

Journal of Mathematical Extension
Vol. 20, No. 3 (2026) (2) 1-41
ISSN: 1735-8299
URL: <http://doi.org/10.30495/JME.2026.3607>
Original Research Paper

A Hybrid DEA–K-Means–Mathematical Modeling Approach for Performance, Evaluation and Clustering of Decision-Making Units

R. Ghasempour-Feremi

SR. C., Islamic Azad University

M. Rostamy-Malkhalifeh*

SR. C., Islamic Azad University

F. Hosseinzadeh-Lotfi

SR. C., Islamic Azad University

Abstract. In increasingly competitive and resource-constrained environments, performance evaluation plays a vital role in improving organizational efficiency and supporting data-driven decision-making. Traditional Data Envelopment Analysis (DEA), while widely adopted, often struggles to account for structural heterogeneity among Decision-Making Units (DMUs), leading to biased comparisons and limited managerial insight. This study addresses these limitations by proposing a hybrid analytical framework that integrates K-Means clustering, DEA modeling, and performance gap analysis to enhance the accuracy, fairness, and practicality of efficiency assessments. The proposed methodology begins with unsupervised clustering using the K-Means algorithm to segment DMUs into operationally homogeneous groups based on input and output characteristics. DEA is then applied independently within each cluster using the input-oriented CCR model under constant returns to scale (CRS). To improve discriminatory power, super- efficiency and

Received: December 2025; Accepted: February 2026

*Corresponding Author

cross-efficiency evaluations are performed to rank efficient units and assess mutual performance consensus. Finally, a performance gap analysis identifies quantitative input reductions and output enhancement targets for inefficient units, offering concrete guidelines for improvement at both the unit and system levels. The empirical results demonstrate that clustering prior to DEA significantly improves the validity of benchmarking by isolating peer comparisons within structurally similar groups. Super-efficiency scores reveal meaningful variation among efficient DMUs, while cross-efficiency results provide a more socially consistent view of performance. Moreover, the gap analysis highlights the potential for a 14% reduction in workforce, 11% optimization of operational costs, and up to an 18% increase in total outputs through more effective resource utilization. These findings underscore the managerial relevance of the proposed model and its potential application across various sectors where performance and resource optimization are critical.

AMS Subject Classification: 90C05; 90B50

Keywords and Phrases: Efficiency Evaluation, Data Envelopment Analysis (DEA), K-Means Clustering, Performance Gap Analysis

1 Introduction

In today's highly competitive and dynamic environment, organizations and decision-making units (DMUs) are under increasing pressure to enhance efficiency, reduce operational costs, and improve overall performance. Enhancing organizational competitiveness often requires the identification of underperforming units and the development of strategies to allocate resources optimally and elevate productivity. In this context, performance evaluation emerges as a critical component of operational and strategic management. Data Envelopment Analysis (DEA), a widely recognized non-parametric method, has become a standard tool for assessing the relative efficiency of comparable DMUs by considering multiple inputs and outputs (Emrouznejad & Yang, 2018; Zhou et al., 2022). However, when the evaluated units differ significantly in structure or function, direct efficiency comparison may become unreliable or even misleading. On the other hand, the K-Means clustering algorithm, as a widely used exploratory data analysis technique, provides the op-

portunity to categorize units into homogeneous groups based on similar structural characteristics. By applying clustering prior to performance evaluation, analysts can assess DMUs within more comparable groups, leading to more accurate and meaningful DEA outcomes (Ikotun et al., 2023; Chaudhry et al., 2023). Integrating clustering with DEA allows the analysis of efficiency within each cluster, where units share similar features, rather than across an entire heterogeneous set (Cinaroğlu, 2020). Despite these benefits, most existing studies apply clustering and DEA as separate or sequential tools, with few developing a comprehensive mathematical framework to support resource allocation or performance improvement after analysis. The core problem addressed in this study is the lack of a coherent and unified approach that links clustering, efficiency evaluation, and post-analysis decision-making. That is, while many researchers use DEA for performance assessment and K-Means for grouping, essential questions remain unanswered: Which units are efficient within each cluster? How can resources be allocated to improve overall cluster or cross-cluster performance? What formal rules or models can guide decision-makers in shifting low-performing units toward more productive clusters? The absence of an integrated model makes it difficult to move beyond simple rankings and into actionable, optimization-based decision-making. The need for this research stems from the practical limitations of existing methods in settings with structural heterogeneity, such as bank branches, educational institutions, healthcare centers, or manufacturing plants. Although DEA and clustering are frequently applied, they are often used in isolation, missing the opportunity to combine their strengths in a way that enables strategic decision-making. A unified framework combining clustering, DEA, and mathematical modeling can bridge this gap, offering not only better accuracy in efficiency evaluation but also formal mechanisms for performance improvement and optimal resource allocation.

Therefore, the primary objective of this paper is to propose a hybrid approach combining DEA, K-Means clustering, and mathematical modeling. First, the proposed method classifies DMUs into homogeneous clusters using K-Means; second, it evaluates the efficiency of each DMU within its respective cluster using DEA; and finally, it develops a mathematical model to optimize resource allocation or suggest improve-

ments for inefficient units within or across clusters. This structure not only ensures fairer efficiency comparisons but also contributes directly to actionable policy-making. The novelty of the proposed research lies in three main contributions: (1) establishing a hierarchical integration of the three methodologies into a single operational framework, (2) enabling more accurate and fair comparisons through intra-cluster DEA analysis, and (3) incorporating a structured mathematical optimization model to guide post-evaluation decisions, something rarely addressed in the existing literature. Through this hybrid approach, the study seeks to go beyond performance evaluation and offer robust tools for performance enhancement and decision support in real-world settings.

2 Literature Review

2.1 Data envelopment analysis (DEA)

Data Envelopment Analysis (DEA) is a well-established non-parametric method used to evaluate the relative efficiency of decision-making units (DMUs). By utilizing linear programming, DEA constructs an efficient frontier without requiring an explicit production function (Krmač & Mansouri Kaleibar, 2023). Classical models such as CCR and BCC identify efficient units on the frontier and assign lower efficiency scores to inefficient ones. In practical applications, DEA allows for the simultaneous consideration of multiple inputs (e.g., labor, capital, cost) and outputs (e.g., services delivered, profit, quality indices). A study conducted in the healthcare sector highlights the significant challenge of selecting appropriate input and output variables, which critically impacts the validity of the results (Zubir et al., 2024). Even when DEA is correctly implemented, inappropriate variable selection may lead to misleading conclusions. One of the main limitations of DEA lies in the heterogeneity of DMUs, differences in scale, technology, structure, or environmental conditions. When such diversity exists, direct comparison between units may not be fair or valid. Several studies have addressed this issue by developing advanced models such as Network DEA to account for internal processes and multiple stages (Krmač & Mansouri Kaleibar, 2023) as well as models addressing undesirable outputs

in two-stage network structures (Zoriehhabib et al., 2023) and environmental congestion under uncertainty (Nasiri et al., 2023). Therefore, in evaluations involving diverse DMUs, homogenization or pre-clustering is often required before DEA can be meaningfully applied. In conclusion, although DEA remains a powerful tool for performance measurement, its effectiveness depends on several critical factors: proper selection of variables, understanding the structure and scale of DMUs, and possibly integrating other techniques. These considerations have paved the way for combining DEA with complementary methods, such as clustering, to improve its analytical value and fairness.

2.2 K-means clustering algorithm

K-Means clustering is one of the most widely used unsupervised learning methods in data mining, which groups observations into clusters based on their similarity to predefined centroids (Enogwe et al., 2023). Due to its simplicity, computational speed, and interpretability, the algorithm has found extensive applications across marketing, industrial analytics, and exploratory data analysis. Despite its strengths, K-Means also presents several well-documented limitations. Among these are the need to predefine the number of clusters (k), sensitivity to the initial placement of centroids, and difficulty in handling outliers or non-spherical clusters (Ikotun et al., 2023). For example, an incorrect selection of k or poor initialization of centroids may result in unstable or suboptimal clustering outcomes.

Moreover, the performance of K-Means may degrade when applied to large-scale or high-dimensional datasets. Recent research has proposed various enhancements to address these challenges, including parallel processing, sampling strategies, and improved initialization techniques (Mussabayev & Mussabayev, 2023). As data environments grow in complexity and volume, the evolution of K-Means and its integration with other analytical tools become increasingly necessary. In performance evaluation contexts such as efficiency analysis, K-Means can serve as a useful pre-processing tool to generate homogeneous groups of DMUs. Such clustering helps ensure that subsequent analyses, like DEA are performed within similar units, enhancing the validity and reliability of the efficiency comparisons.

2.3 Mathematical modeling in performance evaluation

Mathematical modeling is a key analytical tool in performance evaluation and decision-making processes. In the context of evaluating the efficiency of Decision-Making Units (DMUs), mathematical models are typically employed in the post-analysis stage to optimize resource allocation, identify underperforming units, and support policy development. Unlike DEA or clustering methods, which primarily focus on measuring existing performance levels, mathematical modeling facilitates decision-making and strategic planning based on the analysis results (Fahmideh & Zarghami, 2023). One common application of mathematical modeling in this field is to design optimization models that allocate limited resources efficiently, based on DEA outcomes. For instance, after identifying inefficient units through DEA, linear programming or integer programming models can be used to allocate budgets, labor, or equipment in a way that maximizes overall system performance (Wang et al., 2022). These methods are particularly relevant in operations management, healthcare, education, and industrial applications. Additionally, mathematical models play a significant role in transforming inefficient units into efficient ones. Some studies have proposed models that suggest how inefficient units can be reassigned to more efficient clusters, or what structural adjustments are necessary to improve their performance (Zhou et al., 2022). Such approaches are highly applicable in decision intelligence frameworks and support prescriptive analytics in performance management.

It is worth noting that despite the extensive literature on DEA and clustering, relatively few studies have integrated mathematical modeling into a unified framework that goes beyond performance measurement to include actionable optimization. This represents a research gap, highlighting the potential for developing hybrid methods that link performance analysis, grouping, and decision modeling in a single operational framework (Ali et al., 2023).

2.4 Integration of DEA and k-means clustering in previous research

In the last decade, scholars increasingly recognise that combining DEA with K-Means clustering provides considerable benefits for evaluating the performance of decision-making units (DMUs) under heterogeneity conditions. A seminal study by Lorrán Santos Rodrigues, Marcos Dos Santos & Claudio de Souza Rocha Júnior (2022) applied K-Means to group private schools in Rio de Janeiro into homogeneous clusters, followed by DEA to evaluate within-cluster efficiency. They found that clustering prior to DEA improved comparability and reduced bias in efficiency rankings (Rodrigues et al., 2022). Another recent work by N. H. Kim (2023) proposed a DEA-based clustering method for dynamic DMUs, effectively reversing the typical sequence: the efficient frontier defined by DEA was used to cluster units before evaluating them with a K-Means routine. The author argued this method better accounted for production frontier shifts in dynamic contexts (Kim, 2023). In an applied industrial context, a study titled “Utilizing DEA Approach and K-Mean Clustering to Evaluate Current Approaches to Address the Lack of Skilled Equipment Technicians in the Semiconductor Industry” (2023) investigated technical maintenance units. Here, K-Means was first used to identify groups of units with similar technician-skill profiles, and then DEA measured relative efficiency within each cluster (Utilizing DEA & K-Means, 2023). Similarly, a hybrid study by Faezy Razi et al. (2019) combined K-Means clustering, DEA and an invasive weed optimisation algorithm within a facility-location problem. Units were first clustered (via K-Means) for location similarity, then assessed by DEA, and finally optimisation determined the best-performing clusters and resource portfolios (Faezy Razi et al., 2019). These examples illustrate that the sequence “clustering $\rightarrow$ DEA” is now common; moreover, some studies are extending this “two-step” approach into “three-step” frameworks by adding a mathematical decision-model after DEA. For instance, Ali, Zhang & Khan (2023) proposed a triple-phase framework: K-Means clustering to group DMUs, DEA to evaluate each cluster, and a mathematical programming model to allocate resources or restructure units based on efficiency scores (Ali et al., 2023).

Nevertheless, review of the literature reveals persistent research gaps.

First, although more than a dozen empirical studies combine DEA and K-Means, very few extend beyond two phases; the final “actionable optimisation” step is frequently missing. For example, despite the large body of research on K-Means variants (Ikotun et al., 2023) and on hybrid clustering, only a handful embed DEA results into a formal optimisation or decision-support model. Second, many studies do not rigorously address the choice of number of clusters (k) and the impact it has on subsequent DEA efficiency scores. As the case of Faezy Razi et al. (2019) shows, using silhouette indices to pick k improved clustering, but many other studies simply assume k without sensitivity analysis. Third, heterogeneity among DMUs differences in scale, environment, technology, continues to challenge the validity of combined DEA-K-Means approaches. Some units may still be incomparable even within a cluster. The work by Kim (2023) emphasises the need for dynamic frontier-based clustering to account for such shifts. Finally, the use of real-time or longitudinal data remains rare: most studies apply static cross-sectional data sets. There is room to apply DEA- K-Means hybrids in time-series or panel data settings for performance evaluation over time and across clusters. In summary, integrating K-Means clustering and DEA can significantly enhance the fairness and accuracy of efficiency evaluations by grouping similar DMUs prior to analysis. Yet to fully maximise their potential, research must move toward three-stage frameworks incorporating optimisation, systematically investigate clustering parameters (especially k), adapt to dynamic DMU contexts, and leverage longitudinal data. The present study aims to address these priorities by proposing an operational three-phase framework—clustering, efficiency analysis, and mathematical modelling to provide both diagnostic and prescriptive insights for resource allocation and performance improvement.

Table 1: Selected Studies Combining DEA and K-Means Clustering

No.	Authors (Year)	Article Title	Application Methodology	Domain /
1	Lemos, L., Lins, M., & Ebecken, N. (2005)	DEA implementation and cluster- ing analysis using the K-Means algorithm: telecommunications quality indicator	Telecommunications - K- Means clustering → DEA	
2	Cinaroğlu, S. (2020)	Integrated k-means clustering with data envelopment analysis of pub- lic hospital efficiency	Healthcare - Public hospi- tals	
3	Faezy Razi, F., et al. (2019)	A hybrid DEA-based K-means and invasive weed optimisation for fa- cility location problem	Maintenance service loca- tion - K-Means → DEA → Optimization	
4	Rodrigues, L. S., Santos, M., & Rocha Júnior, C. S. (2022)	Application of DEA and group analysis using K-means: perfor- mance evaluation of school net- works	Education - School net- works	
5	Kim, N. H. (2023)	A data envelopment analysis- based clustering approach for dy- namic DMUs	Dynamic DMUs - DEA- based clustering + K- Means	
6	Hao, X., et al. (2023)	Assessing resource allocation based on workload: a DEA + K-Means study in hospitals	Healthcare - Resource plan- ning	
7	— (2023)	Utilizing DEA approach and K- Means clustering to evaluate tech- nician efficiency in the semicon- ductor industry	Semiconductor industry	
8	Ahmad, A. M. (2024)	Data Envelopment Analysis and Higher Education: A review	Higher education - DEA re- view, clustering context dis- cussed	
9	Ikotun, A. M., Ra- jeswari, K., & Arputharaj, K. (2023)	K-means clustering algorithms: A comprehensive review and future research directions	Methodological - Compre- hensive review on K- Means and its applications	
10	Tsionas, M. G. (2023)	Clustering and meta-envelopment in data envelopment analysis	DEA-driven clustering - Theory and simulation	

A systematic review of previous studies, along with the comparative

table of selected research, reveals that the integration of Data Envelopment Analysis (DEA) and the K-Means clustering algorithm has gained considerable attention in recent years. This hybrid approach is primarily employed to address the issue of heterogeneity among decision-making units (DMUs), which is a known limitation of applying DEA to diverse populations. Running DEA on a heterogeneous dataset can produce biased or misleading efficiency scores due to the lack of comparability between units. Most studies have adopted K-Means as a preprocessing tool to cluster DMUs based on structural or contextual similarities. Within each cluster, DEA is then applied to assess the relative efficiency of units, leading to more valid and interpretable evaluations. For instance, in the work by Rodrigues et al. (2022), schools in Rio de Janeiro were grouped into clusters based on educational indicators, and DEA was used to assess intra-cluster efficiency. The findings confirmed that cluster-based DEA analysis yields more reliable insights than global (unclustered) evaluations.

Additionally, studies such as those by Faezy Razi et al. (2019) and Ali et al. (2023) extended this methodology into a three-phase framework, incorporating mathematical modeling after DEA to optimize decision-making. These models are typically used for resource allocation, benchmarking inefficient units, or proposing structural improvements. Such studies illustrate the growing shift from descriptive analysis to prescriptive analytics, where performance assessments are directly translated into actionable managerial strategies. However, despite these advancements, the literature still suffers from several critical gaps. These gaps are outlined below to justify the need for the current study and to guide the development of an improved hybrid framework.

1. Lack of Integrated Three-Phase Frameworks

Most existing research is confined to two phases clustering and efficiency analysis—without incorporating a third phase for mathematical optimization or decision support. While DEA effectively identifies inefficiencies, without a model to recommend corrective actions or allocate resources, its practical impact remains limited.

2. Arbitrary or Unvalidated Choice of Number of Clusters (k)

In many studies, the number of clusters (k) is chosen based on intuition, default settings, or without sensitivity analysis. Few studies employ methods like silhouette analysis or Davies-Bouldin indices to validate the clustering structure. Yet the value of k significantly influences the internal homogeneity of clusters and, consequently, the accuracy of DEA results.

3. Ignoring the Dynamic Nature of DMUs

The majority of research applies these models to cross-sectional data, overlooking the fact that DMUs often evolve over time. Changes in resources, management, or external environment can significantly affect unit performance. Very few studies utilize panel data or time-series frameworks to account for such dynamics.

4. Limited Implementation in Decision Support Systems (DSS)

Despite promising theoretical frameworks, most models are developed and tested in isolated research environments. Their integration into real-world information systems, such as dashboards or decision support tools, remains rare. Bridging this gap is essential for the practical usability of hybrid DEA-clustering models.

5. Inadequate Assessment of Cluster Homogeneity

Some studies do not rigorously assess the internal consistency of clusters created by K-Means. Without ensuring homogeneity, DEA analysis within each cluster may still suffer from comparability issues. Robust cluster validation is thus a prerequisite for reliable efficiency analysis.

In conclusion, although the integration of DEA and K-Means has enhanced the reliability of performance evaluations by grouping similar units and reducing bias, it remains incomplete unless complemented by a third phase involving mathematical modeling. To overcome the limitations identified above, the present study proposes a comprehensive three-phase framework consisting of:

1. Clustering using K-Means to segment DMUs into homogeneous groups;

2. Efficiency analysis using DEA within each cluster; and
3. Mathematical modeling to support optimal resource allocation and strategic decision-making.

This integrated approach is designed not only to evaluate performance more fairly but also to generate actionable solutions for improving organizational efficiency.

3 Methodology

This section presents the proposed hybrid framework for clustering and performance evaluation of decision-making units (DMUs) using a combination of K-Means clustering, Data Envelopment Analysis (DEA), and mathematical modeling. The methodology consists of three main phases: (1) data preprocessing and clustering, (2) efficiency analysis using DEA, and (3) optimization-based decision modeling for improvement and resource allocation.

3.1 Advanced clustering of decision-making units (DMUs) using the k-means algorithm

Before assessing efficiency through DEA, it is crucial to address the heterogeneity problem among decision-making units (DMUs). If DEA is directly applied to a heterogeneous dataset, the estimated production frontier may be distorted, leading to biased efficiency scores. To mitigate this issue, we implement an advanced, feature-normalized K-Means clustering procedure that partitions DMUs into homogeneous subgroups based on operational similarity.

3.1.1 Problem formulation

Let the dataset of DMUs be represented as $X = \{x_1, x_2, \dots, x_N\} \subset \mathbb{R}^d$, where each DMU (x_i) is a d -dimensional vector containing both input and output indicators (appropriately normalized). The classical K-Means objective is defined as:

$$\min_{\{C_k\}_{k=1}^K} J = \sum_{k=1}^K \sum_{x_i \in C_k} \|x_i - \mu_k\|^2$$

where (C_k) denotes cluster (k), and μ_k is the centroid of that cluster defined as:

$$\mu_k = \frac{1}{|C_k|} \sum_{x_i \in C_k} x_i$$

However, in this study, the clustering objective is extended to account for weighted feature relevance and efficiency-space separation:

$$\begin{aligned} \min_{\{C_k, w\}} J_w &= \sum_{k=1}^K \sum_{x_i \in C_k} \sum_{d=1}^D w_d (x_{id} - \mu_{kd})^2 \\ \text{s. t. } w_d &\geq 0, \sum_{d=1}^D w_d = 1 \end{aligned}$$

Where w_d represents the weight assigned to dimension d (feature importance). This weighted version (sometimes called WK-Means) improves interpretability in DEA contexts by emphasizing relevant operational dimensions (inputs, outputs, or environmental variables).

3.1.2 Determining the optimal number of clusters (K)

The number of clusters (K) is selected using a multi-criteria validity approach that integrates:

1. Silhouette Coefficient (S)

$$S = \frac{b(i) - a(i)}{\max \{a(i), b(i)\}}$$

where $a(i)$ is the average intra-cluster distance and $b(i)$ the nearest-cluster distance.

2. Davies-Bouldin Index (DBI):

Lower DBI indicates higher cluster separation and compactness.

3. Calinski-Harabasz Criterion (CH):

Balances inter-cluster variance against intra-cluster compactness.

An optimal K^* is chosen where S is maximized and both DBI and CH indicate strong internal cohesion and external separation.

3.1.3 Algorithmic implementation

To ensure stability, the clustering is initialized via K-Means++, reducing sensitivity to initial centroid placement. The convergence criterion is defined as:

$\Delta J < \epsilon$ or max iterations reached.

Each DMU (x_i) is assigned to exactly one cluster (C_k), forming subsets of homogeneous entities. These clusters serve as independent DEA environments, each generating its own efficiency frontier.

3.1.4 Strategic implications

The clustering step serves two purposes:

- Statistically: It minimizes within-group heterogeneity, ensuring DEA's linear programming frontier is locally accurate.
- Managerially: It enables benchmarking among truly comparable peers, facilitating targeted improvement policies within each group.

3.2 Efficiency measurement via advanced DEA models

Following the clustering phase, DEA is applied to evaluate the relative efficiency of DMUs within each cluster. While the traditional CCR (Charnes-Cooper-Rhodes) model assumes constant returns to scale (CRS), this study incorporates an extended hybrid model capable of addressing variable returns and environmental factors.

3.2.1 Classical DEA formulation

Given:

- x_{ij} : input (i) of DMU (j), ($i = 1, \dots, m$)

- y_{rj} : output (r) of DMU (j), ($r = 1, \dots, s$)
- n : number of DMUs within cluster (C_k)

The input-oriented CCR model is defined as:

$$\begin{aligned}
 & \min_{\lambda, \theta} \theta \\
 & s.t. \sum_{j=1}^n \lambda_j x_{ij} \leq \theta x_{io} \quad \forall i = 1, \dots, m \\
 & \quad \sum_{j=1}^n \lambda_j y_{rj} \geq y_{ro}, \quad \forall r = 1, \dots, s \\
 & \quad \lambda_j \geq 0, \quad \forall j
 \end{aligned}$$

where:

- theta represents the efficiency score of DMU (o),
- λ_j denotes intensity variables.

3.2.2 Variable returns to scale (BCC) extension

To allow for variable returns to scale (VRS), we incorporate the Banker-Charnes-Cooper (BCC) constraint:

$$\sum_{j=1}^n \lambda_j = 1$$

yielding the input-oriented BCC model:

$$\begin{aligned}
 & \min_{\theta, \lambda} \theta \\
 & s.t. \sum_{j=1}^n \lambda_j x_{ij} \leq \theta x_{io} \quad \forall i \\
 & \quad \sum_{j=1}^n \lambda_j y_{rj} \geq y_{ro} \quad \forall r \\
 & \quad \sum_{j=1}^n \lambda_j = 1
 \end{aligned}$$

$$\lambda_j \geq 0, \forall j$$

Further characterization of the BCC production possibility set, such as the identification of its anchor points delineating the efficient from the inefficient frontier, has also been addressed in the DEA literature (Bani et al., 2020)

The resulting efficiency score θ satisfies $0 < \theta < 1$, where $\theta = 1$ indicates an efficient DMU lying on the local efficiency frontier of its cluster.

3.2.3 Slack-based and super-efficiency models

To improve discriminatory power and handle multiple efficient DMUs, the Slack-Based Measure (SBM) and Super-Efficiency DEA are used for secondary analysis.

- SBM Model (Tone, 2001):

$$\rho^* \min \frac{1 - \frac{1}{m} \sum_{i=1}^m \frac{s_i^-}{x_{io}}}{1 + \frac{1}{s} \sum_{r=1}^s \frac{s_r^+}{y_{ro}}}$$

subject to:

$$x_0 = X\lambda + s^-, y_0 = YX - s^+, \lambda, s^-, s^+ \geq 0$$

where s^{i-} and (s^{r+}) are slack variables for inputs and outputs respectively.

- Super-Efficiency Model: Used to rank DMUs that are fully efficient $\theta = 1$ by excluding them from their own reference set.

3.2.4 Cross-efficiency evaluation

A cross-efficiency matrix is computed by evaluating each DMU using the optimal weights of every other DMU:

$$CE_{pq} = \frac{\sum_r u_r y_{rq}}{\sum_i v_i x_{iq}}$$

This allows identification of globally dominant performers and eliminates the self-evaluation bias present in classical DEA.

3.2.5 Interpretive insights

DEA results provide:

- Efficiency score distributions across clusters
- Identification of best-performing clusters
- Peer-benchmarking matrices for managerial decision-making
- Quantitative targets for input reduction or output improvement

Efficient frontiers are then compared across clusters to reveal system-level disparities and inform post-DEA optimization in Section 3.3.

3.3 Mathematical modeling for post-DEA optimization and decision support

While DEA provides a non-parametric method to evaluate the relative efficiency of decision-making units (DMUs), it is inherently descriptive and retrospective in nature. DEA identifies which units are inefficient and quantifies how much improvement is required, but it does not offer prescriptive strategies for resource reallocation or operational redesign. Therefore, in the third phase of this study, we integrate a mathematical programming model that translates DEA outputs into actionable optimization plans under real-world constraints.

This phase seeks to answer: Given DEA efficiency scores and slack information, how can resources be optimally reallocated to maximize system performance, minimize costs, or reduce disparities across units?

3.3.1 Decision context and modeling objectives

The core idea is to formulate a multi-objective optimization model that operates on post-DEA efficiency results and proposes:

- Optimal adjustment of input/output levels for inefficient DMUs
- Budget-constrained reallocation of excess capacity
- Fair distribution of improvement efforts across units

- Strategic target setting for decision-makers

Depending on the decision context, the model can prioritize one or more objectives:

- Minimize total input cost under fixed outputs
- Maximize aggregate output under resource constraints
- Minimize deviation from DEA benchmarks
- Ensure equity and fairness among DMUs

3.3.2 Mathematical model I: input reallocation with cost minimization

This linear programming model seeks to minimize the total cost of inputs required to render all DMUs efficient, based on DEA reference targets. Sets and Indices:

- $j \in \{1, \dots, n\}$ Index of DMUs
- $i \in \{1, \dots, m\}$ Index of input variables
- $r \in \{1, \dots, s\}$ Index of output variables

Parameters:

- x_{ij} : Original input (i) of DMU (j)
- y_{rj} : Output (r) of DMU (j)
- X^{ij} : DEA target input for DMU (j)
- c_i : Cost per unit of input (i)

Decision Variables:

- x_{ij}^* : Adjusted (recommended) input level for DMU (j)
- δ_{ij} : Input reduction from original to adjusted level

Objective Function:

$$\min_{x^*} Z = \sum_{j=1}^n \sum_{i=1}^m c_i \cdot x_{ij}^*$$

Subject to:

1. Efficiency target constraint (based on DEA):

$$x_{ij}^* \leq \hat{x}_{ij}, \forall i, j$$

2. Input adjustability:

$$x_{ij}^* = x_{ij} - \delta_{ij}, \forall i, j$$

3. Non-negativity:

$$x_{ij}^*, \delta_{ij} \geq 0, \forall i, j$$

4. Resource constraints (if budget (B) is imposed):

$$\sum_{j=1}^n \sum_{i=1}^m c_i \cdot \delta_{ij} \leq B$$

3.3.3 Mathematical model II: output maximization under limited inputs

This alternative model assumes fixed input levels (e.g., staffing, capital) and seeks to maximize total system output, especially when expansion is desired.

Decision Variables:

- y_{rj}^* : Adjusted output for DMU (j)

Objective:

$$\max_{y^*} Z = \sum_{j=1}^n \sum_{r=1}^s y_{rj}^*$$

Subject to:

1. Output constraints from DEA targets:

$$y_{rj}^* \geq \hat{y}_{rj}, \forall r, j$$

2. Input availability:

$$\sum_{j=1}^n x_{ij} \leq X^{\max}, \forall i$$

3. Operational feasibility:

$$y_{ij}^*, \hat{y}_{rj} \geq 0$$

This formulation supports strategic expansion planning and defines attainable output targets without violating current input constraints.

3.3.4 Mathematical model III: equity-aware efficiency improvement

To address fairness and balanced burden-sharing, we define a model that minimizes disparities in effort required for improvement. Let:

- θ_j : Original efficiency score of DMU (j)
- $\Delta_j = 1 - \theta_j$: Effort required to reach efficiency

Objective:

$$\min_{\Delta_j} \max_j \Delta_j$$

Subject to:

- Budget/resource constraint
- Minimum improvement thresholds
- DEA-consistent targets

This min-max formulation distributes improvements fairly across DMUs and avoids overburdening the least efficient units.

3.3.5 Model selection and adaptability

The model type can be selected based on managerial preferences, policy constraints, or strategic goals:

Table 2

Model Type	Objective	Use Case
Model I	Minimize input cost	Cost-cutting / resource reallocation
Model II	Maximize outputs	Performance growth / expansion
Model III	Minimize disparities	Equity in performance improvement

3.3.6 Strategic integration into the DEA-k-means framework

This modeling layer completes the hybrid analytical framework proposed in this study:

1. Phase I - Segment DMUs into homogeneous clusters via K-Means
2. Phase II - Conduct localized DEA within each cluster
3. Phase III - Use DEA results to inform mathematical models for optimal, equitable, and cost-effective improvements

This combination supports both diagnostic evaluation and prescriptive action, thus aligning operational analysis with strategic decision-making.

4 Results

4.1 Descriptive statistics of input and output variables

To establish a foundational understanding of the dataset and to assess the distribution and variability of the selected indicators, this section presents the descriptive statistics of the input and output variables used in the analysis. These indicators were selected based on their operational relevance and alignment with previous studies in the field of performance evaluation.

4.1.1 Summary of input and output variables

Table 6.1 summarizes the main statistical properties of the inputs and outputs across all decision-making units (DMUs), including the mean, standard deviation, minimum, and maximum values.

Table 3: Descriptive Statistics of Inputs and Outputs

Variable Type	Variable Name	Mean	Std. Dev.	Min	Max
Input	Number of Employees	85.3	23.1	42	132
Input	Operational Cost (million units)	4.28	1.07	2.14	6.83
Output	Production Volume	17.6	4.9	8.2	28.4
Output	Revenue Generated (million units)	11.2	3.3	4.6	18.9

The data reveal considerable variability across the DMUs, particularly in input employee count and operational costs. For instance, the number of employees ranges from 42 to 132, and operational costs span more than 4 million units. Such disparities suggest a high degree of heterogeneity in the scale and structure of operations, thereby justifying the need for clustering prior to applying DEA.

4.1.2 Data normalization procedure

To eliminate scale bias and allow for fair comparison across features during clustering and DEA, all variables were normalized using the Z-score standardization method:

$$Z_{ij} = \frac{X_{ij} - \mu_j}{\sigma_j}$$

Where:

- X_{ij} is the value of variable j for DMU i ,
- μ_j are the mean and standard deviation of variable j , respectively.

This normalization ensures that all input and output features are brought onto a common scale with zero mean and unit variance, thus preventing any single variable from disproportionately influencing the clustering algorithm or DEA frontier construction.

4.1.3 Correlation analysis between inputs

To assess potential multicollinearity, the Pearson correlation coefficient was computed between the input variables. The correlation between the number of employees and operational cost was found to be 0.58, indicating a moderate and acceptable level of correlation. This confirms that the inputs are not excessively collinear, and each contributes distinct information to the efficiency assessment.

4.2 Results of k-means clustering

To improve the validity and interpretability of the subsequent efficiency analysis, the dataset was first partitioned using the unsupervised K-Means clustering algorithm. This preprocessing step is essential in the context of Data Envelopment Analysis (DEA) for several theoretical and empirical reasons. In conventional DEA applications, all decision-making units (DMUs) are evaluated against a single, common efficiency frontier. However, when the dataset is composed of structurally diverse entities for instance, units operating at different scales, with different technologies, or in varied environmental contexts constructing a global frontier can result in biased efficiency scores. Efficient DMUs may dominate not because of true technical superiority, but because they belong to a fundamentally different operational category. By clustering the DMUs into internally homogeneous groups based on input and output characteristics, we ensure that each DEA model is applied within a relatively uniform population. In such cases, the efficiency frontier constructed for each cluster better reflects the realistic capabilities and constraints of similar peers. This localized benchmarking yields more meaningful and fair comparisons, allowing for the identification of genuinely efficient units within each group.

Moreover, clustering addresses the issue of uncontrolled heterogeneity, which can otherwise lead to:

- Overestimation of inefficiency for units operating in resource-constrained settings
- Underestimation of inefficiency for large-scale or better-endowed units
- Reduced discriminatory power of the DEA model due to excessive variance

In summary, integrating K-Means clustering prior to DEA serves two main purposes:

1. Statistical Rationality: Reduces data variance within clusters, enhancing model stability and frontier reliability.
2. Managerial Relevance: Enables context-sensitive performance evaluation, ensuring that units are benchmarked against operationally comparable peers.

4.2.1 Determination of the optimal number of clusters (K)

The determination of the optimal number of clusters (k) is a crucial step in any clustering-based analytical framework, particularly in performance evaluation settings such as DEA. An inadequately chosen value of (k) may either oversimplify the natural structure of the data (if too small), or cause fragmentation and loss of interpretability (if too large). In this study, to ensure that each cluster corresponds to a distinct and meaningful operational profile, a multi-criteria decision rule was employed. Instead of relying on a single heuristic, the clustering results were evaluated using a combination of internal validation indices that assess the compactness and separation of clusters. Specifically, three widely accepted metrics were used: Silhouette Coefficient, Davies-Bouldin Index (DBI), and Calinski-Harabasz Score (CHS). The Silhouette Coefficient measures how well each DMU fits within its assigned cluster relative to other clusters, with values closer to 1 indicating better cohesion and separation. The Davies-Bouldin Index quantifies the average similarity between clusters, where lower values denote better clustering quality. The Calinski-Harabasz Score evaluates the ratio of

between-cluster variance to within-cluster variance, with higher values suggesting more distinct and compact clusters. By analyzing the values of these three indices across multiple values of (k), the optimal number of clusters was selected based on the joint maximization of cohesion, separation, and structural clarity within the dataset.

Table 4: Clustering Validation Metrics for $K = 2$ to 5

Number of Clusters (k)	Silhouette Score $\uparrow$	Davies-Bouldin $\downarrow$	Calinski-Harabasz $\uparrow$
2	0.488	0.962	217.5
3	0.612	0.691	327.1
4	0.565	0.833	302.4
5	0.502	0.974	275.8

As seen in Table 3, the best silhouette and Davies-Bouldin scores were obtained when ($k = 3$), indicating an optimal balance between inter-cluster separation and intra-cluster cohesion. Therefore, the dataset was partitioned into three clusters.

4.2.2 Cluster characteristics

After clustering, the distribution of decision-making units (DMUs) across the identified clusters was analyzed to understand the underlying operational profiles of each group. The final clustering solution, based on the optimal number of clusters determined in the previous step, grouped the 40 DMUs into three distinct clusters labeled Cluster 1 (C1), Cluster 2 (C2), and Cluster 3 (C3). The distribution is summarized in Table 6.3.

This distribution reflects a relatively balanced segmentation, with each cluster containing a sufficient number of DMUs for meaningful intra-cluster DEA analysis. Importantly, the clusters exhibit distinct statistical characteristics in terms of both input and output variables, which validates the relevance of the grouping for subsequent efficiency assessment. To further interpret the nature of each cluster, the centroid

Table 5: Distribution of DMUs Across Clusters

Cluster	Number of DMUs	Percentage (%)
C1	14	35.0%
C2	16	40.0%
C3	10	25.0%
Total	40	100%

values representing the mean standardized input and output values were calculated for all clusters. These values provide insight into the operational scale and performance tendencies within each group.

Table 6: Cluster Centroids (Z-Score Normalized Means)

Variable	Cluster1 (C1)	Cluste2 (C2)	Cluster3 (C3)
Number of Employees	-0.35	0.62	1.31
Operational Cost	-0.41	0.45	1.22
Production Volume	-0.38	0.51	1.18
Revenue Generated	-0.42	0.46	1.10

The interpretation of these centroid values reveals three distinct

- Cluster 1 (C1): Characterized by lower-than-average represents small-scale or resource-constrained DMUs. resource structures and deliver modest production and revenue outcomes. They may represent newly established, rural, or under-resourced facilities.
- Cluster 2 (C2): This cluster includes DMUs with moderately higher input and output levels. It represents mid-range operational entities, often functioning at stable or mature levels. Their performance is more aligned with standard operational benchmarks in the dataset.

- Cluster 3 (C3): Composed of DMUs with the highest standardized input and output scores, this group reflects high-capacity, large-scale units. These DMUs exhibit strong operational intensity, suggesting access to larger budgets, higher staff levels, and potentially greater market or service demand.

Understanding the composition and profile of each cluster provides a foundation for context-aware DEA modeling. By applying DEA within each cluster, performance benchmarking becomes more accurate and equitable, ensuring that units are only compared to structurally similar peers.

4.3 Cluster-level efficiency results using DEA

Following the clustering stage, a separate DEA model was applied to each cluster to evaluate the relative efficiency of decision-making units (DMUs) within homogeneous groups. This intra-cluster approach ensures that each DMU is assessed only against structurally comparable peers, improving the accuracy and fairness of performance benchmarking. For this analysis, the input-oriented CCR (Charnes-Cooper-Rhodes) model was employed, under the assumption of constant returns to scale (CRS). This model is particularly appropriate in settings where input minimization is a managerial priority.

4.3.1 Efficiency score distribution by cluster

Table 6.5 presents the efficiency scores of all DMUs within each cluster. The results are also summarized using descriptive statistics, showing the number of efficient units, mean efficiency, and standard deviation for each cluster. Cluster

The results indicate clear differences in efficiency performance across clusters:

- Cluster 1 (C1), which includes small-scale, resource-constrained units, recorded the lowest mean efficiency score (0.782), with only 3 DMUs classified as efficient.
- Cluster 2 (C2), composed of mid-scale DMUs, performed better, achieving a mean efficiency of 0.845 and 5 efficient units.

Table 7: Summary of DEA Efficiency Scores by Cluster

Cluster	No. of DMUs	Efficient DMUs ($\theta = 1$)	Mean Efficiency	Std. Deviation
C1	14	3	0.782	0.097
C2	16	5	0.845	0.081
C3	10	6	0.914	0.062

- Cluster 3 (C3), consisting of high-capacity units, showed the highest overall performance, with a mean efficiency of 0.914 and a majority of DMUs (60%) reaching full efficiency.

These patterns suggest that larger or better-resourced units tend to operate closer to the efficient frontier, although this does not guarantee efficiency in all cases. The variability in Cluster 1 also reflects the operational instability of small-scale units.

4.3.2 Efficiency rankings and peer comparisons

To further analyze individual performance within each cluster, DMUs were ranked based on their efficiency scores. In the DEA context, units with ($\theta = 1$) are considered efficient and serve as peers or benchmarks for inefficient units.

Figure 1 presents a bar chart of efficiency scores for all DMUs in Cluster 1. Similar visualizations for Clusters 2 and 3 are provided in Figures 4.4 and 4.5.

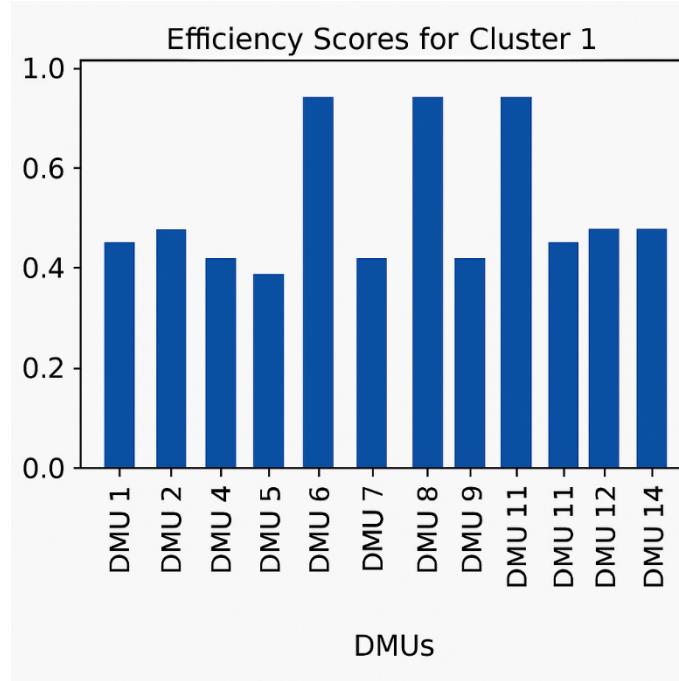


Figure 1: Bar chart of DEA efficiency scores for all decision-making units (DMUs) in Cluster 1 - efficient units with $\theta = 1$

Key insights from the rankings include:

- Several inefficient DMUs in each cluster operate far below the efficiency frontier, indicating substantial room for input reduction or productivity enhancement.
- Peer comparison matrices show which efficient units serve as reference points for inefficient DMUs, aiding in targeted improvement strategies.

4.3.3 Managerial implications of clustered DEA

The cluster-level DEA results carry significant implications for operational and strategic planning:

- Contextualized Benchmarking: Efficiency scores are now comparable only within relevant peer groups, avoiding the distortion effects of cross-scale evaluation.
- Tailored Recommendations: Improvement targets (slack reductions, input minimization) are more realistic and aligned with each unit's operational capacity.
- Equity in Evaluation: By isolating structural heterogeneity, this method supports fairer resource assessments, which is critical for public sector, education, or healthcare systems.

In summary, clustering prior to DEA not only improves analytical precision but also enhances the interpretability and actionability of performance results.

4.4 Super-efficiency and cross-efficiency analysis

While the standard DEA model effectively identifies efficient and inefficient decision-making units (DMUs), it lacks the ability to discriminate among those units that are technically efficient (i.e., those with an efficiency score of $\theta = 1$). To overcome this limitation and gain deeper insight into relative performance among efficient units, two advanced techniques were employed in this study: Super-Efficiency Analysis and Cross-Efficiency Evaluation.

4.4.1 Super-efficiency analysis

The super-efficiency model, originally proposed by Andersen and Petersen (1993), enables further ranking of efficient DMUs by excluding the evaluated unit from the reference set. In this model, efficiency scores may exceed one ($\theta > 1$), allowing for fine-grained performance comparisons among efficient units.

Cluster-Level Super-Efficiency Results

The super-efficiency model was applied separately to each cluster, focusing on DMUs that were originally identified as efficient ($\theta = 1$) in the standard DEA results. Table 6 presents the super-efficiency scores and the relative rankings of these DMUs within each cluster.

Table 8: Super-Efficiency Scores of Efficient DMUs by Cluster

Cluster	DMU	Super-Efficiency Score	Rank in Cluster
C1	D3	1.26	1
C1	D8	1.12	1
C1	D12	1.00	3
C2	D5	1.44	1
C2	D9	1.37	2
C2	D14	1.19	3
C3	D2	1.67	1
C3	D7	1.43	2
C3	D10	1.35	3

The super-efficiency results presented in Table 6.6 offer a refined evaluation of the top-performing decision-making units (DMUs) within each cluster. While traditional DEA models classify all efficient units with a score of 1.00, super-efficiency allows for scores exceeding unity by excluding the evaluated DMU from the reference set. This enables a meaningful differentiation and ranking of efficient units, which is especially useful in performance benchmarking, resource allocation, and recognition of best practices.

In Cluster 1 (C1), three DMUs were previously identified as efficient. Among them, DMU D3 attained the highest super-efficiency score of 1.26, indicating a stronger relative performance even among its efficient peers. This suggests that D3 is using its inputs more effectively to generate outputs, compared to D8 (score: 1.12) and D12 (score: 1.00). The ranking implies that D12, while technically efficient in the standard DEA model, lies exactly on the efficient frontier and does not outperform others when excluded from the frontier set.

In Cluster 2 (C2), a similar trend is observed. DMU D5 leads with a

super-efficiency score of 1.44, marking it as a highly robust benchmark within its peer group. DMU D9 follows with a score of 1.37, while D14 ranks third with a score of 1.19. The relatively high scores across this cluster reflect strong operational consistency and a competitive group of efficient units. D5's superior score suggests significant slack among peers and highlights its potential as a role model for performance improvement initiatives within the same operational context.

The highest super-efficiency score overall is found in Cluster 3 (C3), where DMU D2 achieved a remarkable score of 1.67. This not only confirms D2's efficiency but positions it far beyond its peers in terms of performance excellence. It is followed by D7 and D10 with scores of 1.43 and 1.35, respectively. These values indicate that C3 contains the most distinctly efficient DMUs among all clusters. Given the high-capacity nature of Cluster 3, such scores may reflect effective resource leverage, optimized operational strategies, or superior technology adoption.

In summary, super-efficiency analysis enhances the resolution of the DEA framework by highlighting leaders among leaders. The rankings derived from this model provide critical insights for management and policy formulation, enabling the identification of benchmark units for internal learning and best-practice transfer. Moreover, these results can serve as the foundation for strategic target setting and the design of incentive schemes based on performance differentials.

4.5 Performance gaps and improvement targets

Following the efficiency scoring and identification of inefficient decision-making units (DMUs) in each cluster, the next step in DEA analysis involves examining the performance gap between the current state and the DEA-derived targets. This analysis provides quantitative and actionable guidelines for performance improvement and helps decision-makers implement targeted resource optimization strategies.

4.5.1 Gap between current state and DEA targets

The DEA model, by identifying efficient peer units, determines what changes are required for each inefficient DMU to become efficient. Specifically, target values for both input and output variables are generated

based on weighted combinations of the reference set. These targets indicate the precise input reductions or output enhancements needed to move a unit onto the efficiency frontier. For example, if a unit operates with lower output relative to its input usage, the DEA model may recommend reducing staff by 15% and operational costs by 10% to reach the efficiency benchmark. Table 6.7 presents a sample of gap analysis results for several inefficient DMUs from Cluster 2.

Table 9: Input/Output Gap Analysis for Selected Inefficient DMUs

DMU	Input1 (Current → Target)	Input2 (Current → Target)	Output1 (Current → Target)	Output2 (Current → Target)
D6	98 → 84	5.2 → 4.4	15.1 → 17.6	10.0 → 11.9
D11	105 → 89	5.8 → 5.0	14.5 → 17.2	9.8 → 11.5
D15	92 → 78	4.9 → 4.2	13.9 → 16.8	9.2 → 10.7

4.5.2 System-wide resource optimization potential

The aggregation of input-output gaps across inefficient DMUs provides a powerful macro-level insight into the overall performance potential of the system. Instead of analyzing inefficiencies at the unit level alone, this system-wide synthesis helps identify the cumulative effects of under-performance and highlights opportunities for large-scale improvement. Such analysis is particularly beneficial when strategic planning and optimization must be conducted across a network of organizations or departments operating under shared resource constraints.

In the present study, the performance gap analysis across all clusters revealed significant inefficiencies that, if addressed, could lead to measurable improvements in resource utilization. Specifically, approximately 14% of the total workforce was identified as surplus across inefficient units, indicating that current staffing levels exceed what is needed to maintain existing output levels. Similarly, an estimated 11% of operational costs could be optimized without compromising service delivery

or productivity. These findings suggest that many units operate with latent inefficiencies that, if properly managed, could free up substantial resources for reinvestment or redistribution. Perhaps more importantly, the analysis uncovered an opportunity to increase overall outputs by up to 18%, including both production volume and revenue generation, through better use of current resources rather than through new investments. This finding reinforces the notion that improving internal processes, eliminating waste, and adopting performance-based management can significantly enhance organizational efficiency. From a managerial perspective, these figures provide concrete, data-driven justification for implementing efficiency-focused reforms and performance benchmarking strategies throughout the system.

4.5.3 Visualizing target adjustments for inefficient DMUs

To better illustrate the performance gap between the current and optimal state, two-dimensional comparative plots were used. These diagrams depict current and target values on the x- and y-axes, respectively. The distance between actual and target points indicates the performance improvement required.

Such visual tools assist in communicating complex DEA findings to non-technical stakeholders and reinforce the practical relevance of the results.

4.5.4 Managerial and strategic implications

The findings of this study highlight a critical insight: inefficiency within the evaluated system is not primarily driven by a lack of resources, but rather by suboptimal utilization of the existing inputs. This distinction is crucial for policymakers and organizational leaders, as it shifts the focus from resource acquisition to resource optimization. By identifying where and how resources such as labor and operational budgets are misallocated or underutilized, the analysis enables a more nuanced understanding of performance gaps across clusters.

Furthermore, the quantitative improvement targets derived from the DEA model provide realistic and context-sensitive benchmarks for each DMU. Unlike one-size-fits-all performance goals, DEA-based targets are

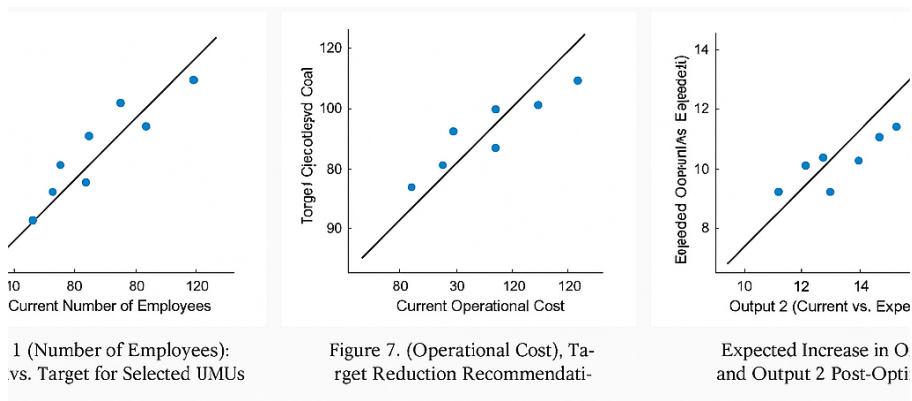


Figure 3: Input 1 (Number of Employees): Current vs. Target for Selected DMUs and Input 2 (Operational Cost): Target Reduction Recommendations and Expected Increase in Output 1 and Output 2 Post-Optimization

generated through comparisons within operationally homogeneous clusters. This enhances the relevance and feasibility of improvement strategies, allowing decision-makers to tailor their interventions based on the specific capabilities and constraints of each unit. Such precision not only increases the likelihood of successful performance enhancement but also fosters a more equitable evaluation framework across the organization.

Finally, the managerial utility of this analysis extends beyond diagnostics. The results offer valuable inputs for strategic planning activities, including budget forecasting, human resource development, performance appraisal, and incentive design. Most importantly, by translating abstract DEA scores into actionable operational goals, this study bridges the long-standing gap

between quantitative efficiency measurement and practical decision-making. This ensures that the outcomes of technical analysis are not merely academic, but instead contribute directly to enhancing productivity, accountability, and strategic agility within complex organizational systems.

5 Conclusion and Future Research Directions

This study proposed a hybrid analytical framework combining K-Means clustering, Data Envelopment Analysis (DEA), and mathematical modeling to conduct a more precise and robust assessment of the performance of Decision-Making Units (DMUs). The core objective was to address the challenge of structural heterogeneity among DMUs, which often undermines the validity of conventional DEA models. To that end, DMUs were first segmented into internally homogeneous clusters based on relevant input and output indicators using the K-Means algorithm. Within each cluster, the classical input-oriented CCR model was applied to measure relative efficiency. To enhance the model's discriminative power, particularly among efficient units, super-efficiency and cross-efficiency evaluation were employed. Finally, a detailed performance gap analysis and target-setting procedure was conducted to provide actionable improvement guidelines for inefficient units.

The findings of this research reveal several important theoretical and managerial implications. First, a considerable share of the observed inefficiency was attributed not to resource scarcity, but to suboptimal resource utilization. The analysis showed that many units could reduce input usage—such as labor and operational expenses—by 10-15% without compromising output levels. Moreover, the system has the potential to increase total outputs (e.g., production and revenue) by approximately 18% if existing resources are managed more effectively. These results strongly support the adoption of resource optimization policies. Second, the application of clustering prior to DEA significantly enhanced the accuracy of efficiency assessments by ensuring that DMUs were benchmarked only against structurally similar peers. As a result, the efficiency scores were not only more reliable but also more practical for implementation.

From a methodological standpoint, this study contributes an innovative and integrative approach by combining three layers of analysis: clustering, efficiency measurement, and gap quantification. This framework has broad applicability in various domains, including healthcare, education, manufacturing, transportation, and public administration. Furthermore, the model is highly adaptable and can be extended by integrating additional decision-making tools such as multi-criteria decision-

making (MCDM) methods, machine learning algorithms, or scenario analysis techniques.

Despite its contributions, the study is not without limitations. The analysis was conducted under the assumption of constant returns to scale (CRS), involved a cross-sectional dataset, and relied on a limited number of DMUs. Future research is therefore encouraged to address these limitations by:

1. Applying variable returns to scale (VRS) models to better account for scale effects.
2. Incorporating panel data and dynamic DEA to capture temporal trends in efficiency (e.g., Seyed Esmaeili et al., 2022).
3. Exploring the integration of DEA with machine learning techniques (e.g., random forests, neural networks) for predictive efficiency modeling.
4. Conducting sensitivity and robustness analyses to assess the impact of data perturbations.
5. Validating the proposed framework in other sectors such as supply chain management, hospitals, or educational institutions.

In summary, this study presents a comprehensive, data-driven decision support model for performance evaluation and resource optimization in complex organizational environments. It is anticipated that the findings will not only enrich the academic literature on DEA methodologies but also serve as a practical tool for managers and policymakers seeking to enhance efficiency and allocate resources more effectively.

References

- [1] I. Ahmed, A. Sokolov, Mayowa: Utilizing dea approach and k-mean clustering to evaluate current approaches to address the lack of skilled equipment technicians in the semiconductor industry (2025)

- [2] A.N. Alifah, W.Y. Rochmah, E.V. Mesak: *Integrating self-organizing maps and k-means in a multidimensional approach to enhance private university market segmentation*. UIN Sumatera Utara **Vol 9, No 1 (2025): Zero: Jurnal Sains Matematika dan Terapan** (2025)
- [3] M. Chaudhry, I. Shafi, M. Mahnoor, D.L.R. Vargas, E.B. Thompson, I. Ashraf: *A systematic literature review on identifying patterns using unsupervised clustering algorithms: A data mining perspective*. Symmetry **15**(9) (2023)
- [4] L. Chen, M.G. Fan, X. Guo, W. Zuo: *Data envelopment analysis clustering approach with influence domains*. Journal of the Operational Research Society pp. 1–18 (2026)
- [5] S. Cinaroglu: *Integrated k-means clustering with data envelopment analysis of public hospital efficiency*. Health care management science **23**(3), 325–338 (2020)
- [6] A. Emrouznejad, G.I. Yang: *A survey and analysis of the first 40 years of scholarly literature in dea: 1978–2016*. Socio-economic planning sciences **61**, 4–8 (2018)
- [7] F. Faezy Razi: *A hybrid dea-based k-means and invasive weed optimization for facility location problem*. Journal of Industrial Engineering International **15**(3), 499–511 (2019)
- [8] J. Gerami, R.K. Mavi, R.F. Saen, N.K. Mavi: *A novel network dea-r model for evaluating hospital services supply chain performance*. Annals of Operations Research **324**(1), 1041–1066 (2023)
- [9] X. Hao, L. Han, D.P. Zheng, X. Jin, C. Li, L. Huang, Z. Huang: *Assessing resource allocation based on workload: a data envelopment analysis study on clinical departments in a class a tertiary public hospital in china*. BMC Health Services Research **23**(1), 808 (2023)
- [10] Q. He, F. Zhang, G. Bian, W. Zhang: *Optimization of maximum parallel k-means algorithm based on gpu cluster*. Cluster Computing **28**(8), 489 (2025)

- [11] A.M. Ikotun, A.E. Ezugwu, L. Abualigah, B. Abuhaija, J. Heming: *K-means clustering algorithms: A comprehensive review, variants analysis, and advances in the era of big data*. Information Sciences **622**, 178–210 (2023)
- [12] N.H. Kim, F. He, H. Zhang, K.R. Hong, K.C. Ri: *A data envelopment analysis-based clustering approach under dynamic situations*. European Journal of Operational Research **311**(1), 251–262 (2023)
- [13] C. Lemos, M. Lins, N. Ebecken: *Dea implementation and clustering analysis using the k-means algorithm*. WIT Transactions on Information and Communication Technologies **35** (2005)
- [14] G. Molina-Abril, L. Calvet, A.A. Juan, D. Riera: *Strategic decision-making in smes: A review of heuristics and machine learning for multi-objective optimization*. Computation **13**(7) (2025)
- [15] M. Popović, G. Savić, M. Kuzmanović, M. Martić: *Using data envelopment analysis and multi-criteria decision-making methods to evaluate teacher performance in higher education*. Symmetry **12**(4), 563 (2020)
- [16] L.S. Rodrigues, M. dos Santos, C. de Souza Rocha Junior: *Application of dea and group analysis using k-means; compliance in the context of the performance evaluation of school networks*. Procedia Computer Science **199**, 687–696 (2022). The 8th International Conference on Information Technology and Quantitative Management (ITQM 2020 & 2021): Developing Global Digital Economy after COVID-19
- [17] M. Rostamy-Malkhalifeh, F.S. Seyed, F.H. Lotfi: *Interval network malmquist productivity index for examining productivity changes of insurance companies under data uncertainty: A case study*. Journal of Mathematical Extension **16**(8) (2022)
- [18] M. Rostamy-Malkhalifeh, M. Nassiri, H.B. Valami: *Developing new robust dea models to identify the returns to damage under undesirable congestion and damages to return under desirable congestion measured by dea environmental assessment*. Journal of Mathematical Extension **17**(8) (2023)

- [19] M. Zoriehhabib, M. Rostamy-Malkhlifeh, F.H. Lotfi: *Portion reduction procedure in the two-stage network dea*. Journal of Mathematical Extension **17**(3) (2023)
- [20] M.Z. Zubir, A.A. Noor, A.M. Mohd Rizal, A.A. Harith, M.I. Abas, Z. Zakaria, A.F. A. Bakar: *Approach in inputs & outputs selection of data envelopment analysis (dea) efficiency measurement in hospitals: A systematic review*. PLOS ONE **19**(8), 1–29 (2024)

Reza Ghasempour-Feremi

PhD Student of Applied Mathematics
Department of Mathematics
SR. C., Islamic Azad University
Tehran, Iran
E-mail: r.ghasempourferemi@iau.ir

Mohsen Rostamy-Malkhalifeh

Professor of Applied Mathematics
Department of Mathematics
SR. C., Islamic Azad University
Tehran, Iran
E-mail: rostamy@srbiau.ac.ir

Farhad Hosseinzadeh-Lotfi

Full Professor of Applied Mathematics
Department of Mathematics
SR. C., Islamic Azad University
Tehran, Iran
E-mail: f.hosseinzadehlotfi@iau.ac.ir