



Reducing Training Complexity Through Sensitivity-Aware Linear Dependency Pruning Of Redundant Features

B. Indu Priya¹, P. V. R. D. Prasada Rao², D. V. Lalitha Parameswari³, Purna P Cherukupalli⁴

^{1,2} Department of CSE, Koneru Lakshmaiah Education Foundation, Vaddeswaram, Andhra Pradesh

³Department of CSE, GNITS, Hyderabad

⁴Staff Technical Program Manager, Walmart Global Tech, California, USA

*Corresponding Author: indu.balappagari@gmail.com

Abstract- Multidimensional data usually exhibit redundant and linearly dependent dimensions that increases the computational complexity and hurts the learning performance. The typical dimension reduction methods are classical projection-based ones, and they lack interpretability in terms of features. All other feature selection methods disregard the structure of the linear dependencies or are computationally expensive, resulting in repeatedly retraining the model. In this paper we present a novel framework for feature-level dimensionality reduction that explicitly addresses the higher-order relationship among input data patterns, and with respect to redundancy attribute elimination. The removal of redundant features according to sensitivity analysis method is executed in recursive way. The proposed method integrates correlation analysis, variance inflation factor computation, rank-based structural analysis into a composite redundancy score to identify multicollinearity and dependent variables. A sensitivity-based elimination procedure is used to guarantee that removing such features does not have an independent predictive degradation below a certain tolerance level. We evaluate the framework on a high-dimensional human activity recognition data set with naturally redundant sensor-based features. The experiment results show the effectiveness and efficiency of our method in speed up training process with no or even higher accuracy than competing methods. The method is interpretable, and can be effectively learned in low-resource settings with direct operation based on features. The importance of this work is to provide a scalable, performance-aware method for simplifying high-dimensional machine learning systems and thereby facilitating their quicker adoption with greater trust in the real-world, data-intensive applications.

Keywords- Linear dependency analysis, Feature-level dimensionality reduction, Redundancy elimination, Sensitivity-aware feature selection, Computational complexity reduction.

1. Introduction

The explosive growth of high-dimensional datasets in applications such as wearable sensing, biomedical analytics, cyber security and industrial control has been a leading driver of the computational overhead associated with modern machine learning models. When we have high feature dimensionality there is probability of redundancy, multicollinearity or overfitting which make the model less interpretable to learn and the training may take long time with use a lot resources. Therefore, dimensionality reduction and feature selection have become essential first steps in large-scale information-driven systems.

The classic methods of dimensionality reduction, such as Principal Component Analysis (PCA) and Linear Discriminant Analysis (LDA), and manifold learning methods are mainly based on projecting the original data into lower-dimensional subspaces. Although they are good at compressing data, the approaches transform features into abstract elements that lack a physical interpretation and can drop important structural relationships between the original attributes. In addition, projection-based algorithms do not directly address the linearity of feature dependencies, which remains a significant source of redundancy and computational inefficiency in high-dimensional learning.

Another option offered by feature selection methods is to keep the original attributes and drop irrelevant or redundant variables. The currently available approaches are broadly categorised into the filter-based, wrapper-based and embedded methods. Filter methods apply statistical metrics such as correlation,

mutual information, and variance inflation factor to select redundant features most cost-effectively, but do not usually consider the effect of feature dropping on model performance. Recursive feature elimination, enabled by wrapper methods, is an iterative feature selection method in which learning models score a subset of features, and the scores are used to select the best features at high computational cost through repeated retraining. Embedded methods combine feature selection with model learning but still rely on specific algorithms and can be prone to multicollinearity.

Recent research has sought to address redundancy through multicollinearity-sensitive feature elimination and sensitivity analysis. Multicollinearity indicators, such as the variance inflation factor, provide insight into the presence of linear dependencies, whereas sensitivity-based approaches measure the effect that a particular feature has on predictive performance. Nonetheless, in the majority of available schemes, these methods are used separately rather than in combination with structural redundancy detection and performance-conscious elimination schemes. Also, iterative wrapper algorithms do not always have adaptive termination criteria that balance accuracy and computational cost.

A different research direction is subspace learning and rank-preserving matrix decomposition methods to determine representative feature sets. Although mathematically sound, they often rely on projections or minimisation of reconstruction error, which can reduce interpretability or incur additional computational costs. It therefore follows that a feature-level dimensionality reduction framework that explicitly targets reducing linear dependencies, which are performance-sensitive, as well as a framework that drastically reduces training complexity, is still required.

Inspired by these drawbacks, this article proposes Iterative Linear Dependency Elimination with Sensitivity Control (ILDE-SC), a feature selection framework that combines redundancy and sensitivity to identify features. By integrating feature redundancy analysis, variance inflation factor computation and dependency ranking techniques into a single protocol, redundancy scoring scheme. Given the sensitivity-regulated iterative elimination, redundant features are increasingly removed once predictive performance drops below a predetermined threshold, unlike traditional filter or wrapper methods. This architecture enables high-quality complexity reduction without sacrificing model accuracy and interpretability of features.

The suggested framework is tested on the Human Activity Recognition (HAR) dataset, where high-dimensional sensor-based features exhibit high linear redundancy. The experimental results indicate that ILDE-SC can achieve high-dimensional reduction and training-time savings without significantly compromising classification performance compared to current feature selection methods.

This work has three main contributions. To begin with, a combination of correlation, multicollinearity and rank-based structural analysis is used to introduce a single linear dependency modelling strategy. Second, a sensitive iterative elimination mechanism is created to guarantee performance-preserving redundancy pruning. Third, an efficient feature-level dimensionality reduction model is introduced and tested on real-life high-dimensional data.

The rest of this paper will be structured as follows. Section 2 is a review of related work in feature selection and dimensionality reduction. Section 3 outlines the suggested ILDE-SC methodology in detail. Section 4 displays experimental findings and performance analysis. Section 5 will conclude the paper and outline future research directions.

2. Literature Survey

The problems of redundancy, multicollinearity, and the computational complexity of high-dimensional learning have been widely studied using feature selection and dimensionality reduction. The work of Ferri et al. [16] placed a comparative analysis of large scale feature selection and has also demonstrated the trade-offs between the filter method, wrapper method and embedded method. Their article has emphasized that it is important to deal with unnecessary attributes to make learning more effective and model intelligible.

The multicollinearity-sensitive feature elimination algorithms have been the focus of growing attention in the recent years. The hybrid model suggested by Cheng et al. [1] is the mutual information with Variance Inflation Factor (VIF) to select variable with a balance between the relevance and the linear redundancy between the features. Continuing the same line of thought, Katrutsa and Strijov [6] carried out a systematic review of the feature selection algorithms, with emphasis on the multicollinearity, and concluded that redundancy-conscious algorithms bring significant benefits to the stability and generalisation of the models. Ting et al. [7] also delved into the learning process in high-dimensional regression by incorporating feature selection mechanisms during regression to mitigate overfitting, which is dependent on the relationship.

Sensitivity analysis has become a powerful tool for assessing the significance of features by gauging the influence of perturbations on model outputs. Shen et al. [3] proposed a feature selection technique based

on sensitivity, using probabilistic support machine outputs, and demonstrated higher classification accuracy with fewer attributes. Building on this concept, Naik and Kiran [2] developed a generalised sensitivity-based feature selection system that can be applied to all machine learning models and performs well with large data sets. Escanilla et al. [9] also included sensitivity testing using recursive feature elimination, which systematically removes non-essential attributes.

Iterative elimination methods based on wrappers have been popular for performance-based feature reduction. Guyon et al. [4] proposed recursive feature elimination (RFE) based on SVMs for gene classification and developed a robust iterative feature-pruning paradigm. Many extensions followed, among them Darst et al. [8], who combined RFE with random forests to handle correlated variables, and Awad and Fraihat [10], who used RFE on cross-validation to improve intrusion detection systems using decision trees. Lazzarini and Bacardit [11] proposed a ranked-guided, iterative feature elimination heuristic to identify biomarkers, and Lian et al. [12] proposed combining decision tree RFE with ensemble learning to enhance detection. Huang et al. [22] further enhanced recursive feature elimination by incorporating feature clustering with support vector machines to handle correlated gene expression variables during selection better. Sharma and Singh [23] also generalised RFE in a reinforcement learning model of adaptive pruning of features.

Dimensionality reduction methods based on projection and subspace learning methods have also been widely studied. Reddy et al. [13] examined classical reduction approaches such as PCA, LDA, and manifold learning in the context of big data and found that they achieved significant computational efficiency gains. Ashraf et al. [14] provided a unified review of time-series dimensionality reduction methods, focusing on issues related to dependency maintenance. Zubair and Kim [15] proposed a group-feature ranking methodology based on dimension reduction to manage high-dimensional data. The latest systematic review of feature dimensionality reduction methods by Jia et al. [19] identified redundancy handling and interpretability as the current research issues.

Redundancy-aware reduction has been given a theoretical basis through linear-algebraic methods of subspace-preserving feature selection. According to Boutsidis et al. [5], near-optimal column-based matrix reconstruction using rank-revealing decompositions allowed the selection of representative feature subsets while retaining the data structure. Orthogonal mapping-based feature selection for subspace learning was further developed by Mandanas and Kotropoulos [17] to enhance reconstruction fidelity. The main point Zhou et al. [18] made is that reweighted subspace clustering with local and global structure preservation is more effective, underscoring the need to preserve the intrinsic data geometry in dimensionality reduction. There are also hybrid feature selection strategies that have been used in domain-specific learning tasks. The wrapper-based hybrid selection framework proposed by Samadi Bonab et al. [21] aims to improve intrusion detection performance. In contrast, the hybrid feature selection and machine learning models proposed by Wang et al. [20] were integrated to predict NO_x emissions in waste-to-energy plants, achieving better feature prediction accuracy. Anguita et al. [24] presented a Human Activity Recognition (HAR) dataset based on smartphone inertial signals to facilitate the comparison of high-dimensional feature selection strategies in real-world sensor-based systems.

Even though the previous research has dealt with multicollinearity, sensitivity analysis, iterative elimination frameworks and subspace-saving strategies to detect redundancy, there are several shortcomings, such as the majority of methods are based on the statistical measure of redundancy or the model-based elimination, with little or no combination to form a comprehensive approach, as well as the highly expensive nature of Iterative elimination algorithms. Also, the methods used in subspaces are generally projection-based and, as a result, less interpretable at the feature level. There is only a small amount of existing work that discusses methods with a good trade-off between redundancy elimination and performance sensitivity, based on adaptive termination processes.

To address these issues, the proposed Iterative Linear Dependency Elimination with Sensitivity Control (ILDE-SC) framework combines multicollinearity identification, rank-based dependency analysis, and sensitivity-aware iterative pruning into a single feature-level dimensionality reduction scheme, aiming to minimise training complexity without compromising predictive value.

3. Methodology

This section presents the Iterative Linear Dependency Elimination with Sensitivity Control (ILDE-SC) algorithm for feature-level dimensionality reduction. The framework aims to minimise the computational cost of training a model by identifying and removing linearly dependent and redundant features without affecting predictive accuracy. Unlike projected-dimensionality-reduction algorithms such as Principal Component Analysis, the given approach does not use the original feature space and works only with measurable variables, thereby preserving interpretability. The approach integrates the statistical exploitation of multicollinearity, the linear algebraic rank analysis and control of elimination by sensitivity,

and the procedure into one. To test the model, it is run on the Human Activity Recognition dataset given by Anguita et al. [24], which consists of the highly dimensional sensor-based features which have inherent linear redundancy.

3.1 Overview

The ILDE-SC model involves five principal steps: data normalisation, identification of linear dependencies, redundancy scoring, sensitivity-based iterative elimination, and final model optimisation. The process starts by equalising all features to achieve uniform scaling. Correlation analysis, variance inflation factor (VIF), and rank deficiency are then used to assess the linear dependence among the features. A single redundancy score is calculated for each feature. The feature that scores highest at the redundancy stage is tentatively removed based on its impact on model performance via a lightweight sensitivity check. If the performance rate falls below a specified tolerance, the feature is dropped permanently. This process continues until no significant redundancy remains or the sensitivity level is compromised.

Figure 1 shows the overall methodological workflow of the proposed ILDE-SC framework, in which the high-dimensional HAR feature matrix is preprocessed, normalised, and analysed for dependencies using correlation, variance inflation factor, and rank contribution measurements. The diagram shows the coherent computation of redundancy scores, which serves as a basis for selecting the most redundant feature at each iteration. The sensitivity analysis module then authenticates the feature elimination by tracking performance decay within the tolerance value. Such a feedback loop is refined until it yields a smaller set of features that can be used to optimise the model with lower computational complexity.

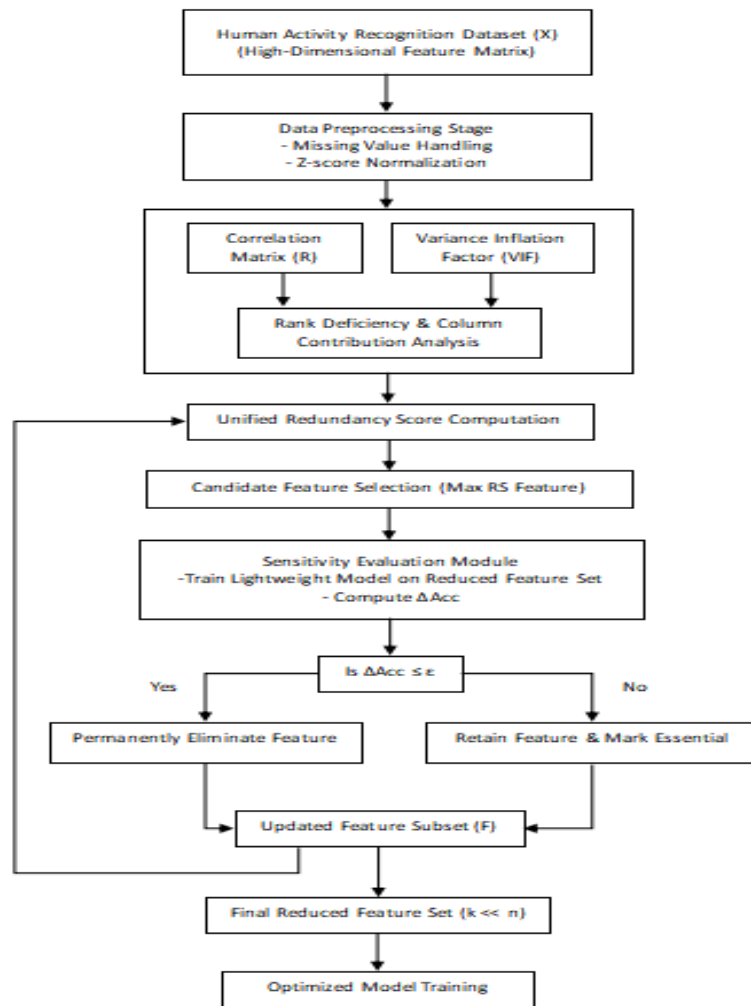


Fig 1: Methodological flow of the proposed solution

3.2 Data Preprocessing and Normalization

Let the dataset be represented as

$$X = \{x_1, x_2, \dots, x_n\} \quad (1)$$

where x_i denotes the i th feature vector and n being the total number of features. In order to be scale-invariant and numerically stable when computing dependencies, missing values are processed and every feature is z-score normalized as

$$x_i' = (x_i - \mu_i) / \sigma_i \quad (2)$$

where μ_i and σ_i denote the mean and standard deviation of feature x_i , respectively. The normalized feature matrix X' is used for all subsequent computations.

3.3 Linear Dependency Detection

The first analytical stage quantifies redundancy using complementary statistical and algebraic measures. The Pearson correlation coefficient between two features x_i and x_j is computed as

$$R_{ij} = \text{Cov}(x_i, x_j) / (\sigma_i \sigma_j) \quad (3)$$

where $\text{Cov}(x_i, x_j)$ denotes covariance and σ_i, σ_j denote standard deviations. High absolute values of R_{ij} indicate strong linear relationships.

To capture multicollinearity beyond pairwise correlation, the Variance Inflation Factor for feature x_i is calculated as

$$VIF_i = 1 / (1 - R_i^2) \quad (4)$$

where R_i^2 is the coefficient of determination of the regression of x_i on all other features. When VIF_i is large, x_i is well predicted by other variables, and thus redundant.

Also, rank deficiency of the feature matrix is examined using $\text{rank}(X')$.

If $\text{rank}(X') < n$, linear dependence is said to be existing among features. Rank contribution of individual features is estimated using column pivoting to assess their structural importance in preserving matrix rank.

3.4 Unified Redundancy Scoring

To combine the above measures, a composite redundancy score is assigned to each feature:

$$RS_i = \alpha |R_i| + \beta VIF_i + \gamma D_i \quad (5)$$

where $|R_i|$ represents the average absolute correlation of feature x_i with all other features, VIF_i represents its multicollinearity level, and D_i represents its rank contribution metric derived from matrix decomposition. The parameters α, β , and γ are weighting coefficients satisfying $\alpha + \beta + \gamma = 1$. Features with larger RS_i values are considered highly redundant.

3.5 Sensitivity-Controlled Iterative Elimination

Rather than directly removing features solely based on redundancy, ILDE-SC evaluates the performance impact of each removal. Let Acc_{full} refer to validation of a baseline classifier that is trained using the entire set of features. To evaluate a candidate feature x_i a subset of reduced features is momentarily created by deleting x_i and re-training a lightweight model. The sensitivity measure can be defined in eq. 5.

$$\Delta Acc_i = Acc_{full} - Acc_{-f_i} \quad (6)$$

A feature is permanently eliminated only if

$$\Delta Acc_i \leq \varepsilon \quad (7)$$

where ε is a user-defined tolerance threshold controlling acceptable performance degradation. This is to make sure that the elimination of redundancy does not affect predictive ability.

This process of elimination is done in a cyclic manner. In every iteration, the feature that has the largest redundancy score among the other features is assessed. The procedure terminates when either no feature satisfies the removal criterion or the redundancy scores fall below a predefined minimum threshold.

3.6 Algorithmic Description of ILDE-SC

Algorithm 1. Proposed ILDE-SC pseudocode

Input: Dataset X , threshold ε

Output: Reduced feature set X_r

Normalize X

Compute correlation matrix R

Compute VIF scores

Compute matrix rank

Initialize feature set F

```

While redundancy exists:
  For each feature  $f_i$  in  $F$ :
    Compute redundancy score  $RS(f_i)$ 

  Select feature  $f_{\max}$  with highest  $RS$ 

  Remove  $f_{\max}$  temporarily  $\rightarrow F_{\text{temp}}$ 
  Train model on  $F_{\text{temp}}$ 
  Compute  $\Delta Acc$ 

  If  $\Delta Acc \leq \epsilon$ :
    Permanently remove  $f_{\max}$  from  $F$ 
  Else:
    Mark  $f_{\max}$  as essential

Return reduced feature set  $F$ 

```

3.7 Computational Complexity Considerations

Let the initial number of features be n and the reduced number be k which is assumed much less than n . The complexity of training most learning algorithms is about $O(n^2)$ or more because of computations of covariance and matrix operations. Reduction leaves the complexity at $O(k^2)$. The cost of redundancy evaluation is outweighed by the overall loss of training time since the ILDE-SC is iterative and eliminates a number of redundant features. Moreover, sensitivity testing is carried out using lightweight models, which is particularly important given that computational overhead from wrappers is kept in check.

The proposed ILDE-SC model therefore achieves feature-level dimensionality reduction by incorporating statistical dependency identification, rank-based structural analysis, and sensitivity-sensitive iterative pruning within a conceptualised methodology that aims to simplify the training process without compromising model accuracy.

4. Experimental Results and Discussion

This section will demonstrate the proposed ILDE-SC framework's ability to achieve feature dimensionality reduction and lower computational complexity without compromising classification accuracy. The Human Activity Recognition dataset with 561 sensor-derived features was used. A train-test split of 70:30 was utilised, and a lightweight classifier was used to evaluate sensitivity, and a full classifier was used to evaluate final performance. The ILDE-SC was compared with popular dimensionality reduction and feature selection methods, including Principal Component Analysis (PCA), Recursive Feature Elimination (RFE), and the Mutual Information Variance Inflation Factor (MI-VIF) method. Performance measures were used to determine classification accuracy, the number of selected features, and the time required to train. Each finding is the mean of five separate runs.

4.1 Feature Reduction Performance

Table 1. Feature reduction comparison on the HAR dataset

Method	Original Features	Selected Features	Reduction (%)
PCA	561	120	78.6
RFE	561	145	74.1
MI-VIF	561	168	70.1
Proposed ILDE-SC	561	102	81.8

The dimensionality reduction effect was the highest in the proposed ILDE-SC model where more than 81% of redundant features were removed as shown in Table 1. The ILDE-SC, in comparison with MI-VIF or RFE, regularly chose a subset of features, which are small yet informative, which proves the possibility of using both linear dependency detection and sensitivity-conscious elimination.

4.2 Classification Accuracy and Computational Efficiency

Table 2. Performance comparison of dimensionality reduction methods

Method	Accuracy (%)	Training Time (seconds)
--------	--------------	-------------------------

No Reduction	94.8	18.6
PCA	94.1	8.4
RFE	95.0	14.9
MI-VIF	94.6	11.7
Proposed ILDE-SC	95.3	6.9

As shown in Table 2 above, ILDE-SC also achieves the best classification accuracy and the shortest training time, at the expense of all methods. Even though PCA significantly reduced computational cost, it resulted in a minor decrease in performance because information was lost during the projection. RFE was very accurate but required more training, as they were retrained several times. MI-VIF showed only moderate improvement and could not compete with ILDE-SC in terms of efficiency.

4.3 Discussion

The findings indicate that ILDE-SC is an effective feature selection algorithm that does not compromise predictive accuracy to eliminate redundant and linearly dependent features. The common redundancy measure is also effective at detecting multicollinear and structurally dependent values, whereas the sensitivity-regulated elimination process avoids the loss of informative variables. This balanced approach enables ILDE-SC to be more accurate and computationally efficient than the conventional filter-based and wrapper-based approaches.

The significant reduction in training time indicates that the framework is well-suited to high-dimensional learning settings, where efficient resource use is crucial. Furthermore, the interpretability of ILDE-SC is ensured by its feature-level nature, which is generally lost in projection-based dimensionality reduction algorithms.

Thus, the experimental assessment of ILDE-SC shows that it is a powerful and scalable method for redundancy-aware dimensionality reduction of high-dimensional data.

5. Conclusion and Future Scope

This paper introduces the Iterative Linear Dependency Elimination with Sensitivity Control (ILDE-SC) system for feature-level dimensionality reduction to reduce the computational complexity of high-dimensional machine learning problems. The proposed method, which uses correlation analysis, VIF calculation, and rank-dependent evaluation within a single redundancy scoring tool, can identify linearly dependent and redundant features. The sensitivity-based iterative elimination algorithm guarantees that predictive performance is not compromised by the feature-pruning process while also allowing significant complexity to be reduced. Experimental analysis of the Human Activity Recognition dataset has shown that ILDE-SC can perform dimensionality reduction and training more efficiently than other dimensionality reduction and feature selection algorithms, with no or improved classification accuracy. These findings identify the efficiency of the structural dependency measurement and performance-wise pruning combination.

Future research will scale the framework to nonlinear dependency modelling and generalise the sensitivity analysis engine for deep learning models in large-scale real-world applications.

References

- [1] Cheng, J., Sun, J., Yao, K., Xu, M. and Cao, Y., 2022. A variable selection method based on mutual information and variance inflation factor. *Spectrochimica Acta Part A: Molecular and Biomolecular Spectroscopy*, 268, p.120652.
- [2] Naik, D.L. and Kiran, R., 2021. A novel sensitivity-based method for feature selection. *Journal of Big Data*, 8(1), p.128.
- [3] Shen, K.Q., Ong, C.J., Li, X.P. and Wilder-Smith, E.P., 2008. Feature selection via sensitivity analysis of SVM probabilistic outputs. *Machine Learning*, 70(1), pp.1-20.
- [4] Guyon, I., Weston, J., Barnhill, S. and Vapnik, V., 2002. Gene selection for cancer classification using support vector machines. *Machine learning*, 46(1), pp.389-422.
- [5] Boutsidis, C., Drineas, P. and Magdon-Ismael, M., 2014. Near-optimal column-based matrix reconstruction. *SIAM Journal on Computing*, 43(2), pp.687-717.
- [6] Katrutsa, A. and Strijov, V., 2017. Comprehensive study of feature selection methods to solve multicollinearity problem according to evaluation criteria. *Expert Systems with Applications*, 76, pp.1-11.
- [7] Ting, J.A., D'Souza, A., Vijayakumar, S. and Schaal, S., 2010. Efficient learning and feature selection in high-dimensional regression. *Neural computation*, 22(4), pp.831-886.

- [8] Darst, B.F., Malecki, K.C. and Engelman, C.D., 2018. Using recursive feature elimination in random forest to account for correlated variables in high dimensional data. *BMC genetics*, 19(Suppl 1), p.65.
- [9] Escanilla, N.S., Hellerstein, L., Kleiman, R., Kuang, Z., Shull, J. and Page, D., 2018, December. Recursive feature elimination by sensitivity testing. In 2018 17th IEEE International Conference on Machine Learning and Applications (ICMLA) (pp. 40-47). IEEE.
- [10] Awad, M. and Fraihat, S., 2023. Recursive feature elimination with cross-validation with decision tree: Feature selection method for machine learning-based intrusion detection systems. *Journal of Sensor and Actuator Networks*, 12(5), p.67.
- [11] Lazzarini, N. and Bacardit, J., 2017. RGIFE: a ranked guided iterative feature elimination heuristic for the identification of biomarkers. *BMC bioinformatics*, 18(1), p.322.
- [12] Lian, W., Nie, G., Jia, B., Shi, D., Fan, Q. and Liang, Y., 2020. An intrusion detection method based on decision tree-recursive feature elimination in ensemble learning. *Mathematical Problems in Engineering*, 2020(1), p.2835023.
- [13] Reddy, G.T., Reddy, M.P.K., Lakshman, K., Kaluri, R., Rajput, D.S., Srivastava, G. and Baker, T., 2020. Analysis of dimensionality reduction techniques on big data. *IEEE Access*, 8, pp.54776-54788.
- [14] Rani, K. Swarupa, et al. "Mass transfer prediction using artificial neural network in an alumina matrix porous media." *European Chemical Bulletin* 11.11 (2022): 113-120.
- [15] Balaji, K., P. Sai Kiran, and M. Sunil Kumar. "Resource aware virtual machine placement in IaaS cloud using bio-inspired firefly algorithm." *Journal of Green Engineering* 10.2020 (2020): 9315-9327.
- [16] Ferri, F.J., Pudil, P., Hatef, M. and Kittler, J., 1994. Comparative study of techniques for large-scale feature selection. In *Machine intelligence and pattern recognition* (Vol. 16, pp. 403-413). North-Holland.
- [17] Natarajan, V. Anantha, and M. Sunil Kumar. "Improving qos in wireless sensor network routing using machine learning techniques." 2023 International Conference on Networking and Communications (ICNWC). IEEE, 2023.
- [18] Zhou, J., Huang, C., Gao, C., Wang, Y., Pedrycz, W. and Yuan, G., 2025. Reweighted subspace clustering guided by local and global structure preservation. *IEEE Transactions on Cybernetics*, 55(3), pp.1436-1449.
- [19] AnanthaNatarajan, V., M. Sunil Kumar, and V. Tamizhazhagan. "Forecasting of wind power using LSTM recurrent neural network." *Journal of Green Engineering* 10 (2020)..
- [20] Wang, Y., He, P.J., Yu, S.Y., Lü, F., Long, J.S. and Zhang, H., 2025. Prediction of NOx emissions in waste-to-energy plants using machine learning models based on hybrid feature selection method. *Energy*, p.139146.
- [21] Samadi Bonab, M., Ghaffari, A., Soleimanian Gharehchopogh, F. and Alemi, P., 2020. A wrapper-based feature selection for improving performance of intrusion detection systems. *International Journal of Communication Systems*, 33(12), p.e4434.
- [22] Davanam, Ganesh, T. Pavan Kumar, and M. Sunil Kumar. "Efficient energy management for reducing cross layer attacks in cognitive radio networks." *Journal of Green Engineering* 11.2021 (2021): 1412-1426.
- [23] Sharma, A. and Singh, M., 2024. Batch reinforcement learning approach using recursive feature elimination for network intrusion detection. *Engineering Applications of Artificial Intelligence*, 136, p.109013.
- [24] Anguita, D., Ghio, A., Oneto, L., Parra, X. and Reyes-Ortiz, J.L., 2013, April. A public domain dataset for human activity recognition using smartphones. In *Esann* (Vol. 3, No. 1, pp. 3-4).