

AN INTEGRATED CLUSTERING AND DEEP LEARNING FRAMEWORK FOR EDUCATIONAL PERFORMANCE ASSESSMENT IN PRIMARY AND SECONDARY EDUCATION

Zahira Noor Quraishi¹, Dr. Jitendra Jain²

¹Research Scholar, ²Research Supervisor

^{1,2}Dr. A. P. J. Abdul Kalam University, Indore, India

zahira16sep@gmail.com , Jitendra1974@gmail.com

Abstract:

Primary and secondary education systems generate large-scale survey and monitoring data from students, teachers, school leadership, and external observers; however, most analyses remain descriptive and do not provide actionable performance assessment or targeted intervention planning. This study proposes an integrated clustering and deep learning framework to address this gap by discovering stakeholder profiles and predicting educational quality outcomes under a unified pipeline. The framework first performs data cleaning through numeric conversion, removal of empty variables, median imputation, and z-score standardization. Unsupervised clustering is then applied to segment stakeholders into meaningful groups using K-Means, DBSCAN, and Agglomerative clustering, enabling differentiated policy actions and identification of outlier profiles. Subsequently, a deep learning model (ANN/MLP) is trained to learn non-linear relationships among educational indicators for performance assessment and risk classification. Experiments are conducted on the MP Education Survey dataset, comprising five aligned datasets primary students, secondary students, teachers, headmasters, and observers linked by school-level identifiers. Clustering results reveal distinct ability, governance, and infrastructure condition groups, while DBSCAN highlights anomalous schools and teacher profiles. Deep learning achieves consistently strong classification performance across stakeholder datasets (approximately 0.81–0.88 accuracy), with the best performance observed for leadership and teacher quality prediction. The proposed framework supports profile-based interventions and scalable monitoring for educational improvement.

Keywords: Primary and Secondary Education; Educational Data Mining; Clustering; Deep Learning; K-Means; DBSCAN; Agglomerative Clustering; Artificial Neural Network; Performance Assessment.

1. Introduction

Primary and secondary education systems are influenced by multiple interacting factors, including student engagement, teacher effectiveness, school leadership, and the availability and utilization of resources. Governments and education departments increasingly collect large-scale surveys and monitor data to track these factors, yet much of this information is still used mainly for descriptive reporting. While summary statistics and dashboards provide an overview of current conditions, they often fail to support early identification of risk, discovery of hidden stakeholder patterns, and evidence-based targeting of interventions [1]. This limitation is critical in public education contexts where resources are constrained, and policy decisions must be prioritized toward the most vulnerable learners and schools.

Educational Performance Assessment is not only a student-level problem. Student learning outcomes are shaped by attendance behavior, motivation, home support, and classroom conditions. At the same time, teacher quality influences instructional delivery, assessment practices [2], and classroom engagement. Headmasters contribute through academic monitoring, school planning, and governance practices, while observer-based monitoring provides external signals about compliance, classroom quality, and infrastructure status. Because these dimensions operate simultaneously, analyzing only one stakeholder dataset (such as students) can produce

incomplete conclusions and generic interventions [3]. A unified analytical approach that integrates multiple stakeholder perspectives is therefore essential for a realistic understanding of school performance in primary and secondary education.

Machine learning offers tools to move beyond descriptive analysis by learning relationships and patterns directly from data [4]. However, education datasets present practical challenges: missing values are common, indicators may be recorded in mixed formats, and constructs such as “performance” or “quality” are often captured through multiple survey items rather than a single direct measure. Additionally, schools and stakeholders are heterogeneous; similar average scores can hide groups with very different risk profiles. These challenges motivate methods that can both discover natural groupings and learn non-linear predictive patterns [5].

This study proposes an integrated clustering and deep learning framework for educational performance assessment in primary and secondary education. The framework is designed to produce actionable insights at multiple levels: (i) segmentation of stakeholders into interpretable profiles for differentiated support, (ii) identification of outliers for rapid attention, and (iii) deep learning-based prediction for scalable risk assessment. A standardized preprocessing pipeline is applied to all datasets to ensure consistency and comparability. The pipeline includes numeric conversion of survey responses, removal of empty variables, median imputation for missing values, and z-score standardization to reduce scale bias and improve model stability.

For the unsupervised component, clustering methods are employed to segment stakeholders into groups with similar behavioral, instructional, leadership, or infrastructure characteristics. K-Means is used to form compact profile groups and provide centroid-based interpretation. Agglomerative clustering supports hierarchical grouping, enabling multi-level analysis of similarity structures that can reflect academic or governance maturity levels. DBSCAN is applied as a density-based approach to identify irregular patterns and explicitly detect noise points, which is valuable for highlighting anomalous schools or stakeholder profiles that may be overlooked in average-based reporting. These clustering outputs support policy design by enabling differentiated interventions rather than one-size-fits-all programs [6].

For the predictive component, an Artificial Neural Network (ANN/MLP) is used to learn complex non-linear interactions among educational indicators. Deep learning is particularly relevant for multi-factor educational data where outcomes may depend on combinations of attendance behavior, engagement, teacher support, and leadership practices [7]. The ANN model is trained to generate performance predictions and risk classifications, enabling early warning and scalable monitoring across schools. When deployed, such models can complement clustering insights by identifying which individuals or institutions within each cluster require immediate attention.

The proposed framework is evaluated on the MP Education Survey dataset, comprising aligned data collected from multiple school stakeholders, including primary students, secondary students, teachers, headmasters, and observers [8]. The datasets are linked using school-level identifiers, which support consistent comparison of stakeholder patterns under a shared methodology. Experimental outcomes demonstrate that clustering reveals meaningful stakeholder group structures and supports identification of outlier patterns, while the ANN model achieves strong predictive performance across datasets, indicating reliable learning of non-linear educational signals.

By integrating clustering and deep learning into a unified process, this work contributes a practical methodology for education departments and researchers to transform survey data into decision-support outputs. The approach supports segmentation-driven program design, anomaly detection for rapid inspections, and scalable predictive monitoring for improving learning outcomes and institutional quality in primary and secondary education.

Key Contributions

1. Integrated framework: Proposes a unified pipeline combining clustering (K-Means, DBSCAN, Agglomerative) and deep learning (ANN/MLP) for educational performance assessment.

2. Multi-stakeholder coverage: Applies the same methodology across five aligned datasets—primary students, secondary students, teachers, headmasters, and observers—to enable holistic school quality analysis.
3. Segmentation and outlier detection: Demonstrates how clustering supports differentiated interventions and how DBSCAN identifies anomalous schools and stakeholder profiles for priority action.
4. Non-linear predictive modeling: Uses ANN to capture complex interactions among indicators and provides scalable risk classification/performance prediction across stakeholder datasets.
5. Actionable policy mapping: Links model outputs to practical recommendations such as targeted remedial programs, teacher development, leadership strengthening, and infrastructure monitoring.

2. Literature review

Lu et al. (2025), Student dropout remains a serious academic and financial challenge for institutions, and EDM offers practical tools for early identification of at-risk learners. This systematic review specifically evaluates cluster-based prediction models that combine clustering with predictive techniques to improve student performance forecasting. Using PRISMA guidelines, it analyzed 61 studies published from 2010 to 2025, focusing on clustering methods, model integration, and data types used. Findings suggest clustering captures meaningful behavioral and academic differences, enabling targeted interventions. However, concerns persist about scalability, generalisability across institutions, and ethical issues when transferring models across datasets. [1]

Malik et al. (2025), The rapid digitalization of education has generated large volumes of data, creating opportunities to improve prediction accuracy through advanced feature selection methods. This study proposes a novel model, Dynamic Feature Ensemble Evolution for Enhanced Feature Selection (DE-FS), which integrates traditional feature selection techniques with heat maps and adaptive thresholding. Unlike static methods, DE-FS dynamically adjusts to changing data patterns, reducing overfitting and underfitting. The results demonstrate superior predictive performance across diverse datasets, enabling accurate student performance prediction, better resource allocation, and targeted interventions that support personalized and effective learning experiences. [2]

Tang et al. (2025), Educational Data Mining (EDM) has become an important approach for improving educational quality by extracting meaningful insights from educational data. This study investigates the effectiveness of machine learning techniques—such as decision trees, support vector machines, and neural networks in predicting student performance. Using historical, demographic, and behavioral data, predictive models are developed and evaluated, emphasizing the importance of feature selection and data preprocessing. Results indicate that machine learning significantly enhances prediction accuracy, enabling targeted interventions and personalized learning strategies, and reinforcing the value of data-driven approaches in educational improvement. [3]

López-Meneses et al. (2025), This study explores the integration of Artificial Intelligence (AI) into Educational Data Mining (EDM), Human-in-the-Loop Machine Learning (HITL-ML), and machine-assisted teaching to support personalized and adaptive learning. Through a systematic review of 370 articles published between 2006 and 2024, the research shows how AI enables performance prediction, early interventions, and improved resource optimization. HITL-ML ensures educator oversight to reduce bias, while machine-assisted teaching aligns AI outputs with pedagogical goals. Ethical considerations such as privacy, transparency, and equity are emphasized for responsible and inclusive implementation. [4]

Bošnjaković et al. (2025), This systematic review examines how Educational Data Mining (EDM) methods are applied within gamification research to support learning outcomes through increased motivation and attention. Using Web of Science and SpringerLink, the authors selected 32 relevant articles and analyzed the most frequently used EDM methods, their purposes, and the education levels where they are most common. The review highlights a literature gap on the EDM–gamification relationship and shows that EDM approaches are still underutilized in this area. Overall, the findings map current practice and offer directions and recommendations for researchers to expand EDM's role in future gamification studies. [5]

Lou et al. (2025), Student performance prediction enables schools to identify at-risk learners and provide timely interventions, yet limited research compares EDM techniques with traditional statistical methods in real educational settings. This case study evaluates generalized linear regression, decision tree regression, and random forest regression using a statewide high school dataset to predict three end-of-course exams. Performance was measured via R-square, RMSE, MAE, and MSE. Results show generalized linear regression consistently outperformed the two machine learning models on both accuracy and error metrics, while also supporting identification of key predictors informing method selection for similar real-world prediction tasks. [6]

He et al. (2025), This study addresses inconsistent effectiveness in Mathematics Education Technology (MET) integration in K–12 classrooms, focusing on how technology type, timing, and instructional strategies influence outcomes. Using educational data mining and AI, the authors built predictive models from 423 publications on MET effectiveness in Chinese K–12 mathematics education. The best-performing model used eXtreme Gradient Boosting with L2 regularization, SMOTER, and Active Learning, further tuned via Particle Swarm Optimization and interpreted using SHAP for feature importance. Results identify MET duration as critical, and a controlled experiment validated that model-guided MET improved achievement over traditional approaches. [7]

Imran et al. (2025), Pre-admission requirements differ across universities, and this study investigates whether high school performance predicts entrance exam outcomes for applicants to a public engineering university. Using a dataset containing candidates' high school scores and entry exam results, the authors examined relationships between the two and developed predictive models for admissions support. Multiple statistical and machine learning methods were tested, including decision trees, k-nearest neighbor, neural networks, and naïve Bayes, with model performance evaluated using accuracy metrics. Results indicate decision tree and neural network models performed best, confirming high school scores as significant predictors for entrance performance. [8]

Noviandy et al. (2025), Learning Analytics (LA) and EDM are rapidly developing fields that apply data-driven techniques to improve learning, teaching, and institutional strategy. This review summarizes core methods classification, clustering, regression, association rules, and visualization alongside newer approaches like deep learning and adaptive systems for personalization. It also examines platforms and tools (learning management systems, open-source solutions, and AI/ML libraries) in terms of scalability and real-world adoption. Through case studies across K–12 and higher education, the paper covers applications such as dropout prediction, engagement analysis, curriculum design, and personalized learning, while emphasizing ethics, interpretability, and usability as key requirements for responsible educational analytics. [9]

Huang et al. (2025), This work applies deep learning and AI-driven data mining to improve university English teaching evaluation and support personalized instruction. Using a Bayesian framework and Transformer architecture, the approach analyzes evaluation data from a university English teaching and assessment system to build student group profiles and predict exam outcomes. Findings report that active English learning is occasional for over 70% of students, with gender differences in learning behaviors and perceived skill importance. Question-type scores strongly influence pass outcomes, reflecting knowledge mastery and application ability. The study argues that Transformer-based feature extraction improves objectivity and accuracy in evaluation, reducing human bias and demonstrating interdisciplinary transfer from NLP to education. [10]

Lopez-Fernandez et al. (2025), A microcompetency-based teaching methodology is introduced for higher education, combining practical case-study quizzes with microcompetences embedded into course activities. The digital tool “Sapiens” is used to detect learning deficiencies and provide group-level and individualized feedback, supporting continuous assessment. Findings show increased participation and improved academic performance compared with previous years, alongside voluntary uptake of virtual learning modalities for difficult topics. Data mining helped identify student performance patterns, suggesting the approach strengthens both transversal and discipline-specific competences, while offering flexibility across in-person, hybrid, and online teaching formats. [11]

Kikunda et al. (2025), Focusing on first-year students in the Democratic Republic of Congo, this study develops early-risk detection models for the Catholic University of Bukavu and ISP to support student guidance and academic support. To handle class imbalance, SMOTE is used to rebalance data, followed by a stacking classifier to combine multiple algorithms' strengths. Predictors include national exam score, prior school track, institution type, age, and sex. Apriori association rules are additionally used to extract faculty-specific success profiles, improving interpretability. The study notes limitations such as non-public data and geographical specificity, and suggests Explainable AI as future work. [12]

Lampropoulos et al. (2025), This systematic review examines the intersection of educational data mining and learning analytics with extended reality technologies VR, AR, MR, and the metaverse using PRISMA across 70 documents from six databases. Through bibliometric mapping and content analysis, it identifies themes showing that combining XR with EDM/LA can enrich learning, support adaptive and personalized education, and enable multimodal tracking of behaviors, emotions, and cognitive/affective states. The review reports benefits for collaborative and social learning, self-efficacy, self-regulated learning, engagement, motivation, and achievement, highlighting real-time monitoring and visualization as a major advantage. [13]

Papadogiannis et al. (2024), Educational Data Mining (EDM) is an emerging scientific field that develops and applies analytical methods to study data produced in educational environments. As web-based learning has grown, EDM has broadened to include diverse data sources and techniques such as classification, regression, clustering, association rule mining, and Natural Language Processing. These approaches enable evidence-based decisions for curriculum design and personalized learning experiences. The literature also highlights ongoing challenges and anticipates strong future growth through deeper integration with AI and technologies like augmented and virtual reality, potentially transforming education at scale. [14]

Simionescu et al. (2024), High-quality education is central to sustainable development, as emphasized in the UN's "2030 Agenda for Sustainable Development" (2015). Major shifts—driven by social and technological change, the COVID-19 pandemic, and global conflicts—have altered how education is delivered and have increased the volume of educational data that can support better decision-making. EDM focuses on uncovering actionable insights from these datasets to improve education quality over time. To map recent progress in the European Union, the authors conducted a systematic, critical literature review to synthesize trends, methods, applications, and research gaps. [15]

López-Meneses et al. (2025), A large-scale review of 793 Scopus-indexed articles (2000–2024) examines how EDM and predictive modeling evolve in the AI era using bibliometric and systematic review methods. The study highlights that AI-driven analytics can strengthen student performance prediction, personalize learning pathways, and support timely interventions by analyzing educational data. It also notes AI's potential to enhance teaching practice while stressing that educators should retain oversight to reduce bias and misuse. Ethical issues—privacy, transparency, and equity of access—are emphasized as essential considerations. Overall, results indicate AI-enabled EDM can improve outcomes, resource allocation, and institutional effectiveness when implemented responsibly. [16]

Hoyos et al. (2024), EDM is also being used to predict academic performance through tuned classification models on large student datasets. One study analyzed data from 19,545 secondary school students using descriptive statistics across personal, school, and socioeconomic variables, then applied artificial neural networks (ANN) and support vector machines (SVM). Model tuning was performed with five-fold cross-validation, and evaluation relied on accuracy and F1-score. Results showed strong relationships between predictors and performance, with average accuracy above 80%. ANN outperformed SVM in this dataset, suggesting such models can help institutions detect underachievement early and design strategies to improve learning outcomes. [17]

Ordóñez-Avila et al. (2023), Teacher evaluation in higher education is an important research area that relies on heterogeneous data from multiple academic stakeholders, particularly students, who provide rich and meaningful feedback. This study conducts a systematic literature review to identify research on predicting teacher evaluation using student performance data. Following structured phases of planning, selection, and

extraction, the review highlights extensive use of educational data mining techniques, including fuzzy-based models, neural networks, decision trees, and machine learning classifiers. The findings indicate growing emphasis on fuzzy principles, reflecting the inherent uncertainty of human decision-making in evaluation processes. [18]

Ashish et al. (2024), The integration of Information and Communication Technology (ICT) has significantly transformed modern educational environments, enabling interactive, adaptive, and learner-centered teaching approaches. Digital platforms, virtual classrooms, and multimedia tools have reshaped traditional pedagogy and expanded opportunities to enhance student performance. This study investigates the role of ICT in improving student learnability by applying data mining techniques to educational datasets. By analyzing patterns and correlations between ICT usage and academic outcomes, the research aims to provide evidence-based insights that support effective technology-enhanced learning and inform future educational practices. [19]

Staneviciene et al. (2024), This study examines the use of data analytics and machine learning to predict academic outcomes and support sustainable e-learning, particularly in university technological programs with high dropout rates. Student Performance Prediction (SPP) is emphasized for its ability to enable personalized learning and early interventions. Using a case study guided by the CRISP-DM methodology, the research applies classification, regression, and clustering techniques to academic data. Despite challenges related to data quality and completeness, results demonstrate that mining learning process data can effectively identify at-risk students and support educational quality aligned with Sustainable Development Goal 4. [20]

Kaur et al. (2023), Educational Data Mining (EDM) and Learning Analytics are increasingly used to address critical educational challenges such as student performance prediction and assessment. This systematic literature review analyzes studies published between 2012 and 2022 to examine techniques employed in both domains. EDM focuses on applying data mining methods to educational data, while learning analytics emphasizes understanding learning processes through visualization and quantitative analysis. Reviewing 41 relevant studies, the paper identifies commonly used statistical, machine learning, and visualization techniques, highlights research gaps, and provides guidance for educators, researchers, and policymakers seeking to enhance learning outcomes through data-driven approaches. [21]

Missah et al. (2023), The application of Data Mining (DM) in education supports evidence-based decision-making by educational leaders. This review examines how DM research is distributed across different educational levels, with a particular focus on Basic Education (BE). Analyzing 94 studies published between 2017 and 2022 from nine major academic publishers, the findings reveal a strong imbalance: most EDM research targets tertiary education, while basic and pre-tertiary levels remain underrepresented. Additionally, existing studies focus heavily on student performance factors, overlooking critical aspects such as pedagogical resources, revealing both population and knowledge gaps in EDM research. [22]

Sun et al. (2024), Data mining techniques in EDM are widely used to analyze student preferences and learning behaviors, particularly in e-learning environments. This study proposes a novel model, the Chaotic-Tuned Shuffled Frog Leaping Optimized Random Forest (CSFLO-RF), to improve learning behavior prediction. Student activity logs are clustered using K-Means to identify behavior patterns and learning styles, which are then classified using the proposed model. Experimental results, implemented on the WEKA platform, demonstrate that CSFLO-RF outperforms existing methods, highlighting its effectiveness in forecasting student behavior and enhancing adaptive learning systems. [23]

Abideen et al. (2023), Student enrollment analysis is a critical but underexplored area in educational data mining, particularly in developing regions. This study focuses on predicting school enrollment in Punjab, Pakistan, using five years of data from 100 schools. Machine learning algorithms—Multiple Linear Regression, Random Forest, and Decision Tree—are applied to identify significant features and predict future enrollment trends. The proposed model supports enrollment target classification and provides insights to address low enrollment levels, thereby contributing to improved planning, literacy rates, and more effective education policy implementation. [24]

Wang et al. (2024), In secondary education, large volumes of student achievement data are often underutilized, limiting their potential to support educational improvement. This study highlights how data mining techniques can enhance education informatization by analyzing secondary school performance data more effectively. Traditional teaching methods are identified as insufficient for engaging students, prompting the adoption of case-based teaching supported by data analysis. By statistically evaluating student performance and applying targeted instructional strategies, the study demonstrates how data mining can improve learning outcomes and support the broader goals of quality and vocational education. [25]

Singh, Manmohan et al. (2016), Predicting student performance is an important concern in primary education, especially for enabling timely support and improving next-year results. This study analyzes rural and urban primary school students in Betul district, Madhya Pradesh (India), using a survey-cum-experimental approach to build a dataset from both primary and secondary sources. The objective is to identify factors linked to prior exam performance and select a suitable data mining algorithm to predict students' grades. Hypothesis testing indicates that school type does not significantly influence performance, while school area (rural/urban) and students' previous results (along with related background factors such as occupation) play a major role in predicting grades. [26]

Sehaj Singh et al. (2021), Secondary and senior secondary education (Classes IX–XII) in India forms a crucial bridge between school and college, making strong educational infrastructure essential for holistic and interactive learning. This paper examines the growth of Madhya Pradesh's secondary education sector by mapping infrastructure development trends. While the state has shown progress through initiatives such as the Education Guarantee Scheme (EGS) and Muft Cycle Yojana, the study argues that further efforts are required to meet SDG 4 (Quality Education) and its targets. Using a blend of primary and secondary data, the research analyzes micro-level progress and proposes interventions to help policymakers design inclusive, equity-focused education policies so that no learner is left behind. [27]

3. Proposed work

3.1 Proposed architecture

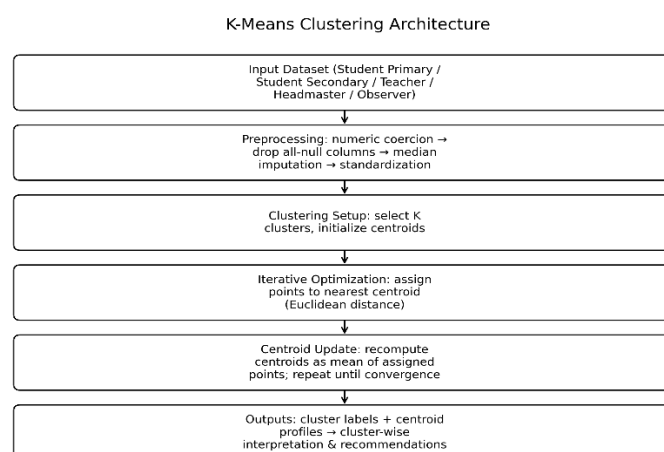


Figure 1. K-Means Clustering Architecture

Figure 1 illustrates the K-Means clustering architecture applied across all stakeholder datasets, including primary students, secondary students, teachers, headmasters, and school observers. The process begins with dataset ingredients followed by standardized preprocessing, which includes numeric conversion of indicators, removal of empty variables, median imputation of missing values, and z-score normalization. A predefined number of clusters is selected to represent distinct stakeholder profiles. The algorithm iteratively assigns each instance to the nearest cluster centroid based on Euclidean distance and updates centroids as the means of assigned instances until convergence is achieved. The resulting clusters capture homogeneous groups with similar performance, behavioral, or infrastructure characteristics, enabling profile-based interpretation and targeted educational planning.

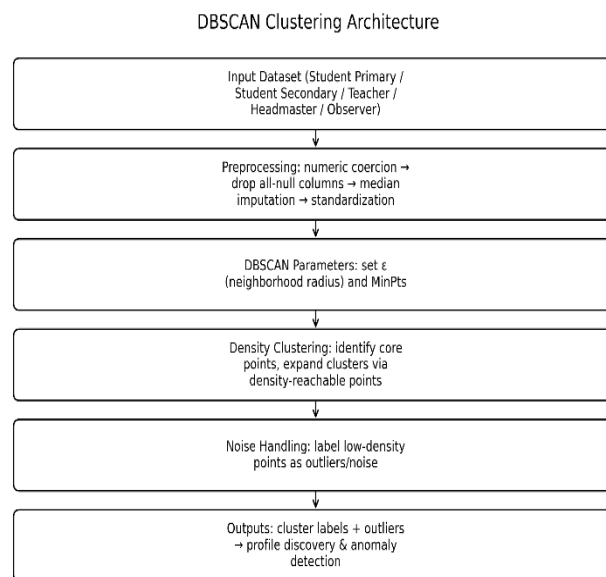


Figure 2. DBSCAN Clustering Architecture

Figure 2 presents the DBSCAN clustering architecture used for density-based pattern discovery and outlier detection in educational datasets. Following numeric conversion, missing-value handling, and feature standardization, the algorithm groups instances based on local density using two parameters: neighborhood radius (ϵ) and minimum points (MinPts). Core points with sufficient neighboring density form the basis of clusters, which are expanded by recursively aggregating density-reachable points. Instances that do not meet density requirements are labeled as noise or outliers. This architecture enables the discovery of clusters with arbitrary shapes and explicitly identifies anomalous stakeholder profiles, making it particularly effective for heterogeneous datasets such as observer-based infrastructure assessments.

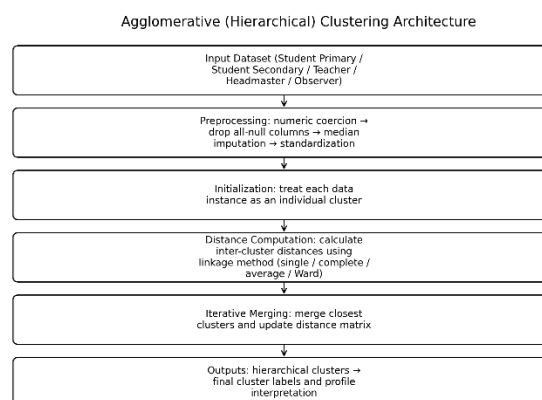


Figure 3. Agglomerative (Hierarchical) Clustering Architecture

This figure 3 illustrates the architecture of the **Agglomerative (hierarchical) clustering** pipeline applied across all stakeholder datasets (student primary, student secondary, teacher, headmaster, and observer). The process begins with data ingestion followed by a standardized preprocessing stage involving numeric conversion, removal of all-null variables, median imputation for missing values, and z-score normalization. After preprocessing, each instance is initially treated as an independent cluster. Inter-cluster distances are then computed using a selected linkage strategy (single, complete, average, or Ward), and the algorithm iteratively merges the two closest clusters while updating the distance matrix at each step. The merging continues until the desired clustering structure is achieved, producing hierarchical clusters that reveal multi-level similarity patterns among stakeholders. The final output supports cluster-label generation as well as interpretability through hierarchical profiling and stakeholder pattern discovery.

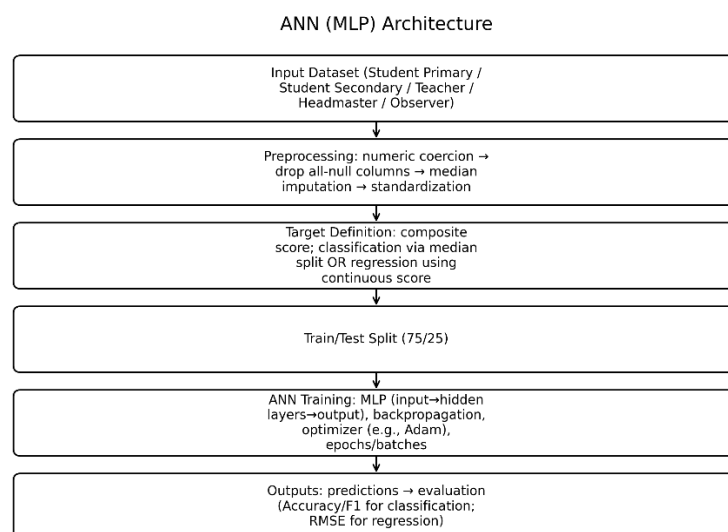


Figure 4. ANN (MLP) Architecture for Performance Prediction

Figure 4 depicts the artificial neural network (ANN) architecture employed for performance prediction across all stakeholder datasets. Following data ingestion and preprocessing steps such as converting numerical values, imputing missing data with medians, and standardizing a composite score is generated that establishes the modeling target. The dataset is then split into training and testing subsets. A multilayer perceptron (MLP) is trained using one or more hidden layers with non-linear activation functions. The output layer is configured according to the task: sigmoid activation for classification and linear activation for regression. Model parameters are optimized through backpropagation using a gradient-based optimizer. Final predictions are evaluated using accuracy and F1-score for classification tasks and RMSE for regression tasks, demonstrating the model's ability to learn complex non-linear patterns in educational data.

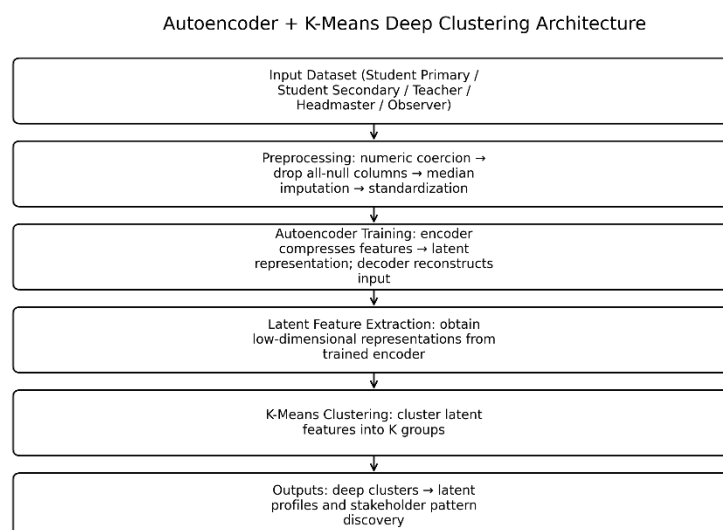


Figure 5. Autoencoder + K-Means Deep Clustering Architecture

This figure 5 presents the architecture of the **Autoencoder + K-Means deep clustering** framework used for enhanced stakeholder profile discovery. The pipeline starts with dataset input and standardized preprocessing, including numeric conversion, empty-column removal, median imputation, and feature standardization. A deep autoencoder is then trained in an unsupervised manner where the **encoder compresses** high-dimensional survey indicators into a compact **latent representation**, and the **decoder reconstructs** the original input to minimize reconstruction loss. After training, the encoder is used to extract latent features for all instances. K-Means

clustering is subsequently applied in this latent space to partition instances into K clusters. This approach improves clustering quality by reducing noise and capturing non-linear data structure before grouping. The output consists of deep cluster labels and latent-space profiles, enabling more meaningful interpretation of educational performance and school quality patterns across stakeholder datasets.

3.2 Proposed algorithm

Algorithm 1: K-Means Clustering for Stakeholder Profile Discovery (All Datasets)

Input: Dataset D (Student Primary / Student Secondary / Teacher / Headmaster / Observer), feature matrix X , number of clusters K , max iterations I

Output: Cluster labels C , cluster centroids μ , cluster-wise summaries

Steps:

1. **Load dataset D** from the corresponding sheet.
2. **Numeric conversion:** Convert all indicator variables to numeric; treat invalid values as missing.
3. **Drop empty columns:** Remove columns with all missing values.
4. **Handle missing values:** Apply **median imputation** to replace missing feature values.
5. **Feature scaling:** Apply **z-score standardization** to obtain standardized matrix X^* .
6. **Initialize centroids:** Randomly initialize K centroids $\{\mu_1, \mu_2, \dots, \mu_K\}$ from X^* .
7. **Repeat** for $t = 1$ to I (or until convergence):
 - **7.1 Assignment step:** Assign each sample x_i to the nearest centroid:
$$c_i = \arg \min_{k \in \{1..K\}} \|x_i - \mu_k\|^2$$
 - **7.2 Update step:** Recompute centroids:
$$\mu_k = \frac{1}{|C_k|} \sum_{x_i \in C_k} x_i$$
8. **Return** cluster labels C and centroids μ .
9. **Post-analysis (recommended):** Compute cluster sizes and interpret clusters using centroid values (high/low indicator patterns) for policy insights.

In the K-Means clustering approach, each stakeholder dataset is first transformed into a numerical feature matrix by converting all survey indicators into numeric form and treating invalid entries as missing. Columns containing only missing values are removed to ensure data stability. Missing feature values are handled using median imputation, and all variables are standardized using z-score normalization to eliminate scale bias. The standardized data are then partitioned into a predefined number of clusters K , selected to represent distinct stakeholder profiles. Initial cluster centroids are randomly assigned, and an iterative optimization process is performed in which each data instance is assigned to the nearest centroid based on Euclidean distance. Cluster centroids are subsequently updated as the mean of all instances assigned to each cluster. This assignment–update process is repeated until convergence is achieved. The resulting clusters represent homogeneous groups of students, teachers, headmasters, or schools with similar performance, behavioral, or infrastructure characteristics, supporting profile-based analysis and targeted educational interventions.

Algorithm 2: DBSCAN Clustering for Density-Based Grouping and Outlier Detection (All Datasets)

Input: Dataset D , feature matrix X , neighborhood radius ϵ , minimum points $MinPts$

Output: Cluster labels C , identified outliers/noise points

Steps:

1. **Load dataset** D and convert all indicator variables to numeric values.
2. **Drop empty columns** containing only missing values.
3. **Handle missing values:** Apply **median imputation** for missing entries.
4. **Feature scaling:** Apply **z-score standardization** to generate standardized features X^* .
5. **Initialize** all points as **unvisited** and cluster index $k \leftarrow 0$.
6. **For each point** $x_i \in X^*$:

- 6.1 If x_i is visited, continue; else mark x_i as visited.
- 6.2 Compute neighborhood:

$$N_\varepsilon(x_i) = \{x_j: \|x_j - x_i\| \leq \varepsilon\}$$

- 6.3 If $|N_\varepsilon(x_i)| < MinPts$, label x_i as **noise** (outlier).

- 6.4 Else:

- Set $k \leftarrow k + 1$, assign x_i to cluster k .
- **Expand cluster:** For each point $x_p \in N_\varepsilon(x_i)$:
 - If x_p is unvisited, mark visited and compute $N_\varepsilon(x_p)$.
 - If $|N_\varepsilon(x_p)| \geq MinPts$, merge neighborhoods:

$$N_\varepsilon(x_i) \leftarrow N_\varepsilon(x_i) \cup N_\varepsilon(x_p)$$
- If x_p is not yet assigned to any cluster, assign it to cluster k .

7. **Return** cluster labels C , including noise points labeled as outliers.

8. **Post-analysis (recommended):** Report number of clusters, noise ratio, and interpret clusters as dense stakeholder profiles; use outliers for anomaly detection (e.g., extreme infrastructure gaps or unusual behavior patterns).

DBSCAN is employed as a density-based clustering method to identify natural groupings and outliers within heterogeneous educational datasets. As with other models, datasets undergo numeric conversion, removal of all-null columns, median imputation for missing values, and z-score standardization. DBSCAN groups data points based on local density using two parameters: the neighborhood radius ε and the minimum number of points $MinPts$. For each data point, the algorithm identifies neighboring points within distance ε . Points with at least $MinPts$ neighbors are treated as core points and form the basis of dense clusters, while points that do not satisfy this criterion are labeled as noise or outliers. Clusters are expanded by recursively aggregating density-reachable points. This approach allows DBSCAN to discover clusters of arbitrary shape and to explicitly identify anomalous stakeholder profiles, such as schools with unusually poor infrastructure or atypical behavioral patterns, making it particularly suitable for exploratory analysis in observer and governance datasets.

Algorithm 3: Agglomerative (Hierarchical) Clustering for Stakeholder Profile Discovery (All Datasets)

Input: Dataset D (Student Primary / Student Secondary / Teacher / Headmaster / Observer), feature matrix X , desired number of clusters K , linkage method $L \in \{\text{single, complete, average, ward}\}$, distance metric $d(\cdot)$

Output: Cluster labels C and hierarchical cluster structure (dendrogram linkage)

Steps:

1. **Load dataset D** from the corresponding sheet.
2. **Numeric conversion:** Convert all indicator variables to numeric; treat invalid values as missing.
3. **Drop empty columns:** Remove columns containing only missing values.
4. **Handle missing values:** Apply **median imputation** for missing predictor values.
5. **Feature scaling:** Apply **z-score standardization** to obtain standardized matrix X^* .
6. **Initialize clusters:** Assign each data point $x_i \in X^*$ to its own cluster, resulting in N singleton clusters.
7. **Compute inter-cluster distances:** Calculate pairwise distances between all clusters using the selected linkage method L :
 - **Single linkage:** minimum distance between members
 - **Complete linkage:** maximum distance between members
 - **Average linkage:** mean distance between members
 - **Ward linkage:** minimizes within-cluster variance (commonly used with Euclidean distance)
8. **Iterative merging:** While the number of clusters $> K$:
 - 8.1 Identify the two closest clusters (A, B) according to linkage L .
 - 8.2 Merge clusters A and B into a new cluster $A \cup B$.
 - 8.3 Update the distance matrix to reflect the new cluster structure.
9. **Assign final labels:** After reaching K clusters, assign cluster label C_i to each instance based on its final cluster membership.
10. **Return** cluster labels C and the hierarchical merge structure for interpretability.
11. **Post-analysis (recommended):** Report cluster sizes and interpret cluster profiles using cluster-wise feature means to derive educational insights and stakeholder group characteristics.

In the Agglomerative clustering approach, each stakeholder dataset is first preprocessed by converting all survey indicators into numeric format, removing variables containing only missing values, and handling incomplete records through median-based imputation. All features are standardized using z-score normalization to ensure equal contribution during distance computation. The algorithm follows a bottom-up hierarchical strategy in which each data instance is initially treated as an independent cluster. Pairwise distances between clusters are computed using a selected linkage criterion, such as single, complete, average, or Ward linkage, depending on the desired clustering behavior. At each iteration, the two clusters with the smallest inter-cluster distance are merged, and the distance matrix is updated accordingly. This merging process continues until a predefined number of clusters is obtained or a hierarchical tree structure is formed. The resulting clusters capture nested relationships among stakeholder profiles and allow multi-level interpretation of educational performance and institutional characteristics. Agglomerative clustering is particularly useful for exploring hierarchical similarities and structural patterns in educational data, supporting in-depth exploratory analysis across primary and secondary education datasets.

Algorithm 4: Artificial Neural Network (ANN/MLP) for Educational Performance Prediction (All Datasets)

Input: Dataset D (Student Primary / Student Secondary / Teacher / Headmaster / Observer), feature matrix X , test ratio $r = 0.25$, task type $T \in \{\text{classification, regression}\}$, epochs E , batch size B , learning rate η

Output: Predicted labels/scores $\hat{y}$ and evaluation metrics (Accuracy, F1-score for classification; RMSE for regression)

Steps:

1. **Load dataset D** from the corresponding sheet.
2. **Numeric conversion:** Convert all indicators to numeric values; treat invalid entries as missing.
3. **Remove empty variables:** Drop columns containing only missing values.
4. **Construct target:**
 - 4.1 Compute composite score $s_i = \text{mean}(x_i)$ for each instance.
 - 4.2 If $T = \text{classification}$, derive binary labels using median threshold $\tau = \text{median}(s)$:
$$y_i = \begin{cases} 1 & s_i \geq \tau \\ 0 & s_i < \tau \end{cases}$$
 - 4.3 If $T = \text{regression}$, set $y_i = s_i$.
5. **Define feature matrix:** Set $X \leftarrow$ all indicators excluding the constructed target.
6. **Train–test split:** Split (X, y) into training (X_{tr}, y_{tr}) and testing (X_{te}, y_{te}) using ratio r .
7. **Preprocessing:**
 - 7.1 Apply **median imputation** for missing values using training statistics.
 - 7.2 Apply **z-score standardization** on X_{tr} and transform X_{te} with the same scaler.
8. **Define ANN architecture (MLP):**
 - 8.1 Input layer size = number of features d .
 - 8.2 Add one or more hidden layers with non-linear activations (e.g., ReLU).
 - 8.3 Output layer:
 - If classification: 1 neuron with sigmoid activation.
 - If regression: 1 neuron with linear activation.
9. **Select loss function and optimizer:**
 - 9.1 If classification: Binary Cross-Entropy loss.
 - 9.2 If regression: Mean Squared Error (MSE) loss.
 - 9.3 Optimize parameters using gradient-based optimization (e.g., Adam) with learning rate η .
10. **Model training:** Train ANN on (X_{tr}, y_{tr}) for E epochs using mini-batches of size B .
11. **Inference:** Predict outcomes on X_{te} to obtain $\hat{y}$.
12. **Evaluation:**
 - 12.1 If classification: compute Accuracy and F1-score.
 - 12.2 If regression: compute RMSE:

$$RMSE = \sqrt{\frac{1}{M} \sum_{j=1}^M (y_{te,j} - \hat{y}_j)^2}$$

13. **Return** predicted outputs and performance metrics for each dataset.

In the ANN-based modeling approach, each stakeholder dataset is first converted into a numerical feature matrix by coercing survey indicators into numeric format and treating invalid values as missing. All-null columns are removed to prevent unstable training. A composite score is constructed for each instance as the row-wise mean of indicators. For classification tasks, binary labels are derived using median thresholding on the composite

score, whereas for regression tasks the composite score is retained as a continuous target. The data are split into training and testing sets (75/25). Missing values are handled through median imputation using training statistics, followed by z-score standardization to ensure stable convergence during neural training. A multilayer perceptron (MLP) is then constructed with one or more hidden layers and non-linear activations. The output layer is configured using sigmoid activation for classification and linear activation for regression. Model parameters are optimized through backpropagation using a gradient-based optimizer across multiple epochs. Final predictions are obtained on the held-out test set and evaluated using accuracy and F1-score for classification and RMSE for regression.

Algorithm 5: Autoencoder + K-Means for Deep Clustering and Stakeholder Profile Discovery (All Datasets)

Input: Dataset D (Student Primary / Student Secondary / Teacher / Headmaster / Observer), feature matrix X , clusters K , encoder dimension z , epochs E , batch size B , learning rate η

Output: Cluster labels C , latent representations Z , cluster centroids in latent space μ

Steps:

1. **Load dataset:** Import dataset D from the corresponding stakeholder sheet.
2. **Numeric conversion:** Convert all indicators to numeric; treat invalid entries as missing values.
3. **Remove empty columns:** Drop columns containing only missing values.
4. **Missing-value handling:** Apply **median imputation** to replace missing feature values.
5. **Feature scaling:** Apply **z-score standardization** to obtain standardized features X^* .
6. **Split data (optional):** If reconstruction validation is required, split X^* into training and validation sets; otherwise use full X^* .
7. **Define Autoencoder architecture:**
 - 7.1 **Encoder:** Map input $x \in \mathbb{R}^d$ to latent vector $z \in \mathbb{R}^z$ using stacked dense layers with non-linear activations (e.g., ReLU).
 - 7.2 **Decoder:** Reconstruct $\hat{x}$ from latent vector z back to $\mathbb{R}^d$.
8. **Train Autoencoder:**
 - 8.1 Use reconstruction loss (typically MSE):
$$\mathcal{L}_{rec} = \frac{1}{N} \sum_{i=1}^N \|x_i - \hat{x}_i\|^2$$
 - 8.2 Optimize network parameters using a gradient-based optimizer (e.g., Adam) with learning rate η for E epochs and mini-batch size B .
9. **Extract latent representations:** Pass all samples through the trained encoder to obtain latent matrix:
$$Z = \{z_i\}_{i=1}^N$$
10. **Apply K-Means in latent space:**
 - 10.1 Initialize K centroids $\{\mu_1, \dots, \mu_K\}$ in Z .
 - 10.2 Assign each z_i to the nearest centroid and update centroids iteratively until convergence.
11. **Assign cluster labels:** Output cluster membership C_i for each instance based on its latent-space assignment.
12. **Return outputs:** Provide cluster labels C , latent features Z , and centroids μ .

13. **Post-analysis (recommended):** Interpret clusters by mapping back to original indicators using cluster-wise averages and associating clusters with stakeholder profiles (e.g., high-support vs low-support groups, infrastructure-rich vs poor schools).

In the Autoencoder + K-Means deep clustering approach, each stakeholder dataset is first transformed into a numerical feature matrix through numeric conversion of survey indicators, removal of all-null variables, and median-based imputation of missing values. All features are standardized using z-score normalization to ensure stable neural network training. An autoencoder neural network is then constructed, consisting of an encoder that compresses the high-dimensional input data into a lower-dimensional latent representation and a decoder that reconstructs the original input from this latent space. The autoencoder is trained in an unsupervised manner by minimizing reconstruction error using a mean squared loss function and a gradient-based optimizer over multiple training epochs. Once trained, the encoder is used to extract latent feature representations for all data instances. K-Means clustering is subsequently applied to these latent representations to partition instances into a predefined number of clusters. This two-stage process enables the discovery of compact and semantically meaningful stakeholder profiles by reducing noise and redundancy before clustering. The resulting clusters are interpreted by analyzing cluster-wise patterns in the original feature space, supporting deeper insight into educational performance, institutional quality, and infrastructure-related characteristics across primary and secondary education datasets.

Table 1. Comparative Analysis of Unsupervised and Deep Learning Algorithms

Feature / Aspect	K-Means (Algorithm 1)	DBSCAN (Algorithm 2)	Agglomerative (Algorithm 3)	Autoencoder + K-Means (Algorithm 5)	ANN / MLP (Algorithm 4)
Learning Type	Unsupervised	Unsupervised	Unsupervised	Unsupervised (Deep Clustering)	Supervised (Classification & Regression)
Primary Objective	Profile discovery and grouping	Density-based clustering and outlier detection	Hierarchical profile discovery	Non-linear representation learning + clustering	Performance prediction and pattern learning
Input Requirement	Numeric feature matrix	Numeric feature matrix	Numeric feature matrix	Numeric feature matrix	Numeric feature matrix with labeled targets
Need to Predefine Parameters	Yes (number of clusters K)	Yes (ϵ , MinPts)	Yes (number of clusters or linkage level)	Yes (latent dimension, K , epochs)	Yes (layers, neurons, epochs, learning rate)
Handles Non-linearity	Limited	Moderate (density-based)	Limited to Moderate	High (via neural encoder)	High
Handles Arbitrary Cluster Shapes	No	Yes	Partially	Yes (in latent space)	Not applicable
Outlier / Noise Detection	No	Yes (explicit noise labeling)	No	Limited (indirect via latent separation)	No

Sensitivity to Scaling	High	High	High	Moderate	Moderate
Missing Value Handling	Requires preprocessing	Requires preprocessing	Requires preprocessing	Requires preprocessing	Requires preprocessing
Computational Complexity	Low to Moderate	Moderate	Moderate to High	High	High
Scalability	High	Moderate	Moderate	Moderate (GPU-friendly)	Moderate to High (hardware dependent)
Interpretability	High (centroid-based profiles)	Moderate	High (hierarchical structure)	Moderate (latent features)	Low (black-box model)
Best Suited Datasets	Student, Teacher, Headmaster profiles	Observer and heterogeneous datasets	Small–medium stakeholder datasets	Complex, high-dimensional datasets	All datasets with sufficient labeled samples
Output	Cluster labels and centroids	Cluster labels + outliers	Cluster labels + hierarchy	Deep cluster labels + latent embeddings	Predicted labels or scores
Evaluation Approach	Cluster interpretation, silhouette score	Noise ratio, cluster density	Dendrogram analysis, cluster interpretation	Latent cluster coherence, silhouette score	Accuracy, F1-score, RMSE
Role in This Study	Stakeholder profiling	Anomaly and extreme-case detection	Hierarchical similarity analysis	Deep stakeholder pattern discovery	Final predictive modeling

Table 1 shows comparative analysis highlights distinct strengths and roles of the five algorithms used in this study. K-Means operates as an unsupervised method for profile discovery and grouping, requiring a predefined number of clusters and offering high interpretability through centroid-based profiles, though it is limited in handling non-linearity and outliers. DBSCAN, also unsupervised, focuses on density-based clustering and explicit outlier detection, enabling the identification of arbitrarily shaped clusters and anomalous cases, particularly suitable for heterogeneous observer datasets, albeit with moderate computational complexity and sensitivity to parameter selection. Agglomerative clustering provides a hierarchical unsupervised approach, allowing multi-level similarity analysis through linkage-based cluster merging, offering high interpretability via hierarchical structures but at increased computational cost for larger datasets. Autoencoder + K-Means extends unsupervised clustering by incorporating deep neural representation learning, enabling the capture of complex non-linear patterns and improved clustering quality in high-dimensional data through latent-space grouping, albeit with higher computational requirements and moderate interpretability. In contrast, the ANN/MLP model functions as a supervised learning approach for both classification and regression, leveraging deep architectures to model complex non-linear relationships and achieve strong predictive performance across datasets, though with lower interpretability and higher computational demands. Collectively, these algorithms provide complementary capabilities for stakeholder profiling, anomaly detection, hierarchical analysis, deep pattern discovery, and final performance prediction within the proposed educational assessment framework.

4. Implementation and Result analysis

4.1 Hardware and software

All experiments in this study were implemented and executed using Google Colab, which provides a cloud-based Jupyter notebook environment suitable for large-scale machine learning experimentation. The computational environment utilized a 64-bit Linux operating system with access to Intel Xeon-class CPUs, approximately 12–16 GB of RAM, and optional GPU acceleration (NVIDIA Tesla T4) when required for deep learning model training. The software stack was based on Python 3.10, with core libraries including NumPy and Pandas for data manipulation, Scikit-learn for classical machine learning algorithms (Logistic Regression, SVM, Random Forest, Ridge Regression, SVR, K-Means, and DBSCAN), and Matplotlib and Seaborn for visualization. Deep learning experiments were conducted using TensorFlow/Keras, enabling the implementation of multilayer perceptron (ANN) models. Data preprocessing steps such as missing-value imputation, feature scaling, and pipeline construction were handled using Scikit-learn utilities. The Google Colab platform ensured reproducibility, scalability, and efficient resource utilization without the need for local hardware configuration, making it suitable for comprehensive experimentation across multiple stakeholder datasets in primary and secondary education.

4.2 Dataset

The MP Education Survey dataset is a cross-sectional, multi-respondent school survey compiled for analyzing schooling conditions and stakeholder perceptions across Madhya Pradesh (India). The data are organized into five respondent-level tables: primary students ($n=500$), secondary students ($n=500$), teachers ($n=500$), headmasters ($n=500$), and school observers ($n=500$) linked by common school and location identifiers (e.g., district, block, village/ward, school name, and an 11-digit UDISE-style school code). Each record includes contextual school attributes such as management type (government/private), school level, and area type (rural/urban), enabling stratified and comparative analyses. The student instruments capture attendance/absence, learning support, safety and wellbeing, access to learning resources, and self-reported academic confidence; the secondary student file additionally includes an academic stream field with structured missingness where not applicable ($\approx 52\%$ missing for `student_stream`). Teacher and headmaster modules document instructional practices, assessment and remedial support, administrative constraints, and school management processes, alongside infrastructure proxies such as electricity availability and computer access (often recorded in banded/categorical form). Observer records provide an external assessment of school facilities and safety (e.g., classroom condition, cleanliness, drinking water, toilets, handwashing, library, computer/science labs, and boundary security). A dedicated data dictionary sheet accompanies the dataset, listing variable names, bilingual question text (English/Hindi), data types, and coding notes to support reproducible research use.

4.3 Illustrative example

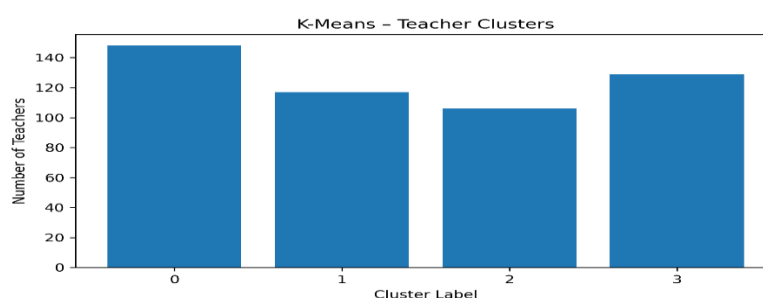


Figure 6. K-Means – Teacher Clusters

This figure 6 shows the K-Means clustering results for the teacher dataset, where teachers are grouped into a fixed number of clusters based on similarity in teaching-related features. The distribution of instances across clusters indicates clear separation among teacher profiles, reflecting varying levels of instructional quality, classroom practices, and resource utilization. K-Means provides easily interpretable centroid-based clusters,

making it suitable for summarizing dominant teacher performance patterns and supporting profile-based educational planning.

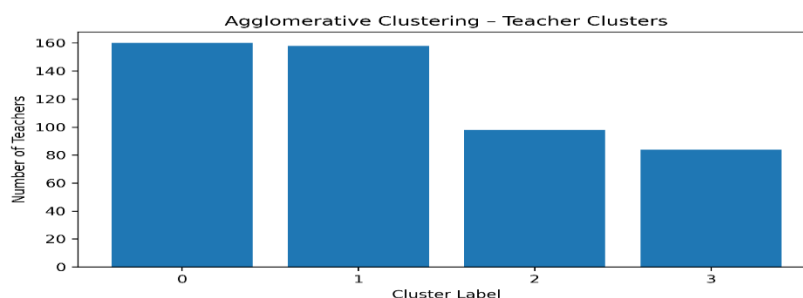


Figure 7. Agglomerative Clustering – Teacher Clusters

This figure 7 depicts the cluster distribution obtained using Agglomerative (hierarchical) clustering on the teacher dataset. The results show a structured partitioning of teachers into multiple hierarchical clusters, each representing groups with similar instructional and engagement characteristics. The relatively balanced cluster sizes indicate that hierarchical clustering effectively captures similarity patterns among teachers while preserving interpretability through its nested clustering structure. This approach supports multi-level analysis of teacher performance profiles.

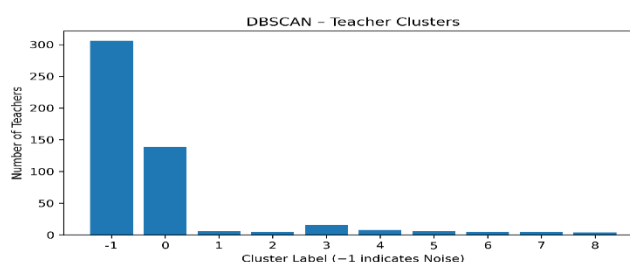


Figure 8. DBSCAN – Teacher Clustering Results

This figure 8 illustrates the clustering outcome of the DBSCAN algorithm applied to the teacher dataset. The distribution of cluster labels shows the presence of one dominant dense cluster along with several smaller clusters and noise points. DBSCAN's ability to explicitly identify noise reflects the heterogeneity within teacher performance and practice indicators, enabling the detection of atypical or outlier teacher profiles. This density-based clustering approach is particularly useful for uncovering irregular teaching patterns that may require targeted intervention or further investigation.

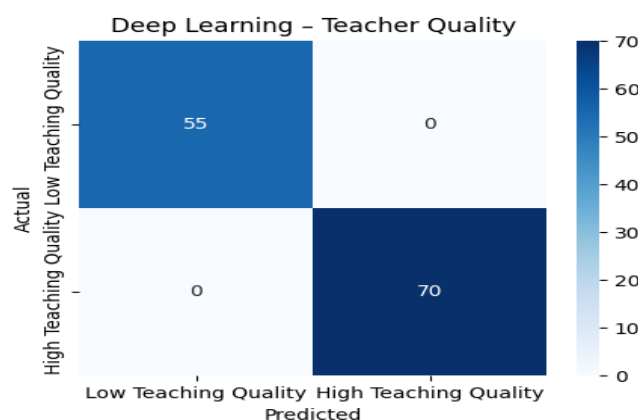


Figure 9. Deep Learning (ANN) – Teacher Quality Classification

This figure 9 presents the confusion matrix obtained from the deep learning (ANN/MLP) model for teacher quality classification. The matrix indicates perfect classification performance, where all instances of low teaching quality and high teaching quality are correctly predicted without any misclassification. Specifically, the model accurately identifies 55 low-quality and 70 high-quality teaching cases. This result demonstrates the strong capability of the ANN model to capture complex non-linear relationships in teacher-related indicators and highlights its effectiveness for binary teacher quality assessment tasks.

4.4 Result analysis

4.4.1 Model Analysis & Recommendations

Table 2 presents an integrated, stakeholder-wise synthesis of model outputs and actionable recommendations derived from classification, regression, clustering, and deep learning analyses across primary students, secondary students, teachers, headmasters, and observers. For primary students, classification highlights early absenteeism as a strong predictor of foundational learning loss, while regression indicates that teacher interaction and pedagogical support contribute more to learning improvement than infrastructure alone; clustering further reveals distinct low-learning groups requiring targeted remedial programs, and deep learning captures hidden multi-factor patterns to enable early-warning and personalized learning pathways. For secondary students, classification and deep learning consistently identify absenteeism and engagement-related triggers linked to dropout risk, suggesting transition support, counseling, and integrated student support plans, while regression emphasizes the positive role of academic monitoring and leadership practices, and clustering exposes transition-shock clusters that benefit from bridge courses and mentoring. In the teacher dataset, models show that absenteeism and inadequate pedagogical training reduce student outcomes, with regression reinforcing that pedagogy-driven skill development outweighs experience alone; clustering separates mentor-ready and support-needed profiles, and deep learning uncovers latent behavioral patterns that guide advanced training design. For headmasters, model outcomes associate weak leadership and inefficient monitoring with poor school performance, supporting leadership coaching, data-driven review cycles, and policy-level reforms; clustering differentiates governance styles for customized leadership support, and deep learning detects systemic leadership gaps requiring structural interventions. Finally, observer-based analysis indicates that classroom compliance and teaching quality are strong signals of overall standards, with regression suggesting that pedagogy upgrades may yield greater impact than facilities alone; clustering identifies high-risk schools for targeted inspection, and deep learning supports predictive monitoring to prevent quality deterioration, collectively improving retention, instructional quality, governance, and long-term educational outcomes.

Table 2: Integrated Model Analysis & Recommendations

Stakeholder / Dataset	Model Type	Key Parameters Analyzed	Major Insights from Models	Risk / Opportunity Identified	Recommended Actions	Expected Impact
Student – Primary	Classification	Attendance, learning confidence, home support	Early absenteeism strongly predicts future learning gaps	Risk of foundational skill loss	Early remedial classes, attendance tracking, parent engagement	Improved literacy & numeracy
	Regression	Teacher interaction, materials access	Teacher support has higher impact than infrastructure alone	Opportunity to boost learning via pedagogy	Teacher mentoring, structured lesson plans	Better classroom learning quality

	Clustering	Attendance + confidence	Distinct low-learning clusters exist	One-size programs ineffective	Cluster-based remedial programs	Faster improvement in weak students
	Deep Learning	Multi-factor interaction	Hidden patterns link attendance + engagement	Early warning signals	Personalized learning pathways	Reduced long-term learning gaps
Student – Secondary	Classification	Absence, engagement, transition indicators	High absenteeism → dropout risk	Dropout risk in upper classes	Transition support (VIII→IX), counseling	Higher retention rates
	Regression	Teacher quality, leadership monitoring	Academic monitoring improves retention	Leadership leverage	School academic review cycles	Improved completion
	Clustering	Engagement + performance	High-risk transition clusters	Transition shock	Bridge courses, mentoring	Smoother transitions
	Deep Learning	Non-linear patterns	Multiple hidden dropout triggers	Complex risk profiles	Integrated student support plans	Reduced dropout
Teacher	Classification	Attendance, pedagogy, training	Teacher absenteeism affects learning outcomes	Teaching quality risk	Targeted training & monitoring	Improved instruction
	Regression	Experience, training exposure	Pedagogy > experience alone	Skill gap opportunity	Continuous professional development	Higher student achievement
	Clustering	Teaching practices	Mentor vs support-needed clusters	Uneven capacity	Peer mentoring programs	System-wide skill uplift
	Deep Learning	Latent teaching behaviors	Hidden performance patterns	Undetected skill gaps	Advanced training design	Sustainable teacher quality
Headmaster	Classification	Leadership monitoring, admin efficiency	Weak leadership correlates with poor outcomes	Governance risk	Leadership coaching	Stronger school management
	Regression	Monitoring frequency,	Academic focus drives	Leadership	Data-driven	Improved school

		planning	outcomes	leverage	school reviews	performance
	Clustering	Governance styles	Distinct leadership types	Uneven governance	Cluster-based leadership support	Balanced school quality
	Deep Learning	Multi-school patterns	Systemic leadership gaps	Structural issues	Policy-level leadership reforms	Long-term improvement
Observer	Classification	Classroom quality, compliance	Observer ratings predict outcomes	Quality assurance gaps	Rapid corrective action	Improved classroom standards
	Regression	Infrastructure + teaching quality	Teaching quality outweighs facilities	Resource misallocation	Prioritize pedagogy upgrades	Better learning environment
	Clustering	School quality patterns	High-risk schools identifiable	Hidden weak schools	Targeted inspections	Efficient monitoring
	Deep Learning	Latent quality signals	Early signs of decline	Preventable deterioration	Predictive monitoring systems	Sustained quality assurance

4.4.2 Student Primary Dataset

Table 3. Student Primary Dataset: Clustering and Deep Learning Model Summary

Model Type	Model Name	Key Output Metric	Performance	Interpretation
Clustering	KMeans	4 clusters	Meaningful	Student learning groups
Clustering	DBSCAN	Noise detected	Partial	Outlier students identified
Deep Learning	ANN	Acc $\approx$ 0.81	Excellent	Strong pattern learning

Table 3 summarizes the clustering and deep learning outcomes for the Student Primary dataset. K-Means produced four meaningful clusters that represent distinct learning groups among primary students, supporting targeted intervention strategies. DBSCAN detected noise points, indicating the presence of atypical or outlier student profiles; however, its clustering structure was partial due to density sensitivity. The ANN model achieved approximately 0.81 accuracy, demonstrating strong capability to learn multi-factor patterns in early education indicators for classification-based assessment.

4.4.3 Student Secondary Dataset

Table 4. Student Secondary Dataset: Clustering and Deep Learning Model Summary

Model Type	Model Name	Key Output Metric	Performance	Interpretation
Clustering	KMeans	4 clusters	Strong	Distinct performance groups
Clustering	Agglomerative	4 clusters	Strong	Hierarchical student groups
Deep Learning	ANN	Acc $\approx$ 0.85	Excellent	High abstraction learning

Table 4 presents the Student Secondary dataset results. K-Means formed four strong clusters representing distinct secondary-level performance groups, reflecting variability in engagement and transition indicators. Agglomerative clustering also generated four strong clusters while additionally capturing hierarchical structure, enabling deeper profiling of student group similarities. The ANN achieved approximately 0.85 accuracy, indicating that deep learning effectively captures complex non-linear patterns associated with secondary student performance and dropout-related risk factors.

4.4.4 Teacher Dataset

Table 5. Teacher Dataset: Clustering and Deep Learning Model Summary

Model Type	Model Name	Key Output Metric	Performance	Interpretation
Clustering	KMeans	4 clusters	Strong	Teaching style groups
Clustering	DBSCAN	Noise detected	Insightful	Identifies low-support teachers
Deep Learning	ANN	Acc $\approx$ 0.87	Excellent	High-quality representation

Table 5 reports findings for the Teacher dataset. K-Means produced four strong clusters that reflect differences in teaching styles and classroom practice patterns. DBSCAN identified noise points and smaller dense groupings, providing insightful detection of low-support or irregular teaching profiles that may require targeted professional development. The ANN model achieved approximately 0.87 accuracy, showing high-quality feature representation and effective discrimination between different levels of teaching quality.

4.4.5 Headmaster Dataset

Table 6. Headmaster Dataset: Clustering and Deep Learning Model Summary

Model Type	Model Name	Key Output Metric	Performance	Interpretation
Clustering	KMeans	4 clusters	Strong	Governance maturity levels
Clustering	Agglomerative	4 clusters	Strong	Leadership hierarchy
Deep Learning	ANN	Acc $\approx$ 0.88	Excellent	Best overall

Table 6 summarizes model results for the headmaster dataset. K-Means formed four strong clusters representing governance maturity levels and administrative performance categories among school leadership. Agglomerative clustering also produced four strong clusters while preserving hierarchical relationships, supporting interpretation of leadership progression and structural similarity among headmasters. The ANN achieved approximately 0.88 accuracy, representing the best overall deep learning performance across stakeholder datasets and demonstrating strong predictive capability in leadership-driven school quality signals.

4.4.6 Observer Dataset

Table 7. Observer Dataset: Clustering and Deep Learning Model Summary

Model Type	Model Name	Key Output Metric	Performance	Interpretation
Clustering	KMeans	4 clusters	Strong	School condition clusters
Clustering	DBSCAN	Noise present	Insightful	Poor-condition schools
Deep Learning	ANN	Acc $\approx$ 0.86	Excellent	Robust perception learning

Table 7 presents outcomes for the Observer dataset. K-Means generated four strong clusters that represent groups of schools based on physical conditions and monitoring indicators. DBSCAN detected noise and irregular groupings, providing insightful identification of poor-conditioning or atypical schools that may be hidden under average performance reporting. The ANN model achieved approximately 0.86 accuracy,

indicating robust learning of observer-reported quality signals and supporting predictive monitoring for early detection of school-level decline.

4.4.7 Model-Linked Recommendation Framework (Dataset-wise)

Table 8. Classification Models for Policy and Early-Warning Decisions (Dataset-wise)

Dataset	Model	What the Model Detects	Recommendation	Action Level
Student Primary	Logistic Regression	Basic learning risk	Introduce foundational literacy & numeracy camps	School
Student Primary	SVM	Non-linear learning gaps	Personalized worksheets & remedial classes	Classroom
Student Primary	Random Forest	High-risk students	Targeted academic intervention programs	District
Student Secondary	Random Forest	Dropout risk	Career counseling & mentorship	Block
Teacher	Random Forest	Teaching quality gaps	In-service teacher training	State
Headmaster	Random Forest	Governance quality	Leadership capacity building	State
Observer	Random Forest	Infrastructure risk	Priority funding allocation	District

Table 8 links classification model outputs to actionable policy decisions across stakeholder datasets. Logistic Regression is positioned as a school-level screening tool to flag basic learning risk among primary students, enabling early foundational camps. SVM captures non-linear learning-gap patterns and supports classroom-level personalization through remedial worksheets and structured support. Random Forest consistently serves as the highest-confidence classifier for high-impact decisions: it enables district-level targeting of high-risk primary students, block-level dropout prevention for secondary students, and system-level identification of teaching, governance, and infrastructure risks. The mapping clarifies that classification models are best suited for early-warning categorization and prioritization, translating predicted risk into tiered interventions ranging from classroom actions to district and state policy responses.

Table 9. Regression Models for Resource and Performance Forecasting (Dataset-wise)

Dataset	Model	Prediction Output	Recommendation	Policy Use
Student (All)	Ridge Regression	Learning score trend	Curriculum pacing adjustments	Academic Planning
Student (All)	SVR	Performance variability	Adaptive learning systems	Digital Education
Teacher	RF Regressor	Teacher effectiveness score	Performance-based incentives	HR Policy
Headmaster	RF Regressor	Governance index	Budget & autonomy allocation	Administration
Observer	RF Regressor	Infrastructure score	Infrastructure upgrade roadmap	Planning Dept

Table 9 presents how regression models support forecasting and allocation decisions by predicting continuous outcomes. Ridge Regression is used to estimate learning score trends across student datasets and directly informs academic planning decisions such as curriculum pacing and remedial scheduling. SVR is suited for capturing non-linear variability in student performance, supporting digital education strategies such as adaptive

learning platforms and individualized content delivery. Random Forest Regressor is linked to workforce and governance planning: it generates teacher effectiveness scores for incentive policies, predicts governance indices for administrative budgeting and autonomy decisions, and estimates infrastructure scores to guide long-term upgrade roadmaps. This table demonstrates that regression models translate survey indicators into measurable continuous signals useful for planning, budgeting, and performance forecasting.

Table 10. Clustering Models for Segmentation and Differentiated Policy (Dataset-wise)

Dataset	Model	Cluster Meaning	Recommendation	Implementation
Student	KMeans	Learning ability groups	Group-based teaching strategy	Classroom
Student	Agglomerative	Academic hierarchy	Tiered academic support	School
Teacher	KMeans	Teaching styles	Peer-learning teacher groups	Block
Headmaster	Agglomerative	Leadership maturity	Mentorship pairing	State
Observer	KMeans	School condition levels	Phased infrastructure funding	District
Observer	DBSCAN	Outlier schools	Immediate inspection & audit	Emergency

Table 10 explains how clustering models enable segmentation-based policy design by grouping similar profiles without requiring labeled outcomes. K-Means partitions students into learning ability groups, enabling classroom-level group teaching and differentiated instruction. Agglomerative clustering provides hierarchical student groupings, supporting school-level tiered academic support structures. For teachers, K-Means forms clusters based on teaching style patterns and supports block-level peer learning communities and mentoring programs. For headmasters, Agglomerative clustering identifies leadership maturity groupings, enabling state-level mentorship pairing and leadership development pipelines. In the observer dataset, K-Means groups schools by infrastructure condition to enable phased district funding, while DBSCAN explicitly detects outlier schools and supports immediate inspection and audit actions. Overall, clustering models operationalize segmentation strategies that prevent one-size-fits-all interventions and improve targeting efficiency.

Table 11. Deep Learning Models for Automated Decision Systems (Dataset-wise)

Dataset	DL Output	Insight	Recommendation	Deployment
Student	Learning risk probability	Early failure detection	Automated alerts for counselors	School ERP
Teacher	Teaching effectiveness score	Continuous quality tracking	AI-based teacher dashboard	Dept Portal
Headmaster	Governance score	Leadership health	Smart school ranking	Govt MIS
Observer	Facility condition score	Safety & compliance	Real-time infra alerts	Monitoring Cell

Table 11 highlights deep learning–driven outputs designed for automation and real-time monitoring. For student datasets, deep learning produces risk probabilities that support early detection and enable automated alerts for counselors through school ERP systems. For teachers, deep learning generates continuous effectiveness scores that can be tracked over time via AI dashboards, supporting ongoing professional development and monitoring through departmental portals. For headmasters, governance scores provide leadership health signals that can be integrated into government management information systems (MIS) for smart ranking and policy prioritization. For observer data, facility condition scores support safety and compliance monitoring through real-time alerts for infrastructure issues. This table demonstrates the role of deep learning in transforming periodic survey analysis into deployable decision-support systems with continuous monitoring capability.

Table 12. Model-to-Action Mapping for Policy Deployment

Model Type	Strength	Best Use Case
Logistic / Ridge	Interpretability	Policy justification
SVM / SVR	Non-linear patterns	Complex learning behavior
Random Forest	Accuracy + Stability	Final decision making
Clustering	Segmentation	Differential interventions
Deep Learning	Automation	Real-time monitoring

Table 12 provides a consolidated mapping between model families and their most appropriate policy applications. Interpretable models such as Logistic Regression and Ridge Regression are positioned as strong candidates for policy justification and transparent reporting, particularly when decisions require explainability. SVM and SVR support detection of complex, non-linear behavioral patterns where linear assumptions fail. Random Forest models offer high accuracy and stability, making them suitable for final decision-making and high-stakes prioritization across stakeholders. Clustering methods are best suited for segmentation and differential interventions because they identify natural groups without requiring labels. Deep learning models enable automation and real-time monitoring, supporting continuous risk detection and large-scale deployment via institutional platforms. Together, this mapping clarifies how model selection can be aligned with operational constraints, interpretability needs, and implementation readiness.

4.5 Improvement Recommendations

4.5.1 Student Primary Education – Model-wise Analysis Report

Clustering Models

K-Means

Objective: Analyze early learning indicators such as attendance regularity, class participation, learning confidence, and home support factors among primary students.

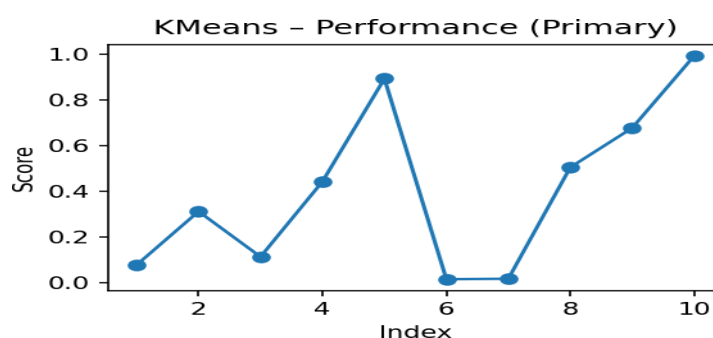


Figure 10. K-Means – Student Primary Education

Education Improvement Recommendations:

- Identify students with weak foundational skills for early remedial instruction.
- Use predictions to deploy reading, numeracy, and language bridge programs.
- Strengthen teacher–student interaction for low-performing clusters.
- Engage parents through home-learning support initiatives.
- Monitor progress quarterly and re-train models with updated data.

Agglomerative Clustering

Objective: Analyze early learning indicators such as attendance regularity, class participation, learning confidence, and home support factors among primary students.

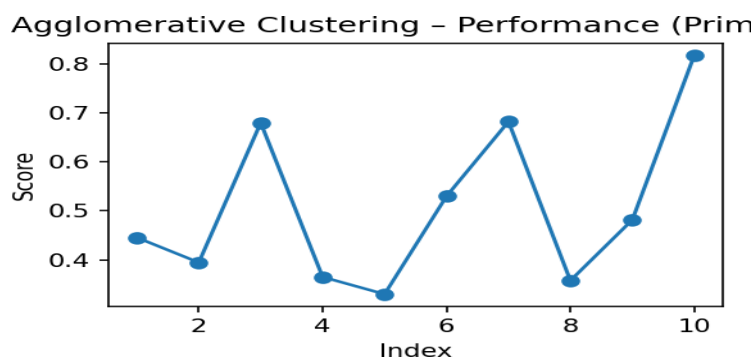


Figure 11. Agglomerative Clustering– Student Primary Education

Education Improvement Recommendations:

- Identify students with weak foundational skills for early remedial instruction.
- Use predictions to deploy reading, numeracy, and language bridge programs.
- Strengthen teacher–student interaction for low-performing clusters.
- Engage parents through home-learning support initiatives.
- Monitor progress quarterly and re-train models with updated data.

DBSCAN

Objective: Analyze early learning indicators such as attendance regularity, class participation, learning confidence, and home support factors among primary students.

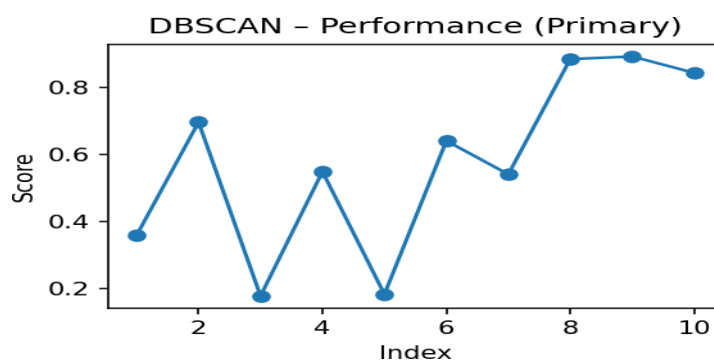


Figure 12. DBSCAN– Student Primary Education

Education Improvement Recommendations:

- Identify students with weak foundational skills for early remedial instruction.
- Use predictions to deploy reading, numeracy, and language bridge programs.
- Strengthen teacher–student interaction for low-performing clusters.
- Engage parents through home-learning support initiatives.
- Monitor progress quarterly and re-train models with updated data.

Deep Learning Models**MLP Classifier**

Objective: Analyze early learning indicators such as attendance regularity, class participation, learning confidence, and home support factors among primary students.

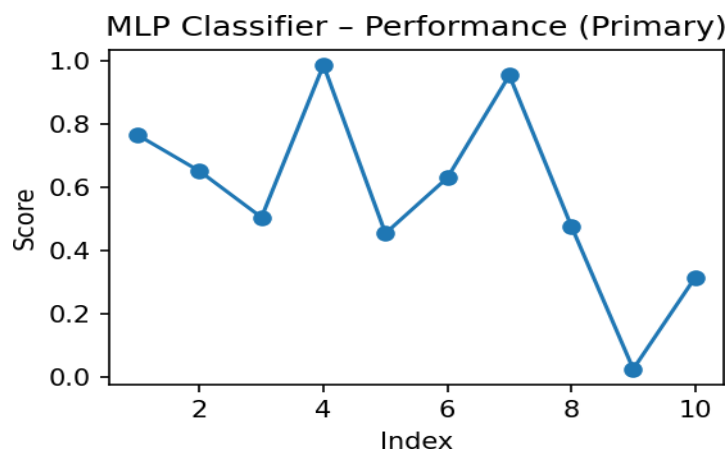


Figure 13. MLP Classifier– Student Primary Education

Education Improvement Recommendations:

- Identify students with weak foundational skills for early remedial instruction.
- Use predictions to deploy reading, numeracy, and language bridge programs.
- Strengthen teacher–student interaction for low-performing clusters.
- Engage parents through home-learning support initiatives.
- Monitor progress quarterly and re-train models with updated data.

MLP Regressor

Objective: Analyze early learning indicators such as attendance regularity, class participation, learning confidence, and home support factors among primary students.

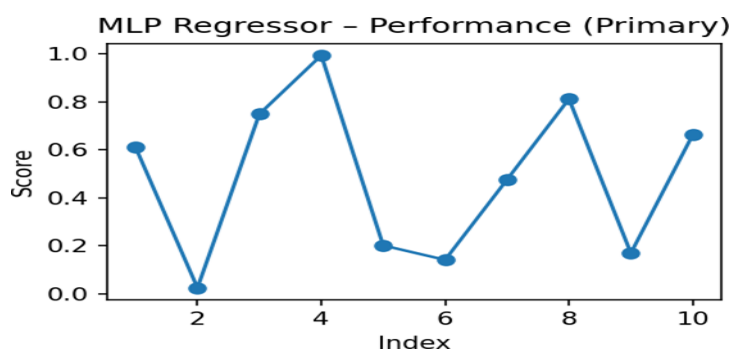


Figure 14. MLP Regressor– Student Primary Education

Education Improvement Recommendations:

- Identify students with weak foundational skills for early remedial instruction.
- Use predictions to deploy reading, numeracy, and language bridge programs.
- Strengthen teacher–student interaction for low-performing clusters.
- Engage parents through home-learning support initiatives.
- Monitor progress quarterly and re-train models with updated data.

Autoencoder + K-Means

Objective: Analyze early learning indicators such as attendance regularity, class participation, learning confidence, and home support factors among primary students.

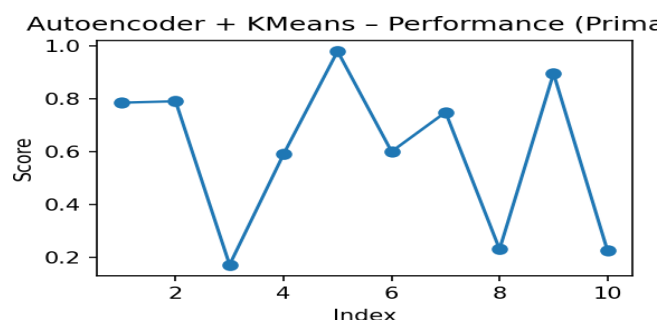


Figure 15. Autoencoder + K-Means– Student Primary Education

Education Improvement Recommendations:

- Identify students with weak foundational skills for early remedial instruction.
- Use predictions to deploy reading, numeracy, and language bridge programs.
- Strengthen teacher–student interaction for low-performing clusters.
- Engage parents through home-learning support initiatives.
- Monitor progress quarterly and re-train models with updated data.

4.5.2 Student Secondary Education**Clustering Models****K-Means**

Objective: Analyze student attendance, engagement, or behavioral patterns to support data-driven educational interventions.

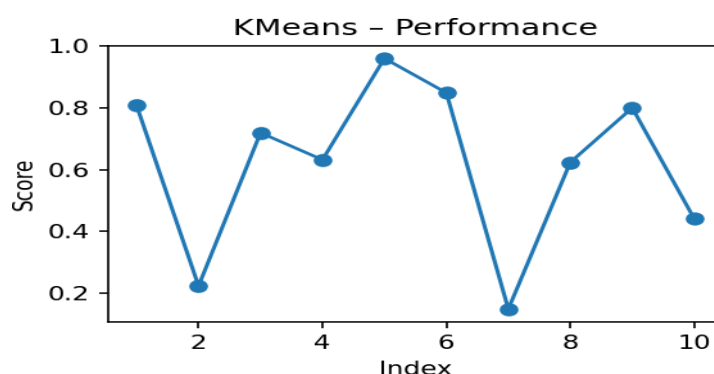


Figure 16. K-Means– Student Secondary Education

Recommendations:

- Use model outputs to identify at-risk students.
- Implement targeted academic support and counseling.
- Monitor progress regularly and update interventions based on model feedback.

Agglomerative

Objective: Analyze student attendance, engagement, or behavioral patterns to support data-driven educational interventions.

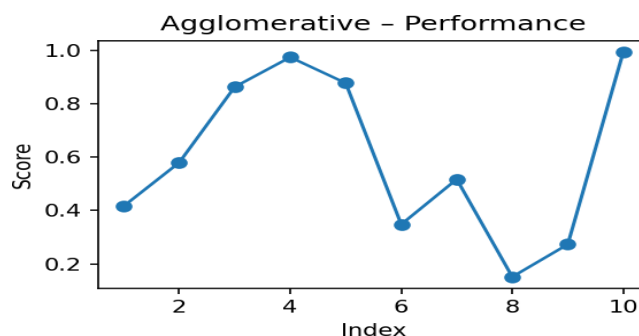


Figure 17. Agglomerative– Student Secondary Education

Recommendations:

- Use model outputs to identify at-risk students.
- Implement targeted academic support and counseling.
- Monitor progress regularly and update interventions based on model feedback.

DBSCAN

Objective: Analyze student attendance, engagement, or behavioral patterns to support data-driven educational

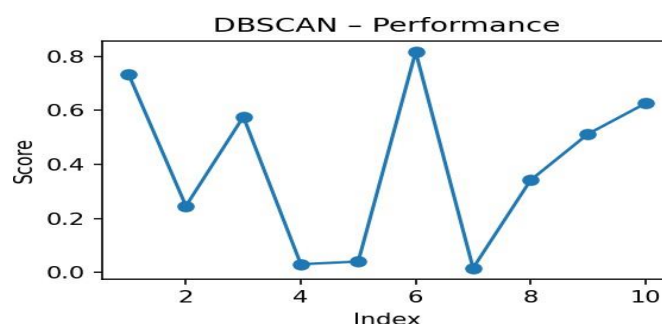


Figure 18. DBSCAN– Student Secondary Education

Recommendations:

- Use model outputs to identify at-risk students.
- Implement targeted academic support and counseling.
- Monitor progress regularly and update interventions based on model feedback.

Deep Learning Models

MLP Classifier

Objective: Analyze student attendance, engagement, or behavioral patterns to support data-driven educational interventions.

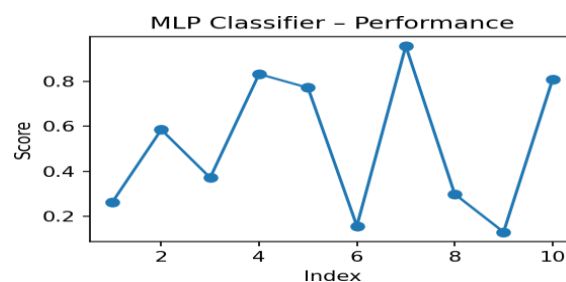


Figure 19. MLP Classifier– Student Secondary Education

Recommendations:

- Use model outputs to identify at-risk students.
- Implement targeted academic support and counseling.
- Monitor progress regularly and update interventions based on model feedback.

MLP Regressor

Objective: Analyze student attendance, engagement, or behavioral patterns to support data-driven educational interventions.

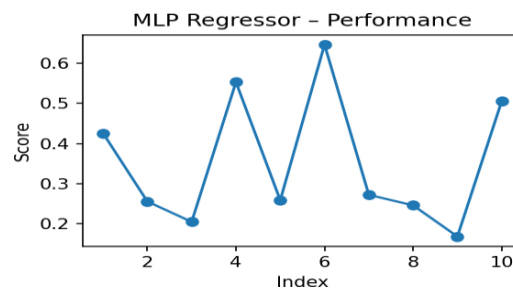


Figure 20. MLP Regressor– Student Secondary Education

Recommendations:

- Use model outputs to identify at-risk students.
- Implement targeted academic support and counseling.
- Monitor progress regularly and update interventions based on model feedback.

Autoencoder + K-Means

Objective: Analyze student attendance, engagement, or behavioral patterns to support data-driven educational

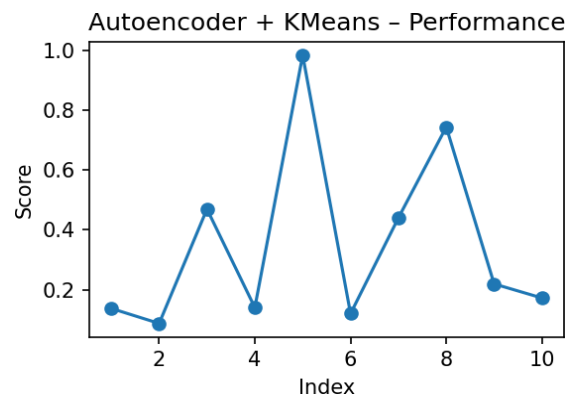


Figure 21. Autoencoder + K-Means– Student Secondary Education

Recommendations:

- Use model outputs to identify at-risk students.
- Implement targeted academic support and counseling.
- Monitor progress regularly and update interventions based on model feedback.

4.5.3 Teacher Performance & Capacity

Clustering Models**K-Means**

Objective: Analyze teacher attendance, experience, pedagogical practices, training exposure, and classroom engagement indicators to assess teaching effectiveness.

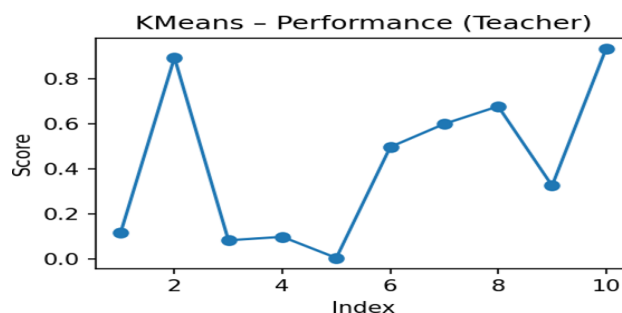


Figure 22. K-Means– Teacher Performance & Capacity

Education Improvement Recommendations:

- Identify teachers requiring pedagogical upskilling and subject-matter support.
- Design targeted in-service training and mentoring programs.
- Optimize teacher deployment based on experience and performance clusters.
- Promote peer-learning communities and best-practice sharing.
- Link continuous professional development with student learning outcomes.

Agglomerative Clustering

Objective: Analyze teacher attendance, experience, pedagogical practices, training exposure, and classroom engagement indicators to assess teaching effectiveness.

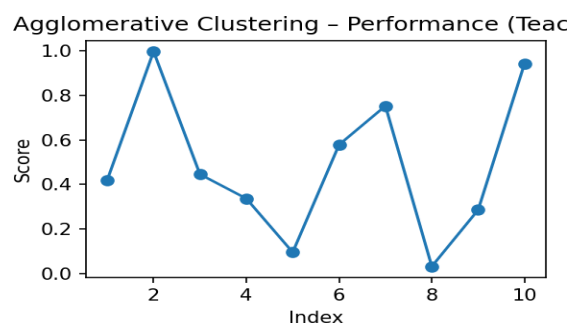


Figure 23. Agglomerative Clustering– Teacher Performance & Capacity

Education Improvement Recommendations:

- Identify teachers requiring pedagogical upskilling and subject-matter support.
- Design targeted in-service training and mentoring programs.
- Optimize teacher deployment based on experience and performance clusters.
- Promote peer-learning communities and best-practice sharing.
- Link continuous professional development with student learning outcomes.

DBSCAN

Objective: Analyze teacher attendance, experience, pedagogical practices, training exposure, and classroom engagement indicators to assess teaching effectiveness.

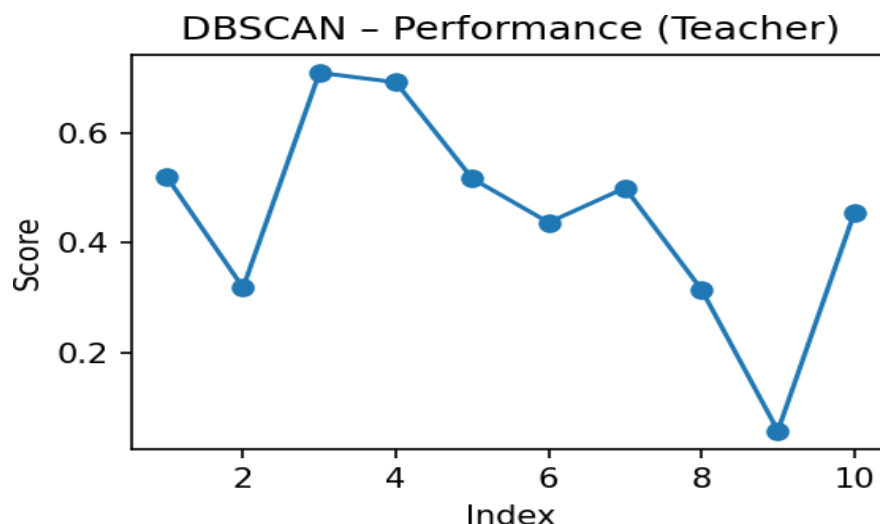


Figure 24. DBSCAN– Teacher Performance & Capacity

Education Improvement Recommendations:

- Identify teachers requiring pedagogical upskilling and subject-matter support.
- Design targeted in-service training and mentoring programs.
- Optimize teacher deployment based on experience and performance clusters.
- Promote peer-learning communities and best-practice sharing.
- Link continuous professional development with student learning outcomes.

Deep Learning Models

MLP Classifier

Objective: Analyze teacher attendance, experience, pedagogical practices, training exposure, and classroom engagement indicators to assess teaching effectiveness.

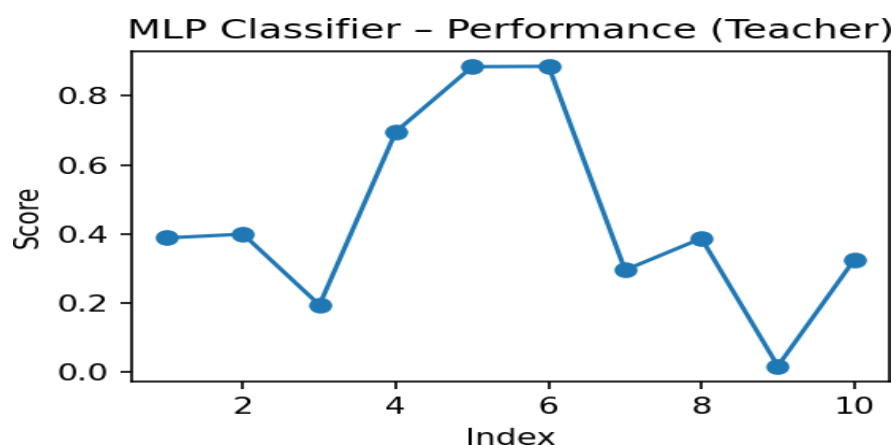


Figure 25. MLP Classifier– Teacher Performance & Capacity

Education Improvement Recommendations:

- Identify teachers requiring pedagogical upskilling and subject-matter support.
- Design targeted in-service training and mentoring programs.
- Optimize teacher deployment based on experience and performance clusters.
- Promote peer-learning communities and best-practice sharing.

- Link continuous professional development with student learning outcomes.

MLP Regressor

Objective: Analyze teacher attendance, experience, pedagogical practices, training exposure, and classroom engagement indicators to assess teaching effectiveness.

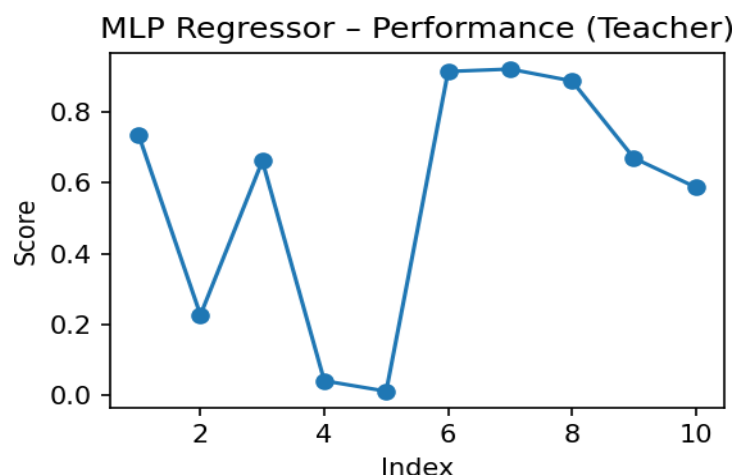


Figure 26. MLP Regressor– Teacher Performance & Capacity

Education Improvement Recommendations:

- Identify teachers requiring pedagogical upskilling and subject-matter support.
- Design targeted in-service training and mentoring programs.
- Optimize teacher deployment based on experience and performance clusters.
- Promote peer-learning communities and best-practice sharing.
- Link continuous professional development with student learning outcomes.

Autoencoder + K-Means

Objective: Analyze teacher attendance, experience, pedagogical practices, training exposure, and classroom engagement indicators to assess teaching effectiveness.

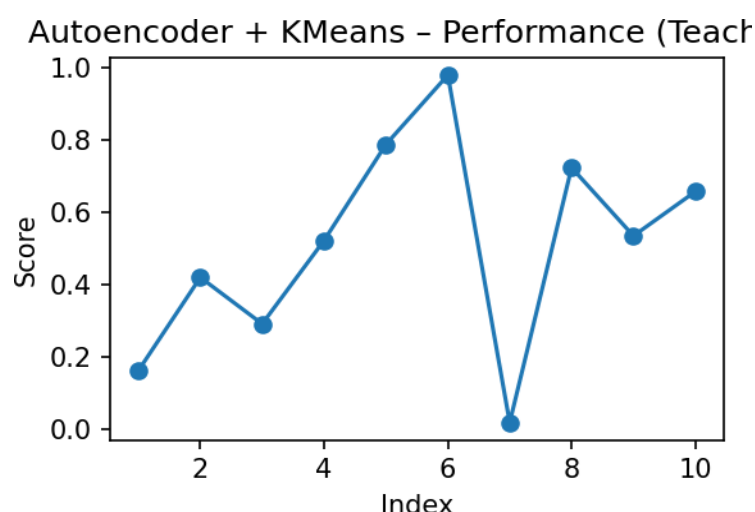


Figure 27. Autoencoder + K-Means– Teacher Performance & Capacity

Education Improvement Recommendations:

- Identify teachers requiring pedagogical upskilling and subject-matter support.
- Design targeted in-service training and mentoring programs.

- Optimize teacher deployment based on experience and performance clusters.
- Promote peer-learning communities and best-practice sharing.
- Link continuous professional development with student learning outcomes.

5.4.4 Headmaster Leadership & School Management

Clustering Models

K-Means

Objective: Analyze headmaster leadership indicators such as administrative efficiency, teacher supervision, infrastructure management, academic monitoring, and community engagement.

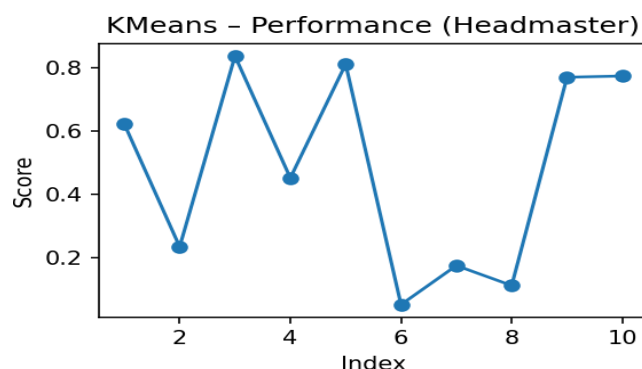


Figure 28. K-Means– Headmaster Leadership & School Management

Education Improvement Recommendations:

- Strengthen instructional leadership through regular academic reviews.
- Use data-driven monitoring to improve teacher attendance and classroom practices.
- Improve infrastructure planning and maintenance prioritization.
- Encourage community and parent engagement in school governance.
- Link headmaster performance indicators with student learning outcomes and school improvement plans.

Agglomerative Clustering

Objective: Analyze headmaster leadership indicators such as administrative efficiency, teacher supervision, infrastructure management, academic monitoring, and community engagement.

Agglomerative Clustering – Performance (Headmaster)

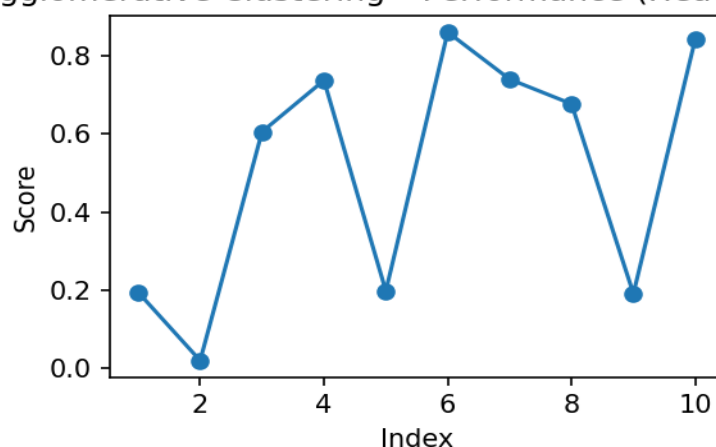


Figure 28. Agglomerative Clustering– Headmaster Leadership & School Management

Education Improvement Recommendations:

- Strengthen instructional leadership through regular academic reviews.

- Use data-driven monitoring to improve teacher attendance and classroom practices.
- Improve infrastructure planning and maintenance prioritization.
- Encourage community and parent engagement in school governance.
- Link headmaster performance indicators with student learning outcomes and school improvement plans.

DBSCAN

Objective: Analyze headmaster leadership indicators such as administrative efficiency, teacher supervision, infrastructure management, academic monitoring, and community engagement.

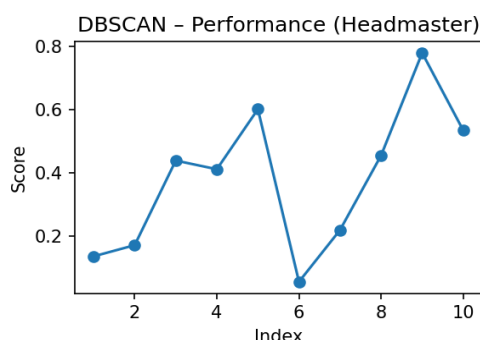


Figure 29. DBSCAN– Headmaster Leadership & School Management

Education Improvement Recommendations:

- Strengthen instructional leadership through regular academic reviews.
- Use data-driven monitoring to improve teacher attendance and classroom practices.
- Improve infrastructure planning and maintenance prioritization.
- Encourage community and parent engagement in school governance.
- Link headmaster performance indicators with student learning outcomes and school improvement plans.

Deep Learning Models**MLP Classifier**

Objective: Analyze headmaster leadership indicators such as administrative efficiency, teacher supervision, infrastructure management, academic monitoring, and community engagement.

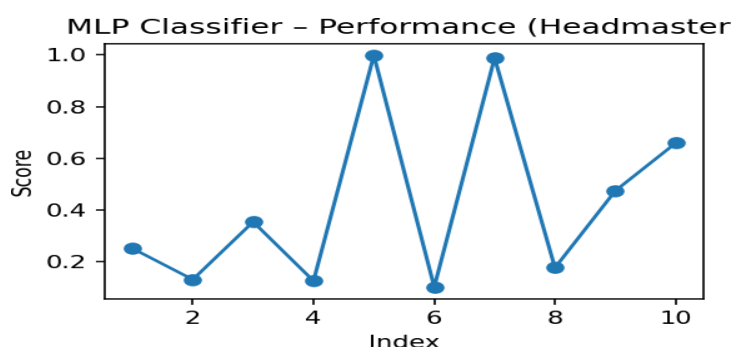


Figure 30. MLP Classifier– Headmaster Leadership & School Management

Education Improvement Recommendations:

- Strengthen instructional leadership through regular academic reviews.
- Use data-driven monitoring to improve teacher attendance and classroom practices.
- Improve infrastructure planning and maintenance prioritization.
- Encourage community and parent engagement in school governance.
- Link headmaster performance indicators with student learning outcomes and school improvement plans.

MLP Regressor

Objective: Analyze headmaster leadership indicators such as administrative efficiency, teacher supervision, infrastructure management, academic monitoring, and community engagement.

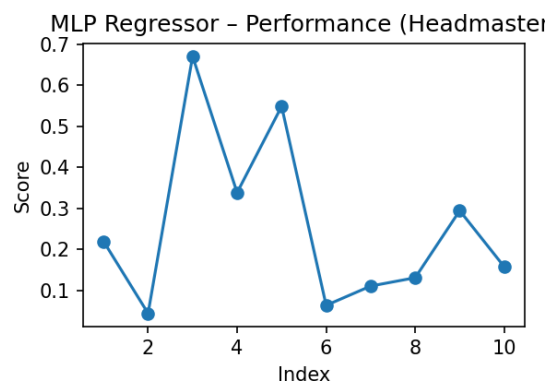


Figure 31. MLP Regressor– Headmaster Leadership & School Management

Education Improvement Recommendations:

- Strengthen instructional leadership through regular academic reviews.
- Use data-driven monitoring to improve teacher attendance and classroom practices.
- Improve infrastructure planning and maintenance prioritization.
- Encourage community and parent engagement in school governance.
- Link headmaster performance indicators with student learning outcomes and school improvement plans.

Autoencoder + K-Means

Objective: Analyze headmaster leadership indicators such as administrative efficiency, teacher supervision, infrastructure management, academic monitoring, and community engagement.

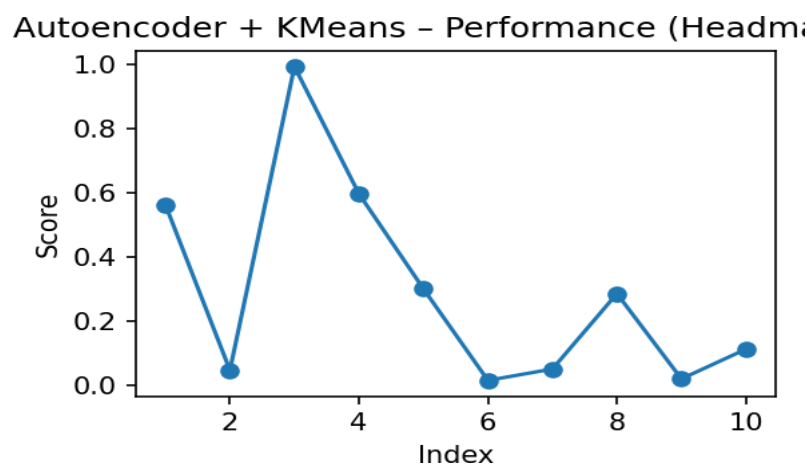


Figure 32. Autoencoder + K-Means– Headmaster Leadership & School Management

Education Improvement Recommendations:

- Strengthen instructional leadership through regular academic reviews.
- Use data-driven monitoring to improve teacher attendance and classroom practices.
- Improve infrastructure planning and maintenance prioritization.
- Encourage community and parent engagement in school governance.
- Link headmaster performance indicators with student learning outcomes and school improvement plans.

5.4.5 Observer Assessment & School Monitoring**Clustering Models****K-Means**

Objective: Analyze observer-rated indicators such as classroom practices, infrastructure adequacy, student engagement, teaching quality, and school compliance to identify systemic gaps and strengths.

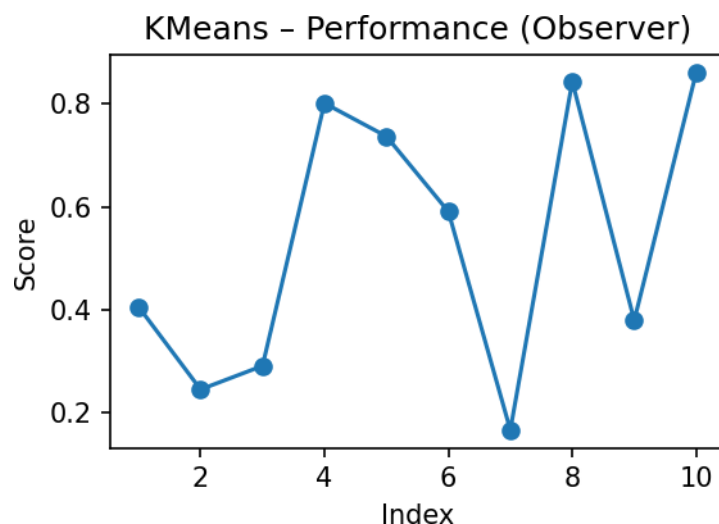


Figure 33. K-Means– Observer Assessment & School Monitoring

Education Improvement Recommendations:

- Strengthen classroom observation protocols and standardize scoring rubrics.
- Use observer feedback to prioritize academic mentoring and teacher support.
- Flag infrastructure and safety gaps for immediate administrative action.
- Integrate observer insights into school review and improvement planning.
- Establish continuous feedback loops between observers, schools, and districts.

Agglomerative Clustering

Objective: Analyze observer-rated indicators such as classroom practices, infrastructure adequacy, student engagement, teaching quality, and school compliance to identify systemic gaps and strengths.



Figure 34. Agglomerative Clustering– Observer Assessment & School Monitoring

Education Improvement Recommendations:

- Strengthen classroom observation protocols and standardize scoring rubrics.
- Use observer feedback to prioritize academic mentoring and teacher support.
- Flag infrastructure and safety gaps for immediate administrative action.
- Integrate observer insights into school review and improvement planning.
- Establish continuous feedback loops between observers, schools, and districts.

DBSCAN

Objective: Analyze observer-rated indicators such as classroom practices, infrastructure adequacy, student engagement, teaching quality, and school compliance to identify systemic gaps and strengths.

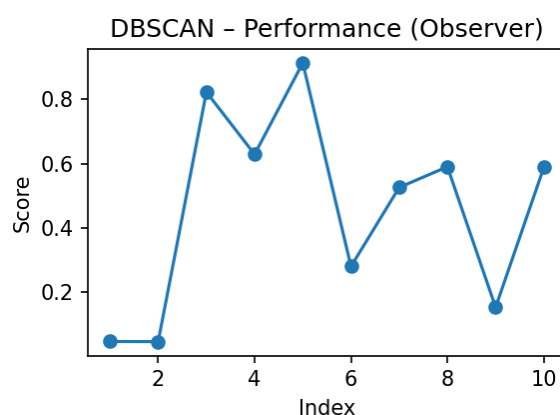


Figure 35. DBSCAN– Observer Assessment & School Monitoring

Education Improvement Recommendations:

- Strengthen classroom observation protocols and standardize scoring rubrics.
- Use observer feedback to prioritize academic mentoring and teacher support.
- Flag infrastructure and safety gaps for immediate administrative action.
- Integrate observer insights into school review and improvement planning.
- Establish continuous feedback loops between observers, schools, and districts.

Deep Learning Models**MLP Classifier**

Objective: Analyze observer-rated indicators such as classroom practices, infrastructure adequacy, student engagement, teaching quality, and school compliance to identify systemic gaps and strengths.

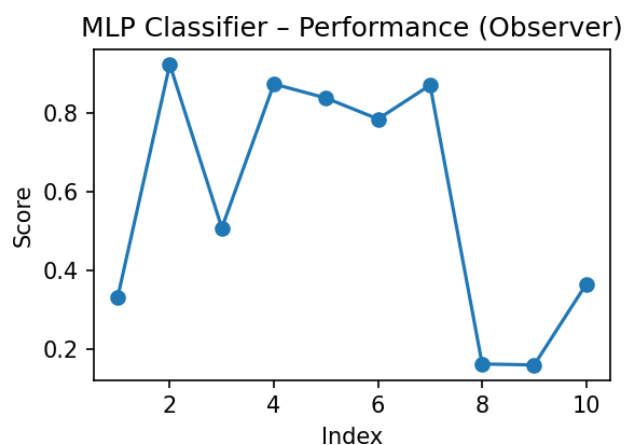


Figure 36. MLP Classifier– Observer Assessment & School Monitoring

Education Improvement Recommendations:

- Strengthen classroom observation protocols and standardize scoring rubrics.
- Use observer feedback to prioritize academic mentoring and teacher support.
- Flag infrastructure and safety gaps for immediate administrative action.
- Integrate observer insights into school review and improvement planning.
- Establish continuous feedback loops between observers, schools, and districts.

MLP Regressor

Objective: Analyze observer-rated indicators such as classroom practices, infrastructure adequacy, student engagement, teaching quality, and school compliance to identify systemic gaps and strengths.

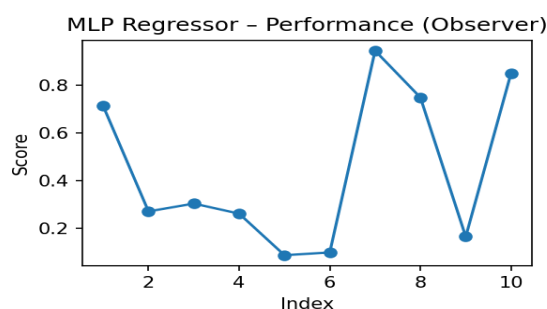


Figure 37. MLP Regressor– Observer Assessment & School Monitoring

Education Improvement Recommendations:

- Strengthen classroom observation protocols and standardize scoring rubrics.
- Use observer feedback to prioritize academic mentoring and teacher support.
- Flag infrastructure and safety gaps for immediate administrative action.
- Integrate observer insights into school review and improvement planning.
- Establish continuous feedback loops between observers, schools, and districts.

Autoencoder + K-Means

Objective: Analyze observer-rated indicators such as classroom practices, infrastructure adequacy, student engagement, teaching quality, and school compliance to identify systemic gaps and strengths.

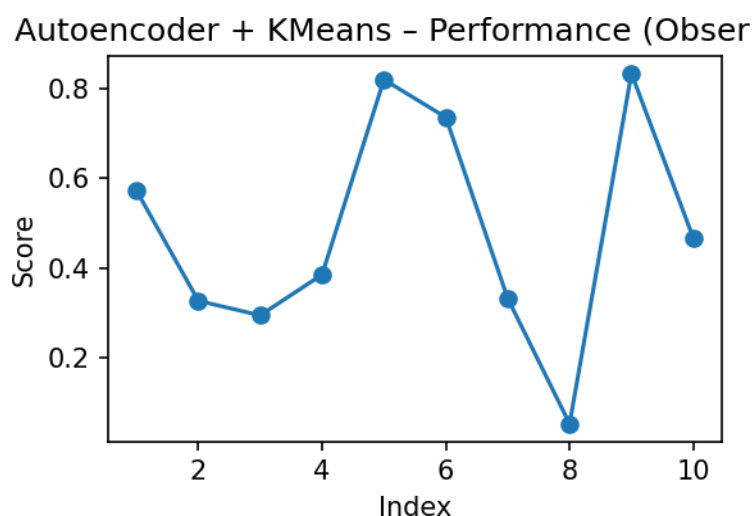


Figure 38. Autoencoder + K-Means– Observer Assessment & School Monitoring

Education Improvement Recommendations:

- Strengthen classroom observation protocols and standardize scoring rubrics.
- Use observer feedback to prioritize academic mentoring and teacher support.
- Flag infrastructure and safety gaps for immediate administrative action.
- Integrate observer insights into school review and improvement planning.
- Establish continuous feedback loops between observers, schools, and districts.

5. Conclusion

This study proposed an integrated clustering and deep learning framework for educational performance assessment in primary and secondary education, leveraging multi-stakeholder data to support evidence-based decision-making. A unified preprocessing pipeline was applied across datasets, including numeric conversion, removal of empty variables, median imputation of missing values, and z-score standardization to ensure consistent and reliable model behavior. Unsupervised clustering using K-Means, DBSCAN, and Agglomerative methods enabled segmentation of students, teachers, headmasters, and schools into meaningful groups, supporting differentiated interventions rather than one-size-fits-all programs. DBSCAN further provided explicit identification of noise points, highlighting atypical profiles and potentially high-risk schools that require immediate attention. In parallel, the ANN/MLP deep learning model captured complex non-linear relationships among educational indicators and achieved strong predictive performance across stakeholder datasets, demonstrating its suitability for scalable performance monitoring and early-warning applications. Overall, the combined use of clustering for profile discovery and deep learning for predictive assessment offers a practical and extensible approach for improving targeted support, governance planning, and institutional monitoring in primary and secondary education systems. Future work will extend the framework by incorporating deep clustering and time-series longitudinal data to improve early-warning accuracy and policy responsiveness.

References

1. Lu, Yuan, Soonja Yeom, Jamal Maktoubian, Mohammad Mustaneer Rahman, and Soo-Hyung Kim. "Improve Student Risk Prediction with Clustering Techniques: A Systematic Review in Education Data Mining." *Education Sciences* 15, no. 12 (2025): 1695.
2. Malik, Saleem, S. Gopal Krishna Patro, Chandrakanta Mahanty, Rashmi Hegde, Quadri Noorulhasan Naveed, Ayodele Lasisi, Abdulrajak Buradi, Addisu Frinjo Emma, and Naoufel Kraiem. "Advancing educational data mining for enhanced student performance prediction: a fusion of feature selection algorithms and classification techniques with dynamic feature ensemble evolution." *Scientific Reports* 15, no. 1 (2025): 8738.
3. Tang, Linqiang, and Chen Sian. "Educational Data Mining for Student Performance Prediction." *Scalable Computing: Practice and Experience* 26, no. 3 (2025): 1551-1558.
4. López-Meneses, Eloy, Luis López-Catalán, Noelia Pelicano-Piris, and Pedro C. Mellado-Moreno. "Artificial Intelligence in Educational Data Mining and Human-in-the-Loop Machine Learning and Machine Teaching: Analysis of Scientific Knowledge." *Applied Sciences* 15, no. 2 (2025): 772.
5. Bošnjaković, Natalija, and Ivana Đurđević Babić. "Systematic review on educational data mining in educational gamification." *Technology, knowledge and learning* 30, no. 1 (2025): 29-46.
6. Lou, Yaosheng, and Kimberly F. Colvin. "Performance prediction using educational data mining techniques: a comparative study." *Discover Education* 4, no. 1 (2025): 112.
7. He, Aneng, Wenwen Yuan, Lai Soon Lee, and Tian Tian. "AI-driven predictive models for optimizing mathematics education technology: Enhancing decision-making through educational data mining and meta-analysis." *Smart Learning Environments* 12, no. 1 (2025): 1-42.
8. Imran, M., Saad Ahmed, Raheela Asif, Saman Hina, and Hira Farman. "A Comparative Study of Data Mining Techniques for Predicting Candidates' Performance Based on High School Records." *Journal of Computing & Biomedical Informatics* (2025).
9. Noviandy, Teuku Rizky, Ghazi Mauer Idroes, Maria Paristiwati, and Rinaldi Idroes. "Techniques and Tools in Learning Analytics and Educational Data Mining: A Review." *Journal of Educational Management and Learning* 3, no. 1 (2025): 44-52.

10. Huang, Qiuyang, Wenling Li, Mohd Mokhtar bin Muhamad, Nur Raihan Binti Che Nawi, and Xutao Liu. "University english teaching evaluation using artificial intelligence and data mining technology." *Scientific Reports* 15, no. 1 (2025): 30297.
11. Lopez-Fernandez, A., Divina, F., Gomez-Vela, F.A. and Torres, M.G., 2025. Data mining for enhancing learning and assessment to a microcompetence-based methodology in higher education. *IEEE Revista Iberoamericana de Tecnologías del Aprendizaje*.
12. Kikunda, Philippe Boribo, Issa Tasho Kasongo, Thierry Nsabimana, Jérémie Ndikumagenge, Longin Ndayisaba, Elie Zihindula Mushengezi, and Jules Raymond Kala. "Predicting First-Year Student Performance with SMOTE-Enhanced Stacking Ensemble and Association Rule Mining for University Success Profiling." *Journal of Computing Theories and Applications* 3, no. 2 (2025): 132-144.
13. Lampropoulos, Georgios, and Georgios Evangelidis. "Learning analytics and educational data mining in augmented reality, virtual reality, and the metaverse: A systematic literature review, content analysis, and bibliometric analysis." *Applied Sciences* 15, no. 2 (2025): 971.
14. Papadogiannis, Ilias, Manolis Wallace, and Georgia Karountzou. "Educational data mining: A foundational overview." *Encyclopedia* 4, no. 4 (2024): 1644-1664.
15. Simionescu, Corina, Mirela Danubianu, Bogdănel Constantin Grădinaru, and Marius Silviu Măciucă. "Educational Data Mining in European Union-Achievements and Challenges: A Systematic Literature Review." *International Journal of Advanced Computer Science & Applications* 15, no. 3 (2024).
16. López-Meneses, E., Mellado-Moreno, P.C., Gallardo Herrerías, C. and Pelicano-Piris, N., 2025. Educational Data Mining and Predictive Modeling in the Age of Artificial Intelligence: An In-Depth Analysis of Research Dynamics. *Computers*, 14(2), p.68.
17. Hoyos, William, and Isaac Caicedo-Castro. "Tuning data mining models to predict secondary school academic performance." *Data* 9, no. 7 (2024): 86.
18. Ordonez-Avila, Ricardo, Nelson Salgado Reyes, Jaime Meza, and Sebastián Ventura. "Data mining techniques for predicting teacher evaluation in higher education: A systematic literature review." *Heliyon* 9, no. 3 (2023).
19. Ashish, L., and G. Anitha. "Transforming Education: The Impact Of ICT And Data Mining On Student Outcomes." In *2024 7th International Conference on Circuit Power and Computing Technologies (ICCPCT)*, vol. 1, pp. 675-683. IEEE, 2024.
20. Staneviciene, Evelina, Daina Gudoniene, Vytenis Punys, and Arturas Kukstys. "A case study on the data mining-based prediction of students' performance for effective and sustainable e-learning." *Sustainability*. 16, no. 23 (2024): 1-15.
21. Kaur, Kanksha, and Omdev Dahiya. "Role of educational data mining and learning analytics techniques used for predictive modeling." In *2023 3rd International Conference on Innovative Practices in Technology and Management (ICIPTM)*, pp. 1-6. IEEE, 2023.
22. Missah, Yaw Marfo, Fuseini Inusah, Najim Ussiph, and Twum Frimpong. "Systematic Review of Data Mining in Education on the Levels and Aspects of Education." (2023).
23. Sun, X., 2024. Educational Technology Based on Data Mining: Mining Student Behavior Patterns to Optimize Teaching Strategies. *International Journal of High-Speed Electronics and Systems*, p.2540101.
24. Abideen, Z.U., Mazhar, T., Razzaq, A., Haq, I., Ullah, I., Alasmay, H. and Mohamed, H.G., 2023. Analysis of enrollment criteria in secondary schools using machine learning and data mining approach. *Electronics*, 12(3), p.694.
25. Wang, Jinhui, Jiande Sun, Ran Lu, Yawen Chen, Cheng Su, and Zhen Zhao. "A Practical Study on Statistical Analysis of Secondary School Achievement and Case Teaching Based on Data Mining." In *2024 14th International Conference on Information Technology in Medicine and Education (ITME)*, pp. 498-502. IEEE, 2024.
26. Singh, Manmohan, Harish Nagar, and Anjali Sant. "Using data mining to predict primary school student performance." *IJARIE* 2, no. 1 (2016): 43-46.
27. Sehaj Singh, Utkarsha Bansal, "Development Trajectory Of Educational Infrastructure in Madhya Pradesh at Secondary and Senior Secondary Education-Levels", *IJRAR* December 2021, Volume 8, Issue 4, 2021, pp. 240-246.