

PARAM: Prediction of Successful Fabrics for Upcoming Seasons in Fashion Industry



ISBN: 978-1-943295-22-7

Savita Kumari
Viswanath Reddy
Shaurya Khanduri
Hakeem Hadi Maqbool
Sargaraj Thaivachathil
ABMCPL, Aditya Birla Group
(savita.kumari@adityabirla.com)
(viswanath.reddy@adityabirla.com)
(shaurya.khanduri@abfirl.adityabirla.com)
(hakeem.hadi-v@adityabirla.com)
(sargaraj.t-v@adityabirla.com)

In the highly competitive and dynamic fashion industry, choosing a high-selling fabric design for the apparel is critical. Conventional approach is time-intensive and dependent on individual expertise. PARAM uses AI to predict the success of a fabric design, well ahead of time. It uses historical data to transform the designing and buying processes into a data-driven decision process. The solution, mentioned in this paper, builds insights on attribute level contribution to success and similarity of fabrics to those in past collection. It minimizes left over inventory and lost opportunity, a challenge that plagues fashion retail.

Keywords: Fashion, Retail, Machine Learning, Artificial Intelligence, Fabric Design, Apparel Industry, Predictive Analytics

1. Introduction

The fashion industry is undergoing rapid transformation, facing considerable pressure to accelerate time-to-market and develop apparel that aligns with consumer demand. This is further complicated by the introduction of new offerings every season and short-lived trends. The challenge now lies in accurately predicting which apparels will be successful. Conventional methods that depend on intuition and past experiences lead to inefficient decisions creating a mismatch between the product offerings and the market demand. Fashion – retail companies currently use fashion trend reports, historical performance, in-store feedback, and the expertise of fashion veterans to make their decisions. While these insights are critical and valuable, they are slow and prone to subjectivity. This leads to an accumulation of unsold products by the end of the season. As a result, retailers are frequently forced to either heavily discount these items or discard them entirely. Thus, forecasting capabilities become essential for maintaining market competitiveness.

This paper explores how data and advanced analytics driven approaches can augment the conventional approaches to provide deep insights to aid the selection of potentially successful fabrics for men's casual and formal shirts. Fashion retail companies can now harness historical performance data to enhance their decision making with actionable insights leading to improved apparel offerings in the market. This will result in driving higher sales performance and lower the risks of accumulating substantial amounts of unsold inventory. The methodology measures how well this approach improves product offerings by examining fabric characteristics and sales performance data from several previous seasons. The results of this study demonstrate how decision-making powered by data and advanced analytics has the potential to fundamentally transform the fashion retail sector.

2. Literature Review

The fashion industry relies heavily on fabric design and selection to drive sales and profitability. Traditional methods of predicting fabric success often involve a designer's intuition and limited data analysis. With the advent of machine learning (ML), there is a potential to enhance these predictions by leveraging data-driven insights.

Research has explored the application of ML in various aspects of the fashion industry, including trend forecasting, customer preferences, and product success. The authors propose applying machine learning models to forecast the demand of new fashion designs [1], which will achieve better accuracy than traditional methods.

The success of fabric designs depends on various factors, including color, season, material, and pattern, etc. Research has examined how these factors impact consumer preferences and sales performance. The authors highlight the role of fabric attributes in shaping brand identity and customer loyalty [2].

Traditionally, predicting the success of fabric designs relies on past sales data and the intuition of designers. However, issues such as uncleaned data, limited access to comprehensive datasets, and reliance on subjective judgment pose challenges. The need for cleaner, more accessible data to improve prediction accuracy is emphasized in [3].

Existing methods for fabric success prediction include statistical analysis and heuristic approaches. However, these methods often can't consider the complicated relationships between different attributes and external factors like upcoming fashion trends.

Advances in Machine learning, such as supervised learning and neural networks, offer promising solutions for predicting success. These models can analyze intricate patterns and relationships in large datasets, providing more accurate predictions than traditional methods. Studies have demonstrated the effectiveness of various Machine learning algorithms in fashion analytics [4].

A significant gap in current research is the lack of high-quality, comprehensive datasets. Many studies rely on limited or uncleaned data, which affects the reliability of predictions. There is a need for improved data collection methods and access to diverse datasets. While research has explored individual attributes, there is less emphasis on understanding the combined impact of multiple attributes on fabric success.

There is a gap in applying Machine learning models to real-world fashion scenarios, particularly in the context of fabric design and sourcing. Integrating Machine learning models into the design and sourcing process can streamline decision-making and improve outcomes. By incorporating data from multiple sources, including project lifecycle management (PLM) systems and sales records, ML engines can provide actionable insights that enhance design choices and procurement strategies.

3. Improvements Over Current Approaches

The PARAM project addresses several gaps identified in the literature by providing an end-to-end ML solution for fabric success prediction. By integrating comprehensive data sources and ML algorithms alongside AI driven descriptive analysis, PARAM aims to enhance the accuracy of predictions and support better decision-making in fabric design and procurement phase. This review highlights the potential of ML to transform fabric success prediction by providing data-driven insights and improving accuracy. This emphasizes the need to improve data quality, attribute granularity and real-world applications for better functioning of ML models for fashion industry.

4. Preliminaries

4.1 Rate of Sale (ROS)

ROS captures the velocity of stock selling and is defined for a product as the sales quantity per week per store. The higher the ROS the faster the stock is selling, thus is an important metric defining to the overall success of a product.

$$\text{Rate of Sale (ROS)} = \frac{\text{Total Sales Quantity}}{\text{Total Stores} * \text{Weighted Shelf - life}}$$

4.1.1 Weighted Shelf-life

Shelf life for apparel is defined as the time the apparel was available in the store for sale. The calculation of weighted shelf life, therefore, is an important step in the calculation of ROS. This is calculated as the weighted average of the number of weeks since inward date using the inward quantity as weights. This is defined in A similar way for style-store level, style level and fabric level.

$$\text{Weighted Shelf life} = \frac{\sum(\text{Individual Store Shelf Life} * \text{Store Inwards})}{\sum(\text{Inwards for Each Store})}$$

4.1.2 ROS – Style Level

Calculation of ROS at Style level, i.e., for each shirt sold is equal to the total number of shirts sold divided by the product of weighted shelf life and the total number of stores the shirt was sent to.

$$\text{Rate of Sale (ROS)} = \frac{\text{Total Sales Quantity Style Level}}{\text{Total Stores Style Level} * \text{Weighted Shelf - life Style Level}}$$

4.1.3 ROS – Fabric Level

Calculation of ROS at Fabric level is slightly different which is because of the one-to-many relation between fabrics and products. One fabric can be used to make multiple shirts; therefore, it is calculated as the total number of shirts sold manufactured using a certain fabric divided the product of the sum of the stores that individual shirts have gone to and the weighted shelf life for the fabric.

$$\text{Rate of Sale (ROS)} = \frac{\text{Total Sales Quantity Fabric Level}}{\text{Total Stores Fabric Level} * \text{Weighted Shelf - life Fabric Level}}$$

4.2 Gross Profit Per Net Sale Value (GP / NSV)

GP/NSV represents the ratio of gross profit (GP) to net sale value (NSV) and assesses the profitability of product. A higher ratio means the product is more profitable because a bigger part of the revenue from sales goes toward profit.

Calculation of GP/NSV: Gross Profit per Net Sale Value is calculated by dividing the total Gross Profit (GP) by total Net Sale value (NSV) of all the shirts sold that were made from the same fabric. Gross Profit is calculated as the difference between Net Sale Value (NSV) and Cost of Goods Sold (COGS).

$$\text{Gross Profit} = \text{Net Sale Value} - \text{Cost of Goods Sold}$$

$$GP/NSV = \text{Total Gross Profit} / \text{Total Net Sale Value}$$

4.3 Success Metric: Combination of ROS and GP/NSV Metric

Once ROS and GP/NSV are calculated for all fabrics the dataset is broken down into training set and test set. The data from 2021 till 2022 is used for training and the data from 2023 is used for testing the models. It is important to split the data before calculating the success metric as we Normalize both the metrics using the Minmax Scaler. We use the training data to fit the scaler and use the same scale to transform the test data. Once normalized we use a weighted sum with 70% weightage given to ROS and 30% weightage given to GP/NSV.

$$\text{Success Metric} = 70\% * \text{ROS} + 30\% * \text{GP/NSV}$$

4.3.1 Categorization of fabrics as successful and unsuccessful:

The success metric forms the basis of evaluation of the performance of a fabric. We split the training data into successful and unsuccessful by using the median value of the success metric as a threshold. All fabrics with the success metric greater than the median value are tagged as successful fabrics and all fabrics with the success metric lesser than the median value are tagged as unsuccessful fabrics. This same median value is used as a threshold to categorize the test set fabrics. This converts the problem into a binary classification problem with a balanced dataset. This methodology was selected with consideration of the relatively small dataset size, ensuring that a balanced dataset provides the models with the maximum number of fabric records for effective learning.

4.4 Gower Distance

Gower Distance is a distance measure that can be used to calculate distance between two entities whose attribute are a mix of categorical and numerical values. The Gower distance d_{ij} between two objects i and j in a dataset with p variables is calculated using the following formula

$$d_{ij} = \frac{\sum_{k=1}^p w_k \cdot d_{ijk}}{\sum_{k=1}^p w_k}$$

Where

- w_k is the weight assigned to variable k ,
- d_{ijk} is the dissimilarity between the k -th variable of objects i and j , and
- d_{ijk} depends on the type of variable:
- For numerical variables

$$d_{ijk} = \frac{|x_{ik} - x_{jk}|}{R_k}$$

Where x_{ik} and x_{jk} are the values of the k -th variable for objects i and j respectively, and R_k is the range of the k -th variable.

- For categorical variables:

$$d_{ijk} = 0 \text{ if } x_{ik} = x_{jk} (\text{objects } i \text{ and } j \text{ have the same category}), \text{ otherwise } d_{ijk} = 1.$$

- For ordinal variables

$$d_{ijk} = \frac{|r(x_{ik}) - r(x_{jk})|}{L_k - 1}$$

where $r(x_{ik})$ and $r(x_{jk})$ are the ranks of the k -th variable for objects i and j respectively, and L_k is the number of levels of the k -th variable.

The Gower distance ranges from 0 (indicating identical objects) to 1 (indicating completely dissimilar objects), and it can handle datasets containing a mix of numerical, categorical, and ordinal variables.

4.5 Visual Similarity – Colorspacious Library

Colorspacious is a powerful and user-friendly library designed for efficient colors space conversions. It supports various standard color spaces, including sRGB, XYZ, xyY, CIE Lab, and CIE LCh. In addition to these common models, Colorspacious offers advanced functionalities such as a complete implementation of the CIECAM02 model, and perceptually uniform spaces like CAM02-UCS, CAM02-LCD, and CAM02-SCD as proposed by Luo et al. (2006) [5].

Conversion of images from their representations in sRGB space to CAM02-UCS space (represented by J, a, b values) allows for an accurate calculation of color difference. This is achieved by calculating the delta-e metric, which determines how different the two colors will appear to an observer [6].

The Delta E (ΔE) equation in the CAM02-UCS colorspace is designed to quantify the difference between two colors based on their J, a, b coordinates. The formula for calculating ΔE in the CAM02-UCS space is as follows:

$$\Delta E_{CAM02-UCS} = \sqrt{(J_1 - J_2)^2 + (a_1 - a_2)^2 + (b_1 - b_2)^2}$$

Where:

- J_1, a_1, b_1 are the J, a, b coordinates of the first color.
- J_2, a_2, b_2 are the J, a, b coordinates of the second color.

4.6 K-Means

K-Means is a clustering algorithm that partitions data points into distinct groups based on their attributes. The primary objective of K-Means is to minimize the distance between points within each cluster while maximizing the distance between different clusters [7].

The K-Means algorithm operates as follows

1. **Initialization:** The user specifies the number of clusters (K) to form. The algorithm randomly selects K initial centroids from the dataset.
2. **Assignment Step:** Each data point is assigned to the nearest centroid based on a distance metric, typically Euclidean distance to form K clusters.
3. **Update Step:** The centroids of each cluster are recalculated by averaging the positions of all data points assigned to that cluster.
4. **Convergence:** Steps 2 and 3 are repeated until the centroids no longer change significantly or until a predetermined number of iterations is reached. At this point, the algorithm has converged.

4.7 Cat Boost

Cat Boost is a state-of-the-art machine learning algorithm developed by Yandex, specifically designed for handling categorical data effectively. It belongs to the family of gradient boosting algorithms, which are known for their high accuracy and robustness in a wide range of predictive tasks. One of the key advantages of Cat Boost is its ability to automatically handle categorical features without the need for extensive pre-processing, such as one-hot encoding or label encoding. This feature makes it particularly well-suited for datasets with a significant number of categorical variables, as it preserves the inherent relationships within the data while reducing the risk of overfitting. Additionally, Cat Boost is recognized for its fast-training speed, even with large datasets, and its capability to produce highly interpretable models. It incorporates a novel technique called Ordered Boosting, which effectively reduces prediction shift and prevents overfitting, making it more reliable in practical applications [8].

4.8 Optuna Library

To optimize the performance of the Cat Boost model, hyperparameter tuning was conducted using the Optuna library, a powerful tool for automated hyperparameter optimization. Optuna's ability to efficiently search the hyperparameter space using techniques like Tree-structured Parzen Estimator (TPE) allowed for the discovery of the optimal set of hyperparameters, enhancing model accuracy and generalization [9].

4.9 Accuracy

Accuracy is a statistical measure used to evaluate the overall performance of a predictive model. It is defined as the ratio of the total number of correct predictions (both true positives and true negatives) to the total number of predictions made by the model. In simpler terms, accuracy reflects how often the model correctly identifies both successful and unsuccessful fabrics.

The formula for accuracy is

$$Accuracy = \frac{True\ Positives\ (TP) + True\ Negatives\ (TN)}{Total\ Predictions\ (TP + TN + FP + FN)}$$

- **True Positives (TP):** The number of correct positive predictions made by the model.
- **True Negatives (TN):** The number of correct negative predictions made by the model.
- **False Positives (FP):** The number of incorrect positive predictions made by the model.
- **False Negatives (FN):** The number of incorrect negative predictions made by the model.

5. Proposed Methodology

This section outlines the methodology employed in the curation of the data that powers the PARAM tool in the backend. The focus of PARAM tool in terms of fashion category is on men's formal and casual shirts. This decision was made for two key reasons: the men's shirt category accounts for the largest share of business revenue, and trend prediction in fashion is particularly challenging in today's fast fashion environment, where trends are short-lived. It was believed that trends in men's shirts tend to be more stable compared to other fashion categories. This section covers the methodologies applied for data preprocessing, predictive analysis, and descriptive analysis in PARAM.

5.1 Data Preprocessing

The initial data pertaining to fabric, product, sales, and inventory inflow is extracted from the business warehouse hosted on Databricks, using SQL queries, and subsequently refined to generate a curated dataset. The fabric related attributes that are captured include fabric pattern, color, fibre composition, weave type, EPI, PPI, fabric finish and business defined assortment groups. An extensive data preprocessing layer is used to transform the raw data into usable format for predictive as well as descriptive analysis. The null value analysis of the showed that several potentially significant columns had a large proportion of missing values and thus were unusable for the analysis. Several other columns (e.g. the weight of the fabric) did contain data, but after consulting with the business, we found the information to be unreliable. As a result, those columns were also excluded. Eventually, filtering out only the pattern, color, fibre composition, weave type, EPI, PPI, fabric finish and business defined assortment groups as the attributes to be considered for the exercise.

Table 1 Summary of Fabric Attributes and Associated Preprocessing Techniques

Attribute Name	Attribute Type	Small Description
Fabric Pattern	Categorical	Essential attribute indicating fabric design (e.g., check, stripe, solid); missing values are dropped.
Fabric Color	Categorical	Contains Pantone codes and business-defined colors; Pantone codes are mapped to RGB values for standardization.
Fibre Constituents	String	Records multiple fiber types in a single string, split into separate columns and combined as needed.
EPI and PPI	Numeric	Ends Per Inch and Picks Per Inch; combined in raw data, split into separate columns, imputed with median values if missing.
Weave Type	Categorical	Describes the fabric structure based on warp and weft interweaving; engineered using a design matrix of 1's and 0's.
Fabric Finish	Categorical	Describes the finish of the fabric, indicating treatments that affect texture, appearance, and performance.
Business Defined Assortment Group	Categorical	A grouping used by the business to categorize fabrics based on specific criteria or market segments.

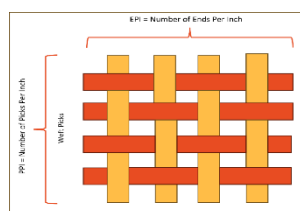


Figure 1 Highlighting EPI (Ends Per Inch) and PPI (Picks Per Inch): Two Key Metrics in Fabric Construction that Influence Texture and Durability.

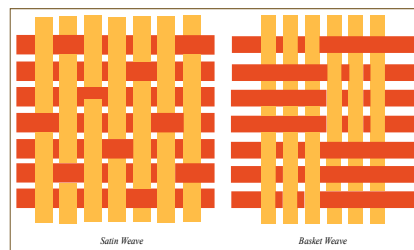


Figure 2 Satin Weave Provides a Glossy Finish, While Basket Weave Showcases a Textured Pattern—Two Examples of the Many Different Weave Type

5.2 Predictive Analysis – PARAM

The predictive analysis in PARAM leverages machine learning to model relationships between historical sales data and corresponding fabric attributes to predict the potential success probability of new season fabrics. The goal of this analysis was

to generate accurate predictions of fabric success, aiding in more data-driven decision-making during the design and buying phases. Initially, we experimented with several models, including logistic regression, decision trees, support vector machines (SVM), random forests, and XGBoost with different loss functions as mentioned in Table.2. However, the performance metrics from these models did not match the accuracy achieved with Cat Boost. In this context, Cat Boost was chosen for its ability to model complex relationships between fabric attributes and their performance outcomes, while efficiently managing the categorical data inherent in the dataset. The algorithm's robust performance and minimal requirement for data preprocessing make it an ideal choice for predictive modelling in the fashion industry, where diverse and complex categorical data is common.

Table 2 Model Specification and Loss Functions for Binary Classification

Model	Criterion / Loss Function
Logistic Regression	Log Loss (Binary Cross-Entropy)
Decision Trees	Gini Impurity / Log Loss
Support Vector Machines (SVM)	Hinge Loss / Squared Hinge Loss
Random Forests	Log Loss (Binary Cross-Entropy)
XG Boost	Log Loss (Binary Cross-Entropy)

We use the data from 2021 and 2022 for training the Cat Boost model. To ensure robust and reliable performance of the predictive models within the PARAM tool, the training process was conducted using the following approach:

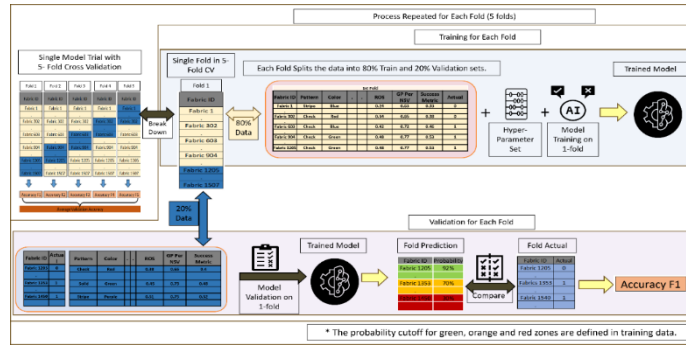


Figure 3 Overview of the Training Process: Implementing 5-fold Cross-Validation with Historical Data and Machine Learning Algorithms to Enhance Fabric Selection and Predict Market Success in Fashion Retail.

5.2.1 Data Splitting and Cross-Validation:

Stratified K-Fold Cross-Validation: Given the balanced nature of the dataset, stratified K-fold cross-validation was employed. This method was chosen to maintain the proportion of the target variable across each fold, ensuring that the model is trained and validated on representative samples of the data. This approach helps in reducing variance in performance metrics and provides a more generalized model evaluation. Stratified 5-fold cross validation was used to train the models to find the best hyperparameters with the help of Optuna (python library for hyperparameter optimization) [10].

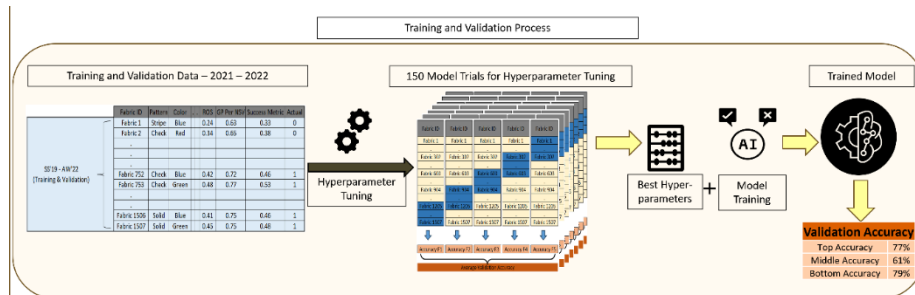


Figure 4 In-Depth Representation of the Process of Training and Validation of each fold in the 5-fold Cross Validation Process.

5.2.2 Hyperparameter Optimization

In our early experiments we noticed that the overall accuracy remained close to 60% on our validation and test sets. Upon further analysis, we identified four distinct clusters of fabrics based on predicted probabilities: those with high probability, above average probability, below average probability, and low probability. The accuracy of classification into successful and unsuccessful was much better, close to 75% for the fabrics with high and low predicted probability. This outcome can be attributed to the inherent complexity of the fashion industry, where certain combinations of attributes have historically led to both successful and unsuccessful fabrics. As such, the models cannot predict that these attribute combinations will consistently

result in successful or unsuccessful performance moving forward; however, for fabrics with a history of clear success or failure, the models reliably assign higher or lower probability scores.

The probability output from the model was, therefore, categorized into four distinct zones: the bottom zone, near bottom zone, near top zone, and top zone. This classification was based on the prediction probabilities that the model assigned to the training data after training. The thresholds for each zone were established as follows:

- **Bottom Zone:** 0 to 25th percentile
- **Near Bottom Zone:** 25th to 50th percentile (median)
- **Near Top Zone:** 50th (median) to 75th percentile
- **Top Zone:** 75th to 100th percentile

These thresholds allowed us to effectively segment the predicted probabilities into meaningful zones for analysis

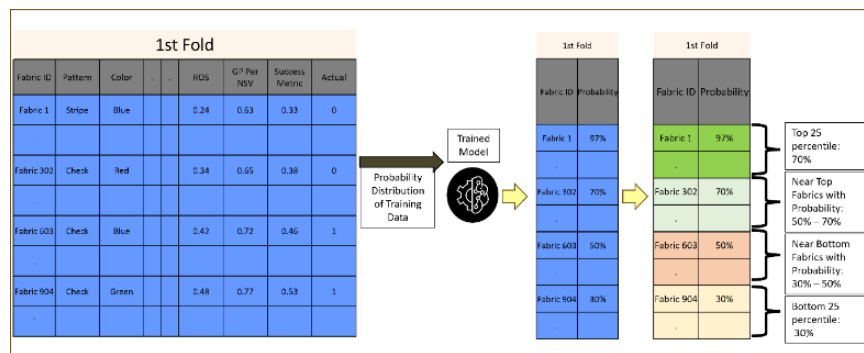


Figure 5 The Process for Splitting the Model Output into 4 Zones: Top, Near To, Near Bottom and Bottom zone.

PARAM recommends that business stakeholders prioritize the purchase of fabrics in the top zone and refrain from selecting those in the bottom zone. In the near bottom and near top zones, PARAM provides valuable analytics to guide decision-making, however, the predictive recommendations are less conclusive. Hence, stakeholders integrate their expertise with our insights to arrive at informed selection choices.

5.2.3 Optimization Objective

We utilized stratified 5-fold cross-validation, which allowed us to obtain a top accuracy and a bottom accuracy from each of the 5-folds. The goal of hyperparameter optimization was to maximize the model's accuracy in both the top and bottom zones. To achieve our objective, we defined a simple objective function: the sum of the average top zone accuracies and the average bottom zone accuracies. Our aim was to maximize this combined accuracy. The goal is to maximize the sum of the top and bottom precision metrics

$$\frac{1}{5} \sum_{i=1}^5 Top\ Fold\ Accuracy_i + Bottom\ Fold\ Accuracy_i$$

- $Top\ Fold\ Accuracy_i$ be the top precision score for fold i .
- $Bottom\ Fold\ Accuracy_i$ be the bottom precision score for fold i .
- 5 be the total number of folds.

Table 3 Defines the Fold Accuracies for Validation

Zone	5-Fold Accuracies	Mean Accuracy
Top Zone	69.76, 77.32, 79.19, 78.11, 77.89	76.45
Bottom Zone	72.73, 72.04, 74.88, 67.82, 71.64	71.82

Once optimal hyperparameters were identified, the model was trained on the complete dataset spanning from 2021 and 2022. The data from 2023 was then used to test the model, with a focus on achieving a business-acceptable accuracy threshold of over 70%. Only after surpassing this benchmark on the test data, we proceed to train a new model on the entire dataset, now including data from 2021 to 2023, to predict the potential performance of fabrics for 2024.

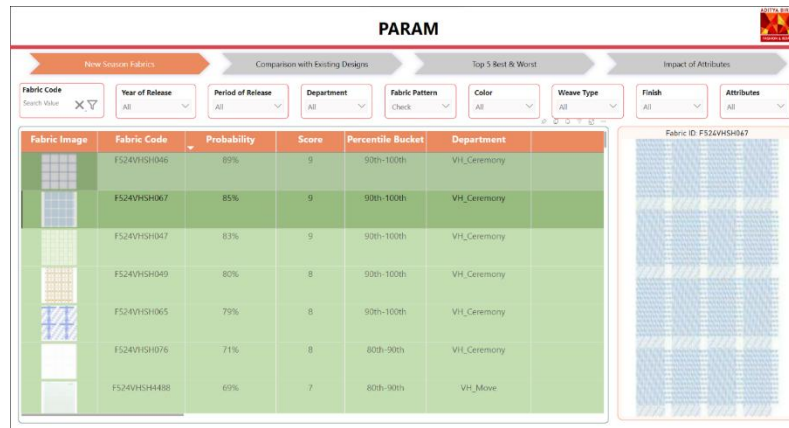


Figure 6 New Season Fabrics view on PARAM Displays all the Newly Designed Fabrics Alongside their Probability Of Success.

The prediction results from the Catboost model are displayed in the new season fabrics view on the tool alongside the fabric images and the fabric attributes as shown in Figure 6.

5.3 Descriptive Analysis – PARAM

Once the fabrics for the upcoming season are designed, analyzing the performance of similar fabrics previously sold provides valuable insights into the potential success of the new season's offerings. Thus, the descriptive analysis feature of PARAM, which identifies the most similar fabrics, is essential for supporting effective decision-making. By streamlining this process, PARAM significantly reduces the effort and time needed for the buying and merchandising team to make well-informed choices. PARAM provides the capability to identify the most similar fabrics based on fabric attributes, images, or a combination of both. This allows end users to find comparable fabrics according to their specific use case and perspective of analysis, offering greater flexibility and precision in fabric selection.

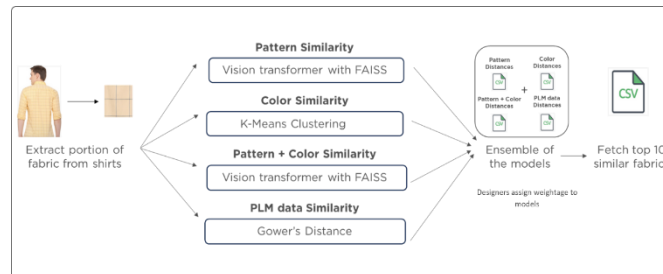


Figure 7 The Overall Process Followed for Identifying Similar Fabrics using an Ensemble of Techniques.

5.3.1 Identification using Fabric Attribute Information and Gower Distance

Identification of similar fabrics is based on the numerical as well as the categorical fabric attributes captured by the design teams. The attributes include business defined assortment group, fabric pattern, season of release, fabric color, cotton%, polyester%, nylon%, bamboo%, viscose fibres%, stretch fibres%, linen fibres%, EPI, PPI and weave type of the fabric. A weighted Gower distance is then utilized as the measurement of similarity between the fabrics. These weights for fabric attributes were decided after multiple rounds of discussions with the business and multiple sets of weights were designed to cater to various kinds of similarity that the end users might be interested in exploring. These similarities can be based of the visual similarity, thereby giving higher weights to fabric color and pattern, or can be based on the structural properties of the fabric thereby giving higher weightages to the weave type and the EPI and PPI of the fabric. Gower distance is calculated for each new fabric from all the previously sold fabrics and 10 fabrics with the least Gower distance are shortlisted to be displayed on the PARAM tool.

5.3.2 Identification using Fabric Images

Identification of similar fabrics based on images begins with the extraction of fabric swatches from garment images. Using the YOLO model for pose estimation, key body landmarks such as the shoulders and hips are detected, enabling precise isolation of the fabric segment for further analysis. Following swatch extraction, two parallel workflows are initiated to capture different aspects of fabric similarity. First, the extracted swatch images undergo color analysis. The images are initially converted from the RGB color space to the CAM02-UCS color space to improve perceptual accuracy in representing color differences [5]. This conversion ensures that color distances reflect human visual perception more accurately. After this transformation, K-means clustering is applied to each image to identify the three most dominant colors, along with their respective proportions within the image.

To calculate the similarity between two images, each with three dominant colors and their associated proportions, the DeltaE color distance metric is employed. The similarity score is derived by comparing each color from one image with every color from the other, resulting in nine total comparisons. The weight of each comparison is determined by multiplying the proportions of the corresponding colors from both images, allowing for a weighted average calculation. This comprehensive approach, which considers both color distribution and perceptual difference, provides a nuanced similarity score that effectively captures the color-based relationship between fabrics.

Table 4 Describe the Delta e Calculations

Proportions	0.75 C _{1a}	0.15 C _{1b}	0.1 C _{1c}
0.5 C _{2a}	$(0.75 \times 0.5) \times \text{Delta-E}(C_{1a}, C_{2a})$	$(0.15 \times 0.5) \times \text{Delta-E}(C_{1b}, C_{2a})$	$(0.1 \times 0.5) \times \text{Delta-E}(C_{1c}, C_{2a})$
0.3 C _{2b}	$(0.75 \times 0.3) \times \text{Delta-E}(C_{1a}, C_{2b})$	$(0.15 \times 0.3) \times \text{Delta-E}(C_{1b}, C_{2b})$	$(0.1 \times 0.3) \times \text{Delta-E}(C_{1c}, C_{2b})$
0.2 C _{2c}	$(0.75 \times 0.2) \times \text{Delta-E}(C_{1a}, C_{2c})$	$(0.15 \times 0.2) \times \text{Delta-E}(C_{1b}, C_{2c})$	$(0.1 \times 0.2) \times \text{Delta-E}(C_{1c}, C_{2c})$

Concurrently, a pattern analysis is conducted by converting the swatch images to grayscale, which isolates the textural and structural features of the fabric without the influence of color. A Vision Transformer (ViT) model, specifically google/vit-base-patch16-224-in21k, is employed to extract feature embeddings from these grayscale images. This model, pre-trained on a diverse set of visual tasks, captures intricate pattern details that are essential for assessing pattern-based similarity. To complement this, the same ViT model is applied to the original color images, thereby generating feature embeddings that incorporate both pattern and color information. This dual approach allows for a comprehensive assessment of fabric similarities, differentiating between purely pattern-based similarities and those influenced by both pattern and color. Output from this was utilized to enhance the result of PARAM fabric similarity.

To harmonize the varied similarity metrics, all calculated distances are normalized using MinMaxScaler [11], ensuring uniformity across different scales. These metrics are then integrated through a weighted ensemble approach, allowing users to assign specific weights to pattern, color, and attribute-based similarities based on their analytical needs. These measures are aggregated to compute a final similarity score for each fabric comparison, reflecting the combined influence of pattern, color, and attribute-based similarities.

Finally, the system ranks fabrics based on their aggregated similarity scores, presenting the top matches for each query. This structured methodology, which combines advanced machine learning models with robust statistical measures, equips stakeholders with precise, data-driven insights for fabric selection and development, ultimately enhancing decision-making and operational efficiency.

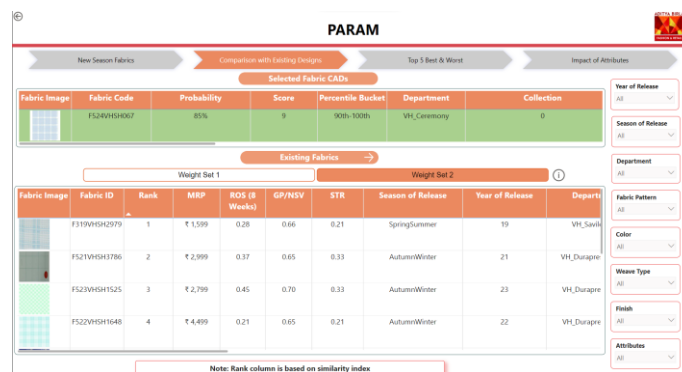


Figure 8 Comparison with Existing Designs view on PARAM Displays the 10 Most Similar Fabrics to the New Season Fabric.

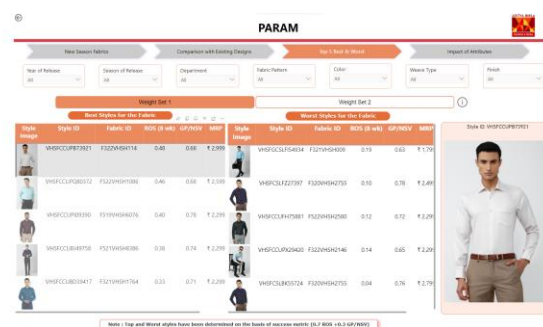


Figure 9 Top 5 Best & Worst view on PARAM Displays the Shirts Made from the most Similar Fabrics Displayed on the Comparison of Existing Design's view.

5.4 Impact of Attributes

The impact of attributes is a descriptive analysis page which helps uncover and compare how different combinations of fabrics attributes have performed in the past. This helps answer questions for designers where they might be contemplating choice between different choices for a certain fabric attribute given that the rest of the attributes remain the same. For example, if a designer is designing a red colored check fabric and wants to understand what weave type have worked best for red colored check fabrics in the past. This is achieved by creating the tree-based drill down visualization using the sales data merged with the fabric attributes data. The page also covers the timeline chart over the first 8-week lifecycle of a product in the store and gets updated every time a certain combination of attributes is chosen.

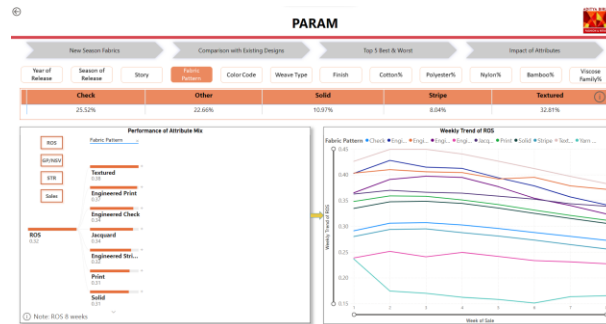


Figure 10 Impact of Attributes view on PARAM displays the overall and the Weekly Performance of Different Combinations of Fabrics Attributes

6. Evaluation

6.1 Predictive Model Performance Criteria

The performance of the Cat Boost model was assessed using accuracy, emphasizing its ability to predict the success or failure of fabrics based on historical data. To evaluate the efficacy of our solution, we focus on the model's ability to accurately predict successful fabrics within the top zone and to identify unsuccessful fabrics within the bottom zone.

In evaluating the performance of the model for the top zone, we consider accuracy because it provides a comprehensive view of how well the model identifies both successful and unsuccessful fabrics. The formula for top zone accuracy is:

$$\text{Top Zone Accuracy} = \frac{\text{No. of successful fabrics predicted}}{\text{No. of successful fabrics predicted} + \text{No. of un - successful fabrics predicted}}$$

Similarly, we define the metric for bottom zone performance evaluation as

$$\text{Bottom Zone Accuracy} = \frac{\text{No. of un - successful fabrics predicted}}{\text{No. of successful fabrics predicted} + \text{No. of un - successful fabrics predicted}}$$

7. Results

7.1 Predictive Model Performance Results

The CatBoost model's performance was optimized through hyperparameter tuning, yielding significant improvements in accuracy. Once the optimal hyperparameters were applied, the model achieved robust results, reducing overfitting and demonstrating strong generalization capabilities across the dataset from 2021 and 2022. Using data from 2021 and 2022, the model achieved a mean 5-fold training accuracy of 76% for the top zone and mean 5-fold accuracy of 72% for the bottom zone. On the test dataset from 2023, the model maintained a top zone accuracy of 71% and the bottom zone accuracy of 72%, surpassing the business threshold of 70%. While many fabrics were accurately classified, a few outliers with unique characteristics presented challenges, highlighting opportunities for further refinement in future iterations. These misclassifications were primarily due to two factors: shifts in market trends and the introduction of entirely new fabric designs that lacked close historical analogues. The retrained model, now incorporating data from 2021 to 2023, generated predictions for over 2500 fabrics designed for 2024. The predictions, supported by high confidence levels, are expected to play a crucial role in prioritizing fabrics for production. Stakeholders can now rely on these insights to make data-backed decisions, enhancing their ability to predict fabric success. By incorporating PARAM's data-driven insights, the fabric selection process has been streamlined, leading to faster decision-making, reduced resource expenditure, and increased market competitiveness. The combination of predictive and descriptive analysis equips stakeholders with a comprehensive tool for both assessing future success and drawing lessons from past performance, ultimately improving product-market alignment.

7.2 Similar Fabrics Performance

The descriptive analysis component, which employs Gower distance, machine learning and a deep learning-based model, effectively identified past fabrics that closely resemble newly designed ones. This comparison offered critical insights into the

historical performance of similar fabrics, enabling stakeholders to draw meaningful parallels between new designs and past successes or failures. By leveraging this historical context, buyers were able to make more informed decisions, refining fabric selection to enhance the likelihood of market success while mitigating risks associated with unsuccessful fabrics.

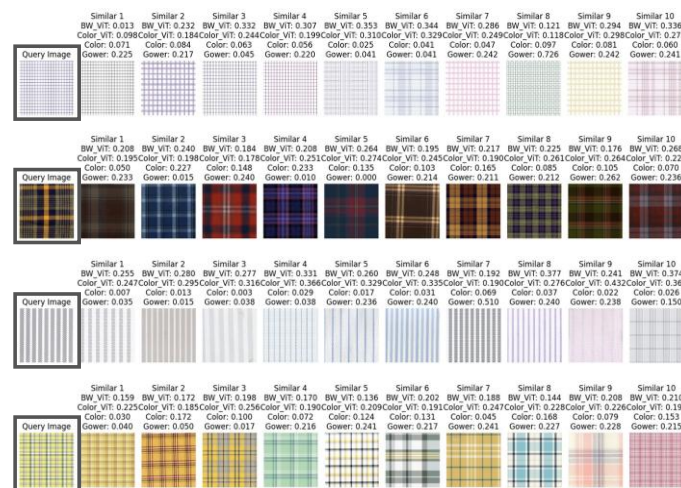


Figure 10 Comparison of the Newly Designed Fabric (Left) with the Ten Most Similar Historical Fabrics (Right). The Similarity Distances Shown Represent: BW_Vit for Grayscale Images, Color_Vit For Colored Images, Color For Distances Based On Color Distributions, And Gower for Distances Based on Fabric Attributes.

8. Conclusion

The implementation of PARAM, an advanced machine learning-based tool, has significantly transformed the fabric selection process in the fashion-retail industry. By combining descriptive and predictive analytics, PARAM empowers stakeholders with actionable insights that drive more informed decision-making, reduce incorrect selections, and increase the overall market competitiveness of new fabric designs.

The descriptive analysis, leveraging Gower distance and deep learning models, offers a robust comparison between newly designed fabrics and their historical counterparts, providing valuable context on past performance. This enables stakeholders to anticipate potential success or failure early in the selection process. The predictive analysis, powered by the CatBoost model, delivers high-accuracy predictions on fabric performance in comparison to traditional methods, enabling businesses to prioritize fabrics with the highest probability of success. The tool's ability to learn from vast historical data and provide real-time predictions has proven instrumental in streamlining the fabric selection process.

However, challenges remain in handling outliers, particularly with fabrics that either deviate from established trends or represent completely new designs without historical analogues. Addressing these limitations will be a key focus in future iterations of PARAM, ensuring that the tool continues to evolve in parallel with emerging trends and innovations in fabric design. Ultimately, PARAM has proven to be a game-changer, allowing the business to optimize fabric selections, and make data-backed decisions. As a result, this tool not only supports more efficient resource utilization but also strengthens the company's ability to maintain a competitive edge in a fast-paced, innovation-driven market. Continued enhancements and iterations of PARAM will further refine its predictive power and broaden its applications, ensuring that the business remains agile and responsive to market dynamics.

9. References

1. Pawan Kumar Singh, Yadunath Gupta, Nilpa Jha, Aruna Rajan, Fashion Retail: Forecasting Demand for New Items, arXiv:1907.01960v1 [cs. OH] 27 Jun 2019.
2. Ahmet Metin and Turgay Tugay Bilgin Bursa Technical University, Bursa, Turkey Automated machine learning for fabric quality prediction: a comparative analysis (23 July 2024). DOI 10.7717
3. Maher Ala'raj1*, Munir Majdalawieh1 and Maysam F. Abbod2(2020), DOI 10.1186, Improving binary classification using filtering based on KNN proximity graphs
4. A.L.D. Loureiro, V.L. Miguéis and Lucas F.M. da Silva, Exploring the use of deep neural networks for sales forecasting in fashion retail, DOI 10.1016, (October 2018)
5. Dosovitskiy, Alexey. "An image is worth 16x16 words: Transformers for image recognition at scale." arXiv preprint arXiv:2010.11929 (2020).
6. Luo, M. R., Cui, G., & Li, C. (2006). Uniform colour spaces based on CIECAM02 colour appearance model. Color Research & Application, 31(4), 320–330. doi:10.1002/col.20227
7. Marco Capó, Aritz Pérez, and Jose A. Lozano, an efficient K-means clustering algorithm for massive data, JOURNAL OF LATEX CLASS FILES, VOL. 14, NO. 8, AUGUST 2015.

8. Liudmila Prokhorenkova^{1,2}, Gleb Gusev^{1,2}, Aleksandr Vorobev¹, Anna Veronika Dorogush¹, Andrey Gulin¹ Yandex, Moscow, Russia, CatBoost: unbiased boosting with categorical features, Moscow Institute of Physics and Technology, Dolgoprudny, Russia, 20 Jan 2019.
9. Optuna: A Next-generation Hyperparameter Optimization Framework Preprint, compiled July 26, 2019, Takuya Akiba¹ *, Shotaro Sano¹, Toshihiko Yanase¹, Takeru Ohta¹, and Masanori Koyama.
10. J.M. Gorriz, F. Segovia, J Ramirez* Department of Signal Theory and Communications University of Granada, Spain, IS K-FOLD CROSS VALIDATION THE BEST MODEL SELECTION METHOD FOR MACHINE LEARNING, January 30, 2024, Department of Communications Engineering University of Malaga, Spain
11. Lucas B.V. de Amorima, b, George D.C. Cavalcantia, Rafael M.O. Cruze, The choice of scaling technique matters for classification performance, arXiv:2212.12343v1 [cs.LG] 23 Dec 2022.
12. Rafiqul Islam, Ronghua Tian, Lynn M. Batten and Steve Versteeg, " Classification of malware based on integrated static and dynamic features". Journal of Network and Computer Applications 36 (2013) 646–656. ELSEVIER.
13. Asaf Shabtai, Robert Moskovitch, Yuval Elovici and Chanan Glezer, " Detection of malicious code by applying machine learning classifiers on static features: A state -of-the-art-survey ", Information Security Technical Report 14 (2009) 16-29, ELSEVIER.
14. Ferreira L, Pilastrri A, Martins CM, Pires PM, Cortez P. 2021. A comparison of AutoML tools for machine learning, deep learning and XGBoost. In: 2021 International joint conference on neural networks (IJCNN). Piscataway: IEEE, 1–8.
15. Feurer M, Klein A, Eggenberger K, Springenberg J, Blum M, Hutter F. 2015. Efficient and robust automated machine learning. Advances in Neural Information Processing Systems 28.
16. Probst P, Boulesteix A-L, Bischl B. 2019. Tunability: importance of hyperparameters of machine learning algorithms. The Journal of Machine Learning Research 20(1):1934–1965.
17. Truong A, Walters A, Goodsitt J, Hines K, Bruss CB, Farivar R. 2019. Towards automated machine learning: evaluation and comparison of AutoML approaches and tools. In: 2019 IEEE 31st international conference on tools with artificial intelligence (ICTAI). Piscataway: IEEE, 1471–1479 DOI 10.1109/ICTAI.2019.00209.
18. Yildirim P, Birant D, Alpyildiz T. 2018. Data mining and machine learning in textile industry. Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery 8(1):1228 DOI 10.1002/widm.1228.
19. Wes McKinney. Pandas: a foundational python library for data analysis and statistics. In SC Workshop on Python for High Performance and Scientific Computing, 2011.
20. Liam Li, Kevin Jamieson, Afshin Rostamizadeh, Ekaterina Gonina, Moritz Hardt, Benjamin Recht, and Ameet Talwalkar. Massively parallel hyperparameter tuning. In NeurIPS Workshop on Machine Learning Systems, 2018.
21. Aaron Klein, Stefan Falkner, Jost Tobias Springenberg, and Frank Hutter. Learning curve prediction with Bayesian neural networks. In ICLR, 2017.
22. Ribeiro R, Pilastrri A, Moura C, Rodrigues F, Rocha R, Cortez P. 2020a. Predicting the tear strength of woven fabrics via automated machine learning: an application of the CRISP-DM methodology. In: 22nd International conference on enterprise information systems (ICEIS) Prague DOI 10.5220/0009411205480555.
23. Ribeiro R, Pilastrri A, Moura C, Rodrigues F, Rocha R, Morgado J, Cortez P. 2020b. Predicting physical properties of woven fabrics via automated machine learning and textile design and finishing features. In: Maglogiannis I, Iliadis L, Pimenidis E, eds. Artificial intelligence applications and innovations. AIAI 2020. IFIP advances in information and communication technology, vol 584. Cham: Springer DOI 10.1007/978-3-030-49186-4_21.
24. Azevedo J, Ribeiro R, Matos LM, Sousa R, Silva JP, Pilastrri A, Cortez P. 2022. Predicting yarn breaks in textile fabrics: a machine learning approach. Procedia Computer Science 207:2301–2310 DOI 10.1016/j.procs.2022.09.289.
25. Bischl B, Binder M, Lang M, Pielok T, Richter J, Coors S, Thomas J, Ullmann T, Becker M, Boulesteix A-L, Deng D, Lindauer M. 2023. Hyperparameter optimization: foundations, algorithms, best practices, and open challenges. Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery 13(2):1484 DOI 10.1002/widm.1484.
26. James Bergstra, Brent Komer, Chris Eliasmith, Dan Yamins, and David D Cox. Hyperopt: A Python library for model selection and hyperparameter optimization. Computational Science & Discovery, 8(1):14008, 2015.
27. James Bergstra, Remi Bardenet, Yoshua Bengio, and Balazs K, egl. Algorithms for hyper-parameter optimization. In NIPS, pages 2546–2554, 2011.
28. BANG LIU, University of Alberta, Canada HANLIN ZHANG, University of Alberta, Canada LINGLONG KONG, University of Alberta, Canada DI NIU, University of Alberta, Canada Fashion Design Classification Based on Machine Learning and Deep Learning Algorithms: A Review, June, 2024) DOI: 10.33022/ijcs.v13i3.3980 License: CC BY-SA 4.0
29. Fashion Design Classification Based on Machine Learning and Deep Learning Algorithms: A Review Ahmed Ali Muhammed¹, Adnan M. Technical College of Informatic, Akre University for Applied Science, Iraq²Duhok Polytechnic University, Duhok, Kurdistan Region, Iraq
30. A Federated Approach for Fine-Grained Classification of Fashion Apparel Authors: Tejaswini Mallavarapu a b, Luke Cranfill a, EunHye Kim a, Reza M. Parizi c, John Morris d, Jung gab Son.