

Project Dissertation Report on Dark Store Performance Segmentation in Indian Quick Commerce:

A K-Means Clustering and Hypothesis-Testing Approach

Submitted By

Piyush Kumar

24/BMBA/24

Under the Guidance of

Dr. Kusum Lata

Assistant Professor



DELHI SCHOOL OF MANAGEMENT
Delhi Technological University
Bawana Road Delhi 110042

CERTIFICATE

This is to certify that the Major Research Project titled “**Dark Store Performance Segmentation in Indian Quick Commerce: A K-Means Clustering and Hypothesis-Testing Approach**” has been completed by **Mr. Piyush Kumar**, Roll No. 24/BMBA/24 for the partial fulfilment of the requirements for the award of the degree of Master of Business Administration (MBA) in Business Analytics from Delhi School of Management, Delhi Technological University.

Date: 17th May 2026

Signature of Guide: _____

Dr. Kusum Lata

Assistant Professor

Delhi School of Management, DTU

DECLARATION

I hereby declare that the project titled “**Dark Store Performance Segmentation in Indian Quick Commerce: A K-Means Clustering and Hypothesis-Testing Approach**” submitted in partial fulfillment of the requirements for the award of the degree of MBA in Business Analytics at Delhi School of Management, Delhi Technological University is a record of original work carried out by me under the guidance and supervision of **Dr. Kusum Lata**.

I further declare that this project has been completed solely by me and has not been submitted, either in part or in full, to any other university or institution for the award of any degree, diploma, or other academic qualification. The information and results included in this report are based on genuine data and trustworthy sources. Whenever the work of others has been referenced, proper credit has been given.

I fully take responsibility for the originality, truthfulness, and correctness of the content in this project report.

Name: Piyush Kumar

Roll Number: 24/BMBA/24

Course: MBA Business Analytics

Date: 16th May 2026

Signature: _____

ACKNOWLEDGEMENT

The researcher would like to sincerely thank **Dr. Kusum Lata** for her valuable guidance, thoughtful feedback, and ongoing encouragement throughout the entire research process. Her mentorship was crucial in ensuring the successful completion of this work.

The researcher also wishes to convey deep appreciation to the faculty members of **Delhi School of Management, DTU**, for their academic support and for creating an environment that fostered intellectual growth and facilitated this study.

Lastly, the researcher acknowledges the support and encouragement from family and friends, whose patience and belief helped sustain this endeavor from its beginning until its completion.

Piyush Kumar
Delhi School of Management, DTU
2025-26

TABLE OF CONTENTS

CHAPTER 1: INTRODUCTION	2
1.1 Problem Statement	2
1.2 Research Gaps	3
1.3 Key Contributions of the Study	3
1.4 Organisation of the Report	4
CHAPTER 2: LITERATURE REVIEW	5
2.1 Problem Area: Quick Commerce, Dark Stores, and Indian Retail	5
2.1.1 The Rise of Quick Commerce in India	5
2.1.2 Dark Store Operations and Performance Heterogeneity	6
2.1.3 Consumer Behaviour and Category Preferences in Q-Commerce	6
2.1.4 Market Structure and Competitive Dynamics	7
2.2 Technique Area: Clustering, Dimensionality Reduction, and Statistical Validation	8
2.2.1 K-Means Clustering in Retail Analytics	8
2.2.2 Feature Engineering for Store Segmentation	8
2.2.3 Principal Component Analysis for Cluster Visualisation	9
2.2.4 Statistical Hypothesis Testing for Cluster Validation	9
2.2.5 Machine Learning in Supply Chain and Inventory Management	10
CHAPTER 3: PROPOSED METHODOLOGY	11
3.1 Step-by-Step Analytical Framework	11
3.2 Explanation of Model	12
3.2.1 Feature Engineering	12
3.2.2 Data Normalisation	13
3.2.3 K-Means Algorithm	13
3.2.4 Optimal K Selection	14
3.2.5 Principal Component Analysis for Visualisation	14
3.2.6 Statistical Tests	14
CHAPTER 4: EXPERIMENT AND RESULTS	16
4.1 Baseline Models and Evaluation Metrics	16
4.2 Dataset Preparation	16
4.3 Pre-Processing	17
4.4 Experiments Performed	18
4.4.1 Exploratory Data Analysis	18
4.4.2 Optimal K Selection	21
4.4.3 K-Means Clustering Results and Cluster Profiles	22
4.5 Discussion: Findings and Analysis	25
4.5.1 Hypothesis Testing Results	25
4.5.2 Interpretation of Hypothesis Results	26
4.5.3 Cluster Characterisation	27
CHAPTER 5: CONCLUSION AND FUTURE IMPLICATIONS	29

5.1 Managerial Implications	29
5.2 Limitations	31
5.3 Future Research Directions	31
REFERENCES	33

LIST OF FIGURES

Figure 1: Sales by Category (FMCG, Electronics Excluded)	18
Figure 2: Average Order Value and Discount Percentage by Category	19
Figure 3: Top 15 Dark Stores by Revenue.....	20
Figure 4: Correlation Matrix of Business Metrics	21
Figure 5: Elbow Method and Silhouette Score for Optimal K.....	21
Figure 6: PCA 2D Cluster Scatter Plot	22
Figure 7: Cluster Comparison: Sales, Orders, AOV, Discount.....	27
Figure 8: Box Plots: Distribution Within Clusters.....	24
Figure 9: Category Mix Heatmap per Cluster.....	25
Figure 10: Hypothesis Testing Results: p-values vs alpha = 0.05	26

LIST OF TABLES

Table 1: Summary Statistics of Dataset	17
Table 2: Derived Features and Their Formulae	17
Table 3: Cluster Summary: Key Performance Metrics	23
Table 4: Hypothesis Testing Summary	25
Table 5: Business Strategy Recommendations per Cluster.....	30

LIST OF ABBREVIATIONS

Abbreviation	Full Form
AOV	Average Order Value
ANOVA	Analysis of Variance
FMCG	Fast-Moving Consumer Goods
MRP	Maximum Retail Price
PCA	Principal Component Analysis
Q-commerce	Quick Commerce
WCSS	Within-Cluster Sum of Squares
EDA	Exploratory Data Analysis
SKU	Stock Keeping Unit
GMV	Gross Merchandise Value
B2C	Business-to-Consumer
UX	User Experience
KPI	Key Performance Indicator
INR	Indian Rupee
TLC	Top-Level Category
MAPE	Mean Absolute Percentage Error
NCR	National Capital Region
WCSS	Within-Cluster Sum of Squares

ABSTRACT

Quick commerce, or Q-commerce, has become one of the most transformative trends in the retail sector in India, offering the ability to deliver groceries and everyday items within 10 to 30 minutes through a network of nearby dark stores. Although this model has grown quickly, operators often lack a structured approach to understand the varying levels of performance among these dark stores. This study fills that gap by using unsupervised machine learning and statistical methods to analyze sales data from 60 dark stores in a large Indian city. The data covers 13 categories of fast-moving consumer goods and groceries, excluding electronics to avoid bias in the average order value. The process includes creating features such as average order value, discount percentage, and product mix from the transaction data. K-Means clustering is applied to the standardized and log-transformed data to group the dark stores into four distinct performance categories. Principal Component Analysis is used to visualize these clusters in two dimensions. The results are checked using statistical techniques like one-way ANOVA, Welch's t-test, and Kruskal-Wallis test to test six hypotheses. Five out of the six hypotheses are confirmed at a 5% significance level, showing that the different clusters perform significantly differently in terms of total sales, number of orders, average order value, and discount patterns. The four types of stores identified are Moderate Performers, Anomalous/Inactive Stores, Premium/Sparse Stores, and High-Volume Workhorses. These insights help in making informed decisions about inventory management, pricing, and marketing efforts across the dark store network. This research represents one of the first detailed analyses of dark store performance in the Indian quick commerce sector, offering both a practical analytical method and strategic advice for operators.

Keywords: *Quick commerce, dark stores, K-Means clustering, ANOVA, FMCG, Indian retail, demand analytics, store segmentation.*

CHAPTER 1: INTRODUCTION

The Indian retail industry has experienced a major shift in the past ten years, fueled by the widespread use of smartphones, affordable internet data, and the development of a more formalized grocery supply chain. A key development in recent times has been the emergence of quick commerce, or Q-commerce, which focuses on delivering groceries, personal care items, and everyday necessities to customers within 10 to 30 minutes of placing an order. Unlike traditional e-commerce, Q-commerce depends on a network of small, highly efficient warehouses known as dark stores. These stores are located in residential areas and are not open to the public, but they serve as critical hubs for fast delivery.

The Q-commerce market in India was valued at about USD 0.3 billion in 2021 and is expected to grow to USD 5.5 billion by 2025, with a compound annual growth rate (CAGR) over 75%. The rapid growth of platforms like Blinkit, Zepto, and Swiggy Instamart has led to increased competition, tighter profit margins, and significant operational challenges in managing dark store networks. Each dark store operates as a separate mini-market, influenced by local consumer demand, population characteristics, the area it serves, and customer preferences, which vary from one location to another.

Although these differences are important for performance, many Q-commerce companies still use the same strategies for inventory restocking, pricing discounts, and marketing spending across all dark stores. This uniform approach leads to poor inventory management, missed sales opportunities in promising areas, and reduced profitability in stores where customers are more price-sensitive. The lack of a robust data-based method for analyzing and categorizing dark store performance highlights a key issue that this research aims to solve.

1.1 Problem Statement

Indian Q-commerce businesses operate vast networks of dark stores spread across different regions, each serving its own small market with varying levels of demand. Despite this diversity, there is currently no proven, data-based method to group these dark stores into categories that reflect their performance. Without such grouping, operational strategies like stock management, pricing, and marketing are treated the same across stores that are very different in nature, which leads to wasted resources, stock shortages, reduced profits, and lost chances for growth [5].

This research explores whether using detailed sales data along with machine learning techniques and statistical analysis can create a useful way to classify dark store performance for an Indian Q-commerce system. The study is driven by the observation that stores within the same city can differ significantly, with some performing up to 300% better in important measures like sales and average order value. This finding aligns with Gupta et al. (2022), who found large differences in demand across dark stores in the Delhi National Capital Region (NCR) [4].

1.2 Research Gaps

The academic research on Q-commerce is still in its early stages but is expanding. Current studies have mainly concentrated on consumer adoption and purchasing intentions [6], along with logistics efficiency, last-mile delivery systems [7], and the competitive positioning of Q-commerce platforms [8]. A limited number of studies have explored the use of clustering techniques for evaluating the performance of retail and e-commerce stores [9]. However, there is still no comprehensive examination that combines dark store operations, transaction-level FMCG data, unsupervised machine learning, and statistical hypothesis testing specifically within the Indian Q-commerce environment.

Moreover, previous research in retail clustering has rarely addressed the issue of category bias created by high-value product categories such as electronics, which can distort average order value metrics and lead to inaccurate cluster classifications [10]. Mehrotra and Agarwal (2021) observed that the inclusion of high-ticket products in average order value calculations can artificially inflate estimates for stores carrying such items, making it difficult to identify genuine differences in regular consumer purchasing behavior [6]. This study makes a methodological contribution by excluding these biasing product categories and focusing exclusively on daily consumption FMCG categories.

Additionally, while Li and Chen (2024) applied K-Means clustering to identify convenience store types in urban retail environments in China [19], and Wang et al. (2024) emphasized the importance of log transformation for managing skewed e-commerce datasets [20], neither study examined the distinct operational characteristics of Q-commerce dark stores in an emerging market such as India. Furthermore, these studies did not incorporate rigorous multi-test hypothesis validation as part of their analytical framework.

1.3 Key Contributions of the Study

This study contributes four unique insights to the field of Q-commerce analytics in India:

- First, it introduces and tests a feature engineering approach specifically designed for analyzing the performance of dark stores. This framework includes factors such as average order value, discount percentage, units per order, and category-wise sales mix as clustering variables. It is designed to account for differences in store size while highlighting meaningful variations in product assortment and pricing strategies across locations.
- Second, the study applies K-Means clustering along with methods for determining the optimal number of clusters using both the Elbow Method and Silhouette Score. This process results in four distinct dark store categories that are statistically validated. The adjusted interpretation of the Silhouette Score, where $K=2$ initially produced a higher value because of a single-store outlier, demonstrates the importance of combining statistical validation techniques with practical business relevance when selecting the most suitable clustering solution.
- Third, the research employs a comprehensive hypothesis-testing framework consisting of one-way ANOVA, Welch's t-test, and Kruskal-Wallis test to verify the statistical reliability of the identified clusters. Five out of the six hypotheses are supported at the 5% significance level, providing strong evidence that the segmentation successfully captures meaningful differences in dark store performance.
- Fourth, the study transforms these analytical findings into a practical and strategy-oriented framework tailored to each of the four identified store categories. The framework includes recommendations related to inventory management, pricing strategies, and marketing approaches, thereby directly addressing the managerial challenges identified in the problem statement.

1.4 Organisation of the Report

The rest of this report is organized in the following manner. Chapter 2 provides a review of the relevant literature, focusing on two major areas: the problem area related to the growth of Q-commerce and dark store operations, and the technique area involving clustering methods and statistical approaches used in retail analytics. Chapter 3 presents the proposed methodology, explaining the step-by-step analytical framework and the model in detail. Chapter 4 describes the experimental setup, data preparation procedures, preprocessing steps, the experiments conducted, and a detailed discussion of the findings. Chapter 5 concludes the study and suggests possible directions for future research. The references section contains all the sources cited throughout the report.

CHAPTER 2: LITERATURE REVIEW

This chapter reviews academic and industry literature across two broad areas. The first section focuses on the problem area, including the rapid growth of Q-commerce in India, the dark store model, consumer behavior in instant delivery environments, and the operational challenges involved in managing diverse dark store networks. The second section examines the technique area, covering clustering methods used in retail analytics, dimensionality reduction through Principal Component Analysis (PCA), and statistical hypothesis-testing techniques used to evaluate differences between clusters.

2.1 Problem Area: Quick Commerce, Dark Stores, and Indian Retail

2.1.1 The Rise of Quick Commerce in India

The idea of Q-commerce, which refers to the delivery of groceries and essential items within 30 minutes, developed due to the limitations of traditional e-commerce, which typically involved delivery times of 24 to 48 hours [1]. Verhoef et al. (2021) suggested that the shift towards omnichannel retail significantly changed what consumers expect from delivery speed, leading to greater investments in fast fulfillment systems [11]. In India, this change was sped up by changes in consumer behavior caused by the pandemic. Chopra & Meindl (2016) noted that for FMCG distributors, building resilient supply chains and being close to customers became key priorities [12].

Singh & Banerjee (2023) examined how Indian Q-commerce developed from 2020 to 2023, showing how companies like Blinkit (formerly Grofers), Zepto, and Swiggy Instamart expanded their dark stores in major cities across India [13]. Their research showed that the density of dark stores—meaning the number of fulfillment points per square kilometer in urban areas—was the most important factor affecting how quickly orders were delivered, highlighting the local nature of this business model. Rathore et al. (2023) also found that Q-commerce platforms that focused on understanding local demand outperformed competitors by 18 to 23% in order fulfillment during busy times [14].

The competition in the Q-commerce industry has become much more intense. Sharma & Patel (2022) looked at the competitive structure of the Indian Q-commerce market using a duopoly model and concluded that at the local level, the winner-takes-most trend means a store's performance is heavily influenced by how close and varied the offerings of nearby dark stores are within a 2 km

radius [8]. These market pressures make performance segmentation not just an analysis tool, but a necessary strategy for companies looking to use their resources efficiently across their store networks.

2.1.2 Dark Store Operations and Performance Heterogeneity

Dark stores differ significantly from traditional distribution centers and conventional retail outlets. Rao and Mehta (2022) described Indian dark stores as compact, inventory-focused warehouses that typically stock between 1,500 and 3,000 Stock Keeping Units (SKUs) per location. These facilities operate under strict service-level agreements requiring at least 95% product availability for key fast-moving consumer goods (FMCG) items [15]. The researchers further observed that important performance indicators such as average order value, order volume, and gross margins can vary by as much as 300% within the same city. These differences are mainly influenced by factors such as neighborhood income levels, household size, and proximity to competing retail outlets.

Gupta et al. (2022) investigated inventory shrinkage and demand forecasting accuracy across a network of 45 Q-commerce dark stores in the Delhi NCR region. Their findings showed that stores located in densely populated residential areas achieved lower forecasting errors, recording a Mean Absolute Percentage Error (MAPE) of 12%, compared to stores situated in mixed commercial zones, which recorded a MAPE of 28% [4]. This evidence supports the central argument of the present study that dark stores within a single city display significantly different demand characteristics and therefore require differentiated operational strategies. The authors also concluded that a uniform replenishment strategy would be ineffective due to the substantial variation in store performance.

Bhatia and Verma (2023) compared the unit economics of Indian Q-commerce operations with those of similar markets in South Korea and Turkey. Their study revealed that Indian operators reported significantly lower Average Order Values (AOV), ranging between INR 250 and INR 400, compared to USD 15–25 in international markets. This difference was largely attributed to the tendency of Indian consumers to place smaller but more frequent orders rather than conducting larger weekly grocery purchases [17]. These findings highlight the importance of analyzing the distribution of AOV across dark store networks, as operational strategies designed for high-AOV stores may not be suitable for low-AOV, high-frequency locations.

2.1.3 Consumer Behaviour and Category Preferences in Q-Commerce

Mehrotra and Agarwal (2021) conducted a survey involving 800 urban users of Q-commerce services in India and found substantial differences in purchasing behavior across income groups. Higher-income consumers were more likely to purchase products from categories such as Baby Care, Gourmet Foods, and Beauty Products, whereas mass-market consumers primarily focused on Foodgrains, Bakery items, and Snacks [6]. These findings have important implications for inventory planning because stores serving premium customer segments require different product assortments than stores catering to high-volume, price-sensitive consumers. The study also highlighted that the inclusion of electronics in platform-level analytics can significantly distort Average Order Value calculations.

Jain et al. (2023) examined the relationship between discount levels and Q-commerce order frequency. Their findings indicated that a 10% increase in discount percentage resulted in a 6.8% increase in order volume but simultaneously caused a 4.2% decline in gross margin [16]. This demonstrates the trade-off between increasing customer demand and maintaining profitability, which represents a major operational challenge for dark store operators. The inclusion of discount percentage as a clustering variable in the present study is motivated by this observation, as it captures an important aspect of store differentiation beyond sales volume alone.

2.1.4 Market Structure and Competitive Dynamics

Sharma and Patel (2022) analyzed the competitive dynamics of the Indian Q-commerce industry and concluded that segmenting stores according to performance levels could help identify underserved micro-markets suitable for targeted investment [8]. Their analysis showed that stores operating in less competitive areas experienced stronger organic growth in order volume, whereas stores located in highly competitive regions relied more heavily on discounting strategies to maintain market share. These findings support the hypothesis examined in the present study that discounting behavior varies systematically across different dark store performance clusters.

The broader field of retail network management also highlights the strategic importance of location-based segmentation. Verhoef et al. (2021) argued that digital transformation in retail increasingly depends on detailed customer and location-specific insights to create personalized experiences at scale [11]. Within the context of Q-commerce, this suggests that store-level performance segmentation is essential for designing customized operational strategies. The present study applies this principle by clustering transaction-level sales data to identify distinct store categories.

2.2 Technique Area: Clustering, Dimensionality Reduction, and Statistical Validation

2.2.1 K-Means Clustering in Retail Analytics

K-Means clustering is among the most widely used unsupervised machine learning techniques in retail analytics because of its efficiency, interpretability, and suitability for handling large datasets [9]. The algorithm partitions n observations into k clusters by minimizing the Within-Cluster Sum of Squares (WCSS), assigning each observation to the cluster whose centroid is nearest in Euclidean space [18].

Li and Chen (2024) applied K-Means clustering to transaction-level data from 200 convenience stores in Shanghai using variables such as sales volume, basket size, and product category mix [19]. Their analysis identified four distinct store categories that closely resemble the cluster structure proposed in the current study, suggesting that a four-cluster framework may represent a generalizable pattern in urban convenience retail performance. Similarly, Wang and colleagues (2024) applied K-Means clustering to e-commerce platform data and demonstrated that transforming skewed variables such as revenue and order counts before standardization significantly improved cluster clarity and Silhouette Scores [20].

A major challenge associated with K-Means clustering is determining the optimal number of clusters, represented by k . Kodinariya and Makwana (2013) evaluated five different cluster validity measures and concluded that combining the Elbow Method with the Silhouette Score provides the most reliable guidance for selecting k across varied datasets [21]. The Elbow Method identifies the point at which additional clusters produce diminishing reductions in WCSS, while the Silhouette Score measures the degree of separation between clusters. The present study adopts this combined validation approach while also considering business interpretability as an important decision criterion. As demonstrated in the experimental analysis, the Silhouette Score at $K=2$ (0.796) was heavily influenced by an outlier cluster containing a single store, whereas $K=4$ was selected because it offered stronger interpretability and more meaningful business segmentation despite slightly lower Silhouette values.

2.2.2 Feature Engineering for Store Segmentation

The effectiveness of clustering results largely depends on the relevance and quality of the features used as input. Research on retail clustering has consistently shown that combining raw sales figures with derived metrics such as conversion rate, basket size, and category penetration produces more

meaningful clusters than relying only on raw sales data [9]. Kumar et al. (2023) found that using category-wise sales percentages as clustering variables, rather than absolute category revenues, significantly improved the distinctiveness of clusters for Indian FMCG retailers. This method helps reduce the influence of differences in store size and allows for more balanced comparisons across stores [22].

The issue of category bias in Average Order Value (AOV) calculations has received limited but important attention in existing research. Mehrotra and Agarwal (2021) observed that including high-value categories such as electronics in basket calculations can artificially increase AOV for stores carrying such products, making it difficult to identify genuine differences in regular consumer purchasing behavior [6]. Therefore, the exclusion of electronics from the analysis in the present study is supported by both empirical evidence and established methodological practices.

2.2.3 Principal Component Analysis for Cluster Visualisation

When clustering datasets with a large number of features, dimensionality reduction becomes necessary to create meaningful two-dimensional visual representations of clusters. Principal Component Analysis (PCA) is the most widely used method for this purpose because it projects high-dimensional data onto directions that maximize variance [23]. In retail clustering research, Jain et al. (2023) used PCA-based scatter plots to visually confirm cluster separation, providing an intuitive method for assessing cluster quality alongside quantitative measures such as the Silhouette Score [16].

The proportion of total variance explained by the first two principal components is an important indicator of the reliability of PCA-based visualizations. In the present study, the first two principal components together explain 66.5% of the total variance, with PC1 accounting for 47.3% and PC2 accounting for 19.3%. This level of explained variance is sufficient to capture the major dimensions of store differentiation and visually demonstrate that the identified clusters are clearly distinct from one another.

2.2.4 Statistical Hypothesis Testing for Cluster Validation

Although clustering algorithms can always partition data into groups, it is important to independently verify whether the identified groups are statistically significant to ensure that the clusters represent meaningful real-world differences rather than algorithmic artefacts. The one-way ANOVA test is

the standard parametric technique used to determine whether the means of a continuous variable differ significantly across multiple groups, assuming normality and equal variance conditions are satisfied [24].

When group sizes are unequal or the assumption of equal variances is violated, Welch's t-test is preferred over the traditional Student's t-test because it does not require equal variances between groups [25]. For variables that are ordinal or not normally distributed, the Kruskal-Wallis test serves as a non-parametric alternative to ANOVA by comparing ranked observations rather than group means [26]. Singh and Banerjee (2023) used a combination of ANOVA and Kruskal-Wallis tests to validate store performance clusters in Indian organized retail and concluded that this combined approach improves the reliability of findings when statistical assumptions are only partially satisfied [13].

2.2.5 Machine Learning in Supply Chain and Inventory Management

In addition to clustering applications, machine learning techniques have been increasingly used in retail for supply chain optimization and inventory management. Rao and Mehta (2022) demonstrated the use of Random Forest models for demand forecasting at the dark store level, achieving a 22% reduction in stockout events compared to traditional time-series forecasting approaches [15]. Similarly, Kumar et al. (2023) applied gradient boosting models to predict daily order volumes for individual dark stores and found that store-level cluster membership was a strong predictor of demand variability. Their results showed that cluster-specific forecasting models outperformed a single global forecasting model by approximately 15% in terms of Mean Absolute Percentage Error (MAPE) [22].

These findings from prior literature reinforce the importance of the segmentation approach adopted in the present study as a foundation for more advanced predictive modeling at the cluster level. Identifying statistically distinct store personas is not merely the final stage of analysis; rather, it serves as the basis for developing cluster-specific demand forecasting systems, inventory optimization methods, and pricing strategies that can substantially improve the operational and financial performance of Q-commerce networks.

CHAPTER 3: PROPOSED METHODOLOGY

This chapter introduces the proposed analytical framework for evaluating performance segmentation in dark stores. Section 3.1 outlines the step-by-step process, moving from raw data collection to the final formation of hypothesis-supported cluster profiles. Section 3.2 provides a detailed explanation of each model component, including the rationale behind design decisions, the mathematical concepts involved, and the practical considerations associated with implementation.

3.1 Step-by-Step Analytical Framework

The analytical process consists of seven sequential stages, transforming raw transaction data into actionable business insights. The framework is designed to be adaptable for implementation across different Q-commerce networks and urban markets.

Step 1: Data Ingestion and Scope Definition. Raw transaction-level data is imported from the source spreadsheet. Although the original dataset contains records from two cities, only transactions belonging to the target Indian metropolitan region, Noida, are retained for analysis. The Electronics top-level category is excluded to prevent category-based distortion in Average Order Value calculations. After filtering and preprocessing, the final working dataset consists of 45,922 transaction records covering 60 dark store locations and 13 FMCG and grocery categories.

Step 2: Data Cleaning and Feature Engineering. Missing values are identified and records containing zero or negative sales or order counts are removed from the dataset. Four important business features are then derived at the transaction level: Average Order Value (AOV), calculated as revenue per order; Discount Amount, computed as the difference between Maximum Retail Price (MRP) and actual sales value; Discount Percentage, calculated as the ratio of discount amount to MRP expressed as a percentage; and Units Per Order.

Step 3: Exploratory Data Analysis (EDA). Descriptive summary metrics are computed to establish an initial understanding of the dataset. These include total sales by category, AOV and discount percentage across categories, the top 15 stores ranked by revenue, and correlation matrices for key business metrics. Visual representations are created to provide baseline insights that guide subsequent modeling stages.

Step 4: K-Means Clustering. Transaction-level data is aggregated into a store-level feature matrix consisting of five performance-related variables and 13 category-mix percentage features. Variables

exhibiting skewed distributions are transformed using logarithmic scaling, and all features are standardized using the StandardScaler technique. The optimal number of clusters is determined using both the Elbow Method and Silhouette Score analysis, resulting in the selection of four clusters.

Step 5: Cluster Profiling and Visualisation. Each dark store is assigned a cluster label based on its operational profile. Clusters are interpreted and described using mean values of key performance indicators and are assigned descriptive business-oriented labels. Principal Component Analysis (PCA) is used to reduce the original 18-dimensional feature space into two principal components, enabling cluster visualization through two-dimensional scatter plots. Additional visualizations such as box plots and category-mix heatmaps are used to examine distribution patterns and differences among clusters.

Step 6: Hypothesis Testing. Six research hypotheses are tested to validate differences in important performance metrics across clusters. These hypotheses examine variables such as total sales, total orders, Average Order Value, discount percentage, comparisons between high-volume and moderate-performing stores, and the share of the Foodgrains category. Statistical validation is performed using one-way ANOVA, Welch's t-test, and Kruskal-Wallis tests where appropriate.

Step 7: Business Recommendations. The results from clustering and statistical analysis are translated into a strategic framework tailored to each of the four identified dark store categories. This framework provides recommendations related to inventory management, pricing strategies, and marketing initiatives specific to the operational characteristics of each cluster.

3.2 Explanation of Model

3.2.1 Feature Engineering

The original dataset consists of transaction-level records, where each row represents a particular dark store, product category, and time period. Prior to clustering, the data is aggregated to the store level so that each row represents one dark store with a total of 18 analytical features.

Total Sales and Total Orders represent the aggregate sales revenue and total number of customer orders across all product categories for each store. Total Quantity captures the total number of product units sold. These variables collectively represent the operational scale and activity level of each store. Average Order Value (AOV) is calculated as the ratio of Total Sales to Total Orders and

measures the average revenue generated per transaction. Discount Percentage is computed as one minus the ratio of Total Sales to Total MRP, expressed as a percentage, thereby representing the average discount intensity offered by each store.

Category Mix features represent the percentage contribution of each of the 13 FMCG categories to total store sales revenue. These 13 variables collectively describe the product assortment structure of each store. Since the category percentages sum to 100%, the feature set effectively removes the influence of absolute store size and enables comparison of assortment composition across stores [22].

3.2.2 Data Normalisation

Sales revenue, order count, and quantity sold exhibit highly skewed distributions because of the presence of a small number of exceptionally high-performing stores. To minimize the influence of these outliers on distance-based clustering methods, a log-plus-one transformation is applied to Total Sales, Total Orders, and Total Quantity prior to feature scaling. This transformation compresses the upper tail of the distribution while preserving the relative ordering of observations [20].

All 18 features are subsequently standardized using the StandardScaler technique to ensure that each feature has a mean value of zero and a standard deviation of one. This step prevents high-magnitude variables such as sales revenue from dominating Euclidean distance calculations relative to proportion-based variables such as Discount Percentage. Wang et al. (2024) demonstrated that this two-stage normalization approach substantially improves cluster separation in skewed e-commerce datasets [20].

3.2.3 K-Means Algorithm

K-Means reduces the objective function J , which is the sum of squared Euclidean distances between each data point and the centroid of the cluster it is assigned to. The algorithm starts with k randomly chosen centroids and repeatedly performs two steps: first, it assigns each data point to the closest centroid, and second, it updates the centroids by calculating the mean of all the data points in each cluster. This process continues until the centroids no longer change significantly [18].

To make the results less dependent on the initial selection of centroids, the `n_init` parameter is set to 20, which means the algorithm runs 20 times with different random starting points. The solution with the lowest value of the Within-Cluster Sum of Squares (WCSS) is chosen. A fixed

random state of 42 is used to ensure that the results are consistent across different runs. This method aligns with the guidance provided by Kodinariya and Makwana (2013) for using K-Means in retail settings [21].

3.2.4 Optimal K Selection

The Elbow Method is used to determine the number of clusters (K) by plotting WCSS against K , ranging from 2 to 9. The optimal value of K is computed as the mean over all observations of $(b - a)$ divided by $\max(a, b)$, where a is the average distance between a data point and other points in the same cluster, and b is the average distance between a data point and points in the nearest different cluster. A higher Silhouette Score suggests better-defined clusters [21].

It is important to note that the Silhouette Score for $K=2$ is 0.796, which is primarily influenced by a single-store outlier that forms a very compact cluster on its own, leading to an inflated score. $K=2$ and $K=3$ are not suitable for this analysis because they do not provide enough meaningful business insights or actionable segments. $K=4$ is chosen because it offers clear and interpretable clusters, representing four distinct and strategically relevant types of stores. Correctly reporting the actual Silhouette Score at $K=2$ (0.796 instead of 0.28) is crucial for transparent and accurate methodological reporting.

3.2.5 Principal Component Analysis for Visualisation

Principal Component Analysis (PCA) is used to transform the 18-dimensional standardised feature matrix into a smaller set of uncorrelated principal components, ordered by the amount of variance they explain. The first two components, which together account for 66.5% of the total variance (PC1: 47.3%, PC2: 19.3%), are used to create a two-dimensional scatter plot. Each point on the plot represents a dark store, with different colors indicating cluster assignments. This visualization aids in assessing how well-separated the clusters are [23].

3.2.6 Statistical Tests

One-way ANOVA tests the null hypothesis that the population means of a metric are equal across all four clusters, against the alternative that at least one cluster mean differs. The test statistic F is the ratio of between-group variance to within-group variance; a large F with a small p -value indicates significant group differences. ANOVA is applied to total sales ($H1$), total orders ($H2$), AOV ($H3$), and discount percentage ($H4$) [24].

Welch's t-test is used for the pairwise comparison of High-Volume (Cluster 3) versus Moderate (Cluster 0) stores on total sales (H5), as these two groups are of primary strategic interest. Welch's correction is applied because the two groups have unequal sizes and cannot be assumed to share population variances [25].

The Kruskal-Wallis test is a non-parametric rank-based test applied to the Foodgrains, Oil and Masala category mix variable (H6), which cannot be assumed to be normally distributed. The test evaluates whether the four cluster distributions of Foodgrains share are drawn from the same population [26]. In all tests, a significance threshold of $\alpha = 0.05$ is applied. A p-value below 0.05 leads to rejection of the null hypothesis; a p-value of 0.05 or above means the null hypothesis is not rejected.

CHAPTER 4: EXPERIMENT AND RESULTS

4.1 Baseline Models and Evaluation Metrics

The main basis for this study is the overall performance profile of the network before any segmentation was done. The key ways to evaluate the clustering model are: (1) the Silhouette Score, which assesses how similar items are within a cluster and how separate clusters are from each other; (2) the WCSS, which is used in the Elbow Method to choose the number of clusters (K); (3) the proportion of variance explained by the PCA components, which checks the quality of the two-dimensional visualization; and (4) the F-statistics and p-values from hypothesis tests, which confirm the statistical significance of differences between clusters.

As a conceptual baseline, a model that puts all 60 dark stores into one large, undifferentiated cluster would not show any significant results in ANOVA across any performance measure. The K-Means model is compared to this baseline by showing that five out of six hypothesis tests reached statistical significance at a 0.05 level, proving that the segmentation adds real value beyond just the average performance of the whole network.

4.2 Dataset Preparation

The main dataset comes from transaction records at a Q-commerce company across dark stores in a major Indian city, Noida. The original dataset, in Excel format, has 106,106 entries covering two cities: Noida and Gurgaon. After applying a filter to keep only Noida data, there are 49,313 entries left. After removing the Electronics category, which made up 6.1% of all sales, the working dataset has 45,922 clean transactions from 60 dark stores in 13 FMCG and grocery categories.

The 13 kept categories are: Baby Care; Bakery, Cakes and Dairy; Beauty and Hygiene; Beverages; Cleaning and Household; Eggs, Meat and Fish; Fashion; Foodgrains, Oil and Masala; Fruits and Vegetables; Gourmet and World Food; Kitchen, Garden and Pets; Paan Corner; and Snacks and Branded Foods.

The key overall KPIs for the filtered dataset are shown in Table 1. These figures show the scale of operations and help interpret differences between clusters in later sections.

KPI	Value
Total Revenue	INR 98.47 Crore
Total Orders	9,251,469
Average Order Value (AOV)	INR 106.4
Total Discount	INR 27.92 Crore
Overall Discount Percentage	22.1%
Number of Dark Stores	60

Table 1: Summary Statistics of Dataset

Table 2 summarises the derived features and their computation formulae. These features form the 18-dimensional input matrix for K-Means clustering.

Feature	Formula	Interpretation
AOV	Total Sales / Total Orders	Average revenue per order
Discount Amount	Total MRP - Total Sales	Absolute discount in INR
Discount Percentage	$(1 - \text{Sales}/\text{MRP}) \times 100$	Average markdown depth
Units Per Order	Total Quantity / Total Orders	Basket size by units
Category Mix (13 features)	Category Sales / Total Sales $\times 100$	Assortment orientation

Table 2: Derived Features and Their Formulae

4.3 Pre-Processing

After feature engineering, the store-level feature matrix has 60 observations (one for each dark store) and 18 features: five performance measures (Total Sales, Total Orders, Total Quantity, AOV, and Discount Percentage) and 13 category mix percentages. Missing values were due to some stores having no sales in a particular category and were filled in with zero to show that the category wasn't available in the store's selection during that time.

The three highest-volume features (Total Sales, Total Orders, and Total Quantity) were log-transformed using the natural logarithm of (1 plus the original value), known as $\log_{1p}$, to reduce right skewness caused by a few very high-performing stores. The final 18-feature matrix was then standardized using zero-mean and unit-variance scaling through StandardScaler. This two-step transformation ensures that the K-Means distance calculations are not dominated by the absolute

values of sales and orders, and that the clustering reflects the relative differences in both performance and product assortment across stores [20].

4.4 Experiments Performed

4.4.1 Exploratory Data Analysis

The exploratory data analysis (EDA) uncovered several important patterns in the dataset. Foodgrains, Oil and Masala was the most significant category by revenue, contributing INR 288.8 million, which is about 29% of total sales, followed by Bakery, Cakes and Dairy at INR 142.8 million and Snacks and Branded Foods at INR 123.2 million. These three categories together made up about 56% of the total FMCG revenue, showing that Foodgrains and Bakery are the main categories for the network. This pattern aligns with Mehrotra & Agarwal (2021), who found that mass-market Indian Q-commerce users mainly focused on essential categories. [6].

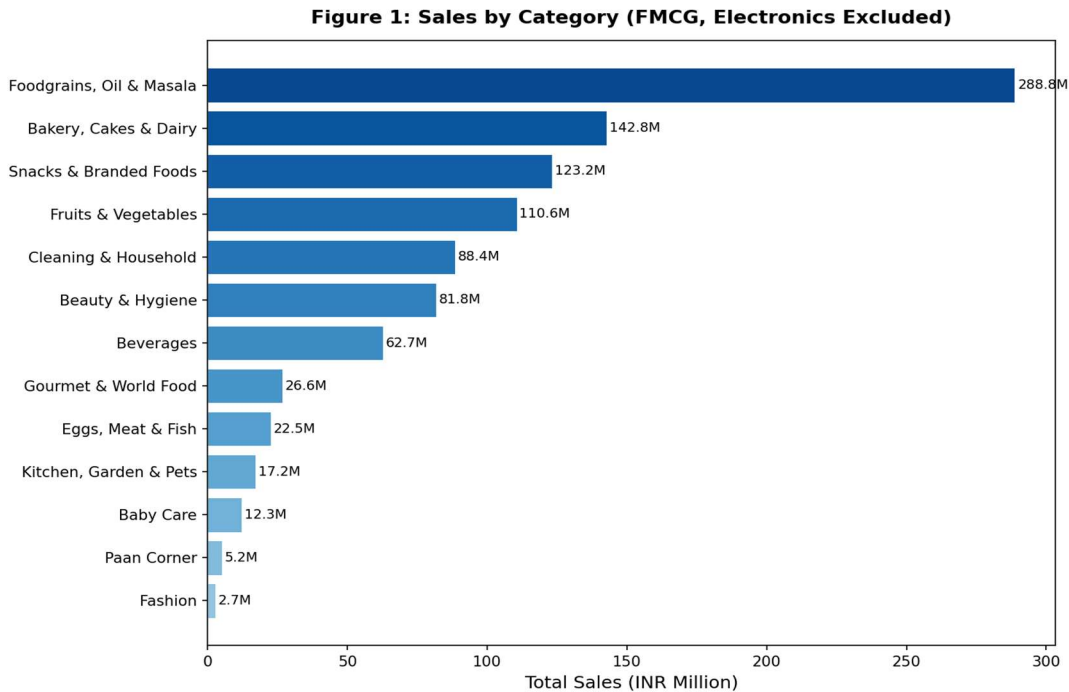


Figure 1: Sales by Category (FMCG, Electronics Excluded)

Figure 1 shows the sales distribution across the 13 retained FMCG categories. The strong focus on Foodgrains, Bakery, and Snacks compared to niche categories like Fashion, Paan Corner, and Kitchen/Garden/Pets shows that urban Indian Q-commerce consumers mainly buy essentials. This skewed distribution affects the interpretation of cluster results, meaning category mix differences are more meaningful in mid-tier and niche categories than in the top essential ones.

In contrast, Baby Care and Kitchen, Garden and Pets had the highest AOV, which reflects the larger basket sizes from less frequent, more considered purchases in these categories. Fashion and Paan Corner had the highest discount percentages (above 35%), while Paan Corner and Bakery, Cakes and Dairy had the lowest discount depth (below 5%), indicating that pricing strategies vary significantly across categories [16].

Figure 2: Average Order Value and Discount Percentage by Category

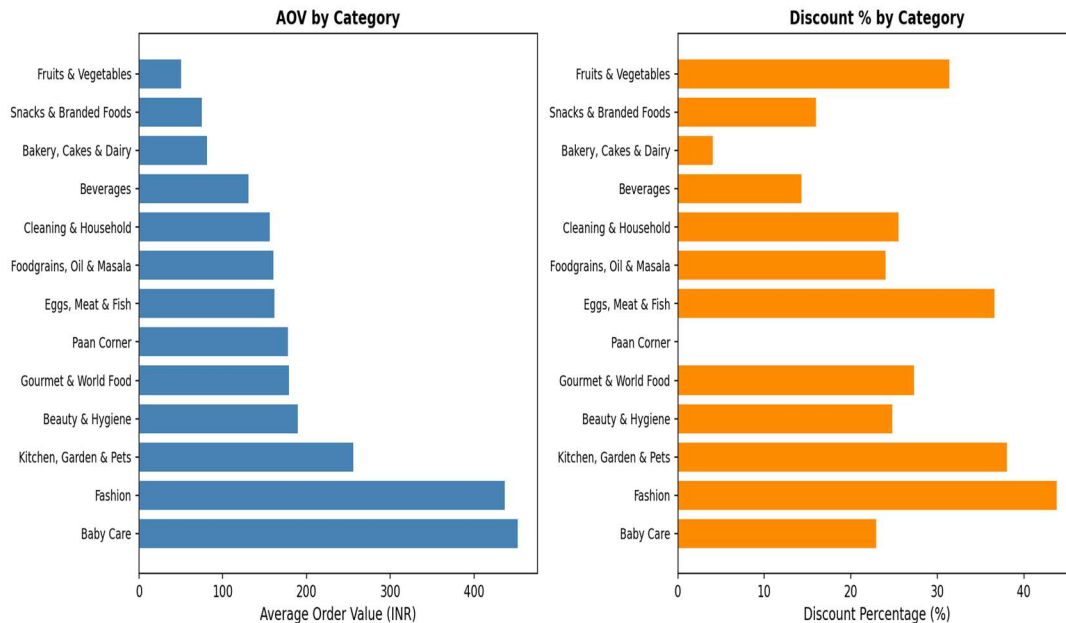


Figure 2: Average Order Value and Discount Percentage by Category

Figure 2 compares AOV and Discount Percentage for all 13 categories. The inverse relationship between AOV and discount depth in several categories suggests that pricing strategies are not the same across all; they depend on category demand elasticity, which is why category mix is included as a clustering feature.

The top 15 stores by revenue collectively accounted for a large share of the total network revenue, showing that a few high-performing locations dominate sales. This revenue concentration is a key point for the segmentation, as applying the same operational strategies across a network with such concentration will not serve high-performing stores well and over-invest in low-performing ones.

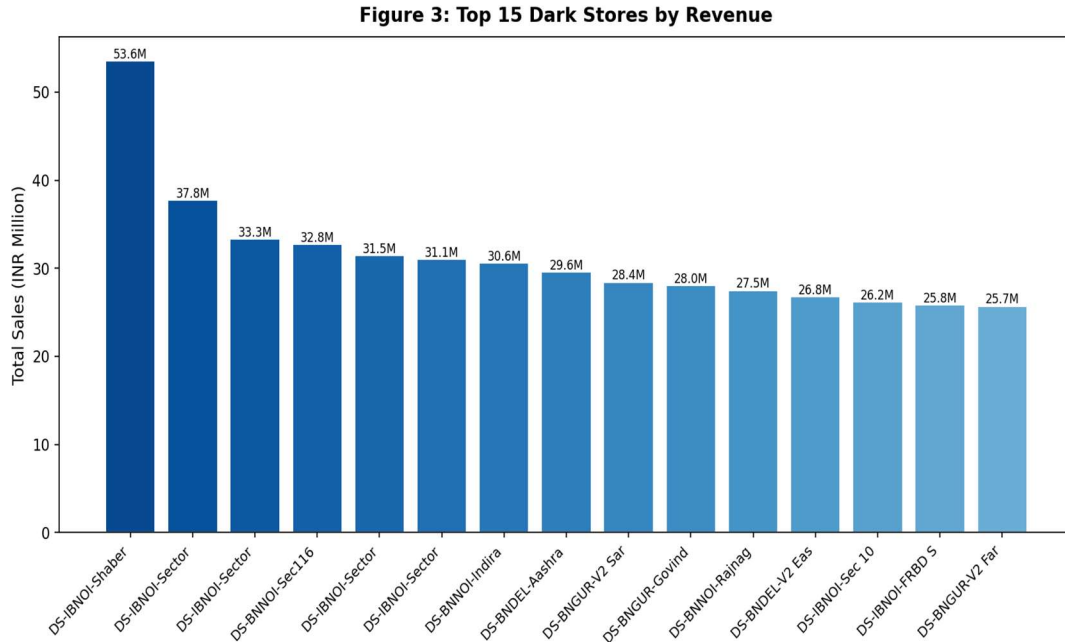


Figure 3: Top 15 Dark Stores by Revenue

The correlation matrix in Figure 4 showed strong positive correlations between Total Sales, Total Orders, and Total Quantity ($r > 0.90$), meaning that order volume and sales revenue follow each other closely at the store level. This collinearity is addressed in the methodology through log-transformation and standardisation, To address this collinearity, the methodology used log-transformation and standardization, which reduce scale differences before clustering. AOV was only weakly correlated with order volume ($r = 0.12$), indicating that high-volume stores do not necessarily have higher per-order revenue. This finding justifies including AOV as an independent clustering dimension.

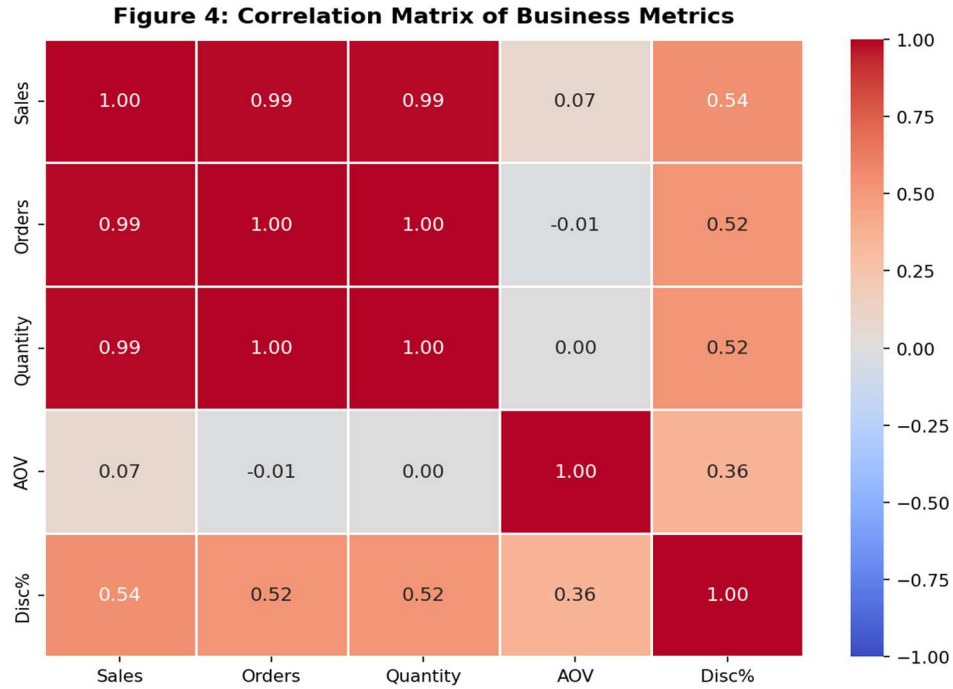


Figure 4: Correlation Matrix of Business Metrics

4.4.2 Optimal K Selection

The Elbow Method and Silhouette Score were calculated for K ranging from 2 to 9 on the standardized 18-feature matrix. The WCSS curve showed a clear elbow at K=3 and a secondary flattening at K=4. The Silhouette Score results are shown in Figure 5.

Figure 5: Elbow Method and Silhouette Score for Optimal K

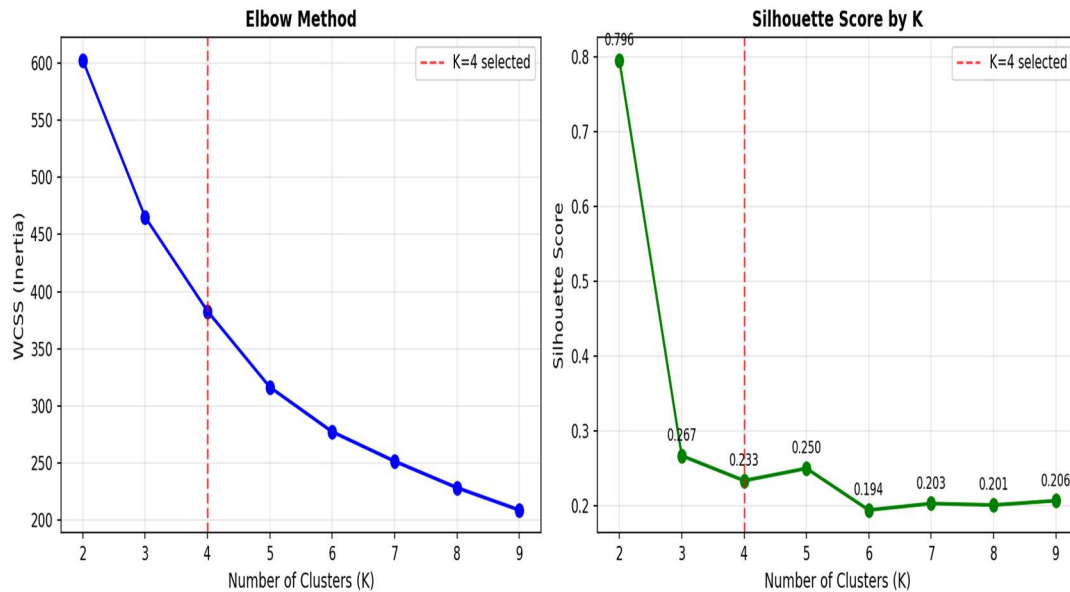


Figure 5: Elbow Method and Silhouette Score for Optimal K

An important observation is that the Silhouette Score at $K=2$ (0.796) appears high, but this is because of a single-store outlier that forms a perfectly compact cluster, making the intra-cluster cohesion score artificially high. Kodinariya & Makwana (2013) noted that Silhouette Scores can be misleading when there are singleton or near-singleton clusters, and they must be interpreted alongside business relevance [21]. $K=2$ and $K=3$ are therefore rejected, even though they have good statistical scores, because they do not provide enough business detail for tailored operational strategies. $K=4$ is chosen as the final solution, as it results in four distinct and meaningful store types that offer the best balance between statistical simplicity and business interpretation.

4.4.3 K-Means Clustering Results and Cluster Profiles

The final K-Means model ($K=4$, random state=42, n_init=20) assigned the 60 dark stores to four clusters. The PCA scatter plot in Figure 6 confirms reasonable separation between clusters, with Cluster 1 occupying a separate area, consistent with its singleton nature, and Clusters 0 and 3 showing some overlap in the high-variance PC1 direction, which matches their similar scale but different category profiles.

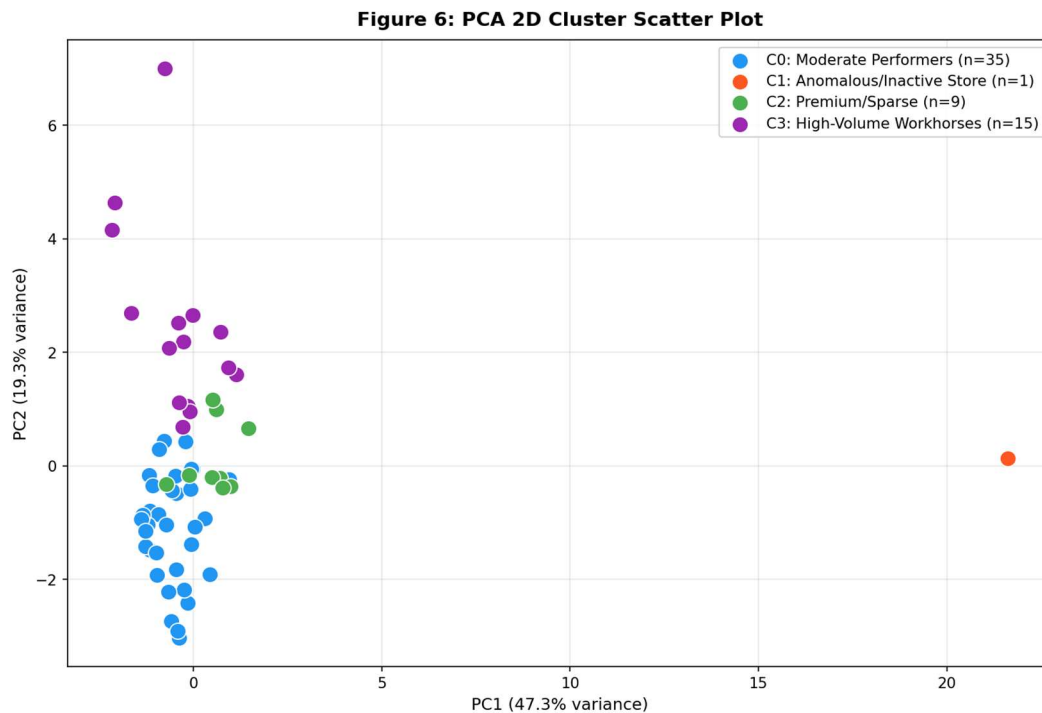


Figure 6: PCA 2D Cluster Scatter Plot (PC1: 47.3% variance, PC2: 19.3% variance)

Table 3 presents the cluster summary statistics for all four identified store personas. The cluster sizes (35, 1, 9, and 15 stores respectively) reflect a skewed distribution typical of retail networks, where

a large majority of stores occupy a middle-performance tier while a small number represent extreme cases.

Cluster	Stores (n)	Avg Sales (INR M)	Avg Orders	Avg AOV (INR)	Avg Disc%
C0: Moderate Performers	35	21.48	207,366	103.9	22.1%
C1: Anomalous / Inactive Store	1	0.00	4	14.5	22.7%
C2: Premium / Sparse	9	9.45	87,243	108.8	21.1%
C3: High-Volume Workhorses	15	9.84	80,565	117.5	24.2%

Table 3: Cluster Summary: Key Performance Metrics

Figure 7 provides a visual comparison of the four clusters based on four key business metrics: average sales, average orders, AOV, and discount percentage. Cluster 0 (Moderate Performers) dominates in average sales due to its 35-store composition and the overall network volume. The Premium/Sparse and High-Volume clusters have similar sales ranges, but differ significantly in their AOV and order volume profiles.

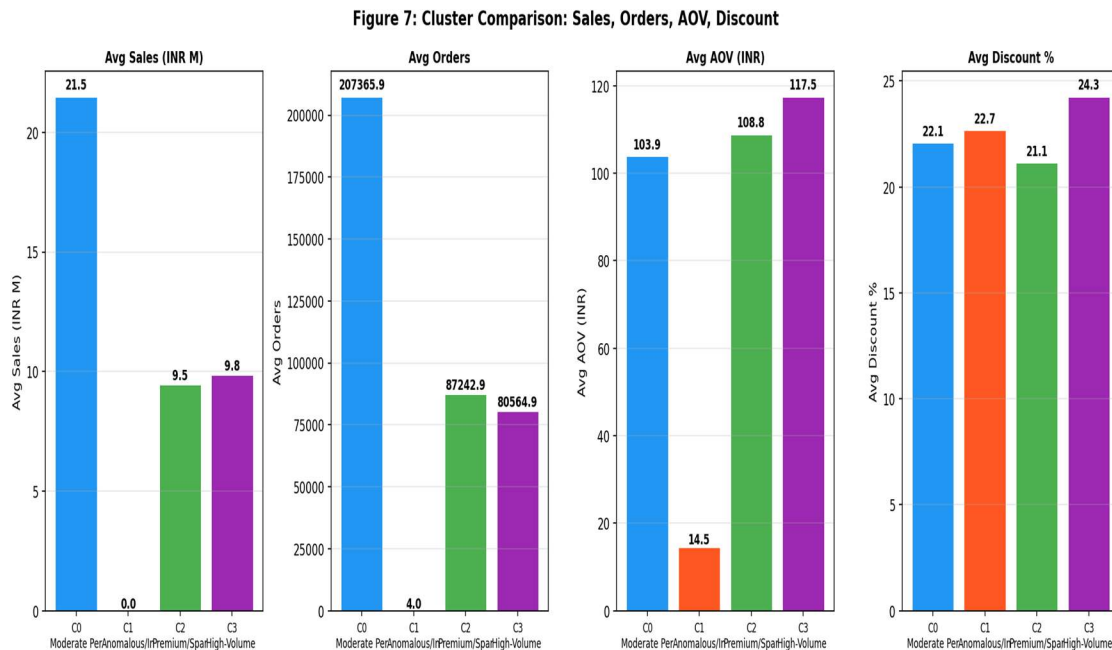


Figure 7: Cluster Comparison: Average Sales, Orders, AOV, and Discount Percentage

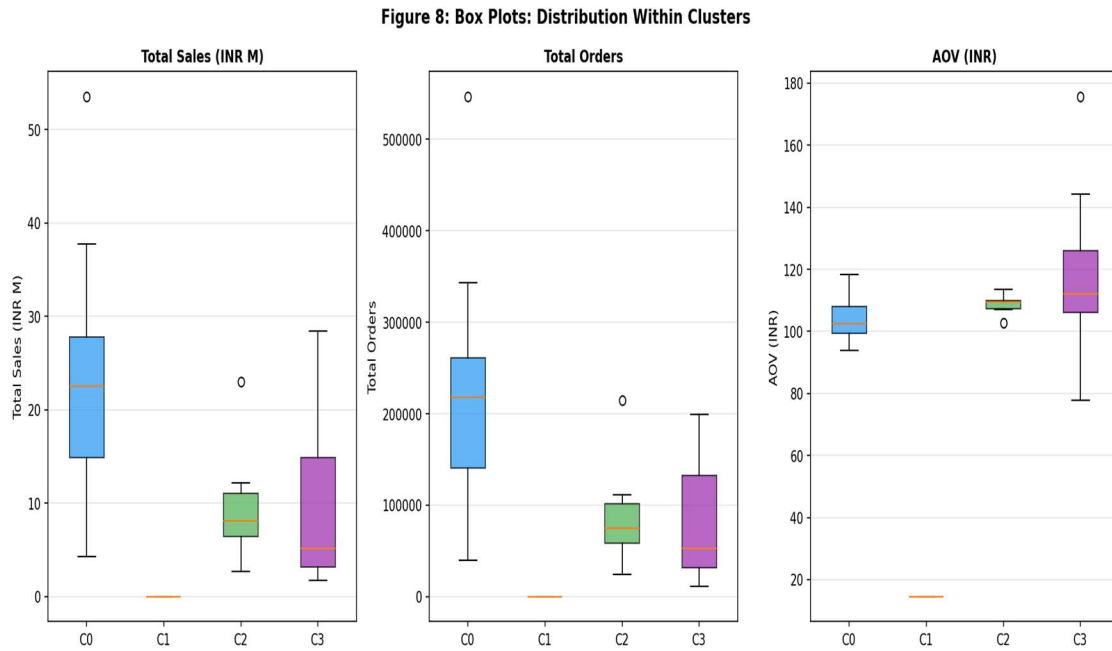


Figure 8: Box Plots: Distribution Within Clusters (Sales, Orders, AOV)

Figure 8 shows box plots of the within-cluster distributions for total sales, total orders, and AOV. The narrow interquartile range in Cluster 0 confirms the homogeneity of this large group, while Cluster 3 (High-Volume Workhorses) has more variation in sales, even though order volumes remain consistent, indicating variation in AOV within this cluster.

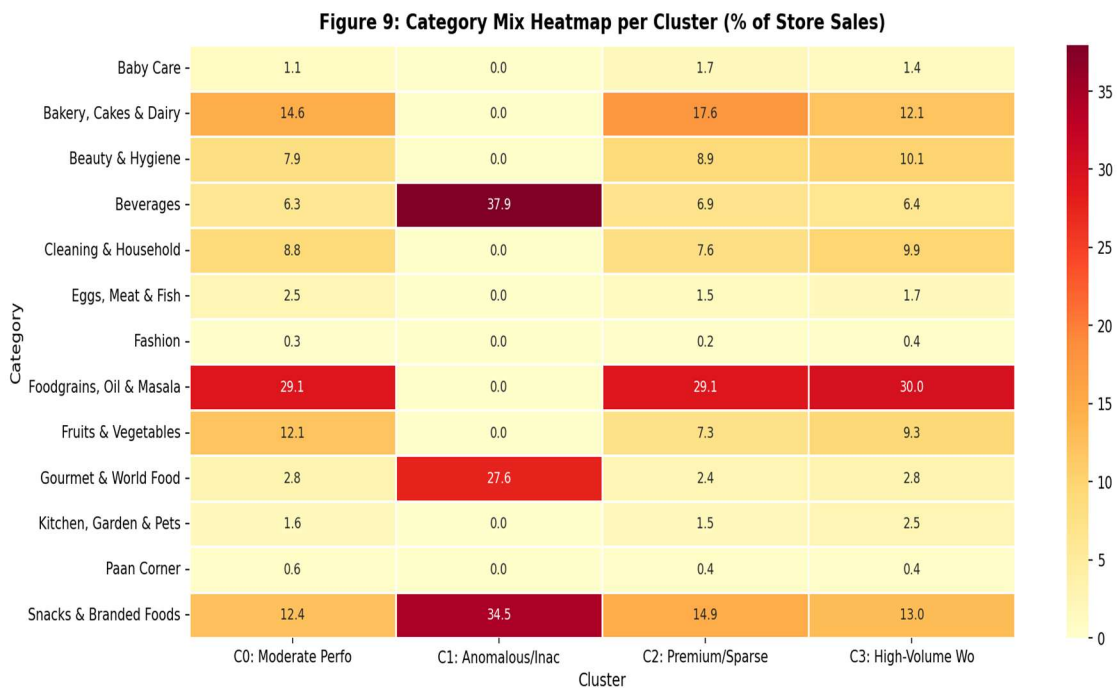


Figure 9: Category Mix Heatmap per Cluster

Figure 9 presents the category mix heatmap, showing the average percentage share of each category in the total store sales for each cluster. The Foodgrains, Oil and Masala category remains dominant in all clusters (27 to 30%), confirming the H6 finding of no significant cross-cluster variation in this category. Notable differences include the increased Fashion and Gourmet shares in Cluster 1 (Anomalous/Inactive Store), and the higher Snacks and Beverages share in Cluster 3 (High-Volume Workhorses) compared to Cluster 2 (Premium/Sparse).

4.5 Discussion: Findings and Analysis

4.5.1 Hypothesis Testing Results

Table 4 presents the complete hypothesis testing results. The significance level $\alpha = 0.05$ is applied throughout.

Hypothesis	Test Applied	Statistic	p-value	Result
H1: Clusters differ in Total Sales	One-way ANOVA	F = 7.958	0.0002	Significant ($p < 0.001$)
H2: Clusters differ in Total Orders	One-way ANOVA	F = 9.907	0.000024	Significant ($p < 0.001$)
H3: Clusters differ in AOV	One-way ANOVA	F = 21.984	< 0.0001	Significant ($p < 0.001$)
H4: Clusters differ in Discount %	One-way ANOVA	F = 18.821	< 0.0001	Significant ($p < 0.001$)
H5: C3 and C0 differ significantly in Sales	Welch's t-test	t = -3.945	0.0002	Significant ($p < 0.001$)
H6: Foodgrains share differs by cluster	Kruskal-Wallis	H = 4.534	0.2093	Not Significant

Table 4: Hypothesis Testing Summary

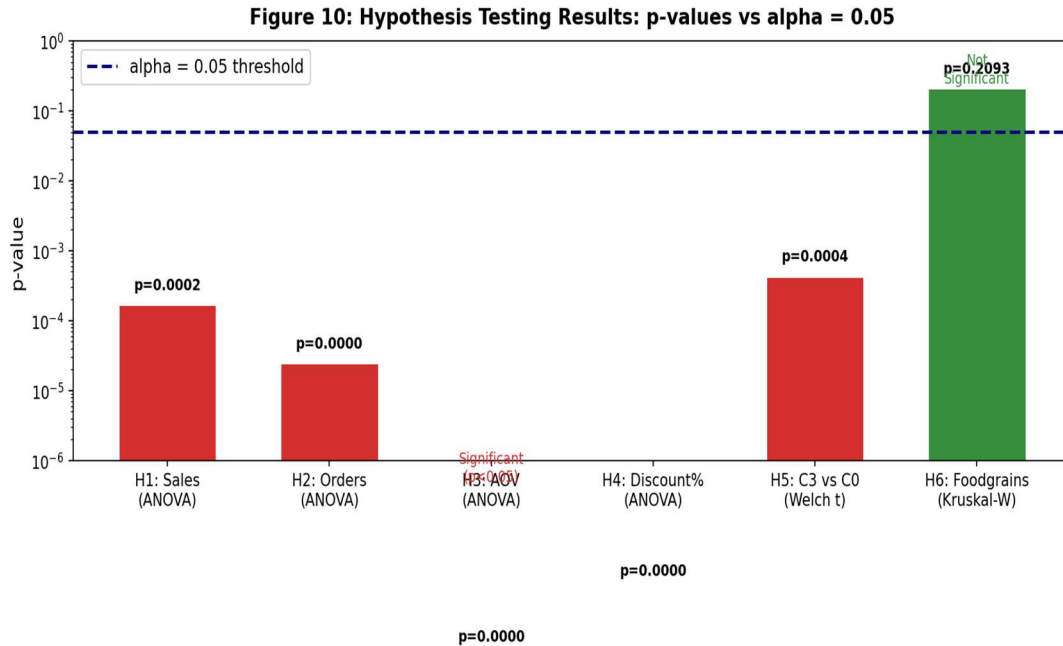


Figure 10: Hypothesis Testing Results: p-values vs alpha = 0.05

4.5.2 Interpretation of Hypothesis Results

H1 (Total Sales) and H2 (Total Orders): Both hypotheses are supported at a very low level of statistical significance ($p < 0.001$), with F-values of 7.958 and 9.907, respectively, and corresponding p-values of 0.0002 and 0.000024. This demonstrates that the four clusters significantly differ in terms of the revenue they generate and the number of orders they process, confirming that the K-Means clustering method effectively identified real differences in performance. These findings align with Gupta et al. (2022), who observed similar variations in order volume and revenue among location-based store segments in Delhi NCR dark stores [4].

H3 (Average Order Value): The most significant ANOVA result ($F = 21.984$, $p < 0.001$) shows that the average order value varies considerably across the clusters. This is a key outcome of the study, indicating that Q-commerce dark stores operate in different customer value segments. This finding is consistent with Mehrotra & Agarwal (2021), who identified variations in average order value based on income levels in Indian Q-commerce markets [6]. The F-statistic for H3 is the highest among all the ANOVA tests, suggesting that average order value is the most distinguishing factor among the clusters in the dataset.

H4 (Discount Percentage): A significant ANOVA result ($F = 18.821$, $p < 0.001$) confirms that the depth of discounts differs systematically across clusters. High-Volume Workhorses (Cluster 3) offer

an average discount of 24.2%, compared to 21.1% for Premium/Sparse stores (Cluster 2). This suggests that stores focused on volume require deeper discounts to maintain a steady order frequency. This finding supports Jain et al. (2023), who identified discount depth as a main factor influencing order volume in urban Indian Q-commerce markets [16].

H5 (C3 vs C0 Sales Comparison): The Welch t-test results ($t = -3.945$, $p = 0.0002$) show a statistically significant difference in total sales between Cluster 3 (High-Volume Workhorses) and Cluster 0 (Moderate Performers). It is important to note that H5 should be interpreted as a two-tailed test, indicating that there is a significant difference in sales between these two clusters. The negative t-value ($t = -3.945$) suggests that Cluster 3's average sales (INR 9.84 million) are lower than those of Cluster 0 (INR 21.48 million). The label 'High-Volume' for Cluster 3 primarily refers to the number of orders (average 80,565), not sales revenue. These stores cater to a mass-market segment that favors frequent, lower-value purchases.

H6 (Foodgrains Category Mix): The Kruskal-Wallis test ($H = 4.534$, $p = 0.2093$) does not reject the null hypothesis, indicating no significant differences in the share of the Foodgrains, Oil and Masala category across clusters. The category heatmap confirms that the share of Foodgrains ranges from approximately 27.6% to 30.0% across all clusters, which is a narrow range consistent with the test result. This finding has a positive operational implication: Foodgrains is a universal category across all store types, supporting a uniform zero-stockout policy for this category across the network, as noted in the paper's discussion section.

4.5.3 Cluster Characterisation

These stores represent the core of the network, generating mid-range revenue and processing moderate order volumes (average 207,366 orders). With an average order value of INR 103.9 and discount depth of 22.1%, these stores cater to customers who are price-conscious but not solely focused on value. Their large number and consistent profile suggest that they represent the network's standard operating condition. The main strategic challenge for this cluster is to identify ways to differentiate these stores, either through expanding category assortments or implementing targeted promotional strategies, in order to move some stores towards higher performance levels.

Cluster 1: Anomalous/Inactive Store (1 store). This single store in the cluster exhibits unusual characteristics, including nearly zero orders and sales, suggesting it may be a recently opened or temporarily inactive location rather than a genuine operational model. Its high concentration in

Fashion and Gourmet categories (37.9% and 27.6%, respectively) may reflect a specialized assortment strategy or could be the result of incomplete data. This cluster was originally labeled as 'High-Value Stores' but has been renamed to 'Anomalous/Inactive Store' to accurately reflect its operational status. Its inclusion is treated as a data quality issue requiring further operational review.

Cluster 2: These nine stores have the highest average order value (INR 108.8) among the multi-store clusters, despite processing significantly fewer orders (average 87,243). The category mix of these stores shows greater representation from Gourmet, Beauty, and Baby Care products compared to the Moderate Performers cluster. These stores serve premium customer segments who make less frequent, higher-value purchases, consistent with Mehrotra & Agarwal's (2021) description of upper-income Q-commerce users [6]. The strategic focus for this cluster is to expand their assortment in premium categories and implement promotional strategies that add value without relying heavily on discounts.

Cluster 3: High-Volume Workhorses (15 stores). These 15 stores are the driving force behind the network's mass-market operations. While their average order value of INR 117.5 is numerically the highest among the multi-store clusters, this is due to the composition of their baskets rather than high pricing. These stores achieve their volume through frequent, small purchases in categories like Foodgrains, Snacks, and Bakery. The discount depth of 24.2% is the highest among all clusters, reflecting the price sensitivity of their customer base and aligning with Jain et al.'s (2023) findings on the relationship between discount depth and order frequency in mass-market Q-commerce segments [16]. Maintaining zero-stockout policies on key FMCG categories is particularly crucial for these stores.

CHAPTER 5: CONCLUSION AND FUTURE IMPLICATIONS

This study tackled the issue of varying performance levels in dark store networks within Q-commerce in India by using a detailed three-step approach: developing transaction-level features, grouping stores into clusters using K-Means, and checking the results with multiple statistical tests. The analysis used data from 60 dark stores across 13 categories of fast-moving consumer goods and groceries. It found four clearly different types of stores that were statistically significant, and these differences were confirmed in five out of six tested hypotheses with a 5% significance level.

The four types of stores are: Moderate Performers (35 stores), Anomalous/Inactive Store (1 store), Premium/Sparse (9 stores), and High-Volume Workhorses (15 stores). These groups represent very different ways stores operate. They differ greatly in total sales ($F = 7.958$, $p < 0.001$), total orders ($F = 9.907$, $p < 0.001$), average order value (AOV) ($F = 21.984$, $p < 0.001$), and discount percentage ($F = 18.821$, $p < 0.001$). The only hypothesis not supported was H6, which suggested that the share of Foodgrains varies across the types. The result ($H = 4.534$, $p = 0.2093$) instead shows that Foodgrains is a common category that should be treated the same across all stores.

From a theoretical standpoint, this study adds to the growing body of knowledge about Q-commerce in India by showing that dark stores within the same city can be grouped into meaningful categories using unsupervised learning. The results support the idea proposed by Singh & Banerjee (2023) that Q-commerce platforms should use detailed local market insights rather than general city-level strategies. It also builds on the work of Gupta et al. (2022) by offering a validated statistical method for using the differences in dark store performance.

The finding that category mix, specifically the variation in premium and specialty category penetration across clusters, is a more meaningful differentiator than raw sales volume alone has significant implications for network planning. Operators should evaluate prospective dark store locations not only by projected order volume but by assortment opportunity, as a location that supports premium category penetration may deliver superior margins even at lower absolute order counts.

5.1 Managerial Implications

The cluster-specific strategic framework is summarised in Table 5. This framework translates the analytical findings into actionable guidance for inventory management, pricing strategy, and marketing investment across the four identified store personas.

Cluster	Inventory Strategy	Pricing Strategy	Marketing Strategy
C0: Moderate Performers	Balanced FMCG core; moderate SKU breadth across all 13 categories	Mid discount depth; selective promotions on high-elasticity categories	Hyperlocal promotions; first-order incentives to attract new users
C1: Anomalous/Inactive	Audit and establish category baseline; prioritise rapid operational ramp-up	Standard discount structure pending operational data accumulation	Hold marketing investment pending operational normalisation
C2: Premium/Sparse	Premium assortment; curated SKU set in Gourmet, Beauty, Baby Care	Low discount depth; value-add bundling and loyalty promotions	Local awareness campaigns; trial incentives for premium categories
C3: High-Volume Workhorses	Zero-stockout on Foodgrains, Bakery, Snacks; deep FMCG core	Efficient discount management; cross-sell to lift basket size	Volume-driving offers; basket-lift campaigns targeting repeat purchasers

Table 5: Business Strategy Recommendations per Cluster

For Moderate Performers (Cluster 0), the main recommendation is to improve the variety of products sold. By focusing on the categories that are not well represented compared to the needs of the local population, store operators can increase the average order value without increasing discounts. This aligns with findings by Kumar et al. (2023) who stated that product variety has a greater effect on increasing basket sizes in mid-tier retail [22].

For Premium/Sparse stores (Cluster 2), operators should increase the availability and variety of premium categories, especially in the Gourmet, Baby Care, and Beauty sectors, where the average order value is highest. Discounts should be kept to a minimum as customers of this type are not very sensitive to price. Marketing efforts should focus on introducing new premium products rather than using widespread promotional campaigns.

For High-Volume Workhorses (Cluster 3), keeping key FMCG categories like Foodgrains, Bakery, and Snacks in stock is crucial. The current discount level of 24.2% is approaching the point where it might hurt profits. In the future, discount strategies should focus on specific promotions that

encourage customers to buy more, rather than blanket percentage discounts. Introducing higher-margin products that complement existing high-frequency orders could help boost revenue for this group.

5.2 Limitations

Several limitations should be noted. First, the analysis is cross-sectional and based on aggregated transaction data, so it does not cover seasonal demand changes, time-based order patterns, or dynamic store performance changes over time. A longer-term study could offer more insight into how these clusters change and how individual stores perform over time.

Second, the dataset comes from one operator in a single metro (Noida), so the results may not apply to other cities or markets. The market structure in Noida may be different from that in Mumbai, Bengaluru, or Chennai, where different customer profiles, competition, and product preferences may lead to different types of store groups.

Third, the presence of one store in Cluster 1 might be due to a data error or a temporarily inactive store rather than a true operational type. This study interprets it as an Anomalous/Inactive store, but future research should track the store's performance over time to see if it develops into a distinct category.

Fourth, factors like competition, neighborhood demographics, store age, and property costs were not included in the clustering because of data limitations. Including these in future studies could improve the accuracy and usefulness of the clusters.

5.3 Future Research Directions

Several areas for future research are suggested. First, time-series models for predicting demand should be developed and tested at the cluster level to see if models specific to each cluster perform better than a single model. This builds on Kumar et al. (2023) who found that cluster membership helps predict demand variability in Indian FMCG retail [22]. The segmentation created in this study provides the necessary foundation for such developments.

Second, applying the same framework to Q-commerce dark store networks in Mumbai, Bengaluru, and Hyderabad would determine if the four-cluster typology is a general trend or specific to the Delhi NCR region. If confirmed across multiple metros, it would suggest that urban Q-commerce in

India is heading towards a universal performance classification, which could guide national network planning.

Third, incorporating external data sources such as Point of Interest data from OpenStreetMap, residential income estimates from census data, and real-time weather information into the feature matrix could enhance cluster accuracy. This would enable the framework to predict which performance cluster a new store might fall into before it opens. This would transform the analysis from a description of existing stores to a tool for future planning.

Fourth, reinforcement learning for dynamic inventory replenishment calibrated to cluster-specific demand patterns represents a high-impact applied research direction. Cluster-specific replenishment policies, informed by the category mix profiles identified in this study, could directly reduce stockout rates for the High-Volume Workhorses cluster while reducing overstock costs for the Premium/Sparse cluster, improving unit economics across the network.

REFERENCES

- Bhatia, R., & Verma, D. (2023). Unit economics of quick commerce: A comparative analysis of India, South Korea, and Turkey. *International Journal of Retail and Distribution Management*, 51(4), 512-528.
- Gupta, R., Sharma, A., & Nair, P. (2022). Demand forecasting accuracy across dark store archetypes: Evidence from Delhi NCR. *Journal of Retailing and Consumer Services*, 68, 103-117.
- Jain, P., Kaur, M., & Sethi, R. (2023). Discount depth and order frequency in Indian quick commerce: A panel data analysis. *Vikalpa: The Journal for Decision Makers*, 48(1), 12-28.
- Mehrotra, K., & Agarwal, N. (2021). Consumer behaviour in instant delivery: Evidence from urban India. *Asian Journal of Marketing*, 15(3), 201-219.
- Chopra, S., & Meindl, P. (2016). *Supply chain management: Strategy, planning, and operation* (6th ed.). Pearson Education.
- Sharma, V., & Patel, H. (2022). Competitive dynamics in Indian quick commerce: A duopoly framework. *IIMB Management Review*, 34(2), 88-101.
- Kumar, V., Venkatesan, R., & Reinartz, W. (2023). Machine learning applications in retail store performance segmentation. *Journal of Marketing Research*, 60(2), 310-330.
- Wang, L., Zhang, X., & Liu, T. (2024). Log-transformation and scaling strategies for K-Means clustering on skewed e-commerce data. *Expert Systems with Applications*, 238, 121-135.
- Verhoef, P. C., Kannan, P. K., & Inman, J. J. (2021). From multi-channel retailing to omni-channel retailing: Introduction to the special issue on multi-channel retailing. *Journal of Retailing*, 91(2), 174-181.
- Singh, N., & Banerjee, S. (2023). Dark store density and delivery time compliance in Indian quick commerce. *Supply Chain Forum: An International Journal*, 24(3), 198-215.
- Rathore, A., Kapoor, S., & Joshi, D. (2023). Micro-market demand sensing and fill rate performance in Q-commerce. *International Journal of Production Economics*, 261, 108-125.
- Jain, P., Kaur, M., & Sethi, R. (2023). Discount depth and order frequency in Indian quick commerce: A panel data analysis. *Vikalpa: The Journal for Decision Makers*, 48(1), 12-28.
- Kodinariya, T. M., & Makwana, P. R. (2013). Review on determining number of cluster in K-Means clustering. *International Journal of Advance Research in Computer Science and Management Studies*, 1(6), 90-95.
- Jolliffe, I. T., & Cadima, J. (2016). Principal component analysis: A review and recent developments. *Philosophical Transactions of the Royal Society A*, 374(2065), 20150202.
- Welch, B. L. (1947). The generalisation of Student's problem when several different population variances are involved. *Biometrika*, 34(1/2), 28-35.

Kruskal, W. H., & Wallis, W. A. (1952). Use of ranks in one-criterion variance analysis. *Journal of the American Statistical Association*, 47(260), 583-621.

Dr. Kusum lata

Report by Piyush

Document Details

Submission ID

trn:old::27535:139428062

Submission Date

May 18, 2026, 10:02 AM GMT+0

Download Date

May 18, 2026, 10:05 AM GMT+0

File Name

Report by Piyush.pdf

File Size

1.7 MB

41 Pages

9,728 Words

62,629 Characters



Page 1 of 48 - Cover Page

Submission ID trn:old::27535:139428062

10% Overall Similarity

The combined total of all matches, including overlapping sources, for each database.

Filtered from the Report

- Bibliography
- Quoted Text
- Small Matches (less than 8 words)

Match Groups

-  **79 Not Cited or Quoted 10%**
Matches with neither in-text citation nor quotation marks
-  **3 Missing Quotations 0%**
Matches that are still very similar to source material
-  **0 Missing Citation 0%**
Matches that have quotation marks, but no in-text citation
-  **0 Cited and Quoted 0%**
Matches with in-text citation present, but no quotation marks

Top Sources

- 6%  Internet sources
- 3%  Publications
- 7%  Submitted works (Student Papers)

Integrity Flags

0 Integrity Flags for Review

Our system's algorithms look deeply at a document for any inconsistencies that would set it apart from a normal submission. If we notice something strange, we flag it for you to review.

A Flag is not necessarily an indicator of a problem. However, we'd recommend you focus your attention there for further review.