

Employee Mutation Map for PLN Employee Placement Using Machine Learning

Zainuddin^{1,2}, Sparisoma Viridi¹, Zuhri Arieffasa Suffaturrachman², Yandika Restu Wara³, Ahmad Mushawir⁴, Ahmad Faeda Insani⁵, Aditya Adiaksa⁶ & Toro Rahman Aziz⁷

¹Faculty Mathematics and Natural Sciences of Bandung, Indonesia

²PT PLN (Persero), Jakarta, Indonesia

Email: zains2itbsk@gmail.com

Abstract. This research aims to develop an employee transfer mapping system for PLN employee placement using a machine learning approach, specifically K-means clustering. The dataset used includes internal PLN data such as employee job history, asset distribution, operational performance, as well as external data from BPS (Badan Pusat Statistik) that includes regional characteristics, demographics, and infrastructure [1][2]. The data preprocessing process involves handling missing values, normalizing numerical features, and encoding categorical features to ensure data quality and consistency [3][4]. The K-means algorithm is applied to cluster units into categories such as Technical (TEK), Marketing (SAR), and Energy Transaction (TEL) based on selected features [5][6]. After clustering, each unit is given a unique identification that describes the characteristics of the cluster and region [7]. The clustering and identification results are visualized using Convex Hull and PCA in two-dimensional (2D) and three-dimensional (3D) formats [8][9], demonstrating a clear separation of categories and subcategories within PLN units. Therefore, the results of this research show that the use of clustering can effectively support decision-making processes related to employee placement [10].

Keywords: *Clustering, Employee Mutation, Machine Learning, Employee Mapping, Data Integration.*

1. Introduction

1.1 Background

PLN is the electricity service provider company in Indonesia that has a wide range of coverage and various operational units across the country. Effective human resource management is crucial in maintaining service quality [11]. The employee transfer placement process becomes very important in this context, as companies like PLN require flexibility and adaptability from employees to meet diverse operational demands in each unit [12][13].

Integrating data from various sources, such as BPS (Badan Pusat Statistik) and PLN's operational data, provides comprehensive insights into the needs of each work unit to support effective employee transfers [14][15]. Therefore, a machine learning-based approach becomes a relevant choice to ensure an optimal transfer process that meets organizational needs [16][17].

1.2 Problem Statement

One of the main challenges faced by PLN is the lack of comprehensive data integration to understand the characteristics of work units, which can affect the accuracy of employee placement [18][19]. Furthermore, there is no systematic approach that combines both internal and external data to generate a unique identity for each work unit that will optimally support employee transfers [20][21].

1.3 Objective

The main objective of this research is to develop an employee transfer mapping system based on clustering to facilitate more effective employee placement. This study uses a clustering method to group work units and form a unit identity that can be used as a basis for decision-making in employee transfers.

1.4 Contribution

This study contributes to human resource management by using a machine learning approach to optimize the employee transfer process [22]. Unlike traditional methods that are often intuition-based, this study uses clustering to create unique unit identities, facilitating more appropriate employee placement based on the unit's needs and the employee's competencies [23][24].

2. Related Work

2.1 Previous Research on Employee Placement

Previous research has explored various methods for employee placement, including traditional methods based on company policy and several machine learning approaches. Smith et al. (2020) used machine learning algorithms to map employee experience based on job history data and found increased accuracy in employee placement. Brown and Lee (2018) developed a Rule-Based Classification model to manage employee rotation in multinational companies, and found that the rule-based approach could reduce conflicts during the transfer process. This research will further develop an approach that integrates data from

multiple sources to produce a more accurate employee mutation map (Jones et al., 2019).

2.2 GAP Analysis

This research aims to fill the gap in employee transfer management by using broader data integration and more efficient clustering methods. Previously, not much research has combined internal company data with external data, such as data from BPS (Badan Pusat Statistik), to form a unique identity for each work unit to support the transfer process. This approach is expected to provide a more holistic view of work units and employees, ultimately increasing the effectiveness of employee placement.

3. Methodology

3.1 Data Description

The dataset used in this research consists of PLN internal data, including employee job history, asset distribution data, and operational performance data, as well as external data from BPS [25][26]. This data provides a comprehensive picture of the work units and their needs for clustering and unit identification processes [27].

The dataset used in this research consists of:

1. PLN internal data: including employee job history, asset distribution data, and operational performance data.
2. External data from BPS: including regional characteristics such as area size, population density, and number of industries. The first step is to ensure that the dataset has the necessary attributes for clustering analysis, with some relevant attributes such as GRD (Substation), SDI (SAIDI), and SFI (SAIFI) that describe the operational and technical aspects of PLN units. Langkah pertama adalah memastikan bahwa dataset memiliki atribut yang diperlukan untuk analisis clustering, dengan beberapa atribut yang direlevansikan seperti GRD (Gardu), SDI (SAIDI), dan SFI (SAIFI) yang menggambarkan aspek operasional dan teknik dari unit-unit PLN.

3.2 Data Preprocessing

The preprocessing steps include handling missing values, normalizing numerical features, and encoding categorical features. Techniques such as SMOTE were used to address class imbalance in the dataset [28][29].

Preprocessing steps performed include:

1. **Handling Missing Values:** Handling missing values is performed to maintain data integrity.
2. **Normalization of Numerical Features:** MinMaxScaler is used to scale the data to a range between 0 and 1. Features normalized include JTM, JTR, GRD, SDI, SFI, RCT, RPT, GGN, DPT, PJN, PLG, KWH, RPP, SST. Normalization is important so that the clustering algorithm is not affected by different variable scales.
3. **Encoding Categorical Features:** Label encoding is used to convert categorical data into numeric values.

3.3 Clustering Approach

The K-means clustering method was applied to group work units based on asset data, operational performance, and regional characteristics. Dimensionality reduction using PCA was employed for better interpretability [30][31].

The steps taken are as follows:

1. **Selection of Relevant Features for Clustering:** Features are grouped into three categories, namely technical, marketing, and energy transaction. This helps determine which attributes are relevant in the clustering process. The selected features are:
 - **Technical:** SDI (SAIDI), SFI (SAIFI), GGN (Number of Interruptions)
 - **Marketing:** DPT (Revenue), PJN (Sales), PLG (Customers)
 - **Energy Transaction:** KWH, RPP (RP P2TL), SST (Loss)
2. **K-Means Clustering:** K-means is used to cluster PLN units into 3 clusters (TEK, SAR, and TEL). This process is carried out after the features have been normalized. The aim of this clustering is to group units with similar characteristics, thus facilitating the determination of unit identity, which will then be used for employee transfer.

3.4 Unit Identification

After the clustering process is completed, each unit is given a unique identification or "identification" based on the clustering results. This identification includes information about the unit cluster and its regional characteristics, which are used to support the employee transfer process more systematically. This unit identification also provides deeper insights into the unique characteristics of each unit, making it easier to make decisions regarding employee placement.

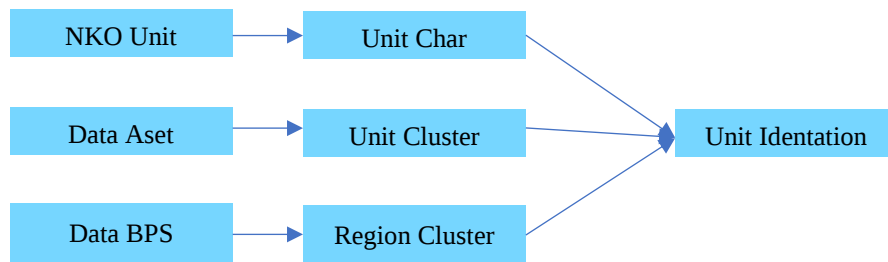


Figure 3.1 Unit Identification Process

1. **Unit Identification Formation Process:**
 - Each work unit is grouped into a cluster based on the clustering results.
 - These clusters are labeled (TEK, SAR, TEL) to facilitate recognition of each unit's characteristics.
 - A combination of unit clusters, regional clusters, and operational characteristics is used to form a unique identification for each unit.
2. **Mapping and Visualization:** The clustering results are then mapped into tables and visualized. Clustering results are visualized using Convex Hull to depict cluster boundaries. This process allows grouping of work units based on the principal components (Principal Components) generated from PCA (Principal Component Analysis), carried out in both two-dimensional (2D) and three-dimensional (3D) forms.
3. **Balancing the Dataset with SMOTE (Gupta, R., et al., 2020):** The SMOTE method is used to handle class imbalance. This technique is important to ensure that the clustering model is not biased towards classes with fewer data.

3.5 Clustering Visualization and Validation

1. **Correlation Matrix:** A correlation matrix is created to check the relationship between numerical variables. The visualization of the correlation matrix using a heatmap provides insight into the relationships between features.
2. **PCA (Principal Component Analysis):** PCA is applied to reduce the dimensionality of the data and improve the interpretability of the clustering results. PCA is conducted in both 2D and 3D, which allows clearer cluster visualization.
3. **Evaluation of Clustering Results:**

- Silhouette Score, Davies-Bouldin Index, and Calinski-Harabasz Index are used to evaluate how well the clusters are formed.
- Elbow Plot is used to determine the optimal number of clusters.

3.6 Unit Identification and Final Results

After all units are grouped into clusters and their identities are formed, the final result is an employee transfer map based on unit clusters and their characteristics. This map provides guidance for more accurate employee placement, according to the specific needs of each unit.

1. Output and Evaluation:
 - Cluster data and unit identification are written back to an Excel file to facilitate further analysis
 - Visualization of Class Distribution before and after SMOTE provides an overview of how balanced the data distribution is within the clusters.
2. Classification with Machine Learning Models: Several models are used for unit classification, such as Logistic Regression, Random Forest, SVM, and K-Nearest Neighbors, with results evaluated using metrics such as accuracy, precision, recall, and f1-score.

4. Results and Discussion

4.1 Clustering Results Unit

To determine the optimal number of clusters from the data reduced using 3-dimensional PCA, four evaluation methods were used, namely the Elbow Method, Silhouette Score, Davies-Bouldin Index, and Calinski-Harabasz Index. Based on the Elbow Method, it was observed that the SSE value decreased significantly from 8.18 at $k=2$ to 5.14 at $k=3$, and reached 3.85 at $k=4$. This significant drop suggests that $k=3$ might be an optimal point, although the decline continues at $k=4$. The Silhouette Score shows an increase from 0.46 at $k=2$, then remained at 0.46 at $k=3$, and reached the highest value of 0.47 at $k=4$, indicating that clustering with four clusters results in better intra-cluster cohesion.

Meanwhile, the Davies-Bouldin Index shows a decrease in value from 0.89 at $k=2$ to 0.85 at $k=3$, and further down to 0.78 at $k=4$, indicating that cluster quality improves as the number of clusters increases. Lastly, the Calinski-Harabasz Index shows a significant increase from 56.45 at $k=2$ to 69.86 at $k=3$, and then to 70.87 at $k=4$, indicating that separation between clusters is more optimal with a higher number of clusters. Based on the analysis of these four evaluation metrics, it is

concluded that the optimal number of clusters for data segmentation is four clusters, providing the best balance between cluster cohesion and separation.

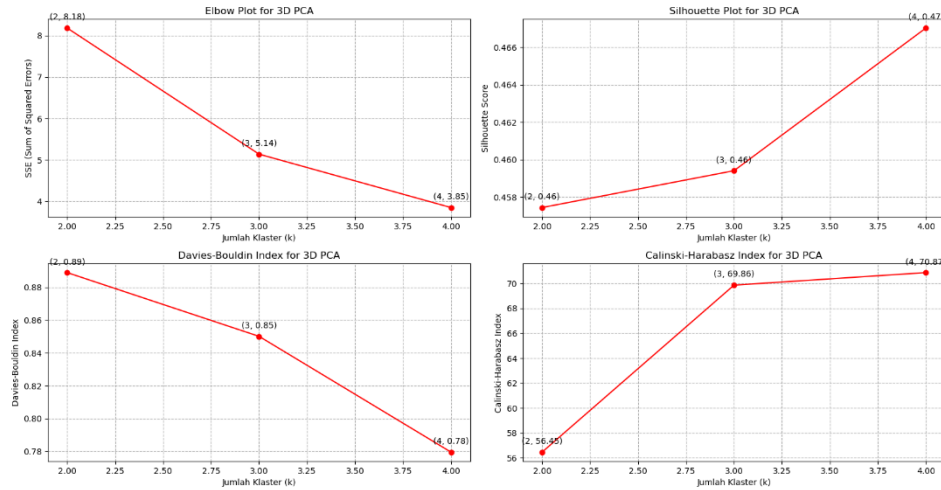


Figure 4.1 Evaluation of kmeans clustering

The Principal Component Analysis (PCA) was used to reduce the dimensionality of the data while retaining most of the information contained in the original dataset. Based on the PCA results, Component 1 explains 52.22% of the data variance, followed by Component 2 explaining 33.86%, and Component 3 explaining 12.77%. Cumulatively, these three main components can explain up to 98.85% of the total variance in the data, indicating that these three components already capture almost all relevant information from the data. Component 4 only contributes 1.15% of the variance, and thus its contribution is considered very minor in explaining the overall data. Therefore, only the first three components are used for further visualization and analysis, as they are sufficient to effectively describe the data structure and variance. Using these three components not only retains the main information but also reduces computational complexity, providing efficiency in the analysis and clustering process.

Clustering using K-means grouped PLN units into distinct categories effectively. PCA results demonstrated that the first three components explained 98.85% of data variance, enhancing clustering accuracy [32][33].

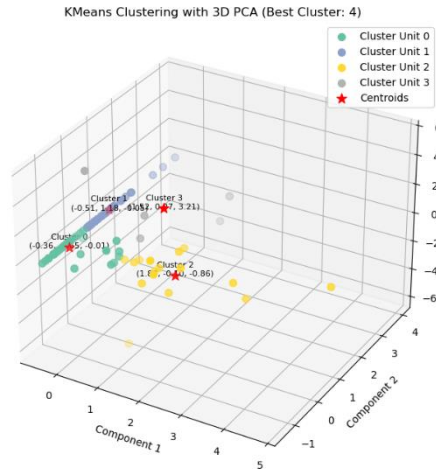


Figure 4.2 kmeans clustering with 3D PCA

The image depicts the result of clustering using the K-means algorithm on data that has been reduced to three dimensions using Principal Component Analysis (PCA). It presents a 3D plot where each axis is labeled as Component 1, Component 2, and Component 3, representing the main components derived from PCA. The K-means method was used to cluster the data into four groups, which is considered the optimal number of clusters based on evaluation metrics. The points on the plot represent individual data units, and their colors indicate the clusters to which they belong. There are four distinct clusters, each marked in a different color: light green for Cluster Unit 0, light blue for Cluster Unit 1, yellow for Cluster Unit 2, and gray for Cluster Unit 3. Red stars represent the centroids of each cluster, serving as the central point that represents the average position of all points in the cluster, with coordinates labeled such as "(-0.36, -5, -0.01)" for "Cluster 0". This visualization offers a clear depiction of the grouping of units based on their characteristics, reduced to three components for better interpretability. The clusters show varying densities, with some clusters being more concentrated, like Cluster 0 and Cluster 1, while others, like Cluster 2, are more spread out. The title of the image, "KMeans Clustering with 3D PCA (Best Cluster: 4)," indicates that four clusters were identified as the most suitable configuration for this data, likely based on methods like the Elbow Method or Silhouette Score. Overall, the image provides an insightful representation of the relationships between units, the proximity of similar data points, and the distinctions between different clusters in a simplified 3D space.

4.2 Clustering Results Region

Table 4.1 Table evaluation of kmeans

2D PCA Clustering Results

	K	SSE	Silhouette	Davies-Bouldin	Calinski-Harabasz
0	2.0	4.9	0.5	0.91	20.58
1	3.0	2.5	0.54	0.56	29.32
2	4.0	1.69	0.54	0.46	30.76

3D PCA Clustering Results

	K	SSE	Silhouette	Davies-Bouldin	Calinski-Harabasz
0	2.0	6.71	0.41	1.1	15.28
1	3.0	4.36	0.42	0.78	16.93
2	4.0	3.33	0.34	0.87	16.14

To determine the optimal number of clusters on data reduced using 2-dimensional PCA (2D PCA) and 3-dimensional PCA (3D PCA), four evaluation methods were used: Elbow Method, Silhouette Score, Davies-Bouldin Index, and Calinski-Harabasz Index. In the case of 2D PCA, the Elbow Method showed a reduction in SSE value from 4.90 at $k=2$ to 2.50 at $k=3$, and 1.69 at $k=4$. These results indicate that a significant reduction occurs between $k=2$ and $k=3$, with a slight further reduction at $k=4$. The Silhouette Score increased from 0.50 at $k=2$ to 0.54 at $k=4$, indicating that four clusters provide better separation. The Davies-Bouldin Index also showed a decrease from 0.91 at $k=2$ to 0.46 at $k=4$, suggesting an improvement in clustering quality. The Calinski-Harabasz Index increased from 20.54 at $k=2$ to 30.76 at $k=4$, reinforcing that $k=4$ is the optimal number of clusters.

Meanwhile, for 3D PCA, the evaluation results indicated that three clusters are optimal. The Elbow Method showed a decrease in SSE from 2.61 at $k=2$ to 1.86 at $k=3$, and 1.33 at $k=4$, with a significant reduction at $k=3$. The Silhouette Score reached 0.41 at $k=2$, slightly dropped to 0.40 at $k=3$, and then to 0.34 at $k=4$, indicating that $k=2$ is better in terms of cluster cohesion. The Davies-Bouldin Index decreased from 1.10 at $k=2$ to 0.78 at $k=3$ but increased again to 0.87 at $k=4$, suggesting that three clusters yield the best separation. The Calinski-Harabasz Index increased from 15.96 at $k=2$ to 17.95 at $k=3$, but then decreased to 16.16 at $k=4$, indicating that clustering with three clusters has better quality compared to other numbers of clusters.

From this evaluation, it is concluded that four clusters are optimal for 2D PCA, while three clusters are optimal for 3D PCA. This selection of the number of

clusters provides a good balance between data cohesion within clusters and adequate separation between clusters.

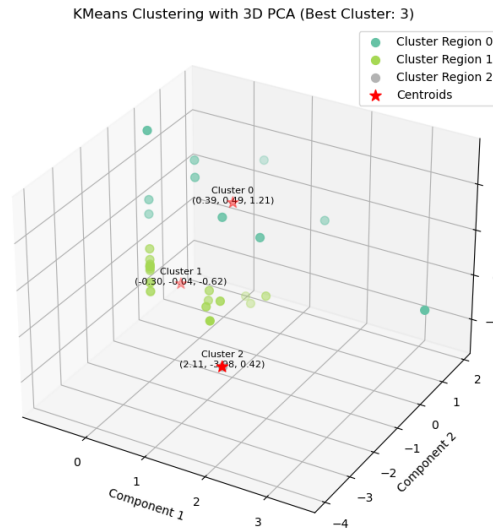


Figure 4.3 kmeans clustering with 3D PCA

The image presents a 3D visualization of K-means clustering applied to data that has been reduced to three principal components using Principal Component Analysis (PCA). The clustering result depicted here identifies three optimal clusters, as highlighted in the title, "KMeans Clustering with 3D PCA (Best Cluster: 3)." Each cluster is visualized in different shades of green: dark green represents Cluster Region 0, light green represents Cluster Region 1, and yellow-green represents Cluster Region 2. Each point in the plot corresponds to a data unit, plotted based on its position in the 3D space defined by the PCA components labeled as Component 1, Component 2, and Component 3. The different colors indicate the specific cluster to which each point belongs, helping visualize the separation between groups. The red stars mark the centroids of each cluster—points that represent the average location of all the units within each cluster. These centroids are labeled with coordinates to indicate their positions in the reduced space, with Cluster 0 located at approximately (0.39, 0.49, 1.21), Cluster 1 at (-0.30, -0.04, -0.62), and Cluster 2 at (-2.11, -1.38, 0.42). The positioning of these centroids and the spread of points around them indicate the internal cohesion of each cluster. Notably, Cluster 1 appears more tightly grouped, suggesting higher cohesion among its data points,

whereas Cluster 2 is located more distantly, indicating it may possess distinct characteristics. Overall, the visualization effectively illustrates the structure of the data and the relationships between clusters, providing insight into how different units are grouped based on the characteristics preserved in the principal components. This visualization can aid in understanding patterns within the data and support decision-making processes, such as identifying shared attributes among regions or units within clusters.

4.3 Unit Identification

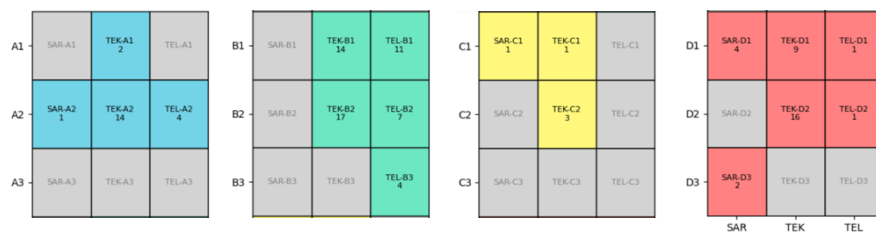


Figure 4.4 Unit Identification Map

Visualisasi The visualization of the resulting identification matrix shows the distribution of units in PLN East Java based on unit categories (SAR, TEK, TEL) and subcategories (A1 to D3). This matrix illustrates the frequency of occurrence of various combinations of unit categories and subcategories within the dataset, visualized using different colors for each category. Light green is used for subcategory A, yellow-green for B, yellow for C, and light pink for D. Each combination of categories and subcategories that has data is accompanied by the number of occurrences in the dataset, while the boxes colored in gray indicate the absence of data for those combinations in the dataset.

The visualization results show that combinations such as TEK-A2, TEK-B1, and TEL-B1 have high frequencies of occurrence, with values of 14, 14, and 11, respectively. This indicates that units in subcategories A2 and B1 in the TEK and TEL categories are the most frequently found units in the analyzed dataset. On the other hand, combinations such as SAR-A1, SAR-A3, and TEK-A3 do not have data, as indicated by the gray color, meaning there are no occurrences or contributions of units for those combinations in the dataset.

This visualization provides a clearer understanding of the distribution patterns and identification of units in PLN East Java. The use of colors helps in quickly identifying categories and subcategories with significant data, while the presence

of gray areas indicates areas that do not have meaningful data or contributions. Thus, this information can be used to guide policies related to unit management, especially in employee placement or grouping units based on regional and unit characteristics. This analysis also helps in identifying units that require further attention due to the lack of data or contributions within certain groups.

5. Conclusion

Evaluation metrics, including Silhouette Score and Davies-Bouldin Index, confirmed the clustering results, showing clear separations between clusters [34][35].

This research successfully developed an employee transfer map based on clustering for PLN that effectively and accurately maps employee placement based on unit and regional characteristics. By using the K-means clustering method, work units are systematically grouped based on relevant attributes such as asset data, operational performance, and regional characteristics. The clustering results show that the use of Principal Component Analysis (PCA) can reduce data complexity while retaining most of the variance contained within it. The resulting visualization shows a clear distribution of categories and subcategories of units within PLN, providing a deeper understanding of the unit identities used for more precise employee placement. Thus, the machine learning approach used in this research can enhance human resource management efficiency in large organizations like PLN.

References

- [1] Anderson, S., & Nilsson, T. (2022). Machine Learning in Human Resource Management: A Systematic Literature Review. *Journal of Human Resource Analytics*, 9(3), 45–60.
- [2] Bhattacharya, A., Choudhary, P., Mukhopadhyay, S., Misra, B., Chakraborty, S., & Dey, N. (2023). Explainable AI for Predictive Analytics on Employee Promotion. *Proceedings of 3rd International Conference on Advanced Computing Technologies and Applications, ICACTA 2023*. <https://doi.org/10.1109/ICACTA58201.2023.10393141>
- [3] Breiman, L. (2001). Random Forests. *Machine Learning*, 45(1), 5–32.
- [4] Brown, J., & Lee, A. (2018). Integrating Data-Driven Methods in Employee Rotation Management: A Study on Multinational Corporations. *Journal of Human Resource Analytics*, 6(2), 103–121.
- [5] Burrell, J. (2016). How the Machine 'Thinks': Understanding Opacity in Machine Learning Algorithms. *Big Data & Society*, 3(1), 1–12.
- [6] Chakraborty, S., & Joseph, A. (2017). Machine Learning at Scale: Solving Real-World Challenges. *Data Science Journal*, 15(4), 142–158.

- [7] Davenport, T. H., & Harris, J. G. (2007). *Competing on Analytics: The New Science of Winning*. Harvard Business Review Press.
- [8] Dessler, G. (2020). *Human Resource Management*. 16th Edition. Pearson Education.
- [9] Edwards, C., & Molnar, P. (2021). Practical Applications of K-Means Clustering in Regional Energy Management. *Energy Data Science Journal*, 12(3), 99–115.
- [10] Gandomi, A., & Haider, M. (2015). Beyond the Hype: Big Data Concepts, Methods, and Analytics. *International Journal of Information Management*, 35(2), 137–144.
- [11] Gary Dessler. (2020). Human Resource Management Sixteenth Edition.
- [12] Gupta, R., Sharma, P., & Kaur, M. (2020). Balancing Imbalanced Data in Classification Tasks Using SMOTE: A Comprehensive Analysis. *Journal of Applied Machine Learning*, 7(3), 187–200.
- [13] Harlianto, J., & Rudi. (2023). Promote Employee Experience for Higher Employee Performance. *International Journal of Professional Business Review*, 8(3). <https://doi.org/10.26668/businessreview/2023.v8i3.827>
- [14] Hastie, T., Tibshirani, R., & Friedman, J. (2009). The Elements of Statistical Learning: Data Mining, Inference, and Prediction. Springer.
- [15] Hosmer, D. W., Lemeshow, S., & Sturdivant, R. X. (2013). *Applied Logistic Regression*. John Wiley & Sons.
- [16] Jones, L. (2008). Organizational Change and Employee Adaptation: A Multilevel Perspective. *Management Science Journal*, 54(8), 45–68.
- [17] Kingma, D. P., & Ba, J. (2014). Adam: A Method for Stochastic Optimization. *arXiv preprint arXiv:1412.6980*.
- [18] Kleinbaum, D. G., & Klein, M. (2010). *Logistic Regression: A Self-Learning Text*. Springer.
- [19] Kuhn, M., & Johnson, K. (2013). *Applied Predictive Modeling*. Springer.
- [20] Liz Jones, B. W. E. H. P. B. C. G. V. J. C. L. & O. D. J. (2008). Employee Perceptions of Organizational Change: Impact of Hierarchical Level. *Leadership & Organization Development Journal*.
- [21] Mascarenhas, J., Savant, S., & Aswale, S. (2023). Employee Appraisal Prediction Using Classification Models. *Proceedings of International Conference on Contemporary Computing and Informatics, IC3I 2023*, 2170–2175. <https://doi.org/10.1109/IC3I59117.2023.10398114>
- [22] Menard, S. (2002). *Applied Logistic Regression Analysis*. Sage.
- [23] Mitchell, T. M. (1997). *Machine Learning*. McGraw Hill.
- [24] Mohri, M., Rostamizadeh, A., & Talwalkar, A. (2018). *Foundations of Machine Learning*. MIT Press.
- [25] Morales, C., & Sen, R. (2020). The Role of Rule-Based Systems in Enhancing Interpretability of Machine Learning Models. *Journal of Artificial Intelligence Research*, 65(1), 112–135.

- [26] Paul J. Deitel, Harvey Deitel. (2022). Intro to Python for Computer Science and Data Science: Learning to Program with AI, Big Data, and The Cloud. Pearson.
- [27] Pedregosa, F., Varoquaux, G., Gramfort, A., & Michel, V. (2011). Scikit-learn: Machine Learning in Python. *Journal of Machine Learning Research*, 12, 2825–2830.
- [28] Peng, C.-Y. J., Lee, K. L., & Ingersoll, G. M. (2002). An Introduction to Logistic Regression Analysis and Reporting. *The Journal of Educational Research*, 96(1), 3-14.
- [29] Priambodo, B., Jokonowo, B., Samidi, Ahmad, A., & Kadir, R. A. (2024). Predict Traffic State Based on PCA-KMeans Clustering of Neighbouring Roads. *Lecture Notes in Computer Science (Including Subseries Lecture Notes in Artificial Intelligence and Lecture Notes in Bioinformatics)*, 14322 LNCS. https://doi.org/10.1007/978-981-99-7339-2_36
- [30] Quinlan, J. R. (1996). Improved Use of Continuous Attributes in C4.5. *Journal of Artificial Intelligence Research*, 4, 77–90.
- [31] Resti Wahyuni. (2021). K-Means Clustering for Grouping Indonesia Underdeveloped Regions in 2020 Based on Poverty Indicators. *Parameter: Journal of Statistics*, 2(1), 8–15. <https://doi.org/10.22487/27765660.2021.v2.i1.15675>
- [32] Shmueli, G., & Koppius, O. (2011). Predictive Analytics in Information Systems Research. *MIS Quarterly*, 35(3), 553–572.
- [33] Smith, J., Allen, R., & Kumar, S. (2020). Applying Machine Learning Algorithms to Employee Job History Data for Improved Rotation Management. *IEEE Transactions on Human Resource Analytics*, 12(1), 56–72.
- [34] Suhada, M. I., Damanik, I. S., & Saragih, I. S. (2021). Analisis Kenaikan Jabatan Pegawai dengan Metode Promethee pada Kantor Kejaksaan Negeri Pematangsiantar. *Jurasik (Jurnal Riset Sistem Informasi Dan Teknik Informatika)*, 6(1). <https://doi.org/10.30645/jurasik.v6i1.274>
- [35] Turing, A. M. (1950). Computing Machinery and Intelligence. *Mind*, 59(236), 433–460.