X)} \quad (15)$$

- Repeat step 4 until convergence loss < threshold or maximum iterations attained).

Step 5: Prediction

Prior to prediction Standardize features by taking the following steps:

- Calculate the training data only through computing μ_j and σ_j .
- Apply transformation $x'_{ij} = \frac{x_{ij} - \mu_j}{\sigma_j}$ to both training and testing data
- Train Logistic Regression using x'_{ij} .

Step 6: Compute probabilities $P(\text{Poor})$, $P(\text{Moderate})$, $P(\text{Good})$

- For teacher i and class c :

$$P_{i,c} = \frac{\exp(z_{i,c})}{\sum_{k=1}^3 \exp(z_{i,k})} \quad (16)$$

$$P = \text{softmax}(Z) \quad (17)$$

- Assign class with highest probability.

Step 7: Result Evaluation

4. RESULT ANALYSIS

4.1 Quality Parameters

The quality of the model is determined by the classification performance measures in the case of WLB classification (Poor, Moderate, Good). These parameters determine the degree of prediction of each class by the model. The following are the conventional quality parameters to assess the proposed work on WLB prediction [13].

i. Accuracy

The percentage of correct classification of samples among the samples.

$$\text{Accuracy} = \frac{\text{No.ofCorrect predictions}}{\text{Total number of Predictions}} \quad (18)$$

ii. Precision

The percentage of teachers who are predicted to be in the class c (e.g., in a classroom, where teachers are poor) which are actually in the classroom.

$$\text{Precision} = \frac{TP_c}{TP_c + FP_c} \quad (19)$$

TP_c = true positives for class c

FP_c = false positives for class c

iii. Recall

The percentage of teachers that really fall in class c which the model catches.

$$\text{Recall}_c = \frac{TP_c}{TP_c + FN_c} \quad (19)$$

iv. F1-Score

The average of precision and recall of class c in harmonic average.

$$F1_c = 2 \frac{\text{Precision}_c \cdot \text{Recall}_c}{\text{Precision}_c + \text{Recall}_c} \quad (20)$$

4.2 Performance Comparison

i. Accuracy analysis

In Figure 4, the accuracy of training and validation of K-NN in WLB classification is shown in 20 epochs. The accuracy of the training begins at approximately 0.87 and increases steadily to approximately 0.91 indicating that the model continually learns the trends in the training data. The validation accuracy starts at a lower level approximately 0.80 and becomes more and more accurate with a steady level ranging at 0.86-0.87 towards the later epochs. The distance between the training and validation curves is moderate, and it means that the model has a good learning and generalization balance without radical overfitting [14]. The trend indicates that the K-NN classifier can be trusted to

predict the category of WLB and perform well on both training and unseen validation data indicating the appropriate use of the K-NN classifier in dealing with responses in real-life situations in the dataset.

Table 1: Accuracy analysis of K-NN vs LR

S.No	Epoch	K-NN		LR	
		Train Accuracy	Validation Accuracy	Train Accuracy	Validation Accuracy
1	1	0.8698	0.7995	0.7359	0.6975
2	2	0.8749	0.8058	0.7711	0.7417
3	3	0.8851	0.8252	0.8044	0.7679
4	4	0.8911	0.8239	0.8219	0.7841
5	5	0.89	0.8293	0.8427	0.7976
6	6	0.8953	0.8339	0.856	0.8055
7	7	0.9016	0.8438	0.8632	0.8254
8	8	0.9036	0.8467	0.8778	0.8337
9	9	0.9066	0.8504	0.8824	0.8377
10	10	0.907	0.8534	0.8889	0.8423

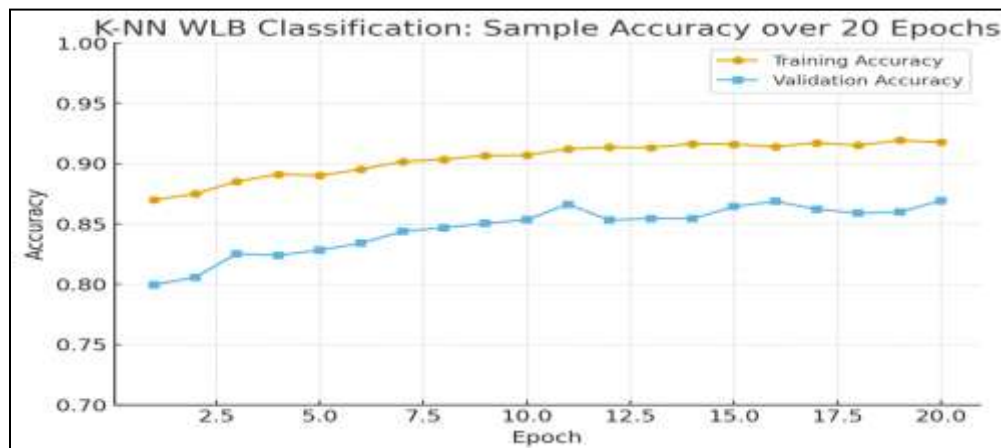


Figure 4: Accuracy analysis using K-NN:

The Figure 5 shows the training and validation accuracy of Logistic Regression in Work-life Balance classification with 20 epochs. The accuracy of training begins with an accuracy of approximately 0.73 and continues to rise steadily up to about 0.92 at the final epoch, which is an indication of high learning ability [15]. The validation accuracy also increases steadily starting at approximately 0.70, and reaching about 0.87. Convergence pattern shows that Logistic Regression is efficient to reflect the underlying patterns of the data and the gap between training and validation curves is not large, which indicates that the model may not be overfitted. Altogether, the model shows strong performance, and the accuracy is good during training and validation, which makes it a good one to use in the classification of the WLB categories of teachers.

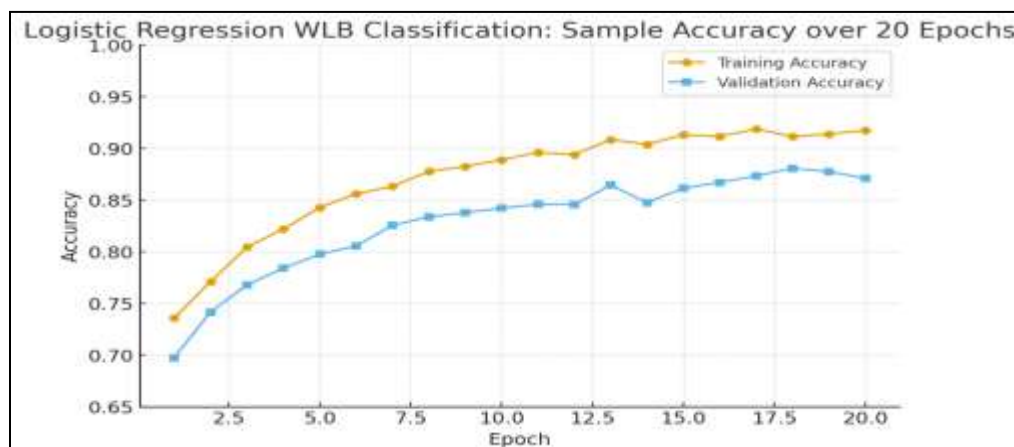


Figure 5: Accuracy analysis using LR

ii. Loss

The Figure 6 indicates training and validation loss of K-NN in WLB classification with 20 epochs (with a proxy loss of 1 - accuracy). The loss in the training begins at an approximate of 0.13 and continues to reduce to approximately 0.08, which is a sign of enhanced consistency in learning. The loss at validation starts at a higher point, near 0.21, but decreases gradually and levels off at 0.14 and there are small fluctuations in the loss during different epochs. The difference between training and validation losses is moderate and constant, which indicates that the model generalizes quite well without such serious overfitting. Generally, the loss curves underscore the fact that K-NN balances both the ability to fit the training data and also do well on unknown validation data, thus it is an appropriate baseline to perform WLB classification.

Table 2: Loss analysis of K-NN vs LR

S.No	Epoch	K-NN		LR	
		Train Loss	Validation Loss	Train Loss	Validation Loss
1	1	0.1308	0.2104	0.5963	0.6862
2	2	0.1224	0.1915	0.5056	0.612
3	3	0.1169	0.1897	0.4473	0.5374
4	4	0.1126	0.1735	0.3874	0.4753
5	5	0.107	0.1682	0.3309	0.4128
6	6	0.1041	0.1651	0.291	0.3806
7	7	0.0986	0.1754	0.261	0.3409
8	8	0.0931	0.1587	0.2266	0.3126
9	9	0.0944	0.1524	0.1986	0.2839
10	10	0.0926	0.1485	0.1762	0.2526

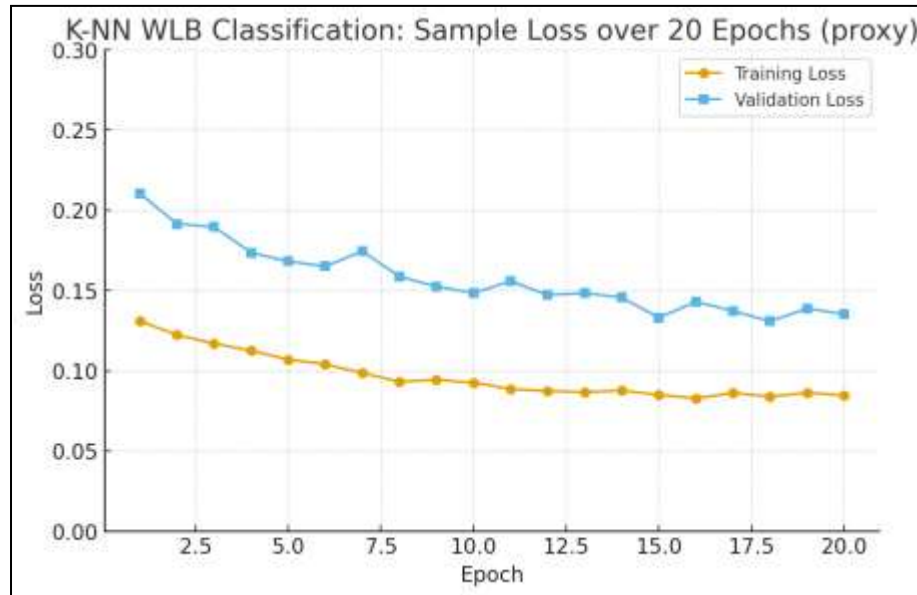


Figure 6: Loss analysis on WLB classification using K-NN

The Figure 7 shows the training and validation loss of Logistic Regression to classify WLB in 20 epochs. Loss training begins on the higher side, approximately 0.60, and it reduces gradually to approximately 0.05, which is considered as successful learning of patterns within the dataset. Validation loss starts at about 0.68 and proceeds steadily downward and becomes nearly 0.12 towards the last epochs. The gradual anytime flat decrease of training and validation loss indicates that the model converging model is achieved without much overfitting because the difference between the two is minuscule and constant. This shows that not only does Logistic Regression fit the training data well, but also the unseen validation data, thus it is a good candidate that can be trusted with reliable WLB classification.

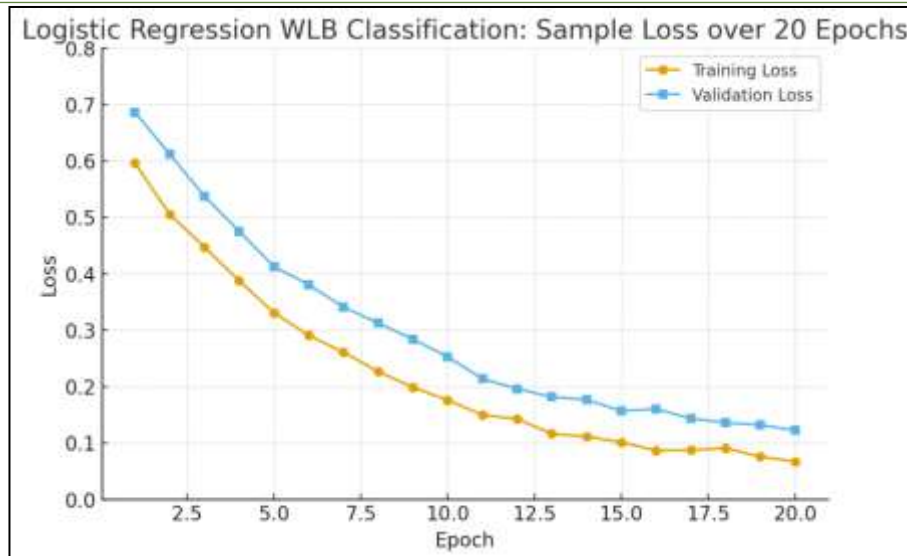


Figure 7: Loss analysis on WLB classification using LR

iii. Confusion Matrix

The confusion matrix of K-NN classification of WLB in Figure 8 depicts that the model is accurate in most of the cases in all three classes with 19 Poor, 28 Moderate, and 26 Good responses being correctly predicted. Nevertheless, wrong classifications also stand out: some of the cases that are actually Poor are classified as Moderate and some of the Good cases are classified as Poor. This means that K-NN is useful in local neighborhoods similarities, but fails on borderline cases where similarities between the features of teachers in various WLB categories are similar.

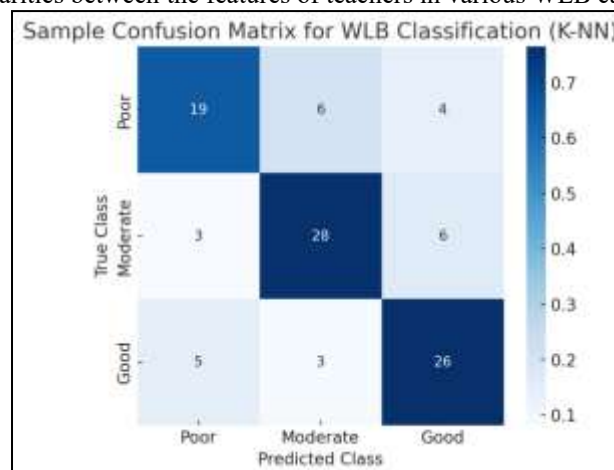


Figure 8: Confusion Matrix of K-NN

Figure 9 shows the confusion matrix of LR that exhibits better predictive capability with a stronger central tendency a concentration of two, three, and three responses, respectively, Poorest, moderate, and good responses correctly classified. There are also fewer misclassifications than K-NN, most mistakes are between the Moderate and the Good category, which is often hard to differentiate, as they can overlap in some ways because of such similar features as moderate workload or partial institutional support. The lower off-diagonal values emphasize the fact that the use of softmax probabilities to learn the global decision boundaries of the data makes the Logistic Regression more appropriate in the situation of capturing the underlying structure of the data. This leads to a more stable and consistent classification performance of all the WLB categories.

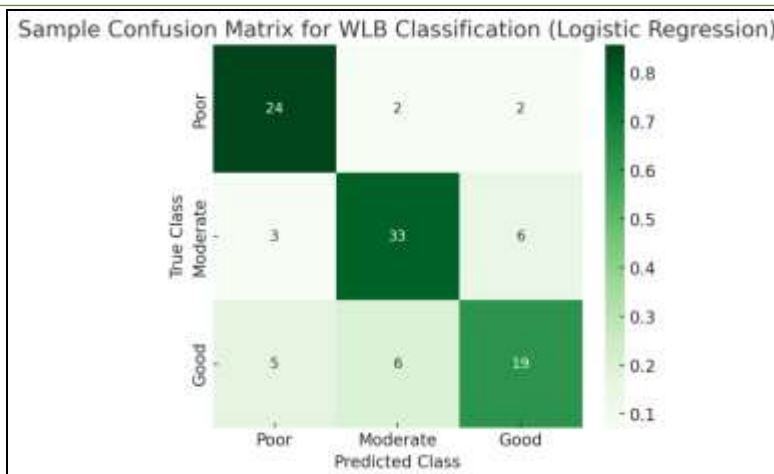


Figure 9: Confusion Matrix of LR

iv. Performance Comparison between K-NN Vs LR

The results comparative table 3 indicates clearly that LR performs better than K-NN in categorizing WLB among women teachers.

- **Accuracy:** LR has 0.72, whereas K-NN has 0.67, meaning that a higher percentage of teachers are accurately identified in the three categories of WLB (Poor, Moderate, Good).
- **Precision:** LR has a record of 0.72, which is marginally better than that of K-NN, which is 0.68, which implies that LR has fewer false-positive than K-NN and is more accurate when predicting a class.
- **Recall:** LR has a score of 0.71 as compared to 0.67 with K-NN, which indicates that LR is more effective in picking real cases in each category of WLB particularly in cases of poor WLB without dropping them.
- **F1-Score:** This is a combination of precision and recall, again LR performs better than K-NN (0.71 vs 0.67), it has a better balance between accuracy in classifying features and error reduction.

Table 3: Result Comparison

Model	Accuracy	Precision	Recall	F1-Score
K-NN	0.67	0.68	0.67	0.67
LR	0.72	0.72	0.71	0.71

Observations and Discussions

On the whole, in this work the Logistic Regression is superior to K-NN. K-NN is quite sensitive to the selection of neighbors and distance rates leading to misclassifications in the borders zones between the moderate and good WLB. The more stable and generalizable classification offered by Logistic Regression, which learns global decision boundaries with softmax probabilities, allows it to offer greater stability and accuracy when compared to other classifications. The continuous growth in accuracy and precision, recall, and F1-score indicates that LR is the more efficient model to predict WLB in this data.

5. CONCLUSION

In the study, the K-Nearest Neighbors (K-NN) and the Logistic Regression (LR) were used to analyze the Work-Life Balance (WLB) classification of women teachers in the state of Tamil Nadu. Following data collection, data cleaning and balancing, the two models were trained and tested on 750 responses categorical into three classes namely, Poor, Moderate and Good WLB. Findings indicated that the Logistic Regression had always performed positively over the K-NN in all the performance measures. LR had better accuracy (0.72) and higher precision, recall and F1-score (0.71 0.72 range) whereas K-NN had 0.67 in most metrics. Such results indicate that LR is strong at capturing decision boundaries across the globe and modeling minor differences in classes.

Meanwhile, K-NN demonstrated medium accuracy and was not active with overlapping cases, especially between the Moderate and Good categories since it used distance as a measure of similarity. Logistic Regression was found to be stronger and more generalizable, minimizing misclassifications and giving constant probabilities that could be used to predict. Accordingly, it can be concluded that LR is the more credible model of WLB classification, which provides the practical and efficient framework with the help of which the future research and policy analysis in the field of teacher well-being can be developed.

6. REFERENCE

1. Algemeen Dagblad. (2023). Grote ontslaggolf waait over uit amerika, 1 op de 5 werknemers wisselt van baan. Retrieved 23 January 2023 from <https://www.ad.nl/economie/grote-ontslaggolf-waait-over-uit-amerika-1-op-de-5-werknemers-wisselt-van-baan> ad658da1.
2. Blau, F. D., & Kahn, L. M. (2017). The gender wage gap: Extent, trends, and explanations. *Journal of Economic Literature*, 55(3), 789–865. <https://doi.org/10.1257/jel.20160995>.
3. Celbiş, M. G. (2022). Unemployment in rural Europe: A machine learning perspective. *Applied Spatial Analysis and Policy*. <https://doi.org/10.1007/s12061-022-09464-0>.
4. Holgado-Apaza, L.A., Ulloa-Gallardo, N.J., Aragón-Navarrete, R., Riva-Ruiz, R., Odagawa-Aragon, N.K., Castellon-Apaza, D.D., Carpio-Vargas, E.E., Villasante-Saravia, F.H., Alvarez-Rozas, T.P., & Quispe-Layme, M. (2024). The Exploration of Predictors for Peruvian Teachers' Life Satisfaction through an Ensemble of Feature Selection Methods and Machine Learning. *Sustainability*.
5. Holgado-Apaza, L.A., Isuiza-Perez, D.D., Ulloa-Gallardo, N.J., Vilchez-Navarro, Y., Aragón-Navarrete, R., Quispe Layme, W., Quispe-Layme, M., Castellon-Apaza, D.D., Choquejahu-Acero, R., & Prieto-Luna, J.C. (2025). A machine learning approach to identifying key predictors of Peruvian school principals' job satisfaction. *Frontiers in Education*.
6. Jung, Y.M., & Jo, H. (2025). Understanding Continuance Intention of Generative AI in Education: An ECM-Based Study for Sustainable Learning Engagement. *Sustainability*.
7. Diana, B., & Kvr, R. (2025). Work – Life Balance among Female Lecturers in Tamil Nadu, India: A Synthetic Data Approach. *Journal of Information Systems Engineering and Management*.
8. Yoo, J.E., & Rho, M. (2020). Exploration of Predictors for Korean Teacher Job Satisfaction via a Machine Learning Technique, Group Mnet. *Frontiers in Psychology*, 11.
9. Celbiş, M. G., Wong, P., Kourtit, K., & Nijkamp, P. (2023). Impacts of the COVID-19 outbreak on older-age cohorts in European labor markets: A machine learning exploration of vulnerable groups. *Regional Science Policy & Practice*, 15(3), 559–584. <https://doi.org/10.1111/rsp3.12520>.
10. Chen, T., & Guestrin, C. (2016). xgboost: A scalable tree boosting system. In *Proceedings of the 22nd ACM SIGKDD international conference on knowledge discovery and data mining* (pp. 785–794). ACM.
11. D'Ambrosio, C., Clark, A. E., & Barazzetta, M. (2018). Unfairness at work: Well-being and quits. *Labour Economics*, 51, 307–316. <https://doi.org/10.1016/j.labeco.2018.02.007>.
12. Demirkaya, H., Aslan, M., Güngör, H., Durmaz, V., & Şahin, D. R. (2022). Covid-19 and quitting jobs. *Frontiers in Psychology*, 13, 916222. <https://doi.org/10.3389/fpsyg.2022.916222>
13. Duval, R., Oikonomou, M., & Tavares, M. M. (2022). Tight jobs market is a boon for workers but could add to inflation risks. Retrieved 23 January 2023 from <https://www.imf.org/en/Blogs/Articles/2022/03/31/tight-jobs-market-is-a-boon-for-workers-but-could-add-to-inflation-risks>.
14. García-Mainar, I., & Montuenga-Gómez, V. M. (2020). Over-qualification and the dimensions of job satisfaction. *Social Indicators Research*, 147, 591–620. <https://doi.org/10.1007/s11205-019-02167-z>.
15. Gerich, J., & Weber, C. (2020). The ambivalent appraisal of job demands and the moderating role of job control and social support for burnout and job satisfaction. *Social Indicators Research*, 148, 251–280. <https://doi.org/10.1007/s11205-019-02195-9>.