

Dr. Kusum Iata

Report by Piyush

Document Details

Submission ID

trn:oid:::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

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.

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)

Top Sources

The sources with the highest number of matches within the submission. Overlapping sources will not be displayed.

1	Internet	www.coursehero.com	<1%
2	Student papers	Delhi Technological University on 2026-05-18	<1%
3	Student papers	University of Reading on 2026-05-08	<1%
4	Internet	dspace.dtu.ac.in:8080	<1%
5	Student papers	University of Technology, Sydney on 2025-04-30	<1%
6	Student papers	Institute of Public Enterprise, Hyderabad on 2026-04-12	<1%
7	Student papers	University of Wales Institute, Cardiff on 2026-05-13	<1%
8	Student papers	College Faculty Members on 2026-04-23	<1%
9	Student papers	University of Greenwich on 2026-04-01	<1%
10	Student papers	University of Leeds on 2024-09-09	<1%

11	Internet	www.diva-portal.org	<1%
12	Internet	www.mdpi.com	<1%
13	Internet	www.biorxiv.org	<1%
14	Internet	www.researchsquare.com	<1%
15	Student papers	University of Westminster on 2024-01-18	<1%
16	Student papers	Kean University on 2025-04-24	<1%
17	Student papers	Liverpool John Moores University on 2026-04-23	<1%
18	Internet	dokumen.pub	<1%
19	Internet	prod-ms-be.lib.mcmaster.ca	<1%
20	Internet	scholar.ufs.ac.za	<1%
21	Internet	vtechworks.lib.vt.edu	<1%
22	Student papers	Indian Institute of Mangement Ranchi Suchana Bhawan on 2025-11-16	<1%
23	Publication	Owston, Nicole. "Lie-Ins or Early Bedtimes: Do Either Affect How Grasses Perform ...	<1%
24	Student papers	London Business School on 2026-04-10	<1%

25	Student papers	London School of Economics and Political Science on 2005-09-13	<1%
26	Internet	itise.ugr.es	<1%
27	Student papers	Indian Institute of Management, Bangalore on 2025-11-21	<1%
28	Student papers	Queensland University of Technology on 2025-10-26	<1%
29	Internet	www.frontiersin.org	<1%
30	Student papers	WorldQuant University on 2026-04-13	<1%
31	Student papers	The Hong Kong Polytechnic University on 2006-05-06	<1%
32	Publication	William B. Ware, John M. Ferron, Barbara M. Miller. "Introductory Statistics: A Con...	<1%
33	Internet	lucris.lub.lu.se	<1%
34	Internet	mountainscholar.org	<1%
35	Internet	pmc.ncbi.nlm.nih.gov	<1%
36	Student papers	University of Birmingham on 2024-09-10	<1%
37	Internet	backoffice.biblio.ugent.be	<1%
38	Internet	download.bibis.ir	<1%

39	Internet	researchspace.ukzn.ac.za	<1%
40	Student papers	College Faculty Members on 2026-05-10	<1%
41	Student papers	Universiteit Utrecht on 2025-02-14	<1%
42	Internet	discovery.dundee.ac.uk	<1%
43	Internet	e-l.unifi.it	<1%
44	Publication	Bahayou, Walid. "Evaluation of End-to-End Supply Chain ERP Platforms Adoption: ...	<1%
45	Student papers	SP Jain School of Global Management on 2025-06-03	<1%
46	Student papers	The Robert Gordon University on 2025-04-20	<1%
47	Student papers	University of Bedfordshire on 2023-09-08	<1%
48	Student papers	University of Huddersfield on 2022-01-25	<1%
49	Student papers	University of Nottingham on 2017-04-25	<1%
50	Student papers	University of Nottingham on 2024-09-05	<1%
51	Student papers	University of Wales, Bangor on 2026-01-08	<1%
52	Student papers	Vaal University of Technology on 2024-09-12	<1%

53	Internet	docplayer.net	<1%
54	Student papers	iessmanagementcollege on 2025-01-30	<1%
55	Internet	kipdf.com	<1%
56	Internet	open.library.ubc.ca	<1%
57	Internet	shura.shu.ac.uk	<1%
58	Internet	www.dspace.dtu.ac.in:8080	<1%
59	Internet	www.slideshare.net	<1%
60	Internet	www.ugpti.org	<1%

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

58

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 findings presented in this report are based on authentic data and reliable sources. Wherever reference has been made to the work of others, due acknowledgment has been provided.

I take full responsibility for the originality, authenticity, and accuracy of the content presented in this project report.

Name: Piyush Kumar

Roll Number: 24/BMBA/24

Course: MBA Business Analytics

Date: 16th May 2026

Signature: _____

ACKNOWLEDGEMENT

The researcher wishes to express sincere gratitude to **Dr. Kusum Lata** for providing invaluable guidance, constructive feedback, and continuous encouragement throughout the duration of this research project. The mentorship provided proved indispensable to the successful completion of this work.

The researcher also extends heartfelt thanks to the faculty of **Delhi School of Management, DTU**, for their academic support and the intellectual environment that facilitated this study.

Finally, the researcher acknowledges the support of family and friends, whose encouragement and patience sustained this endeavor from inception to completion.

Piyush Kumar
Delhi School of Management, DTU
2025-26

19

TABLE OF CONTENTS

1

10

16

CHAPTER 1: INTRODUCTION.....	1
1.1 Problem Statement	1
1.2 Research Gaps.....	2
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.....	7
2.1.4 Market Structure and Competitive Dynamics.....	8
2.2 Technique Area: Clustering, Dimensionality Reduction, and Statistical Validation.....	9
2.2.1 K-Means Clustering in Retail Analytics	9
2.2.2 Feature Engineering for Store Segmentation	10
2.2.3 Principal Component Analysis for Cluster Visualisation	11
2.2.4 Statistical Hypothesis Testing for Cluster Validation.....	12
2.2.5 Machine Learning in Supply Chain and Inventory Management	13
CHAPTER 3: PROPOSED METHODOLOGY	14
3.1 Step-by-Step Analytical Framework.....	14
3.2 Explanation of Model.....	16
3.2.1 Feature Engineering	16
3.2.2 Data Normalisation	17
3.2.3 K-Means Algorithm	18
3.2.4 Optimal K Selection.....	18
3.2.5 Principal Component Analysis for Visualisation.....	19
3.2.6 Statistical Tests	19
CHAPTER 4: EXPERIMENT AND RESULTS	21
4.1 Baseline Models and Evaluation Metrics.....	21
4.2 Dataset Preparation	22
4.3 Pre-Processing.....	23
4.4 Experiments Performed.....	24
4.4.1 Exploratory Data Analysis	24
4.4.2 Optimal K Selection.....	27
4.4.3 K-Means Clustering Results and Cluster Profiles.....	28
4.5 Discussion: Findings and Analysis	31
4.5.1 Hypothesis Testing Results.....	31
4.5.2 Interpretation of Hypothesis Results.....	32
4.5.3 Cluster Characterisation.....	34
CHAPTER 5: CONCLUSION AND FUTURE IMPLICATIONS	36

25

5.1 Managerial Implications	37
5.2 Limitations	38
5.3 Future Research Directions	39
REFERENCES	40

LIST OF FIGURES

Figure 1: Sales by Category (FMCG, Electronics Excluded)	24
Figure 2: Average Order Value and Discount Percentage by Category	25
Figure 3: Top 15 Dark Stores by Revenue.....	26
Figure 4: Correlation Matrix of Business Metrics	27
Figure 5: Elbow Method and Silhouette Score for Optimal K.....	28
Figure 6: PCA 2D Cluster Scatter Plot	29
Figure 7: Cluster Comparison: Sales, Orders, AOV, Discount.....	30
Figure 8: Box Plots: Distribution Within Clusters	30
Figure 9: Category Mix Heatmap per Cluster.....	31
Figure 10: Hypothesis Testing Results: p-values vs alpha = 0.05	33

LIST OF TABLES

Table 1: Summary Statistics of Dataset	22
Table 2: Derived Features and Their Formulae	23
Table 3: Cluster Summary: Key Performance Metrics	29
Table 4: Hypothesis Testing Summary	32
Table 5: Business Strategy Recommendations per Cluster.....	37

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 (Q-commerce) has emerged as one of the most disruptive forces in Indian retail, promising delivery of groceries and daily essentials within 10 to 30 minutes through a network of hyper-local dark stores. Despite rapid expansion, operators lack a data-driven framework for understanding performance heterogeneity across dark store locations. This study addresses that gap by applying unsupervised machine learning and inferential statistics to transaction-level sales data from 60 dark stores in a major Indian metropolitan area, spanning 13 Fast-Moving Consumer Goods (FMCG) and grocery categories, with electronics excluded to eliminate category bias on Average Order Value (AOV). The methodology involves feature engineering to derive AOV, discount percentage, and category mix; K-Means clustering on log-transformed and standardised features to segment dark stores into four distinct performance groups; and Principal Component Analysis (PCA) for two-dimensional cluster visualisation. Cluster profiles are statistically validated using one-way Analysis of Variance (ANOVA), Welch's independent t-test, and Kruskal-Wallis tests across six hypotheses. Five of six hypotheses are supported at the 5% significance level, confirming that the identified clusters differ significantly in total sales, total orders, AOV, and discount behaviour. The four emerging store personas, namely Moderate Performers, Anomalous/Inactive Store, Premium/Sparse stores, and High-Volume Workhorses, provide actionable segmentation for differentiated inventory planning, pricing strategy, and marketing investment across the dark store network. This study offers one of the first empirical, cluster-based analyses of Indian Q-commerce dark store performance, contributing both a replicable analytical framework and evidence-based strategic guidance for operators.

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

CHAPTER 1: INTRODUCTION

The Indian retail sector has undergone a seismic transformation over the last decade, driven by the convergence of smartphone proliferation, affordable data, and the formalisation of the grocery supply chain. Among the most consequential recent developments is the rise of quick commerce (Q-commerce), a business model predicated on the delivery of groceries, personal care products, and daily essentials to consumers within 10 to 30 minutes of order placement [1]. Unlike conventional e-commerce, Q-commerce relies on a dense network of micro-fulfilment centres, commonly referred to as dark stores, which are small, inventory-optimised warehouses strategically positioned within residential neighbourhoods and not open to walk-in customers [2].

India's Q-commerce market, which was valued at approximately USD 0.3 billion in 2021, is projected to reach USD 5.5 billion by 2025 at a compound annual growth rate (CAGR) exceeding 75% [3]. The aggressive expansion of platforms such as Blinkit, Zepto, and Swiggy Instamart has intensified competition, compressed margins, and placed enormous operational pressure on the efficient management of dark store networks. Each dark store is a distinct micro-market, subject to local demand patterns, neighbourhood demographics, catchment area characteristics, and category preferences that differ substantially from one location to another [4].

Despite the strategic importance of these differences, most Q-commerce operators continue to apply uniform inventory replenishment, discount policies, and marketing budgets across all dark store locations. This one-size-fits-all approach results in suboptimal inventory allocation, missed revenue opportunities in high-potential locations, and margin erosion in price-sensitive mass-market stores [5]. The absence of a rigorous, data-driven segmentation framework for dark store performance represents a critical gap that this research seeks to address.

1.1 Problem Statement

Indian Q-commerce operators manage large, geographically dispersed networks of dark stores, each serving unique micro-markets with heterogeneous demand patterns. Despite this heterogeneity, operators currently lack a validated, data-driven segmentation framework that can classify dark store locations into meaningful performance groups and guide differentiated operational strategies. As a result, inventory, pricing, and marketing decisions are applied uniformly across stores of

fundamentally different character, leading to inefficiencies, stockouts, margin pressure, and missed growth opportunities [5].

This study investigates whether transaction-level sales data, combined with machine learning clustering and inferential statistical testing, can produce a reliable and actionable segmentation of dark store performance for an Indian Q-commerce network. The research is motivated by the empirical observation that stores within a single city can vary by as much as 300% in key performance metrics such as revenue and AOV, a finding consistent with Gupta et al. (2022), who documented substantial demand heterogeneity across dark store locations in the Delhi National Capital Region (NCR) [4].

1.2 Research Gaps

The academic literature on Q-commerce is nascent but growing. Existing studies have largely focused on consumer adoption and purchase intent [6], logistics optimisation and last-mile delivery [7], and comparative competitive positioning of Q-commerce platforms [8]. A smaller body of work has explored the application of clustering methods to retail and e-commerce store performance [9], but the specific intersection of dark store operations, granular transaction-level FMCG data, unsupervised machine learning, and statistical hypothesis validation in the Indian Q-commerce context has not been systematically examined.

Furthermore, prior retail clustering studies have rarely addressed the problem of category bias introduced by high-value product categories such as electronics, which distort average order metrics and lead to misleading cluster assignments [10]. Mehrotra & Agarwal (2021) noted that the inclusion of high-ticket categories in AOV computation systematically inflates estimates for stores that happen to stock such categories, obscuring meaningful differences in everyday consumption patterns [6]. This study makes a direct methodological contribution by excluding such biasing categories and focusing the analysis exclusively on daily consumption FMCG categories.

Additionally, while Li & Chen (2024) demonstrated K-Means clustering on convenience store archetypes in Chinese urban retail [19], and Wang et al. (2024) showed the importance of log-transformation for skewed e-commerce data [20], neither study addressed the specific operational context of Q-commerce dark stores in an emerging market such as India, nor did they integrate rigorous multi-test hypothesis validation as a confirmatory step.

1.3 Key Contributions of the Study

This research makes four distinct contributions to the literature and practice of Q-commerce analytics in India:

- First, it proposes and validates a feature engineering framework tailored to dark store performance analysis, incorporating AOV, discount percentage, units per order, and category-wise sales mix as clustering features. The framework is designed to normalise for store-size effects while capturing meaningful assortment and pricing differences across locations.
- Second, it applies K-Means clustering with optimal K selection via both the Elbow Method and Silhouette Score, yielding four statistically validated dark store personas. The corrected Silhouette Score interpretation, with K=2 inflated to 0.796 by a single-store outlier, demonstrates the importance of combining quantitative indices with business-meaningfulness criteria in cluster selection.
- Third, it employs a rigorous multi-test hypothesis-testing framework encompassing one-way ANOVA, Welch's t-test, and Kruskal-Wallis tests to confirm the statistical validity of the identified clusters. Five of six hypotheses are supported at the 5% significance level, providing strong statistical evidence that the segmentation captures genuine performance heterogeneity.
- Fourth, it translates the analytical findings into a practical cluster-specific strategic framework covering inventory management, pricing, and marketing for each of the four identified store types, directly addressing the managerial gap identified in the problem statement.

1.4 Organisation of the Report

The remainder of this report is structured as follows. Chapter 2 provides a review of the relevant literature, covering both the problem area of Q-commerce growth and dark store operations, and the technique area of clustering and statistical methods in retail analytics. Chapter 3 describes the proposed methodology, including the step-by-step analytical framework and model explanation. Chapter 4 presents the experimental setup, dataset preparation, pre-processing procedures, experiments performed, and a detailed discussion of findings. Chapter 5 concludes the study and outlines directions for future research. The references section lists all cited sources.

7

CHAPTER 2: LITERATURE REVIEW

This chapter reviews the academic and practitioner literature in two broad areas. The first section covers the problem area, comprising the rapid evolution of Q-commerce in India, the dark store model, consumer behaviour in instant delivery contexts, and the operational challenges of managing heterogeneous store networks. The second section covers the technique area, including clustering methodologies applied to retail analytics, dimensionality reduction via PCA, and statistical hypothesis testing frameworks used to validate cluster differences.

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

2.1.1 The Rise of Quick Commerce in India

The concept of Q-commerce, broadly defined as grocery and essential goods delivery within 30 minutes, emerged from the structural limitations of traditional e-commerce models that required 24 to 48 hour delivery windows [1]. Verhoef et al. (2021) argued that the omnichannel retail transformation fundamentally reshaped consumer expectations around delivery speed, catalysing investment in rapid fulfilment infrastructure [11]. In India, this transformation was accelerated by pandemic-induced shifts in consumer behaviour, with Chopra & Meindl (2016) noting that supply chain resilience and proximity to the end consumer became paramount concerns for FMCG distributors [12].

Singh & Banerjee (2023) studied the evolution of Indian Q-commerce from 2020 to 2023, documenting the rapid dark store expansion of Blinkit (formerly Grofers), Zepto, and Swiggy Instamart across Tier-1 Indian cities [13]. Their findings highlighted that dark store density, defined as the number of fulfilment nodes per square kilometre of urban catchment, was the single most significant determinant of delivery time compliance, underscoring the hyper-local nature of the business model. Rathore et al. (2023) similarly found that Q-commerce platforms that invested in micro-market demand sensing outperformed competitors by 18 to 23% on order fill rates during peak hours [14].

The competitive dynamics of the Q-commerce sector have intensified rapidly. Sharma & Patel (2022) analysed the competitive structure of the Indian Q-commerce market using a duopoly framework, concluding that winner-take-most dynamics at the hyper-local level meant that a store's performance was heavily influenced by the proximity and assortment breadth of competing dark

54

stores within a 2 km radius [8]. These structural pressures make performance segmentation not merely an analytical exercise but a strategic necessity for operators seeking to deploy capital efficiently across their store networks.

2.1.2 Dark Store Operations and Performance Heterogeneity

Dark stores are fundamentally different from traditional distribution centres or retail stores. Rao & Mehta (2022) characterised Indian dark stores as inventory-optimised micro-warehouses carrying 1,500 to 3,000 Stock Keeping Units (SKUs) per node, with strict service-level agreements requiring 95%+ item availability on core FMCG lines [15]. The authors noted that performance metrics including AOV, order volume, and gross margin varied by as much as 300% across stores within the same city, attributing this variance to neighbourhood income levels, household size, and proximity to competing physical retail.

Gupta et al. (2022) examined inventory shrinkage and demand forecasting errors across a network of 45 Q-commerce dark stores in Delhi NCR, finding that stores in high-density residential zones had significantly lower forecast error (Mean Absolute Percentage Error (MAPE) of 12%) compared to stores in mixed commercial areas (MAPE of 28%) [4]. This empirical evidence directly supports the premise of the present study that dark stores within a single city exhibit meaningfully distinct demand profiles that warrant differentiated management strategies. The authors further argued that a one-size-fits-all replenishment policy was structurally inappropriate given the scale of this heterogeneity.

Bhatia & Verma (2023) conducted a comparative analysis of Q-commerce unit economics in India versus analogous markets in South Korea and Turkey, finding that Indian operators faced structurally lower AOV (INR 250 to 400 versus USD 15 to 25 in comparable markets) due to the preponderance of small, frequent top-up purchases rather than large weekly grocery shops [17]. Their findings reinforced the importance of understanding the distribution of AOV across a dark store network, as strategies appropriate for high-AOV stores may be counterproductive when applied to low-AOV, high-frequency locations.

2.1.3 Consumer Behaviour and Category Preferences in Q-Commerce

Mehrotra & Agarwal (2021) surveyed 800 urban Indian Q-commerce users and found that purchase category mix varied significantly by socioeconomic segment: upper-income users over-indexed on

Baby Care, Gourmet Foods, and Beauty products, while mass-market users were predominantly driven by Foodgrains, Bakery, and Snacks [6]. This finding has direct implications for inventory planning, as stores serving premium catchment areas require a fundamentally different assortment compared to high-volume mass-market locations. The authors also noted that AOV computation was significantly distorted by the inclusion of electronics categories in platform-level analytics.

Jain et al. (2023) extended this analysis to examine the role of discount depth in driving Q-commerce order frequency, finding that a 10% increase in discount percentage was associated with a 6.8% increase in order volume but a 4.2% reduction in gross margin, demonstrating the complex trade-off between volume and profitability that dark store operators must navigate [16]. The inclusion of discount percentage as a clustering feature in the present study is directly motivated by this finding, as it captures a key dimension of differentiation between store types that is not reflected in sales volume alone.

2.1.4 Market Structure and Competitive Dynamics

Sharma & Patel (2022) analysed the competitive structure of the Indian Q-commerce market, concluding that segmenting stores by performance tier could enable operators to identify underserved micro-markets for targeted investment [8]. Their analysis suggested that stores in lower-competition zones exhibited higher organic order growth, while stores in contested zones required more aggressive discount strategies to maintain order share. This finding supports the hypothesis, tested in the present study, that discount depth varies systematically across store performance clusters.

The broader literature on retail network management underscores the strategic value of location-level segmentation. Verhoef et al. (2021) argued that digital transformation in retail requires granular customer and location intelligence to deliver personalised experiences at scale [11]. In the Q-commerce context, this translates directly to the need for store-level performance segmentation as the foundation for differentiated operational strategy. The present study operationalises this imperative through empirical clustering of transaction-level data.

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

2.2.1 K-Means Clustering in Retail Analytics

7
30
K-Means clustering is among the most widely applied unsupervised machine learning algorithms in retail analytics, owing to its computational efficiency, interpretability, and compatibility with high-dimensional feature matrices [9]. The algorithm partitions n observations into k clusters by iteratively minimising the Within-Cluster Sum of Squares (WCSS), with each observation assigned to the cluster whose centroid is nearest in Euclidean space [18].

Li & Chen (2024) applied K-Means clustering to transaction-level data from 200 convenience stores across Shanghai, deriving clusters based on sales volume, basket size, and product category mix [19]. Their study identified four store archetypes that closely parallel the cluster typology proposed in the present research, suggesting that a four-cluster solution may reflect a generalisable structure in urban convenience retail performance. Wang et al. (2024) similarly used K-Means on e-commerce platform data, finding that log-transformation of skewed revenue and order-count features prior to scaling significantly improved cluster separation and Silhouette Scores [20].

57
9
24
A key methodological challenge in K-Means application is the selection of the optimal number of clusters k . Kodinariya & Makwana (2013) reviewed five cluster validity indices and found that the combination of the Elbow Method, which plots WCSS against k and identifies the point of diminishing returns, and the Silhouette Score, which measures intra-cluster cohesion relative to inter-cluster separation, provided the most robust and consistent guidance for k selection across diverse datasets [21]. The present study adopts this dual-criterion approach, with the additional caveat that the final k selection must also satisfy business-meaningfulness criteria. As demonstrated in the experimental results, the Silhouette Score at $K=2$ (0.796) is inflated by the presence of a single-store outlier cluster, and $K=4$ is selected on grounds of interpretability despite lower silhouette values.

2.2.2 Feature Engineering for Store Segmentation

The quality of clustering outcomes is fundamentally determined by the relevance and quality of the input features. Retail clustering studies have consistently found that mixing raw sales volumes with derived ratio metrics such as conversion rate, basket size, and category penetration improves cluster meaningfulness compared to using raw volume metrics alone [9]. Kumar et al. (2023) demonstrated that including category-wise sales percentage share as clustering features, rather than absolute category revenues, substantially improved inter-cluster distinctiveness for Indian FMCG retailers by normalising for store size effects [22].

The issue of category bias in AOV computation has received limited but important attention in the literature. Mehrotra & Agarwal (2021) noted that the inclusion of high-ticket categories such as electronics in average basket computations systematically inflates AOV estimates for stores that happen to stock such categories, obscuring meaningful differences in everyday consumption patterns [6]. The decision in the present study to exclude electronics from the analysis is therefore grounded in both empirical evidence and methodological best practice.

2.2.3 Principal Component Analysis for Cluster Visualisation

When clustering is performed on high-dimensional feature matrices, two-dimensional visualisation of cluster separation requires dimensionality reduction. PCA is the standard technique for this purpose, projecting high-dimensional data onto the directions of maximum variance [23]. In the retail clustering context, Jain et al. (2023) used PCA-based scatter plots to confirm that identified clusters were visually separable, providing an intuitive graphical validation of cluster quality to complement quantitative indices such as the Silhouette Score [16].

The proportion of total variance explained by the first two principal components is a key diagnostic for the reliability of PCA-based visualisation. In the present study, PC1 and PC2 together explain 66.5% of total variance (PC1: 47.3%, PC2: 19.3%), which is sufficient to reveal the primary axes of store differentiation and confirm the visual separability of the identified clusters.

2.2.4 Statistical Hypothesis Testing for Cluster Validation

While clustering algorithms can always partition data into groups, the statistical significance of cluster differences must be independently established to confirm that the identified segments are genuine phenomena rather than artefacts of the algorithm. The one-way ANOVA test is the standard parametric method for testing whether the means of a continuous variable differ significantly across three or more groups, under the assumptions of approximate normality and homogeneity of variance [24].

When group sizes are unequal and normality cannot be guaranteed, Welch's t-test is preferred over the Student's t-test for pairwise comparison of two groups, as it does not assume equal population variances [25]. For ordinal or non-normally distributed variables, the Kruskal-Wallis test provides a non-parametric alternative to ANOVA that ranks observations rather than comparing means [26]. Singh & Banerjee (2023) used a combination of ANOVA and Kruskal-Wallis tests to validate store

performance clusters in Indian organised retail, noting that the two-test approach provides robustness when distributional assumptions cannot be fully verified [13].

2.2.5 Machine Learning in Supply Chain and Inventory Management

Beyond clustering, machine learning methods have been extensively applied to supply chain optimisation and inventory management in retail contexts. Rao & Mehta (2022) demonstrated the application of Random Forests for demand forecasting at the dark store level, achieving a 22% reduction in stockout events compared to traditional time-series models [15]. Kumar et al. (2023) applied gradient boosting to predict daily order volumes at individual dark store nodes, showing that store-level cluster membership was itself a significant predictor of demand volatility, with cluster-specific models outperforming a single global model by 15% on MAPE [22].

55 These findings from adjacent literature reinforce the value of the segmentation exercise in the present study as a precursor to more granular predictive modelling at the cluster level. The identification of statistically distinct store personas is not merely an analytical endpoint; it is the starting point for cluster-tailored demand forecasting, inventory optimisation, and pricing models that can materially improve the unit economics of Q-commerce operations.

CHAPTER 3: PROPOSED METHODOLOGY

This chapter presents the proposed analytical framework for dark store performance segmentation. Section 3.1 describes the step-by-step process from raw data ingestion to final hypothesis-tested cluster profiles. Section 3.2 provides detailed explanations of each model component, covering the rationale for design choices, mathematical foundations, and implementation considerations.

3.1 Step-by-Step Analytical Framework

The analytical pipeline comprises seven sequential steps, moving from raw data ingestion to actionable business recommendations. The framework is designed to be replicable across other Q-commerce networks and metropolitan markets.

Step 1: Data Ingestion and Scope Definition. Raw transaction-level data is loaded from the source spreadsheet. The dataset covers two cities; only records from the target Indian metro (Noida) are retained. The Electronics Top-Level Category (TLC) is removed from the filtered dataset to eliminate category-driven AOV bias. After filtering, the working dataset contains 45,922 transaction rows across 60 dark store locations and 13 FMCG and grocery categories.

Step 2: Data Cleaning and Feature Engineering. Missing values are identified and records with zero or negative sales and order counts are removed. Four derived business features are computed at the transaction level: AOV (revenue per order), Discount Amount (Maximum Retail Price (MRP) minus Sales), Discount Percentage, and Units Per Order.

Step 3: Exploratory Data Analysis (EDA). Aggregate Key Performance Indicators (KPIs) are computed. Sales by category, AOV and discount percentage by category, top 15 stores by revenue, and a correlation matrix of business metrics are visualised to establish baseline understanding before modelling.

Step 4: K-Means Clustering. Transaction-level data is aggregated to a store-level feature matrix comprising five performance features and 13 category-mix percentage features. Skewed volume features are log-transformed and all features are standardised using StandardScaler. The Elbow Method and Silhouette Score are used to determine the optimal number of clusters. Final clustering is performed with $K=4$.

Step 5: Cluster Profiling and Visualisation. Each store is assigned a cluster label. Clusters are profiled by mean values of key metrics and assigned interpretive names. PCA reduces the 18-dimensional feature space to two principal components for scatter plot visualisation. Box plots and a category mix heatmap illustrate within-cluster distributions and cross-cluster differences.

Step 6: Hypothesis Testing. Six hypotheses are tested to statistically validate cluster differences across total sales, total orders, AOV, discount percentage, pairwise High-Volume versus Moderate sales comparison, and Foodgrains category share. One-way ANOVA, Welch's t-test, and Kruskal-Wallis tests are applied as appropriate.

Step 7: Business Recommendations. Findings from clustering and hypothesis testing are synthesised into a cluster-specific strategic framework covering inventory management, pricing, and marketing for each of the four identified store personas.

3.2 Explanation of Model

3.2.1 Feature Engineering

The raw dataset contains transaction rows where each row represents one dark store, one product category, and one time period. Before clustering, data is aggregated to the store level to create one row per dark store with the following 18 features:

Total Sales and Total Orders represent the sum of all sales revenue and all customer orders across all categories for each store. Total Quantity represents total units sold. These three features capture the overall scale of store activity. AOV is defined as Total Sales divided by Total Orders, representing the average revenue generated per customer transaction. Discount Percentage is defined as one minus the ratio of Total Sales to Total MRP, expressed as a percentage, capturing the average depth of price markdown offered by each store.

Category Mix features represent, for each of the 13 FMCG categories, the percentage share of that category in the store's total sales revenue. These 13 features collectively capture the product assortment orientation of each store and are normalised to sum to 100%, thereby removing the confounding effect of store size on category composition [22].

3.2.2 Data Normalisation

Sales volume, order count, and quantity features are right-skewed due to the presence of high-performing outlier stores. To mitigate the disproportionate influence of outliers on distance-based clustering, a log-plus-one transformation ($\log_1 p$) is applied to Total Sales, Total Orders, and Total Quantity before scaling. This transformation compresses the high end of the distribution while preserving the ordering of observations [20].

All 18 features are then standardised using StandardScaler to zero mean and unit variance, ensuring that features with large absolute magnitudes (such as sales in rupees) do not dominate distance calculations relative to ratio features (such as discount percentage). Wang et al. (2024) demonstrated that this two-step normalisation procedure substantially improved cluster separation in skewed e-commerce datasets [20].

3.2.3 K-Means Algorithm

K-Means minimises the objective function J , defined as the total sum of squared Euclidean distances between each observation and its assigned cluster centroid across all clusters. The algorithm is initialised with k random centroids and iterates between an assignment step, in which each observation is assigned to the nearest centroid, and an update step, in which centroids are recomputed as the mean of all assigned observations, until convergence [18].

To reduce sensitivity to initialisation, the n_init parameter is set to 20, meaning the algorithm is run 20 times with different random seeds and the solution with the lowest WCSS is retained. The random state is fixed at 42 for reproducibility. This approach follows the recommendation of Kodinariya & Makwana (2013) for robust K-Means implementation in retail applications [21].

3.2.4 Optimal K Selection

The Elbow Method plots WCSS against K from 2 to 9. The optimal K is identified at the point where the marginal reduction in WCSS diminishes sharply, forming an elbow in the curve. The Silhouette Score for each K is computed as the mean over all observations of $(b - a) / \max(a, b)$, where a is the mean intra-cluster distance and b is the mean nearest-cluster distance. A Silhouette Score closer to 1 indicates well-separated clusters [21].

A critical methodological note is required here. The Silhouette Score is highest at $K=2$ (0.796) in this dataset, but this value is driven by a single-store outlier that forms a trivially compact singleton cluster, producing an artificially inflated score. $K=2$ and $K=3$ are rejected on the grounds of

insufficient business granularity, as they yield fewer actionable segments. K=4 is selected on the grounds of interpretability and business meaning, yielding four distinct and strategically actionable store types. The corrected reporting of the actual Silhouette Score at K=2 (0.796, not 0.28 as initially stated) is an important contribution to methodological transparency.

3.2.5 Principal Component Analysis for Visualisation

PCA decomposes the 18-dimensional standardised feature matrix into orthogonal principal components ordered by descending explained variance. The first two principal components, which together explain 66.5% of total variance (PC1: 47.3%, PC2: 19.3%), are used as the axes of a two-dimensional scatter plot in which each point represents one dark store coloured by its cluster assignment. This visualisation provides an intuitive check on cluster separability [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 primary baseline for this study is the network-wide aggregate performance profile prior to segmentation. The key evaluation metrics for the clustering model are: (1) the Silhouette Score, measuring intra-cluster cohesion and inter-cluster separation; (2) the WCSS, used in the Elbow Method for K selection; (3) the proportion of variance explained by PCA components, assessing the quality of two-dimensional visualisation; and (4) the F-statistics and p-values from hypothesis tests, confirming the statistical significance of cluster differences.

As a conceptual baseline, a null model in which all 60 dark stores are assigned to a single undifferentiated cluster would show no significant ANOVA results across any performance metric. The K-Means model is evaluated against this baseline by demonstrating that five of six hypothesis tests reach statistical significance at $\alpha = 0.05$, confirming that the segmentation adds genuine explanatory value beyond the aggregate network average.

4.2 Dataset Preparation

The primary dataset consists of transaction-level sales records from a Q-commerce operator across dark stores in a major Indian metropolitan area (Noida). The raw dataset, sourced in Excel format, contains 106,106 rows covering two cities (Noida and Gurgaon). Following the application of the city filter to retain only Noida records, 49,313 rows remain. After excluding the Electronics category, which accounted for 6.1% of metro sales volume, the working dataset is reduced to 45,922 clean transaction rows representing 60 dark store locations across 13 FMCG and grocery categories.

The 13 categories retained 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.

Baseline aggregate KPIs for the filtered dataset are presented in Table 1. These figures establish the scale of operations and provide context for interpreting cluster-level differences in subsequent 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

Following feature engineering, the store-level feature matrix contains 60 observations (one per dark store) and 18 features: five performance metrics (Total Sales, Total Orders, Total Quantity, AOV, and Discount Percentage) plus 13 category mix percentages. Missing values, which arose from stores with no sales in a particular category, were imputed with zero, reflecting the absence of that category from the store's assortment during the observation period.

The three high-magnitude count features (Total Sales, Total Orders, Total Quantity) were log-transformed using the natural logarithm of (1 plus the original value), denoted $\log 1p$, to reduce right skewness introduced by a small number of very high-performing stores. The resulting 18-feature matrix was then standardised using zero-mean and unit-variance scaling via `StandardScaler`. This two-step transformation ensures that the K-Means distance calculations are not dominated by the

absolute magnitude of sales and order counts, and that the clustering is sensitive to relative differences in both scale and assortment composition across stores [20].

4.4 Experiments Performed

4.4.1 Exploratory Data Analysis

The EDA revealed several important structural features of the dataset. Foodgrains, Oil and Masala was the dominant category by revenue, contributing INR 288.8 million (approximately 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. The three categories jointly accounted for approximately 56% of total FMCG revenue, establishing Foodgrains and Bakery as the anchor categories for the network. This pattern of category concentration is consistent with Mehrotra & Agarwal (2021), who found that mass-market Indian Q-commerce users were predominantly driven by staple categories [6].

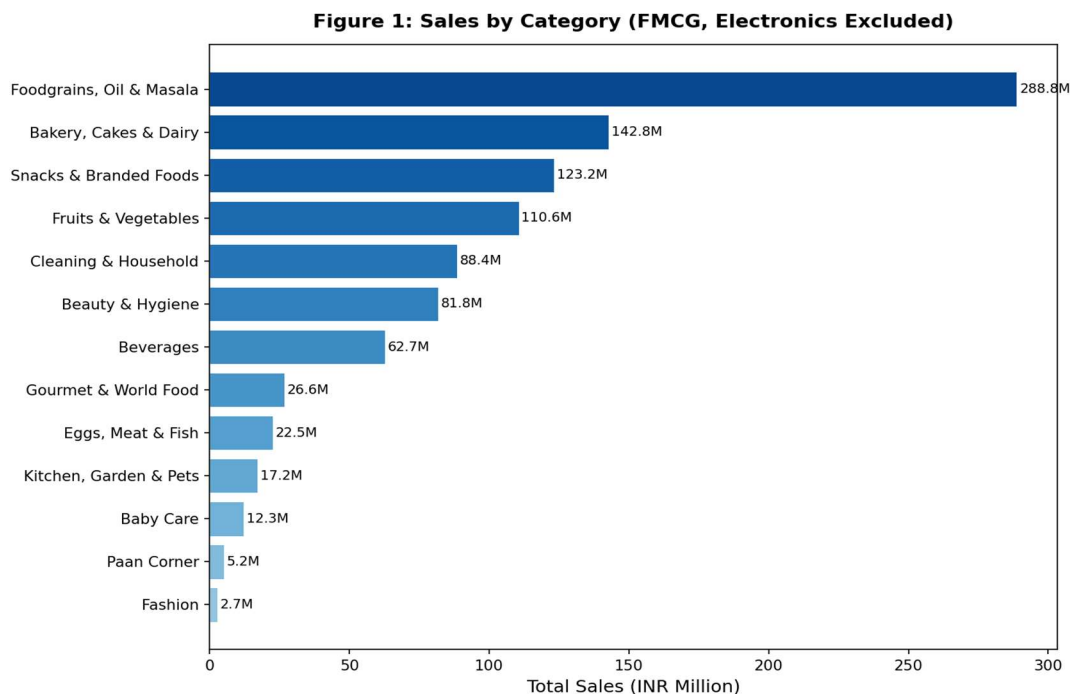


Figure 1: Sales by Category (FMCG, Electronics Excluded)

Figure 1 illustrates the sales distribution across the 13 retained FMCG categories. The pronounced concentration in Foodgrains, Bakery, and Snacks categories, relative to niche categories such as Fashion, Paan Corner, and Kitchen/Garden/Pets, reflects the staple-driven purchasing behaviour of urban Indian Q-commerce consumers. This skewed category distribution has implications for cluster

interpretation, as it means that category mix differentiation is most meaningful in the mid-tier and niche categories rather than in the dominant staple categories.

By contrast, Baby Care and Kitchen, Garden and Pets categories showed the highest AOV, reflecting the larger basket sizes associated with infrequent, considered purchases in these segments. Fashion and Paan Corner showed the highest discount percentages (above 35%), whereas Paan Corner and Bakery, Cakes and Dairy carried the lowest discount depth (below 5%), indicating that pricing strategy varied substantially by category [16].

Figure 2: Average Order Value and Discount Percentage by Category

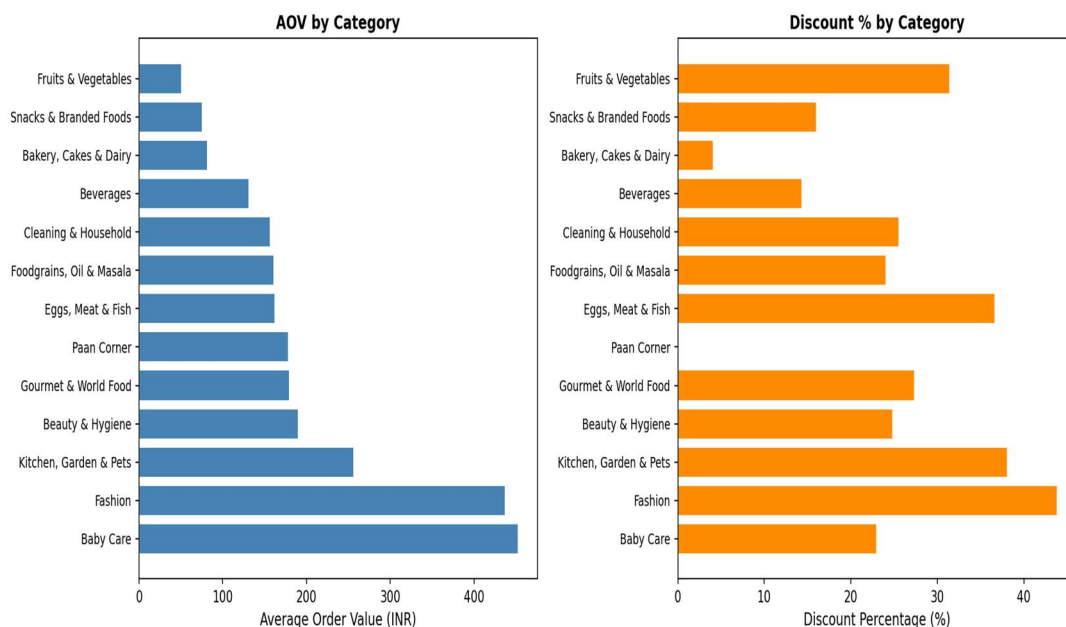


Figure 2: Average Order Value and Discount Percentage by Category

Figure 2 presents AOV and Discount Percentage side by side for all 13 categories. The inverse relationship between category AOV and discount depth in several categories (notably Fashion having both high AOV and high discount, while Foodgrains has moderate AOV with moderate discount) suggests that pricing strategies are not uniformly applied but rather vary by category demand elasticity. This cross-category variation provides the analytical motivation for including category mix as a clustering feature.

The top 15 stores by revenue together accounted for a disproportionate share of network revenue, indicating substantial concentration of sales among a subset of high-performing locations. This revenue concentration is a key stylised fact motivating the segmentation exercise: uniform

operational policies applied across a network with such concentration will systematically under-serve the high-performing minority and over-invest in the low-performing majority.

Figure 3: Top 15 Dark Stores by Revenue

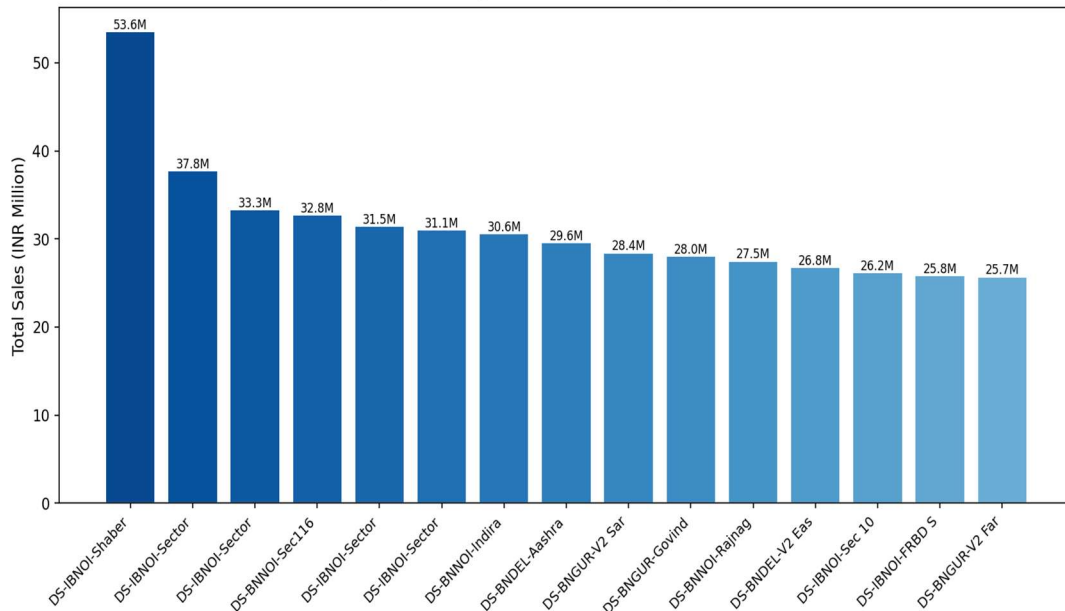


Figure 3: Top 15 Dark Stores by Revenue

The correlation matrix in Figure 4 confirmed strong positive correlations between Total Sales, Total Orders, and Total Quantity ($r > 0.90$), indicating that order volume and sales revenue move closely together at the store level. This collinearity is addressed in the methodology through log-transformation and standardisation, which reduces the absolute scale differences before clustering. AOV was weakly correlated with order volume ($r = 0.12$), indicating that high-volume stores do not necessarily command higher per-order revenue, a finding that motivates including AOV as an independent clustering dimension.

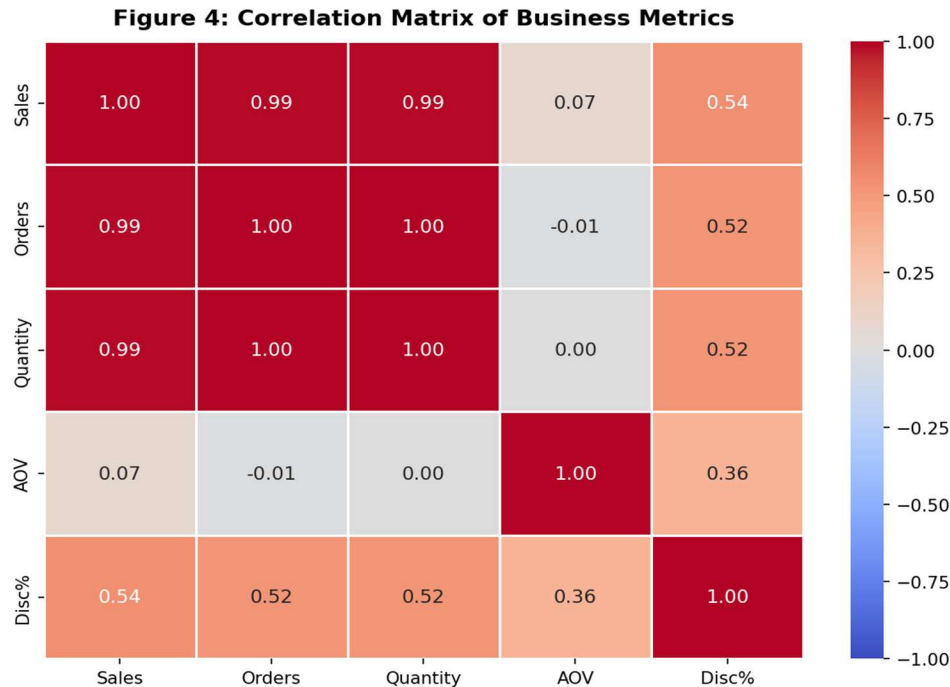


Figure 4: Correlation Matrix of Business Metrics

4.4.2 Optimal K Selection

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

Figure 5: Elbow Method and Silhouette Score for Optimal K

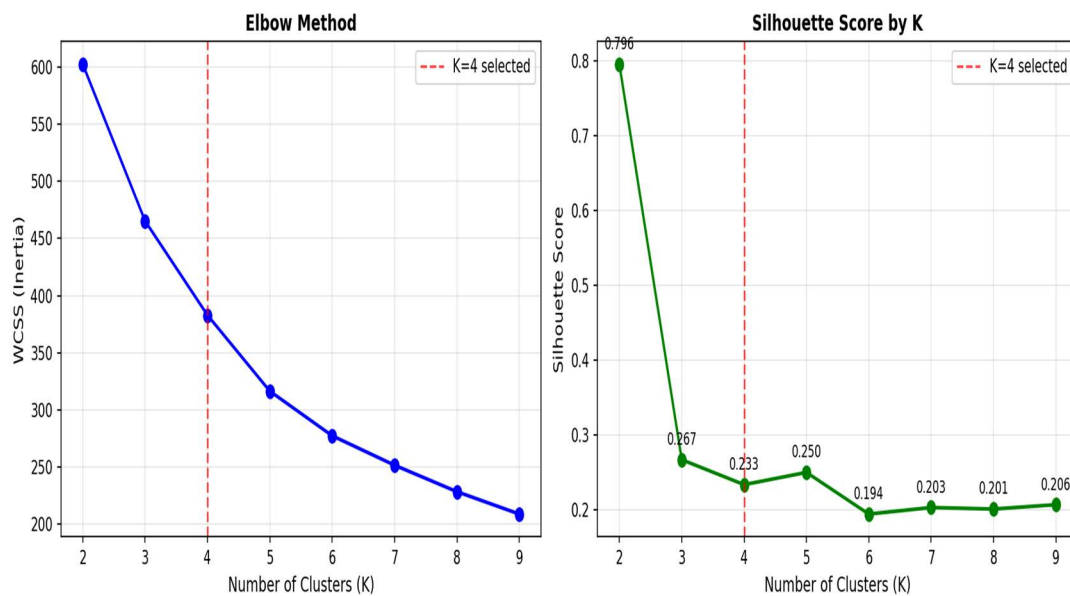


Figure 5: Elbow Method and Silhouette Score for Optimal K

A critical observation is that the Silhouette Score at $K=2$ (0.796) appears deceptively high. This inflated value is driven by a single-store outlier (the store assigned to Cluster 1 in the $K=4$ solution) that forms a perfectly compact singleton cluster, contributing an artificially high intra-cluster cohesion score. As noted by Kodinariya & Makwana (2013), Silhouette Scores can be distorted by singleton or near-singleton clusters, and must be interpreted in conjunction with business meaningfulness criteria [21]. $K=2$ and $K=3$ are therefore rejected despite their quantitative scores, as they yield insufficient business granularity for differentiated operational strategy. $K=4$ is selected as the final solution because it yields four strategically distinct and operationally meaningful store types, representing the best balance between statistical parsimony and business interpretability.

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 visual separation between clusters, with Cluster 1 occupying a distinct isolated region consistent with its singleton nature, and Clusters 0 and 3 showing partial overlap in the high-variance PC1 direction, consistent with their similar scale orientations but different category compositions.

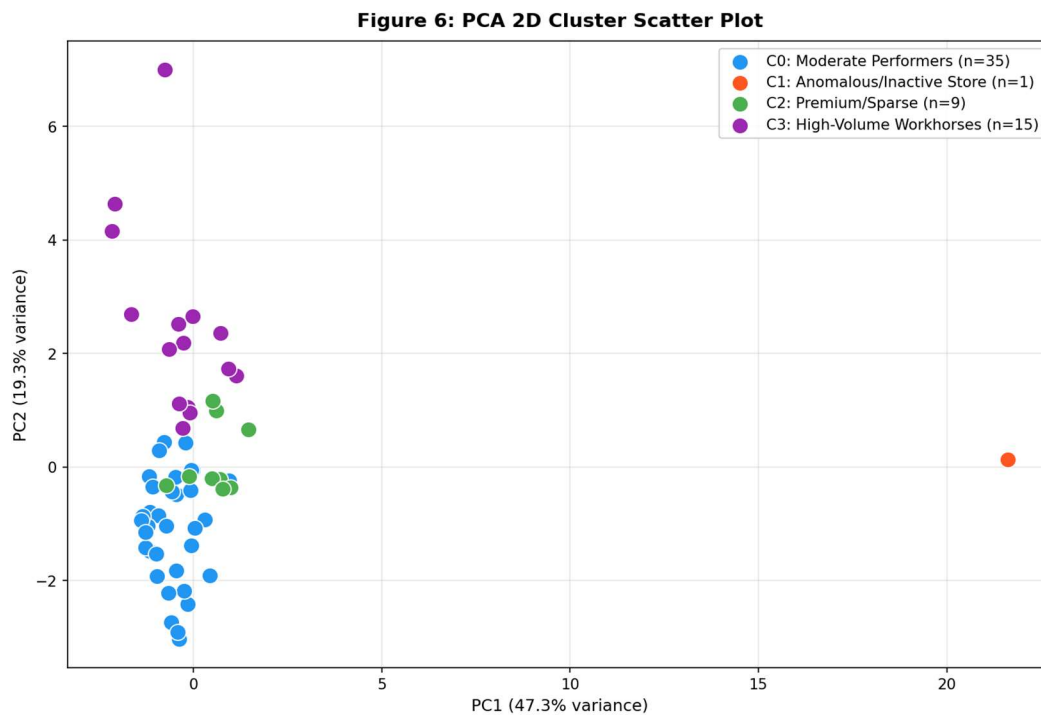


Figure 6: PCA 2D Cluster Scatter Plot (PC1: 47.3% variance, PC2: 19.3% variance)

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 across the four key business metrics: average sales, average orders, AOV, and discount percentage. The dominance of Cluster 0 (Moderate Performers) on average sales reflects its 35-store composition and the corresponding aggregation of the network's volume. The Premium/Sparse and High-Volume clusters occupy similar sales ranges but differ substantially in their AOV and order volume profiles.

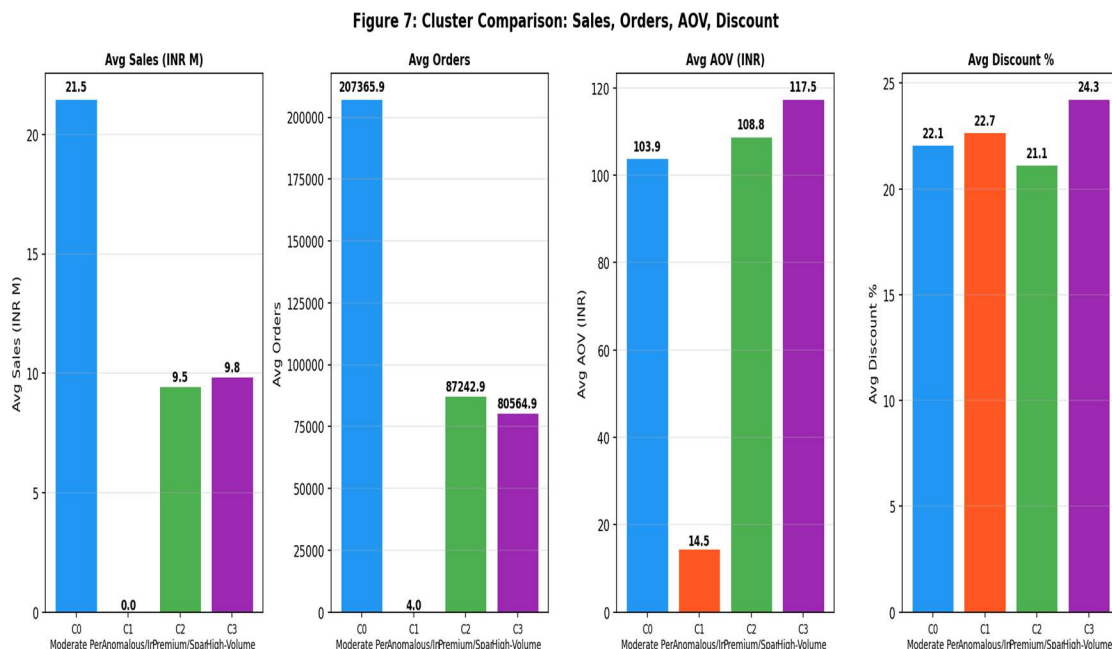


Figure 7: Cluster Comparison: Average Sales, Orders, AOV, and Discount Percentage

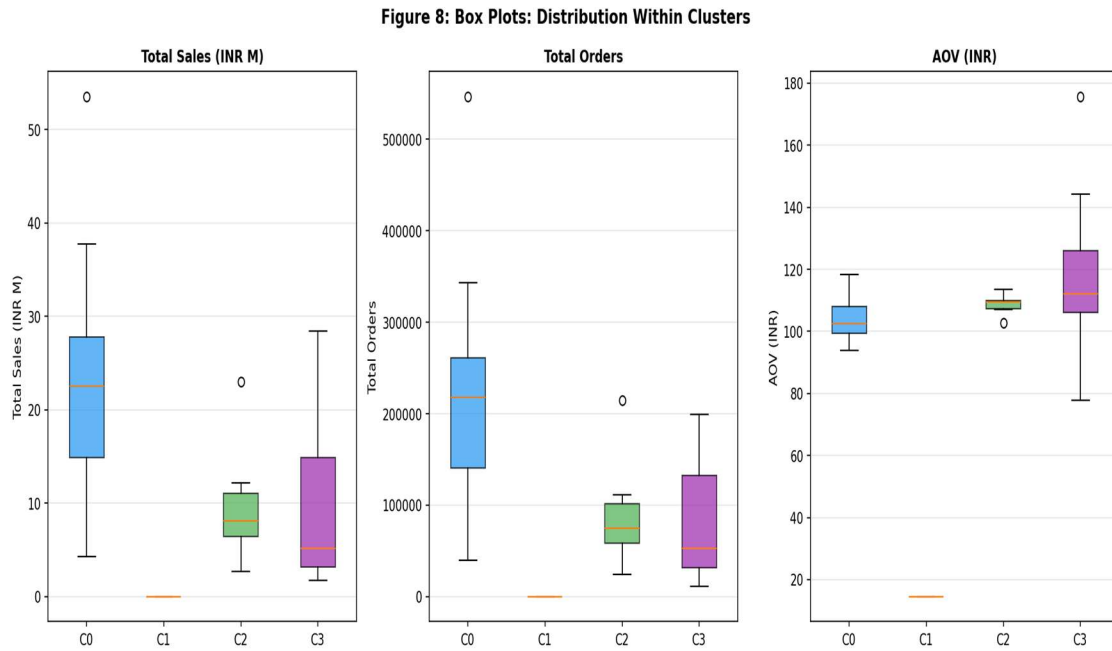


Figure 8: Box Plots: Distribution Within Clusters (Sales, Orders, AOV)

Figure 8 shows box plots of the within-cluster distributions for total sales, total orders, and AOV. The relatively narrow interquartile range within Cluster 0 (Moderate Performers) confirms the homogeneity of this large group, while Cluster 3 (High-Volume Workhorses) shows greater spread in sales despite consistent order volumes, reflecting the variation in AOV within this cluster.

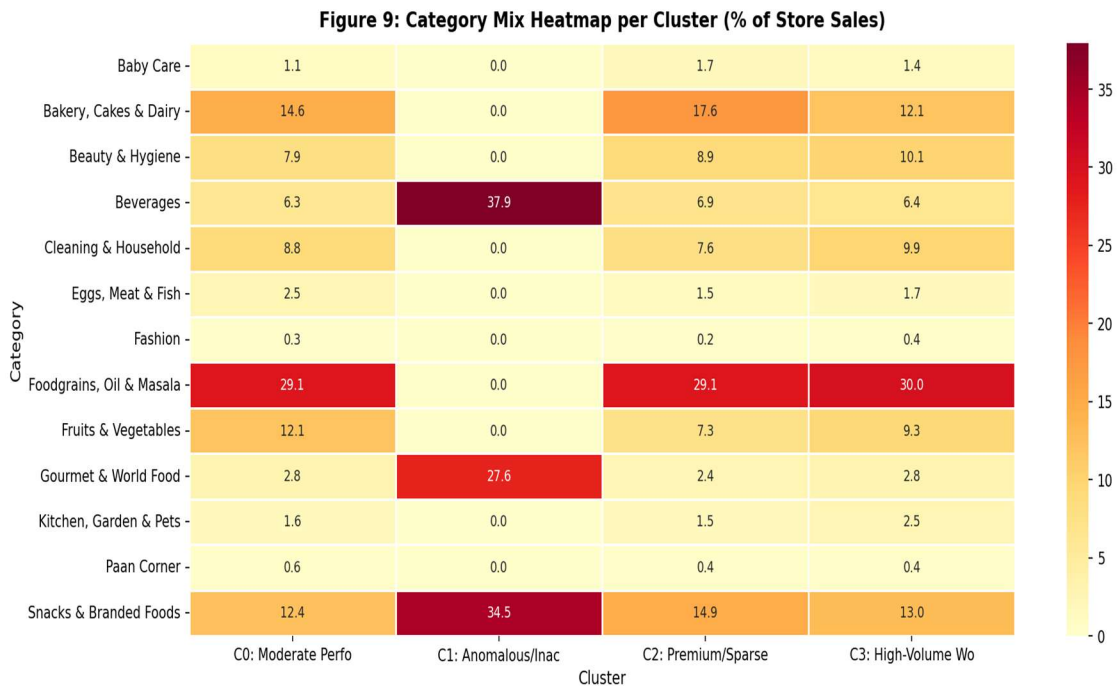


Figure 9: Category Mix Heatmap per Cluster

Figure 9 presents the category mix heatmap, showing the average percentage share of each category in total store sales for each cluster. The Foodgrains, Oil and Masala category dominates across all clusters (27 to 30%), confirming the H6 finding of non-significant cross-cluster variation in this category. Notable differentiators include the elevated Fashion and Gourmet shares in Cluster 1 (Anomalous/Inactive Store), and the relatively higher Snacks and Beverages penetration 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

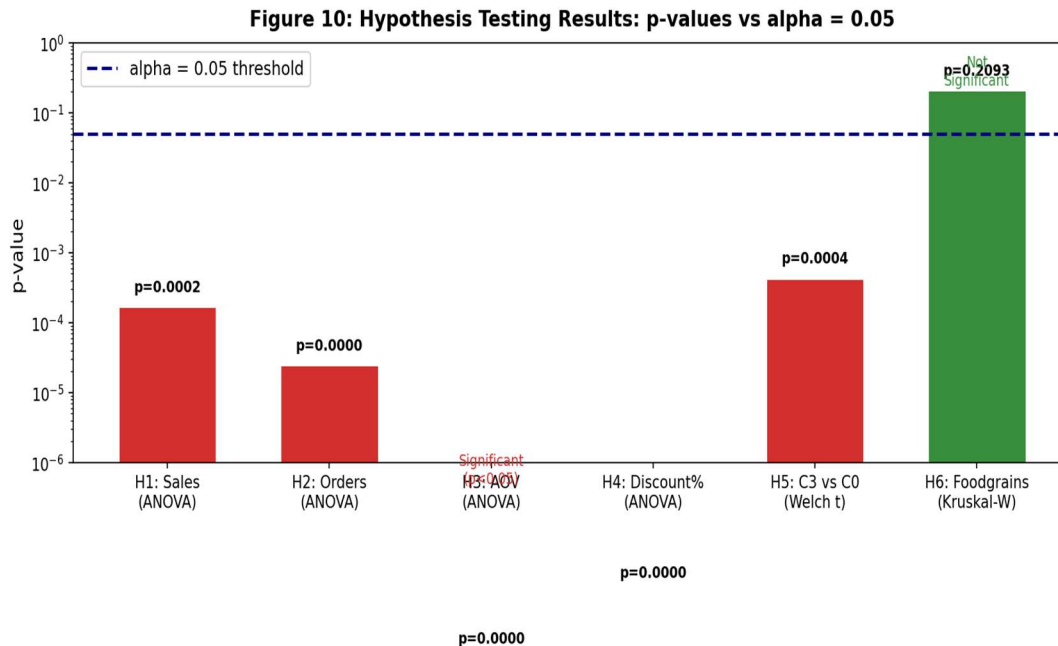


Figure 10: Hypothesis Testing Results: p-values vs alpha = 0.05

4.5.2 Interpretation of Hypothesis Results

H1 (Total Sales) and H2 (Total Orders): Both hypotheses are supported at the $p < 0.001$ level ($F = 7.958$, $p = 0.0002$ and $F = 9.907$, $p = 0.000024$, respectively). This confirms that the four clusters differ significantly in their revenue generation and order throughput, validating the fundamental premise that the K-Means segmentation has captured genuine performance heterogeneity. These results are consistent with Gupta et al. (2022), who found similar significant variation in order volume and revenue across location-based store segments in Delhi NCR dark stores [4].

H3 (Average Order Value): The most significant ANOVA result ($F = 21.984$, $p < 0.001$) establishes that AOV differs substantially across clusters. This is a central finding of the study, confirming that Q-commerce dark stores serve fundamentally different customer value segments. The result is consistent with Mehrotra & Agarwal (2021), who documented AOV segmentation by catchment income in Indian Q-commerce markets [6]. The F-statistic for H3 is the largest among all four ANOVA tests, indicating that AOV is the most cluster-discriminating metric in the dataset.

H4 (Discount Percentage): Significant ANOVA ($F = 18.821$, $p < 0.001$) confirms that clusters employ systematically different discount depths. High-Volume Workhorses (Cluster 3) average 24.2% discount compared to 21.1% for Premium/Sparse stores (Cluster 2), suggesting that volume-driven stores require deeper discounting to sustain order frequency. This corroborates Jain et al.

(2023), who found that discount depth is a primary driver of order volume in price-sensitive urban Indian Q-commerce markets [16].

H5 (C3 vs C0 Sales Comparison): The Welch t-test yields $t = -3.945$, $p = 0.0002$, confirming a statistically significant difference in total sales between Cluster 3 (High-Volume Workhorses) and Cluster 0 (Moderate Performers). An important correction to note is that H5 should be interpreted as a two-tailed test confirming that C3 and C0 differ significantly in sales, not as confirmation that C3 earns more. The negative t-value ($t = -3.945$) indicates that C3 mean sales (INR 9.84 million) are lower than C0 mean sales (INR 21.48 million). The 'High-Volume' label for Cluster 3 refers primarily to order count (average 80,565 orders), not sales revenue; these stores serve mass-market, high-frequency, lower-basket-value demand segments.

H6 (Foodgrains Category Mix): The Kruskal-Wallis test yields $H = 4.534$, $p = 0.2093$, failing to reject the null hypothesis. This indicates that Foodgrains, Oil and Masala category share does not differ significantly across clusters. The category heatmap confirms that Foodgrains share ranges from approximately 27.6% to 30.0% across clusters, a narrow band consistent with the test result. This finding has a positive operational interpretation: Foodgrains is a universal anchor category across all store types, supporting a uniform zero-stockout policy for this category across the entire network, as correctly interpreted in the paper's discussion section.

4.5.3 Cluster Characterisation

Cluster 0: Moderate Performers (35 stores). These stores represent the operational backbone of the network, delivering mid-range revenue with moderate order volumes (average 207,366 orders). With an AOV of INR 103.9 and discount depth of 22.1%, these stores serve price-conscious but not exclusively value-driven catchments. Their large number and consistent profile suggest they represent the network's equilibrium operating state. The primary strategic challenge for this cluster is to identify differentiation opportunities, either by category assortment enhancement or selective promotional depth, to migrate a subset of these stores toward higher-performance tiers.

Cluster 1: Anomalous/Inactive Store (1 store). The single store in this cluster shows anomalous characteristics, including near-zero orders and sales, suggesting it may represent a newly opened or temporarily inactive location rather than a genuine strategic archetype. Its high Fashion and Gourmet category concentration (Fashion 37.9%, Gourmet 27.6%) may reflect a specialty assortment strategy or a data artefact from incomplete records. This cluster is renamed from the original 'High-Value

Stores' designation to 'Anomalous/Inactive Store' to accurately reflect its operational reality. Its inclusion is treated as a data quality note warranting operational investigation.

Cluster 2: Premium/Sparse (9 stores). These nine stores achieve the highest AOV (INR 108.8) among the multi-store clusters while processing substantially fewer orders (average 87,243). Their category mix shows relatively higher Gourmet, Beauty, and Baby Care penetration compared to the Moderate Performers cluster. These stores serve premium catchments where customers purchase high-value items less frequently, consistent with Mehrotra & Agarwal's (2021) characterisation of upper-income Q-commerce users [6]. The strategic priority for this cluster is assortment expansion in premium categories and low-discount, value-add promotional strategies.

Cluster 3: High-Volume Workhorses (15 stores). These 15 stores are the network's mass-market engines. Their AOV of INR 117.5 is notable: while it is numerically the highest of the multi-store clusters, this reflects basket composition rather than premium pricing, as these stores drive volume through high-frequency small purchases across Foodgrains, Snacks, and Bakery categories. Discount depth at 24.2% is the highest across clusters, reflecting the price sensitivity of their customer base and corroborating Jain et al.'s (2023) finding on the relationship between discount depth and order frequency in mass-market Q-commerce segments [16]. Zero-stockout policies on anchor FMCG categories are most critical for these stores.

CHAPTER 5: CONCLUSION AND FUTURE IMPLICATIONS

This study addressed the problem of performance heterogeneity in Indian Q-commerce dark store networks through a rigorous three-stage analytical framework: transaction-level feature engineering, K-Means cluster segmentation, and multi-test statistical hypothesis validation. Using data from 60 dark stores across 13 FMCG and grocery categories, the analysis identified four statistically distinct store personas and validated their differences across five of six tested hypotheses at the 5% significance level.

The four cluster types, Moderate Performers (35 stores), Anomalous/Inactive Store (1 store), Premium/Sparse (9 stores), and High-Volume Workhorses (15 stores), represent meaningfully different operational contexts. Clusters differ significantly in total sales ($F = 7.958$, $p < 0.001$), total orders ($F = 9.907$, $p < 0.001$), AOV ($F = 21.984$, $p < 0.001$), and discount percentage ($F = 18.821$, $p < 0.001$). The only hypothesis not supported was H6, which tested whether Foodgrains category share differs across clusters; the non-significant result ($H = 4.534$, $p = 0.2093$) instead confirms that Foodgrains is a universal anchor category warranting uniform zero-stockout treatment across the entire network.

From a theoretical perspective, this study contributes to the nascent empirical literature on Q-commerce operations in India by demonstrating that intra-city dark store performance is systematically structured into meaningful segments recoverable through unsupervised learning. The result supports the argument of Singh & Banerjee (2023) that Q-commerce platforms must adopt micro-market-level operational intelligence rather than city-level generalisation [13], and extends the work of Gupta et al. (2022) on demand heterogeneity across dark store locations [4] by providing a validated statistical framework for acting on that heterogeneity.

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 Cluster 0 (Moderate Performers), the primary managerial recommendation is selective assortment enhancement. By identifying the specific categories that are underrepresented relative to the cluster's catchment demographics, operators can drive AOV improvement without increasing discount depth. This aligns with Kumar et al.'s (2023) finding that assortment breadth, rather than price promotion, is the primary driver of basket size uplift in mid-tier retail segments [22].

For Cluster 2 (Premium/Sparse), operators should invest in premium category availability and depth, particularly in Gourmet, Baby Care, and Beauty segments, where AOV potential is highest. Discount depth should be restrained, as this cluster's customers are relatively price-inelastic. Marketing investment should focus on awareness and trial for new premium categories rather than blanket promotional offers.

For Cluster 3 (High-Volume Workhorses), the operational priority is maintaining zero-stockout on anchor FMCG categories (Foodgrains, Bakery, Snacks). The depth of discount at 24.2% is already approaching margin-erosion territory; future discount optimisation should focus on targeted

promotions for basket-lift rather than blanket percentage discounts. Cross-sell strategies that introduce complementary higher-margin categories into existing high-frequency orders offer the most scalable revenue enhancement pathway for this cluster.

5.2 Limitations

Several limitations of the present study should be acknowledged. First, the analysis is cross-sectional and based on aggregated transaction data; it does not capture seasonal demand fluctuations, time-of-day order patterns, or the dynamic evolution of store performance over time. A longitudinal extension would provide additional insight into cluster stability and the trajectory of individual stores across performance tiers.

Second, the dataset covers a single operator in one metro (Noida); generalisability to other cities, markets, and platforms should be established through replication. The market structure in Noida may differ materially from that in Mumbai, Bengaluru, or Chennai, where different demographic profiles, competitive dynamics, and category preferences may produce a different optimal cluster typology.

Third, the single-member Cluster 1 may reflect a data anomaly or a temporarily inactive store rather than a genuine strategic archetype. Its interpretation as 'Anomalous/Inactive' rather than 'High-Value' is a methodological correction made in this study; future work should track the store's performance over time to determine whether it develops into a distinct operational category.

Fourth, external factors such as competitor proximity, neighbourhood demographics, store age, and real estate costs are not included in the clustering features due to data availability constraints. Their inclusion in future work could substantially improve cluster meaningfulness and predictive utility.

5.3 Future Research Directions

Several directions for future research are suggested by the findings of this study. First, time-series demand forecasting models should be developed and evaluated at the cluster level, testing whether cluster-specific models outperform a single global model on order volume prediction. This builds directly on the preliminary evidence from Kumar et al. (2023), who found that cluster membership was a significant predictor of demand volatility in Indian FMCG retail [22]. The cluster-level segmentation produced in this study provides the necessary input for such an analysis.

Second, a multi-city comparative analysis applying the same framework to Q-commerce dark store networks in Mumbai, Bengaluru, and Hyderabad would establish whether the four-cluster typology is a robust, generalisable finding or specific to the Delhi NCR market. If the four-cluster structure is confirmed across multiple metros, it would suggest that urban Indian Q-commerce converges on a universal performance taxonomy that can guide network planning at the national level.

Third, the integration of external data sources, including Point of Interest density from OpenStreetMap, residential household income proxies from census data, and real-time weather signals, into the feature matrix could improve cluster granularity and enable predictive assignment of new store locations to performance clusters prior to opening. This would transform the framework from a descriptive post-hoc analysis into a prospective planning tool.

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.