



Sharif University of Technology

<https://sjie.journals.sharif.edu/>



Research Article

Topic Modeling and Sentiment Analysis of Customers Using Natural Language Processing and Machine Learning Techniques

Razieh Rahimi¹ and Reza Ghasemy Yaghin^{*2}

¹Department of Textile Engineering, Amirkabir University of Technology, Tehran, Iran.

²Department of Industrial Engineering and Management Systems, Amirkabir University of Technology.

* corresponding author: (yaghin@aut.ac.ir)

Article Info

Article history:

Received:15 October 2024

Revised:15 January 2025

Accepted:10 March 2025

Keywords:

Sentiment analysis,
machine learning,
natural language processing,
text mining,
customer behavior.

Abstract

Given that modeling and predicting customer behavior using data science helps companies gain a better understanding of customer behavior, this research focuses on analyzing customer reviews in the women's clothing domain within e-commerce. We employ machine learning techniques and natural language processing (NLP) to achieve this goal. The machine learning models used include Support Vector Machine, Logistic Regression, Decision Tree, Random Forest, Multinomial Naive Bayes, Complement Naive Bayes, XGBoost, and LightGBM. To extract and vectorize text features from the reviews, we utilize the TF-IDF and Word2vec algorithms. We employ Topic Modeling using the Latent Dirichlet Allocation (LDA) method and k-means clustering. The dataset consists of women's clothing reviews, with the target variable being customer ratings in those reviews. The study is conducted in binary, three-class, and five-class scenarios. The target variable, which originally has five classes (scores 1 to 5), is categorized into two-class and three-class modes. In the two-class mode, scores below 3 are class zero, while scores of 3 and above are class one. In the three-class mode, scores below 3 are class zero, scores equal to 3 are class one, and scores above 3 are class two. In all three cases, the Random Forest model performs best, achieving an accuracy of 0.98 in the binary case, 0.95 in the three-class case, and 0.91 in the five-class case. After performing the required preprocessing and feature engineering, principal component analysis (PCA) and T SNE are applied. After that, the scatter diagram of the data is drawn and the optimal number of clusters, 3, is estimated using the elbow and silhouette diagrams. In the next step, by removing punctuation marks, stop words and words with fewer than three letters, converting the first letter of the words to lowercase and lemmatization, data cleaning was done. After that, topic modeling was done and each of the topics and words related to them were examined. In the next step, the topics were examined in different clusters. These analyses provide a comprehensive understanding of the key themes and concerns customers have when considering womenswear items in each of the four topics.

Funding: This research did not receive any specific grant from funding agencies in the public, commercial, or not-for-profit sectors.

Conflicts of Interest: The authors declare no conflict of interest.

Author Contributions: All authors contributed to the reported work for conceptualization, methodology, validation, writing-review and editing, and supervision.

To Cite this article:

Rahimi, R. and Ghasemy Yaghin, R. 2026. Topic modeling and sentiment analysis of customers using natural language processing and machine learning techniques, Sharif Industrial Engineering and Management Journal, 41(2), pp.45-61. <https://doi.org/10.24200/ij65.2025.65318.2418>



مدل سازی موضوعی و تحلیل احساس های مشتریان با استفاده از روش های پردازش زبان طبیعی و یادگیری ماشین

راضیه رحیمی^۱ و رضا قاسمی یقین^{۲*}

^۱ دانشکده مهندسی نساجی، دانشگاه صنعتی امیرکبیر، تهران، ایران

^۲ دانشکده مهندسی صنایع و سیستم های مدیریت، دانشگاه صنعتی امیرکبیر، تهران، ایران.

*نویسنده مسئول (yaghin@aut.ac.ir)

چکیده

در نوشتار حاضر، به بررسی و تجزیه و تحلیل نظرهای مشتریان پوشاک زنان با استفاده از روش های یادگیری ماشین و پردازش زبان طبیعی پرداخته شده است. مدل های یادگیری ماشین استفاده شده، شامل: ماشین بردار پشتیبان، رگرسیون لجستیک، درخت تصمیم، جنگل تصادفی، بیز ساده و چندجمله ای، بیز ساده مکمل، XGBoost، و LightGBM بوده اند. به منظور برداری کردن و استخراج ویژگی متن نظرها از الگوریتم های TF-IDF و Word2vec استفاده شده است. مدل سازی موضوعی با استفاده از روش تخصیص دیریکله ی پنهان و خوشه بندی K-Means انجام شده است. مجموعه ی داده ی استفاده شده مربوط به نظرهای پوشاک زنان و متغیر هدف، امتیازهای مشتریان در نظرها بوده است. مطالعه ی حاضر، در حالت های: دو کلاسه، سه کلاسه، و پنج کلاسه انجام و در هر سه حالت، بهترین عملکرد مربوط به مدل جنگل تصادفی با دقت ۰/۹۸ در حالت دو کلاسه، ۰/۹۵ در حالت سه کلاسه، و ۰/۹۱ در حالت پنج کلاسه برآورد شده است.

اطلاعات مقاله

تاریخ دریافت ۱۴۰۳/۰۷/۲۴

تاریخ اصلاحیه: ۱۴۰۳/۱۰/۲۶

تاریخ پذیرش: ۱۴۰۳/۱۲/۲۰

واژگان کلیدی:

تجزیه و تحلیل احساس ها،

یادگیری ماشین،

پردازش زبان طبیعی،

متن کاوی،

رفتار مشتری،

مدل سازی موضوعی.

۱. مقدمه

کاربران اینترنت با نوشتن نظرهای خود و رتبه بندی محصول ها از گیرنده های اطلاعات سنتی به ناشران اطلاعات تغییر پیدا کرده اند؛ که منجر به ذخیره ی حجم قابل توجهی از اطلاعات ارزشمند توسط وبسایت های تجارت الکترونیک می شود، و تجزیه و تحلیل آن ها به منظور به کارگیری روش های بازاریابی مناسب، از توانایی های پردازش شناختی انسان ها پیشی می گیرد.^[۱ و ۲] رتبه بندی و نظرها، نقش مهمی در درک احساس های مشتریان دارند. تحلیل احساس ها در تجارت الکترونیک بسیار مهم است و به کسب و کارها کمک می کند تا بازخورد مشتری را تجزیه و تحلیل و روندها را شناسایی کنند. تحلیل احساس ها، راهبرد استفاده شده برای استخراج و ارزیابی احساس های بیان شده در داده های متنی است،^[۳] و هدف آن استخراج مزایا و معایب محصول از نظرهای مشتریان، پرداختن به تجزیه و تحلیل نظرهای مثبت و منفی، و استخراج اطلاعات برجسته به منظور بهبود محصول، انتخاب برای تولید، و افزایش نرخ خرید است.^[۴] برای تجزیه و تحلیل اطلاعات اخیر و درک احساس های مشتری، می توان از روش های مبتنی بر یادگیری ماشین و پردازش زبان طبیعی استفاده کرد.

به همین منظور در پژوهش حاضر، رویکردی مبتنی بر روش های پردازش زبان

طبیعی و مدل های یادگیری ماشین براساس نظرها و رتبه بندی های مشتریان ارائه شده است، که به کمک آن می توان با آگاهی از احساس های مشتریان خود، خدمات و راهبردهای بازاریابی مناسب را به کار گرفت. در نوشتار حاضر، با استفاده از مجموعه ی داده ی نظرهای پوشاک زنان در خریدهای الکترونیکی به بررسی رضایت یا ناراضیاتی از خرید مشتریان، شناسایی عوامل مؤثر در رفتار آن ها از طریق شناسایی احساس های متن نظرها، مدل سازی موضوعی براساس دسته بندی محصول ها به منظور شناسایی موضوع های کلیدی به همراه خوشه بندی مشتریان و در نهایت، ارائه ی مدل های کلاس بندی احساس ها با استفاده از روش های پردازش زبان طبیعی و یادگیری ماشین پرداخته شده است. درک احساس های نهفته در نظرهای مشتریان اهمیت فراوانی دارد و اطلاعات ارزشمندی در مورد رضایت مشتریان نسبت به محصول ها یا خدمات شرکت ارائه می دهد. احساس های مثبت، نشان دهنده ی خوشحالی و رضایت مشتریان و احساس های منفی، نشان دهنده ی زمینه های نیازمند بهبود هستند. با شناسایی احساس های مذکور، کسب و کارها می توانند راهبردهای خود را برای رسیدگی به نیازها و نگرانی های خاص مشتریان خود تنظیم کنند. به همین منظور هدف کلی نوشتار حاضر، ارائه ی روش ها و مدل هایی است که با ترکیب پردازش زبان طبیعی و یادگیری ماشین، اطلاعات و داده های شرکت های

استناد به این مقاله:

رحیمی، راضیه و قاسمی یقین، رضا، ۱۴۰۴. مدل سازی موضوعی و تحلیل احساس های مشتریان با استفاده از روش های پردازش زبان طبیعی و یادگیری ماشین. مهندسی صنایع و مدیریت شریف، ۴۱ (۲)، صص. ۴۵-۶۱. <https://doi.org/10.24200/j65.2025.65318.2418>



غیرقابل‌اعتماد بودن نتایج و معیارهای ارزیابی مدل‌های پیاده‌سازی شده، موضوع متعادل کردن مجموعه‌ی داده مهم است و باید به آن توجه کرد.

چهار مطالعه در میان کارهای مرتبط انجام شده، بحث متعادل کردن مجموعه‌ی داده را مطرح و بررسی کرده‌اند. کوبروسلی^۱ و همکاران (۲۰۲۲)،^[۶] به منظور انجام این کار به صورت تصادفی، مجموعه‌های متعادل آموزش و آزمایش را انتخاب کرده و توضیح دقیق تری از روش انجام کار ارائه نداده‌اند. سونیل و شیرازی (۲۰۲۳)،^[۸] به منظور متعادل کردن مجموعه‌ی داده از روش افزایش تعداد داده‌ها در کلاس اقلیت استفاده کرده‌اند، اما به روش انجام آن به صورت صریح اشاره‌ای نکرده‌اند. لوکیلی^۲ (۲۰۲۳)،^[۹] در بخش تجزیه و تحلیل اکتشافی، به موضوع متعادل‌نبودن مجموعه‌ی داده با متغیر هدف شاخص پیشنهاد محصول پی برده است. با این حال در ادامه‌ی نوشتارش، توضیحی مبنی بر استفاده از روش‌هایی برای رفع نامتعادل بودن مجموعه‌ی داده در نظر گرفته نشده است. شتی^۳ و همکاران (۲۰۲۳)،^[۱۱] به نامتعادل بودن مجموعه‌ی داده اشاره کرده‌اند؛ اما راهکاری به منظور برطرف کردن آن ارائه نکرده‌اند. با این حال، ایشان بیان کرده‌اند که به دلیل نامتعادل بودن مجموعه‌ی داده، معیار درستی قابل اعتماد نیست و بهتر است از معیارهای دقت و پوشش برای ارزیابی عملکرد مجموعه‌ی داده‌های نامتعادل استفاده شود.

در مطالعه‌ی حاضر، با در نظر گرفتن امتیاز محصولات به عنوان متغیر هدف و به کارگیری روش SMOTE برای رفع مشکل نامتعادل بودن مجموعه‌ی داده و به کمک روش‌های یادگیری ماشین و پردازش زبان طبیعی، به تجزیه و تحلیل احساس‌ها در حالت‌های دو کلاسه، سه کلاسه، و پنج کلاسه پرداخته شده است. روش SMOTE، با استفاده از الگوریتم k -نزدیکترین همسایه به تولید داده‌های مصنوعی جدید می‌پردازد،^[۱۸] و در مقایسه با روش‌های متداول، افزایش تعداد داده‌ها در کلاس اقلیت مانند OverSampling، که ممکن است داده‌های تکراری زیادی به مجموعه اضافه شوند و نتایج را غیرقابل اعتماد کنند، انتخاب بهتری است. همچنین انجام نوشتار حاضر در حالت‌های دو کلاسه، سه کلاسه، و پنج کلاسه به بررسی دقیق تر و با جزئیات بیشتری به رفتار و احساس‌های مشتریان پرداخته است. علاوه بر موضوعاتی که اشاره شد، در نوشتار کنونی، مدل‌سازی موضوعی و برآورد موضوع‌های نهفته در متن نظرها با استفاده از روش تخصیص دیریکله‌ی پنهان نیز انجام شده است. نتایج حاصل از همه‌ی موارد ذکر شده، اطلاعات مهمی در رابطه با شناخت عوامل مؤثر در رفتار مشتریان و تحلیل احساس‌های آن‌ها در اختیار پژوهشگران قرار می‌دهد.

۳. روش‌شناسی

در شکل ۱، مراحل مطالعه‌ی حاضر مشاهده می‌شود؛ که مطابق آن، پس از انتخاب مجموعه‌ی داده‌ی مربوط به نظرهای پوشاک زنان، به انجام تمیز کردن داده‌ها و پیش‌پردازش‌هایی چون مدیریت مقادیر گمشده و حذف ستون مربوط به شماره‌ی ردیف داده‌ها پرداخته شده است.

در مرحله‌ی بعد، ابتدا به منظور شناخت و درک مجموعه‌ی داده به تجزیه و تحلیل اکتشافی پرداخته شده است، که در ادامه در مورد آن بحث شده است. در بخش بعد، پس از انجام پیش‌پردازش به انتخاب ویژگی^۴، استخراج^۵، و مهندسی ویژگی^۶ توسط الگوریتم TF-IDF، متعادل کردن مجموعه‌ی داده با استفاده از

پوشاک را به بینش‌های عمیق و عملی برای رشد و موفقیت آن‌ها تبدیل و از قدرت بازخوردهای مشتریان برای بهبود محصول‌ها، خدمات، و تجربه‌ی کلی مشتری از خرید خود استفاده کنند.

ساختار ادامه‌ی نوشتار به این صورت است، که در بخش ۲، به مرور ادبیات موضوع و بیان کارهای مرتبط؛ در بخش ۳، به ارائه‌ی روش‌شناسی و مبانی نظری؛ در بخش ۴، به ارزیابی مدل‌های پیاده‌سازی شده و ارائه‌ی نتایج به دست آمده؛ و در بخش ۵، به ارائه‌ی نتیجه‌گیری و پیشنهاد‌های آینده پرداخته شده است.

۲. مرور ادبیات موضوع

خلاصه‌ی کارهای مرتبط در جدول ۱ ارائه شده است؛ که مطابق آن، مطالعات پیشین به منظور بررسی نظرهای مشتریان پوشاک در خریدهای الکترونیک با مجموعه داده‌های متفاوت، متغیر هدف، و مدل‌های یادگیری ماشین متفاوت انجام شده‌اند. از آنجایی که مجموعه داده‌ی استفاده شده در نوشتار حاضر با عنوان نظرهای تجارت الکترونیک پوشاک زنانه با برخی نوشتارهای پیشین،^[۳] و^[۱۴-۵] یکسان بوده است؛ در ادامه، به بررسی بیشتر آن‌ها پرداخته شده است. مجموعه داده‌ی مذکور در مطالعات متعددی استفاده شده است، که بیشتر بر تحلیل احساس‌های مشتریان و پیش‌بینی رفتارهای آن‌ها متمرکز بوده‌اند. برخی از آن‌ها، شامل تحلیل اکتشافی داده‌ها و پیش‌بینی احساس‌ها با استفاده از مدل‌های یادگیری ماشین، مانند ماشین بردار پشتیبان و درخت تصمیم، همچنین مقایسه‌ی مدل‌های مختلف یادگیری ماشین برای تحلیل نظرهای مشتریان بوده‌اند. در مقایسه با پژوهش‌های ذکر شده، پژوهش حاضر با استفاده از مدل‌های مختلف یادگیری ماشین تلاش کرده است تا دقت مدل‌ها را در تحلیل و طبقه‌بندی نظرهای مشتریان بهبود بخشد. علاوه بر این، استفاده از روش‌های پردازش زبان طبیعی مانند TF-IDF و Word2Vec در استخراج ویژگی‌ها و مدل‌سازی موضوعی با روش تخصیص دیریکله‌ی پنهان، از ویژگی‌های خاص پژوهش حاضر است، که به دقت بیشتر در پیش‌بینی احساس‌ها و دسته‌بندی نظرها کمک می‌کند. لذا، با استفاده از ترکیب روش‌های مختلف و ارزیابی دقیق تر مدل‌ها، به نتایج متفاوتی نسبت به سایر مطالعات مشابه دست یافته است. در برخی مطالعات،^[۳، ۹-۱۱] متغیر هدف در نظر گرفته شده به منظور پیاده‌سازی مدل‌های یادگیری ماشین، ستون ویژگی شاخص پیشنهاد محصول در مجموعه‌ی داده و در حالت دو کلاسه (۰ برای پیشنهاد نکردن محصول توسط مشتری و ۱ برای پیشنهاد محصول توسط مشتری) انجام شده است. همچنین، در مطالعه‌های دیگری،^[۱۲ و ۱۳] متغیر هدف در نظر گرفته شده، ستون ویژگی امتیاز محصول‌ها و در حالت دو کلاسه با مرز امتیاز ۳ (امتیازهای برابر ۳ و بالاتر به عنوان کلاس ۱ و امتیازهای کمتر از ۳ به عنوان کلاس ۰) صورت گرفته‌اند. موضوع قابل توجه این است که مجموعه داده‌ی استفاده شده در همه‌ی مطالعات یکسان بوده و همان‌طور که در ادامه با تجزیه و تحلیل داده‌های اکتشافی نشان داده شده است، مجموعه داده‌ی مذکور برای متغیرهای هدف شاخص، پیشنهاد محصول، و امتیاز مشتریان متعادل نبوده است. از آنجایی که بسیاری از مدل‌های یادگیری ماشین، مانند رگرسیون لجستیک، هر زمان که اختلاف نسبت بین تعداد نقاط داده نامتعادل باشد، با مشکلات عملکردی مواجه می‌شوند،^[۸] و به منظور جلوگیری از

⁴ Feature Selection

⁵ Feature Extraction

⁶ Feature Engineering

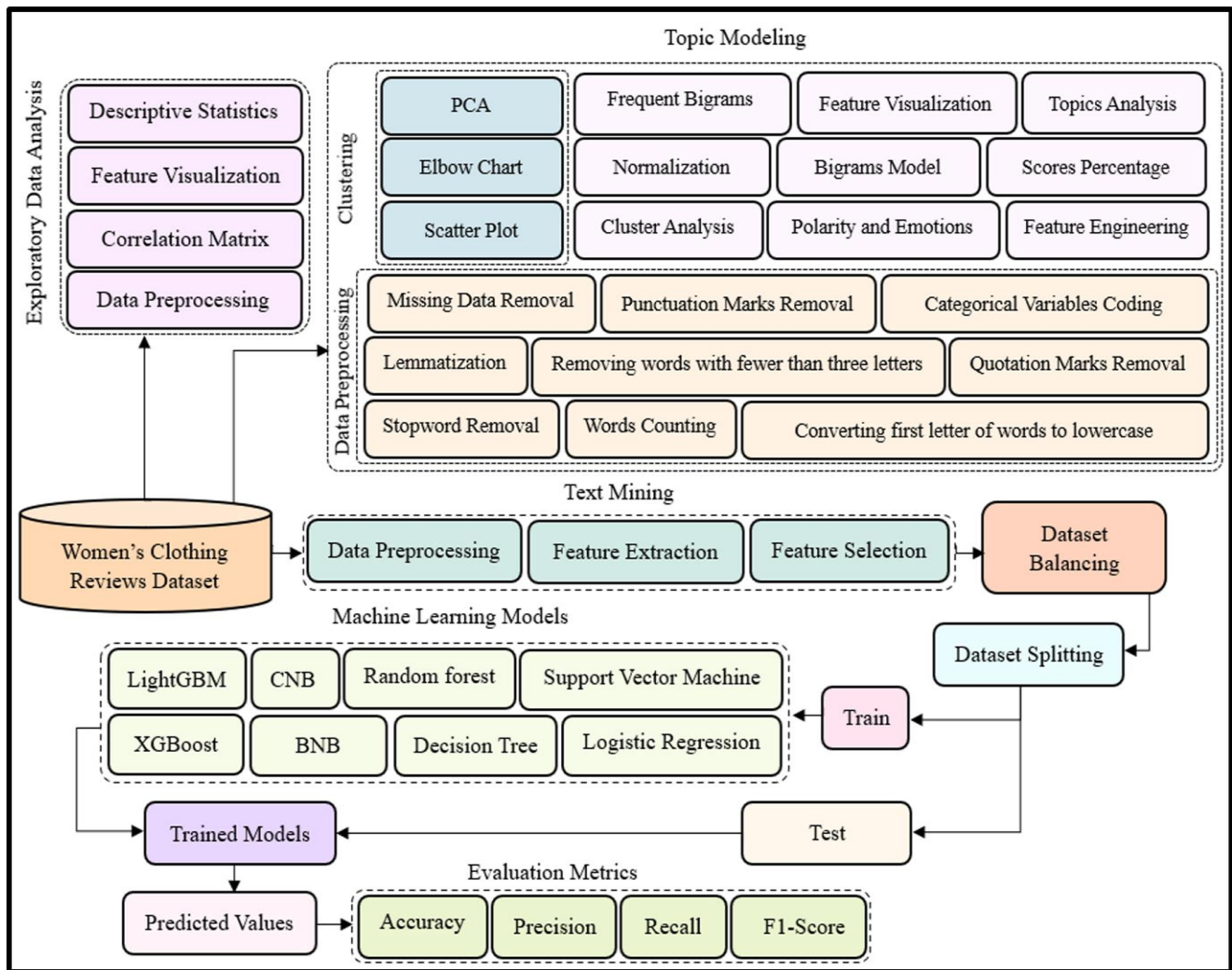
¹ Kubrusly

² Loukili

³ Shetty

جدول ۱. خلاصه‌ی کارهای مرتبط.

Title and publication year	Dataset	Target variable	Dataset balancing	ML models
Sentiment analysis of e-commerce customer reviews based on natural language processing 2020 ^[5]	Women's e-commerce clothing reviews	Recommended Index	-	LR, SVM, RF, XGBoost, LightGBM
A statistical analysis of textual e-commerce reviews using tree-based methods 2022 ^[6]	Women's e-commerce clothing reviews	Recommended Index	Random selection of balanced training (5,840) and test (2,504) Sets	RF, XGBoost, DT, GB
Analyzing online reviews of customers using machine learning techniques 2022 ^[7]	Women's e-commerce clothing reviews	Recommended Index	-	NB, LR, SVM, NN
Customer review classification using machine learning and deep learning techniques 2023 ^[8]	Women's e-commerce clothing reviews	Recommended Index	Enhancing the minority class sample size	LR, Linear SVM, DT, RF, AdaBoost
Explainable sentiment analysis for textile personalized marketing 2023 ^[3]	Women's e-commerce clothing reviews	Recommended Index	-	LR
Sentiment analysis of product reviews for e-commerce recommendation based on machine learning 2023 ^[9]	Women's e-commerce clothing reviews	Recommended Index	-	LR, RF, KNN, CatBoost
Exploring commonly used terms from online reviews in the fashion field to predict review helpfulness 2023 ^[16]	Customer reviews on amazon's fashion platform	Positive feedback count	-	RF, SVM, LR, SGB, TM, LDA
It is about inclusion! Mining online reviews to understand the needs of adaptive clothing customers 2023 ^[17]	Customer reviews from Amazon, Silverts, IZ-Adaptive, and ResellerRatings	-	-	TM, LDA
Using Topic Modeling for Extracting Customers' Expectations: a case of women apparel 2023 ^[10]	Women's e-commerce clothing reviews	-	-	TM, LDA
Sentiment forecasting in women's fashion e-Commerce: a machine learning perspective 2024 ^[18]	Reviews on Women's Apparel	Sentiment	-	NB, SVM, NN, LR
Hyperparameter optimization of machine learning models using grid search for amazon review sentiment analysis 2024 ^[11]	Women's e-commerce clothing reviews	Recommended Index	-	DT, LR, SVM, NB, XGBoost, RF
Sentiment exploring on feedback of e-commerce data using machine learning algorithms 2024 ^[12]	Women's e-commerce clothing reviews	Customer Ratings (Binarized with a Threshold of 3)	-	LR, NB, MNB, BNB, SVM, RF, AdaBoost
Factors influencing recommendations for women's clothing satisfaction: a latent dirichlet allocation approach using online reviews 2024 ^[14]	Women's e-commerce clothing reviews	-	-	TP, LDA
This study: Topic Modeling and Sentiment Analysis of Customers Using Natural Language Processing and Machine Learning Techniques	Women's e-commerce clothing reviews	Customer Ratings in Binary, Three-class, and Five-class Settings	SMOTE	LR, SVM, DT, RF, CNB, MNB, XGBoost, LightGBM, TM, LDA, NLP



شکل ۱. روش‌شناسی مطالعه‌ی حاضر.

۲.۳. تجزیه و تحلیل داده‌های اکتشافی

به منظور انجام تجزیه و تحلیل داده‌های اکتشافی (EDA)^{۱۴}، پس از بررسی و حذف مقادیر داده‌ی گمشده^{۱۵}، ابتدا محاسبه‌ی پارامترهای آمار توصیفی برای ویژگی‌های عددی انجام و سپس ماتریس همبستگی ویژگی‌های مذکور ترسیم شده است. در ادامه، به تحلیل ویژگی‌ها پرداخته شده است. پس از حذف مقادیر داده‌های گمشده، تعداد ردیف‌های مورد بررسی مجموعه‌ی داده برابر با ۱۹۶۶۲ در نظر گرفته شده است. میانگین گروه سنی افرادی که نظرهای خود را در رابطه با پوشاک خریداری‌شده‌ی زنان به اشتراک گذاشته‌اند، حدود ۴۳ سال و در بازه‌ی ۱۸ تا ۹۹ سال برآورد شده است. میانگین امتیازها برابر با ۴/۱ و میانگین شاخص پیشنهاد محصول ۰/۸۱ محاسبه شده و نشان‌دهنده‌ی آن است که در حالت کلی، مشتریان از خرید خود رضایت داشته‌اند. تعداد مشتریانی که نظرهای ثبت‌شده را مفید و کمک‌کننده دانسته‌اند، نیز از ۰ تا ۱۲۲ عدد متغیر بوده است. در شکل ۲، ماتریس همبستگی ویژگی‌های عددی مشاهده می‌شود؛ که مطابق آن، بیشترین مقدار ضریب همبستگی مربوط به ویژگی‌های امتیاز و شاخص پیشنهاد محصول و برابر با ۰/۷۹ برآورد شده است. شکل ۳ (a)، نحوه‌ی

روش SMOTE، و تفکیک آن به دو دسته‌ی مجموعه‌ی داده‌های آموزش و آزمایش صورت گرفته است. در نهایت، مدل‌های یادگیری ماشین مورد نظر، پیاده‌سازی، و عملکرد آن‌ها ارزیابی شده است.

۳.۱. جمع‌آوری داده

مجموعه‌ی داده‌ی استفاده‌شده در مطالعه‌ی حاضر، شامل داده‌های تجاری واقعی و با عنوان نظرهای تجارت الکترونیک پوشاک زنانه^۷، از وبسایت Kaggle دانلود شده است.^[۱۹] مجموعه‌ی داده شامل ۲۳۴۸۶ ردیف و ۱۰ ستون ویژگی بوده است و هر ردیف آن بیان‌کننده‌ی شناسه‌ی قطعه‌ی در حال بررسی، سن مشتری صاحب‌نظر، عنوان نظر، متن نظر، امتیاز مشتری به محصول مورد نظر از ۱ (بدترین) تا ۵ (بهترین)، شاخص پیشنهاد محصول توسط مشتری، تعداد مشتریانی که این نظر را مثبت و کمک‌کننده دانسته‌اند، نام بخش محصول سطح بالا^۸، نام بخش محصول، و نام کلاس محصول است. ستون‌های ویژگی مجموعه‌ی داده‌های اخیر، شامل: متغیرهای عددی^۹، رشته-ای^{۱۰}، عددی ترتیبی^{۱۱}، عددی دودویی^{۱۲}، و طبقه‌ای^{۱۳} می‌شود.

^{۱۲} Binary variable

^{۱۳} Categorical

^{۱۴} Exploratory Data Analysis

^{۱۵} Missing values

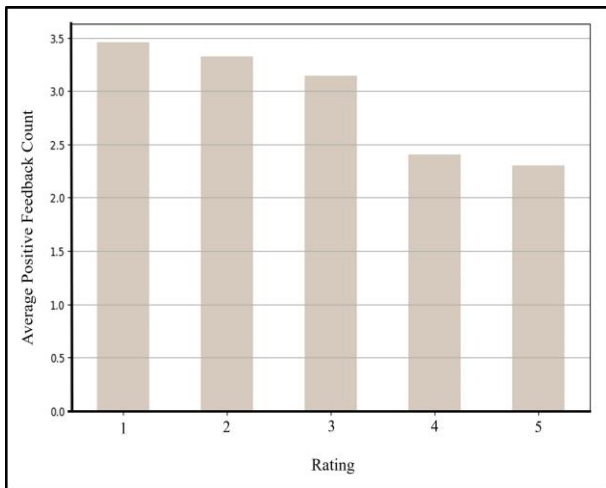
^۷ Women's clothing e-commerce reviews

^۸ Product high level division

^۹ Numerical

^{۱۰} String

^{۱۱} Ordinal variable



شکل ۵. میانگین تعداد مفیدبودن نظرها برای هر امتیاز.

تعداد کل مجموعه‌ای داده داشته و پس از آن به ترتیب امتیازهای ۴، ۳، ۲، و ۱ قرار گرفته‌اند. بیشتر امتیازها توسط افراد در محدوده سنی ۲۰ تا ۷۵ سال ثبت شده و بیشترین نظرها، امتیاز ۳ و بالاتر داشته‌اند. در شکل ۵، میانگین تعداد مفیدبودن نظرها برای هر امتیاز مشاهده می‌شود؛ که مطابق آن، بیشترین مقدار میانگین تعداد مفیدبودن نظرها، مربوط به امتیازهای ۱ تا ۳ بوده است.

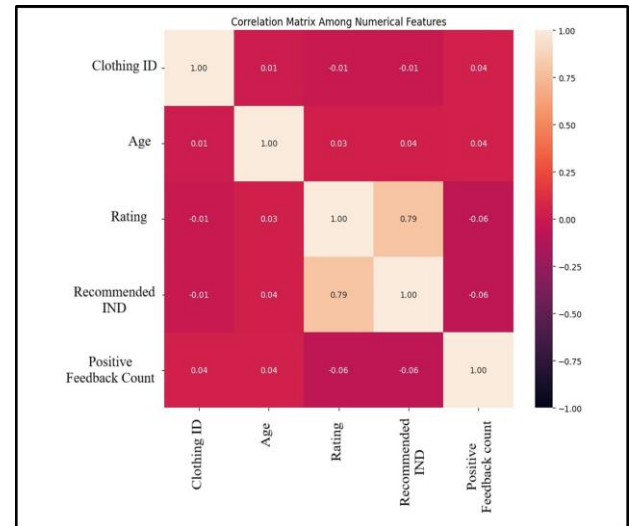
۳.۳. مدل‌سازی موضوعی

در بخش کنونی، ابتدا به بیان توضیح مختصری از روش تخصیص پنهان دیریکله (LDA)^{۱۶} در مدل‌سازی موضوعی^{۱۷}، خوشه‌بندی^{۱۸}، و چرایی استفاده از دو روش مذکور در نوشتار حاضر پرداخته شده است. در ادامه، روش شناسی بخش حاضر، بیان و نتایج برآوردشده گزارش شده است.

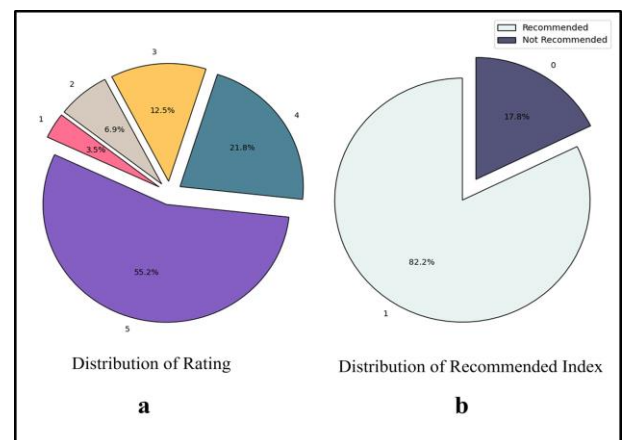
- **تخصیص پنهان دیریکله**، شکلی از یادگیری ماشین بدون نظارت است، که برای کشف موضوع‌های پنهان از متن نظرها استفاده می‌شود و هدف آن، یافتن گروهی از موضوع‌هاست که کلمات مشاهده‌شده را در تمام اسناد به بهترین شکل توصیف کنند. اطلاعات تولیدشده از مدل ذکرشده، شامل واژگان کلیدی مرتبط با هر یک از موضوع‌ها و احتمال مرتبط شدن هر یک از نظرها متن با هر موضوع است.^[۲۰]

- **خوشه‌بندی K-means**، از الگوریتم‌های یادگیری بدون نظارت است، که برای طبقه‌بندی مجموعه‌ای از داده‌ها به کلاس‌هایی از داده‌های مشابه استفاده می‌شود. طبقه‌بندی مجموعه‌ای داده‌ی N مورد براساس ویژگی‌ها به k زیرمجموعه‌ی مجزا با کمینه‌سازی فاصله بین آیت‌م داده و مرکز خوشه‌ی مرتبط انجام می‌شود.^[۲۱]

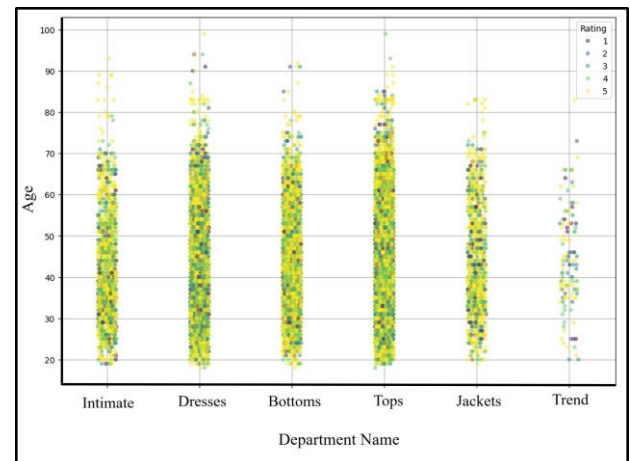
در نوشتار حاضر، برای انجام مدل‌سازی موضوعی از روش تخصیص دیریکله‌ی پنهان و خوشه‌بندی K-Means به‌طور ترکیبی استفاده شده است. تخصیص دیریکله‌ی پنهان، به‌عنوان یک روش مبتنی بر احتمال، توانایی شبیه‌سازی توزیع موضوع‌های پنهان در اسناد مختلف را دارد و می‌تواند به‌طور مؤثری، موضوع‌های غالب در متن‌های مشتریان را شناسایی کند. روش تخصیص دیریکله‌ی پنهان به تحلیل ساختار معنایی داده‌ها کمک می‌کند و امکان استخراج ویژگی‌های دقیق موضوعی از متون نظرها مشتریان را فراهم



شکل ۲. ماتریس همبستگی ویژگی‌های عددی.



شکل ۳. (a) نحوه‌ی توزیع امتیازها، (b) توزیع شاخص پیشنهاد محصول.



شکل ۴. امتیازهای مشتریان با توجه به گروه سنی و نام بخش.

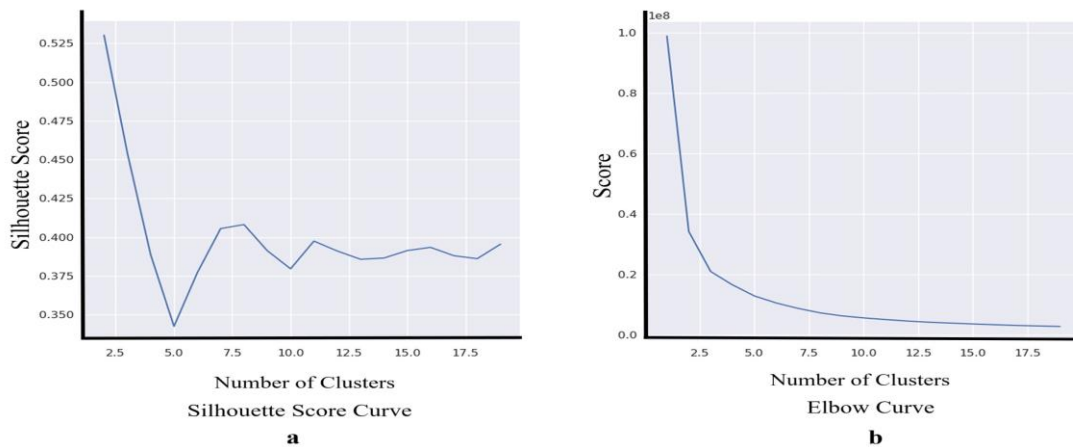
توزیع امتیازهای موجود در مجموعه‌ی داده، شکل ۳ (b)، نحوه‌ی توزیع شاخص پیشنهاد محصول توسط مشتریان، و شکل ۴، امتیازهای مشتریان را با توجه به گروه سنی و نام بخش نشان می‌دهند.

همان‌طور که مشاهده می‌شود، ۸۲/۲٪ از مشتریان، محصول خریداری‌شده را به دیگران پیشنهاد می‌کنند. امتیاز ۵ با مقدار ۵۵/۲٪، بیشترین سهم را در

¹⁸ Clustering

¹⁶ Latent Dirichlet Allocation

¹⁷ Topic Modeling



شکل ۶. (a) نمودار میانگین سیلوئت، (b) نمودار Elbow.

مرحله‌ی بعد، میزان قطبیت و احساس‌های متن‌نظرها به همراه عنوان آن‌ها، برآورد و در ستونی با عنوان امتیاز احساس‌های کل به مجموعه‌ی داده افزوده شده است. به‌منظور بهبود عملکرد الگوریتم خوشه‌بندی K-Means، کاهش ابعاد و رسم مجموعه‌ی داده، داده‌ها با استفاده از تجزیه و تحلیل مؤلفه‌های اصلی (PCA)^{۱۹} و سپس تحلیل T-SNE^{۲۰} به ابعاد پایین‌تر کاهش یافته و به عنوان ورودی الگوریتم K-Means در نظر گرفته شده‌اند. T-SNE یک روش کاهش ابعاد غیرخطی است، که برای نمایش ساختارهای پیچیده‌ی داده‌ها در فضای دو یا سه بعدی بسیار مؤثر است و امکان خوشه‌بندی مؤثرتر داده‌های پیچیده را با حفظ روابط غیرخطی میان نقاط داده فراهم می‌کند.^[۲۳] پس از آن، با استفاده از نمودار Elbow و معیار سیلوئت، تعداد بهینه‌ی خوشه‌ها برآورد شده است. روش elbow برای تعیین تعداد بهینه‌ی خوشه‌ها در الگوریتم K-Means براساس معیار مجموع مربعات خطا (SSE)^{۲۱} عمل می‌کند. معیار SSE، مجموع فاصله‌های مربعات داده‌ها تا نزدیک‌ترین مرکز خوشه را اندازه‌گیری و به عنوان شاخصی برای پراکندگی داده‌ها درون خوشه‌ها عمل می‌کند. برای این منظور، SSE برای تعداد خوشه‌های مختلف محاسبه و مقادیر به‌دست‌آمده در قالب یک نمودار رسم و نقطه‌ای که در آن کاهش SSE روند کندتری پیدا کرده است، به عنوان مقدار بهینه‌ی تعداد خوشه‌ها انتخاب شده است. کاهش SSE به این معناست که داده‌ها به خوشه‌های مجزا و منسجم‌تری تقسیم می‌شوند.^[۲۴] میانگین سیلوئت نیز به منظور ارزیابی کیفیت خوشه‌بندی داده‌ها استفاده شده است، که نشانگر انسجام درون خوشه‌ای و جداسازی بین خوشه‌هاست؛ که برای محاسبه‌ی آن، فاصله‌ی هر نقطه از نقاط دیگر در خوشه‌ی خود و نزدیک‌ترین خوشه محاسبه می‌شود.^[۲۵] در مطالعه‌ی حاضر، مقدار سیلوئت برای تعداد دو خوشه برابر با ۰/۵۳ و برای تعداد سه خوشه برابر با ۰/۴۵ به‌دست آمده است. این تفاوت نشان می‌دهد که هر دو تعداد خوشه (دو و سه خوشه) کیفیت مشابهی در خوشه‌بندی داده‌ها داشته‌اند. از طرفی، با توجه به نمودار elbow، تعداد سه خوشه به عنوان تعداد بهینه‌ی خوشه‌ها با کمترین مجموع مربعات خطا و تقسیم‌بندی بهتر داده‌ها شناسایی شده است. با توجه به اینکه نتایج روش elbow و امتیاز سیلوئت به‌طور کلی هم‌راستا هستند و تعداد سه خوشه، هم‌زمان کیفیت بالاتری در خوشه‌بندی و انسجام بهتر داده‌ها ارائه می‌دهد، تعداد سه خوشه به عنوان تعداد بهینه‌ی خوشه‌ها در نظر گرفته شده است. در نهایت، عملیات خوشه‌بندی با استفاده از الگوریتم K-means انجام شده است. در شکل ۶ (a)، نمودار میانگین سیلوئت و در شکل

می‌آورد. به‌علاوه، به دلیل سادگی در پیاده‌سازی، قابلیت تفسیر بالا، و کارایی آن در تحلیل داده‌های متنی با حجم بالا به عنوان یکی از روش‌های محبوب در پردازش زبان طبیعی شناخته می‌شود. به‌ویژه برای مسائلی که نیاز به شناسایی و استخراج موضوع‌های پنهان از مجموعه‌های بزرگ داده دارند، انتخابی مناسب و کارآمد است.^[۲۶] در کنار آن، الگوریتم خوشه‌بندی K-Means برای تقسیم‌بندی اسناد به خوشه‌های مختلف، براساس توزیع‌های موضوعی به‌دست‌آمده از تخصیص پنهان دیریکله استفاده شده است. ترکیب مذکور، داده‌ها را براساس شباهت‌های موضوعی به‌طور مؤثری گروه‌بندی می‌کند و تحلیل‌های دقیق‌تری را از نظرهای مشتریان امکان‌پذیر می‌سازد. علاوه بر این، در مطالعه‌ی حاضر، ارتباط موضوع‌های غالب در هر خوشه با سن مشتریان و تعداد کلمات در متن نظرهای مشتریان بررسی شده است. تحلیل‌های انجام‌شده نشان می‌دهند که چگونه سن و طول متن نظرها می‌تواند بر تمایل به صحبت در مورد موضوع‌های خاص تأثیر بگذارد و به درک بهتر الگوهای رفتاری و اولویت‌های مشتریان کمک کند.

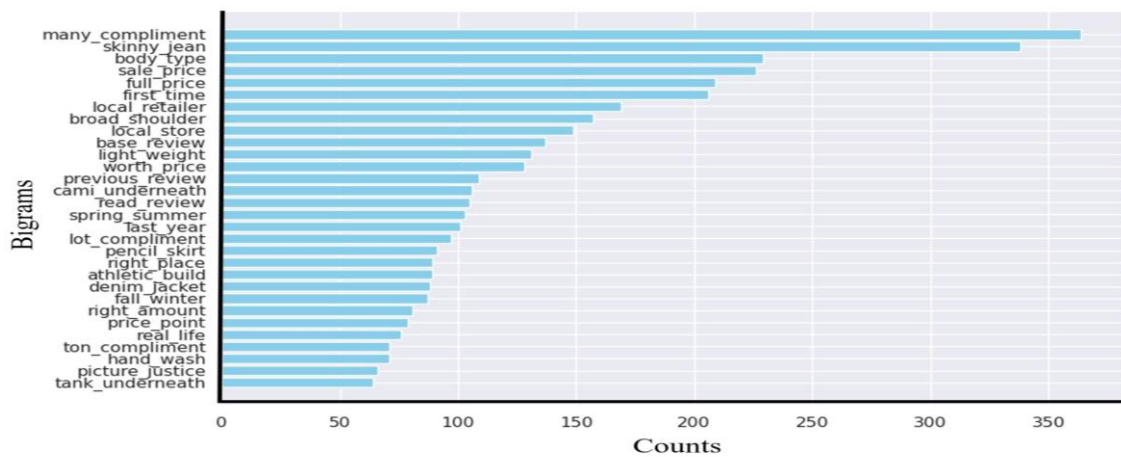
بررسی هر خوشه به‌طور ویژه مزایای زیادی از منظر راهبردهای بازاریابی و توسعه‌ی کسب و کار در بر دارد. با شناسایی ویژگی‌های خاص هر خوشه، کسب‌وکارها می‌توانند راهبردهای بازاریابی متناسب با نیازها و خواسته‌های خاص مشتریان در هر گروه را طراحی کنند. برای مثال، یک خوشه ممکن است تمایل زیادی به ویژگی‌های خاص محصول نشان دهد، که نیاز به تبلیغات هدفمند و تشویق به خرید دارد؛ در حالی که خوشه‌ای دیگر ممکن است بیشتر به خدمات پس از فروش یا تجربه‌ی مشتری علاقه‌مند باشد. به‌علاوه، تحلیل موضوع‌های موجود در هر خوشه می‌تواند به کسب‌وکارها کمک کند تا پیشنهادها، محصولات و خدمات را براساس اولویت‌های خاص مشتریان هر خوشه سفارشی کند و تجربه‌ی مشتری را بهبود بخشد. در نتیجه، روش مذکور نه فقط به تحلیل دقیق‌تر رفتار مشتریان کمک می‌کند، بلکه می‌تواند راهکارهای مؤثری برای رشد کسب‌وکار و جذب مشتریان بیشتر ارائه دهد.

پس از انجام پیش‌پردازش‌هایی مانند مدیریت مقادیر گمشده، کدبندی ویژگی‌های طبقه‌ای، شمارش تعداد کلمات، و تعداد کلمات یکتا، شناسه‌های لباس با فراوانی کمتر از ۱٪ تعداد کل، شناسایی و شناسه‌ی یکسانی به آن‌ها نسبت داده شده است. سپس تعداد تکرار هر یک از شناسه‌های لباس، نرمال و در ستونی با نام شناسه‌ی لباس جدید به مجموعه‌ی داده اضافه شده است. در

²¹ Sum of Squared Errors

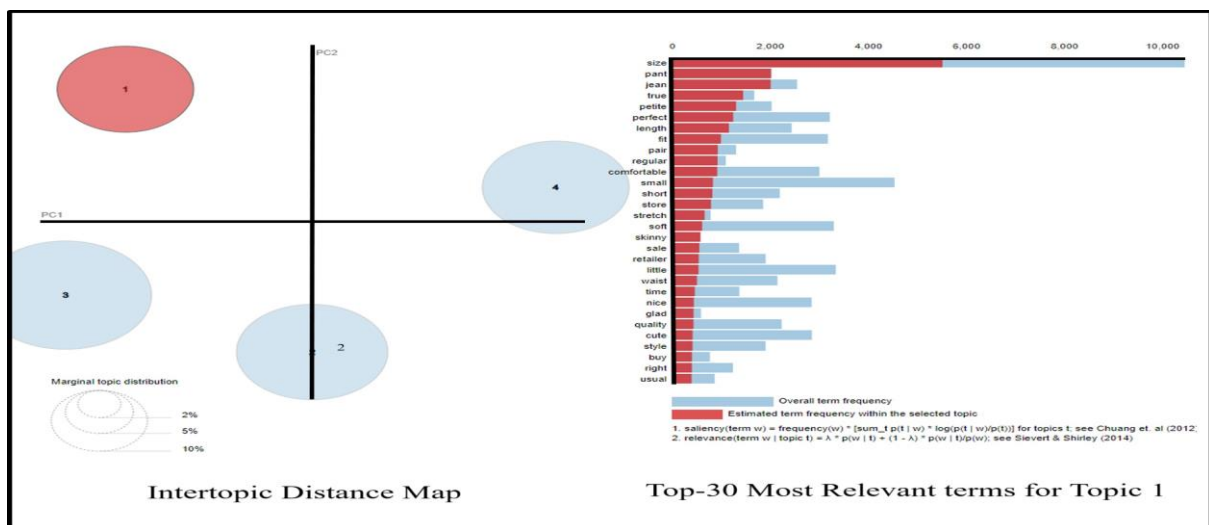
¹⁹ Principal component analysis

²⁰ T-distributed Stochastic Neighbor Embedding



Top 30 Bigrams and their Counts

شکل ۷. سی بایگرام پرتکرار.



شکل ۸. نمودار فاصله‌ی بین موضوعی و کلمات پرتکرار موضوع یک.

مانند: شلوار جین، دامن، ژاکت، و غیره، تناسب و اندازه‌ی لباس‌ها و همچنین فصل‌های سال نیز از موضوعات مهم در نظر گرفته شده‌اند. در شکل ۸، نمودار فاصله بین موضوعی و کلمات پرتکرار مربوط به موضوع یک؛ در شکل ۹، نمودار فاصله بین موضوعی و کلمات پرتکرار مربوط به موضوع دو؛ در شکل ۱۰، نمودار فاصله بین موضوعی و کلمات پرتکرار موضوع سه؛ و در شکل ۱۱، نمودار فاصله بین موضوعی و کلمات پرتکرار مربوط به موضوع چهار مشاهده می‌شوند. همچنین، در جدول ۲، بررسی هر یک از موضوع‌ها و کلمات مرتبط با آن‌ها به صورت مختصر ارائه شده است. تحلیل‌های صورت گرفته، درک جامعی از موضوع‌های کلیدی و نگرانی‌هایی که مشتریان هنگام بررسی اقلام لباس زنانه در هر یک از چهار موضوع دارند، ارائه می‌دهند. موضوع قابل توجه این است که در جدول اخیر، ترتیب دسته‌بندی هر یک از موضوع‌ها براساس تعداد کلمات به کاررفته در آن دسته است. بنابراین در موضوع یک، پراهمیت‌ترین مورد: تناسب و اندازه؛ در موضوع دو: پارچه و راحتی؛ در موضوع سه: تناسب و اندازه؛ و در موضوع چهار، تناسب و راحتی هستند. پس از آن، تخصیص دیریکله‌ی پنهان برای حالت‌های ۲ تا ۲۳ موضوعی آزمایش و با توجه به معیار coherence، بهترین تعداد موضوع‌ها، چهار در نظر گرفته شده است.

۶ (b)، نمودار Elbow برای محاسبه‌ی تعداد خوشه‌های بهینه مشاهده می‌شوند.

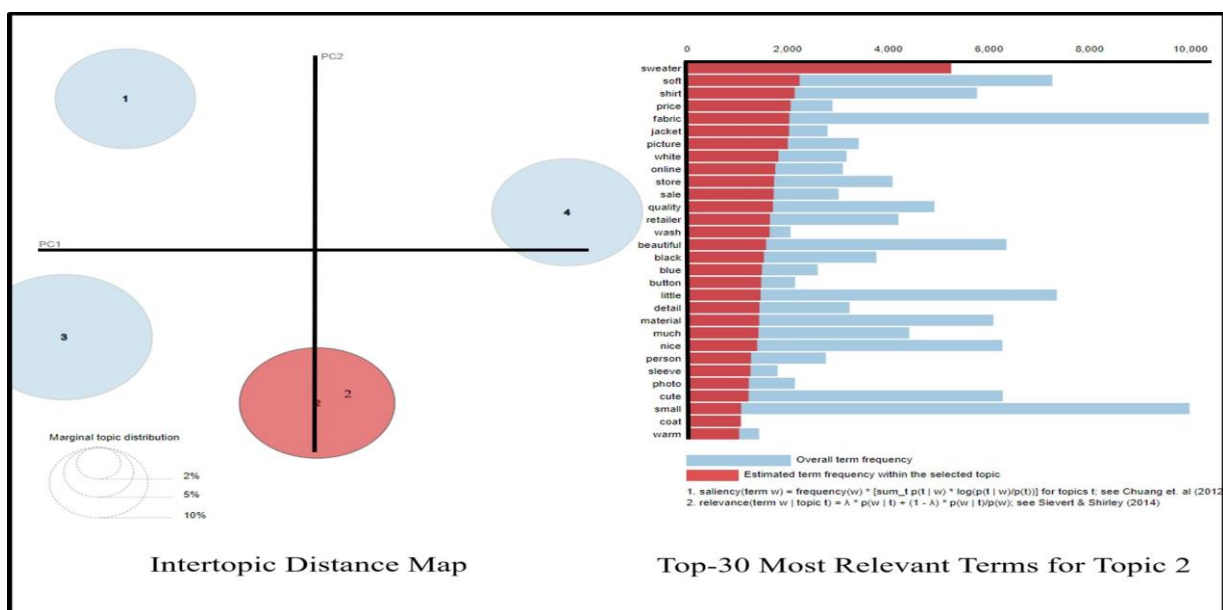
در مرحله‌ی بعد، برای اطمینان حاصل کردن از اینکه کلمات به‌درستی براساس بخش گفتارشان^{۲۲} به صورت کلمه بیان می‌شوند، با استفاده از یک تابع، اسم، صفت، قید، و یا فعل بودنشان مشخص شده است. پس از آن با حذف علامت‌های نگارشی، کلمات توقف و کلمات دارای کمتر از سه حرف، تبدیل حرف اول کلمات به حرف کوچک و ریشه‌یابی^{۲۳} تمیزسازی داده‌ها انجام شده است. سپس با استفاده از کلمات اسم و صفت به ساخت مدل بایگرام^{۲۴} به کمک روش Word2vec (برای استخراج ویژگی) و برآورد ۳۰ بایگرام پرتکرار پرداخته شده است. بایگرام‌های مذکور در شکل ۷ مشاهده می‌شوند؛ که مطابق آن، عبارت‌هایی مانند تمجید فراوان نشان می‌دهد که در حالت کلی، مشتریان از خرید خود رضایت دارند.

از طرفی، نوع اندام، قیمت فروش، ارزش کالا با توجه به قیمت، مقایسه با فروشنده‌های محلی، وزن سبک، و نحوه‌ی شست‌وشو از عبارت‌های پرتکرار و در نتیجه از موضوعاتی هستند که باید به آن‌ها توجه کرد. دسته‌بندی لباس‌ها،

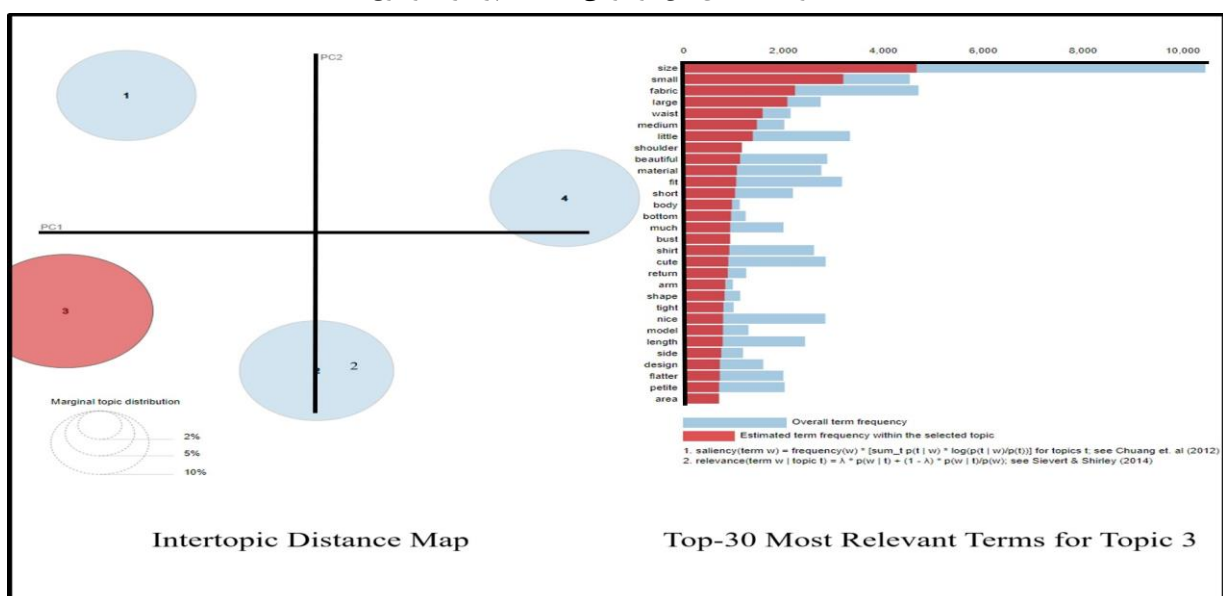
²⁴ Bigram

²² Part-of-speech

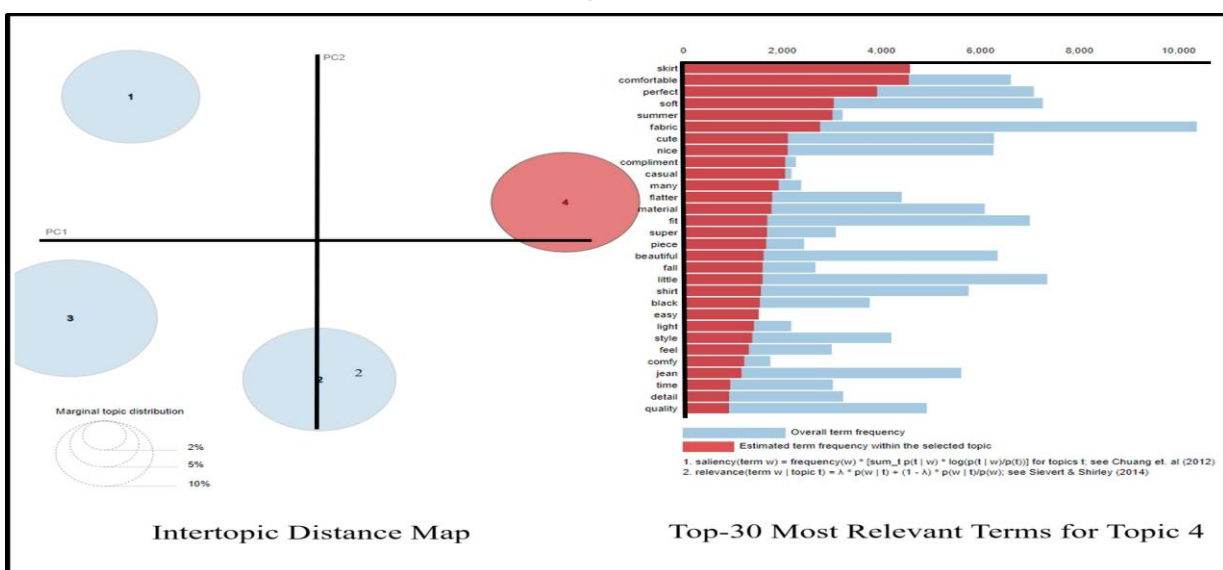
²³ Lemmatization



شکل ۹. نمودار فاصله‌ی بین موضوعی و کلمات پرتکرار موضوع دو.



شکل ۱۰. نمودار فاصله‌ی بین موضوعی و کلمات پرتکرار موضوع سه.



شکل ۱۱. نمودار فاصله‌ی بین موضوعی و کلمات پرتکرار موضوع چهار.

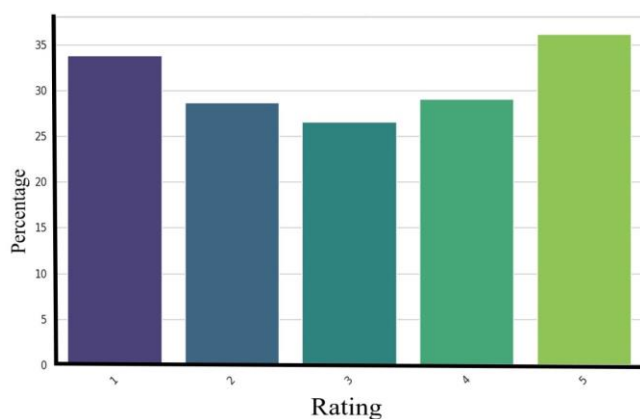
جدول ۲. تجزیه و تحلیل نتایج مدل سازی موضوعی.

	Items	Frequent words	Analysis
Topics	Topic one (focusing on pants and jeans)	Fit	"size," "petite," "regular," "small," "waist," "short," "usual"
		Comfort	"comfortable," "stretch," "soft," "skinny"
		Style	"pair," "perfect," "nice," "cute," "style," "quality"
		Shopping experience	"store," "sale," and "retailer"
		Occasions and applications	"time," "buy," "right," "glad"
	Topic two (focusing on jacket and shirt)	Fabric and comfort	"soft," "fabric," "material," "warm," "nice"
		Style	"beautiful," "cute," "black," "blue," "button," and "detail"
		Price and value	"price," "sale," and "retailer"
		Fit	"small" and "little"
		Shopping experience	"store," "online," "photo," and "picture"
	Topic three (focusing on shirts and tops)	Fit	"size," "small," "large," "medium," "shoulder," "waist," "body," "bust," and "arm"
		Fabric and material	"fabric," "material," and "soft"
		Style	"beautiful," "cute," "nice," "design," "flatter," "shape," and "model"
		Occasions and applications	"return," "design"
		Details	"short," "bottom," "tight," "length," "side," "area"
	Topic four (focusing on skirts and casual)	Fit and comfort	"comfortable," "soft," "fit," "comfy," "perfect"
		Style	"cute," "nice," "beautiful," "flatter," "style," "detail," and "quality"
		Occasions and applications	"casual," "summer," "fall," "piece," "easy,"
		Fabric and material	"fabric," "material," "light," and "feel"
		Shopping experience and satisfaction	"compliment," "time," and "glad"

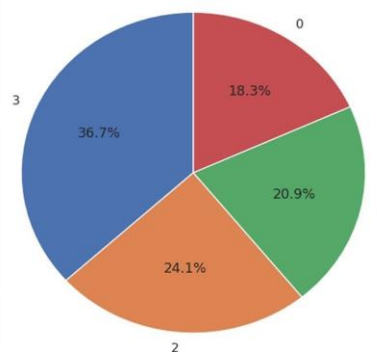
بر تناسب و راحتی، پارچه و جنس، ظاهر و استایل، و جزئیات پیراهن و تاپ است. خوشه‌ی دو، دربرگیرنده‌ی نظرهای زنان میانسال و سالمند با گروه سنی ۶۰-۳۵ سال است؛ که نظرهای آن‌ها به‌صورت توصیفی با نزدیک به ۴۰-۹۰ کلمه و تقریباً ۳۳٪ آن‌ها مربوط به موضوع ۲ است. همان‌طور که گفته شد، این موضوع بیشتر روی ژاکت و پیراهن، راحتی، پارچه و ارزش، و قیمت متمرکز شده است. خوشه‌ی سه، شامل نظرهای همه‌ی گروه‌های سنی از ۲۵ تا ۶۰ سال است، که نظرهای بسیار دقیق با تعداد ۸۰-۱۱۰ کلمه نوشته‌اند. نزدیک به ۴۲٪ از نظرهای اخیر مربوط به موضوع ۲ است.

درنهایت، تجزیه و تحلیل مربوط به هر خوشه با توجه به درصد امتیازها، مصورسازی ویژگی‌های عددی، و موضوع‌های اصلی در آن خوشه انجام شده است. در شکل‌های ۱۲ الی ۱۴، به ترتیب درصد موضوع‌ها و امتیازها در خوشه‌های یک، دو، و سه مشاهده می‌شوند.

در شکل ۱۵، نحوه‌ی توزیع برخی از ویژگی‌های عددی در هر سه خوشه مشاهده می‌شود. با توجه به نمودارهای گزارش‌شده، خوشه‌ی یک، شامل زنان جوان و میانسال با گروه سنی ۲۵-۴۰ است، که نظرهای نسبتاً دقیقی (۱۰-۵۰ کلمه) نوشته‌اند و حدود ۴۷٪ از نظرهای آن‌ها مربوط به موضوع ۳ و متمرکز

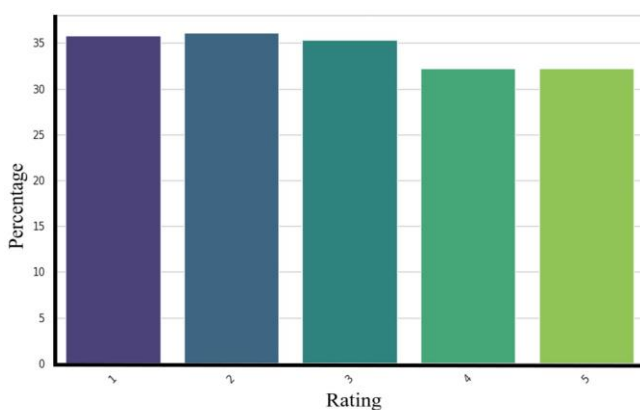


a

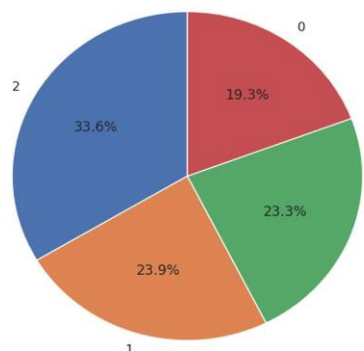


b

شکل ۱۲. (a) درصد امتیازها، (b) درصد موضوع‌ها (در خوشه‌ی یک).

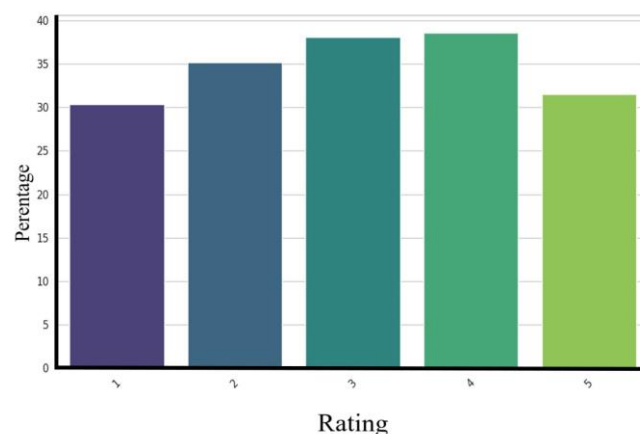


a

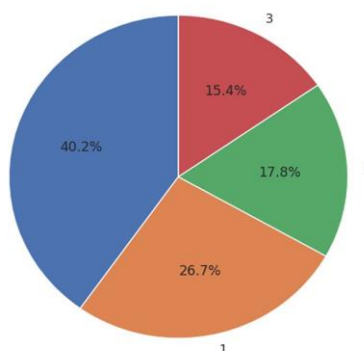


b

شکل ۱۳. (a) درصد امتیازها، (b) درصد موضوع‌ها (در خوشه‌ی دو).

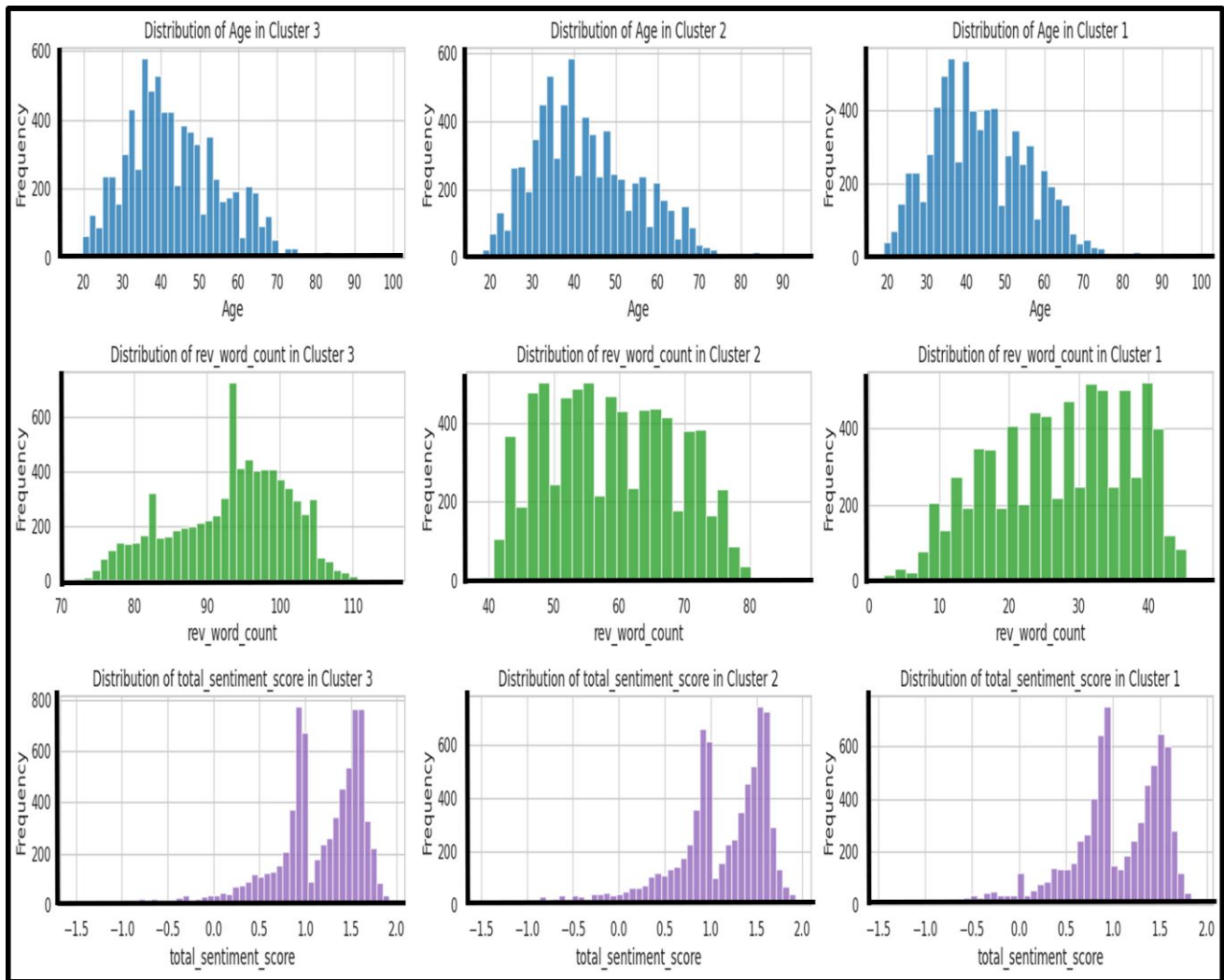


a



b

شکل ۱۴. (a) درصد موضوع‌ها، (b) درصد امتیازها (در خوشه‌ی سه).



شکل ۱۵. نحوه توزیع ویژگی‌های عددی در خوشه‌ها.

۴.۳. انتخاب ویژگی

در سند، ابزار TfidfVectorizer از کتابخانهی sklearn استفاده می‌شود، که برای انجام عملیات اخیر از الگوریتم TF-IDF استفاده می‌شود؛ و در نهایت، کلمات موجود در متن نظر را به یک ماتریس برداری تبدیل می‌کند.^{[۸] و [۱۱]}

۴.۳.۶. متعادل کردن مجموعه داده

همان‌طور که در قسمت تجزیه و تحلیل داده‌های اکتشافی در بخش بررسی ویژگی امتیاز مشاهده شد، مجموعه داده دارای توزیع نامتعادل کلاس امتیازهاست. بنابراین، به منظور اعمال مدل‌های یادگیری ماشین، ابتدا باید مجموعه داده متعادل شود. از آنجایی که در روش‌های معمول، افزایش داده‌های اقلیت، مانند OverSampling، ممکن است داده‌های تکراری زیادی به مجموعه اضافه و نتایج را غیرقابل اعتماد کنند، روش انتخابی برای متعادل کردن مجموعه داده در مطالعه حاضر، روش SMOTE^{۲۵} انتخاب شده است؛ که با استفاده از الگوریتم k-نزدیک‌ترین همسایه به منظور متعادل کردن مجموعه داده، نمونه‌های مصنوعی جدید تولید می‌کند.^[۱۵] از آنجایی که متغیر هدف امتیاز معمولاً پنج کلاس (امتیاز ۱ تا ۵) دارد، به منظور انجام مطالعه و اجرای مدل‌ها در حالت‌های دوکلاسه و سه‌کلاسه، امتیازها دسته‌بندی شده‌اند. در حالت دوکلاسه، امتیازهای قبل از ۳ به عنوان کلاس ۰ و امتیاز ۳ و بالاتر به عنوان کلاس ۱ تعریف شده‌اند. در حالت سه‌کلاسه، امتیازهای قبل از

پس از بررسی میزان همبستگی ویژگی‌ها با ویژگی امتیاز به عنوان متغیر هدف و با توجه به درستی به دست آمده از اعمال مدل‌های یادگیری ماشین بر مجموعه داده، تصمیم گرفته شد که به منظور انجام مطالعه حاضر از ویژگی متن نظرها و امتیاز استفاده شود.

۴.۳.۵. پیش‌پردازش داده‌ها و استخراج ویژگی

در مطالعه حاضر، به منظور انجام پیش‌پردازش و استخراج ویژگی از الگوریتم TFIDF به همراه حذف کلمات توقف استفاده شده است؛ که در ادامه، به بیان توضیح مختصری از آن‌ها پرداخته شده است.

• **کلمات توقف:** رایج‌ترین اصطلاح‌ها در یک زبان که ارزش معنایی ندارند و فقط به درک کلی متن کمک می‌کنند، را کلمات توقف می‌گویند؛ که با حروف تعریف، حروف اضافه، علائم نگارشی، حروف ربط، و ضمایر مشخص می‌شوند. حذف آن‌ها مقدار توکن‌ها را کاهش می‌دهد و تحلیل را بهبود می‌بخشد.^[۶]

• **الگوریتم TFIDF برای استخراج ویژگی:** به منظور تبدیل متن به بردار و همچنین امکان محاسبه وزن نسبت داده شده به هر عبارت موجود

²⁵ Synthetic Minority Oversampling Technique

۳ به عنوان کلاس ۰، امتیاز ۳ به عنوان کلاس ۱، و امتیازهای بالاتر از آن به عنوان کلاس ۲ در نظر گرفته شده اند. در مرحله ی بعد، به اجرای مدل های یادگیری ماشین پرداخته شده است. پس از بررسی کارهای مرتبط و مطالعات پیشین و همچنین اعمال برخی از مدل های یادگیری ماشین و با توجه به درستی حاصل، مدل های یادگیری ماشین در نظر گرفته شده برای مطالعه ی حاضر عبارت از: رگرسیون لجستیک، ماشین بردار پشتیبان، جنگل تصادفی، LightGBM، درخت تصمیم، XGBoost، بیز ساده ی چندجمله ای، و بیز ساده ی مکمل (CNB)^{۲۶} بوده اند، که در ادامه، به بیان توضیح مختصری از آن ها پرداخته شده است.

۳.۷. مدل های یادگیری ماشین

- **رگرسیون لجستیک**، از مدل های یادگیری ماشین تحت نظارت است، که در کلاس بندی، تحلیل، و پیش بینی با استفاده از مشاهده های قبلی یک مجموعه ی داده استفاده می شود. رگرسیون لجستیک، از یک تابع لجستیک به نام تابع سیگموئید که احتمال رخ دادن متغیر وابسته را بین ۰ و ۱ محدود می کند، برای تخمین احتمال ها استفاده می کند. [۱۲ و ۲۶]
- **ماشین بردار پشتیبان**، از الگوریتم های یادگیری ماشین تحت نظارت به منظور انجام وظایف کلاس بندی است و هدف آن یافتن ابرصفحه ای است که بیشترین حاشیه بین کلاس های داده های آموزشی را داشته باشد. [۵ و ۲۶]
- **درخت تصمیم**، روش یادگیری ماشین تحت نظارت است، که در مسائل رگرسیون و همچنین کلاس بندی استفاده می شود و با خلاصه کردن ویژگی های داده از داده های آموزشی به پیش بینی مقدار متغیر هدف کمک می کند. [۲۶ و ۲۷]
- **جنگل تصادفی**، از روش های کلاس بندی مبتنی بر یادگیری گروهی^{۲۷} است، که با استفاده از روش گروه بندی موازی^{۲۸} و ترکیب چندین درخت تصمیم و انتخاب تصادفی داده ها و ویژگی ها، عملکرد مدل را بهبود می بخشد. بنابراین، مشکل بیش برآزش را به میزان کمینه می رساند و درستی پیش بینی را افزایش می دهد. [۱۲ و ۲۶]

- **LightGBM**^{۲۹}، الگوریتمی مبتنی بر یادگیری گروهی و درخت های تصمیم است، که با استفاده از چارچوب GBDT^{۳۰} طراحی شده است. این مدل با استفاده از روش های نمونه برداری یک طرفه ی مبتنی بر گرادین (GOSS)^{۳۱} و دسته بندی ویژگی های انحصاری (EFB)^{۳۲}، در حالی که زمان آموزش و استفاده از حافظه را کاهش می دهد، عملکرد مدل را بهبود می بخشد. [۲۶ و ۵]

- **بیز ساده ی چندجمله ای و مکمل**: بیز ساده از الگوریتم های یادگیری ماشین تحت نظارت مبتنی بر احتمال و قضیه ی بیز با فرض استقلال بین هر جفت ویژگی است و هدف آن بهبود احتمال پسین با استفاده از داده های آموزشی است، که در نهایت منجر به یک قانون تصمیم گیری برای داده های جدید می شود. [۲۶ و ۲۸] بیز ساده ی چندجمله ای براساس تعداد دفعاتی که یک کلمه در یک سند ظاهر می شود، عمل می کند. الگوریتم بیز ساده ی مکمل، نیز اقتباسی از بیز ساده ی چندجمله ای است، که به جای محاسبه ی احتمال یک نمونه متعلق به یک کلاس خاص، احتمال تعلق نمونه به همه ی کلاس ها را محاسبه می کند. [۱۲ و ۲۹]
- **XGBoost**^{۳۳} (تقویت گرادین شدید، شکلی از الگوریتم های تقویت گرادین است که تقریب های دقیق تری را هنگام تعیین بهترین مدل در نظر می گیرد و مانند جنگل های تصادفی، یک الگوریتم یادگیری مجموعه ای است که یک مدل نهایی را براساس یک سری مدل های جداگانه، معمولاً درخت های تصمیم، تولید می کند. [۲۶ و ۲۷]

۳.۸. معیارهای ارزیابی

معیارهای ارزیابی استفاده شده در مطالعه ی حاضر عبارت اند از: درستی، دقت، پوشش، و امتیاز F1، که خلاصه ای از نحوه ی محاسبه و توضیحات آن ها در جدول ۳ ارائه شده است.

۳.۴. ارزیابی و نتایج

در بخش کنونی، به بیان معیارهای ارزیابی برآورد شده توسط مدل های یادگیری

جدول ۳. معیارهای ارزیابی و نحوه ی محاسبه. [۹، ۱۱]

Name	Formula	Definition
Recall	$TPR = \frac{TP + TN}{FN + TP}$	The proportion of true positive samples that are correctly classified as positive
Accuracy	$A = \frac{TP + TN}{TN + FP + FN + TP}$	The proportion of correct predictions among the total number of predictions
Precision	$P = \frac{TP}{FP + TP}$	The proportion of predicted positive instances that are actually positive
F1-Score	$F1 = 2 \frac{P \times TPR}{P + TPR}$	To measure the harmonic mean of precision and recall, providing a balanced evaluation by considering both false positives and false negatives

³⁰ Gradient Boosted Decision Trees

³¹ Gradient-based One-Side Sampling

³² Exclusive Feature Bundling

³³ Extreme Gradient Boosting (XGBoost)

²⁶ Complement Naive Bayes

²⁷ Ensemble learning

²⁸ Parallel ensembling

²⁹ Light Gradient Boosting Machine

جدول ۴. مقایسه‌ی نتایج حالت دوکلاسه با کارهای مشابه.

		shetty, et.al (2023)						This study	
		Evaluation Metrics							
		Accuracy	Precision	Recall	F1-score	Accuracy	Precision	Recall	F1-score
Machine Learning Models	Support Vector Machine	0.94	0.94	0.99	0.96	0.95	0.96	0.95	0.95
	Logistic regression	0.93	0.94	0.99	0.97	0.90	0.91	0.92	0.91
	Decision Tree	-	-	-	-	0.91	0.92	0.92	0.92
	Random forest	0.90	0.90	1.00	0.95	0.98	0.96	0.99	0.98
	AdaBoosting	0.91	0.93	0.97	0.95	-	-	-	-
	Bernoulli Naive Bayes	0.92	0.95	0.97	0.96	-	-	-	-
	Multinomial Naive Bayes	0.89	0.89	1.00	0.94	0.87	0.86	0.94	0.89
	Complement Naive Bayes	-	-	-	-	0.87	0.92	0.86	0.89
	XGBoost	-	-	-	-	0.91	0.91	0.94	0.93
	LightGBM	-	-	-	-	0.90	0.91	0.92	0.91

جدول ۵. نتایج حالت‌های سه کلاسه و پنج کلاسه.

		Three-class						Five-class	
		Evaluation Metrics							
		Accuracy	Precision	Recall	F1-score	Accuracy	Precision	Recall	F1-score
Machine Learning Models	Support Vector Machine	0.90	0.90	0.88	0.89	0.83	0.83	0.83	0.83
	Logistic regression	0.84	0.83	0.78	0.79	0.75	0.74	0.74	0.75
	Decision Tree	0.83	0.81	0.81	0.81	0.72	0.72	0.72	0.72
	Random forest	0.95	0.96	0.93	0.94	0.91	0.92	0.91	0.91
	Multinomial Naive Bayes	0.78	0.83	0.68	0.67	0.73	0.73	0.73	0.73
	Complement Naive Bayes	0.81	0.81	0.75	0.76	0.72	0.73	0.72	0.72
	XGBoost	0.86	0.87	0.82	0.83	0.79	0.80	0.79	0.79
	LightGBM	0.85	0.85	0.81	0.82	0.81	0.81	0.81	0.80

درستی ۰/۹۸ و پس از آن درستی مدل ماشین بردار پشتیبان برابر با ۰/۹۵ در نظر گرفته شده است. موضوع قابل توجه این است که در مطالعه‌ی شتی و همکاران (۲۰۲۳)^[۱۲]، به منظور رفع نامتعادل بودن مجموعه‌ی داده، راهکاری ارائه نشده است.

مقادیر درستی، دقت، پوشش، و امتیاز F1 برای مدل‌های یادگیری ماشین اجرا شده در حالت‌های سه کلاسه و پنج کلاسه در جدول ۵ مشاهده می‌شود. در حالت سه کلاسه، بهترین عملکرد مربوط به مدل جنگل تصادفی با درستی ۰/۹۵، دقت ۰/۹۶، پوشش ۰/۹۳، و امتیاز F1 ۰/۹۴ بوده است. پس از آن مدل ماشین بردار پشتیبان با داشتن مقادیر درستی ۰/۹۰، دقت ۰/۹۰، پوشش ۰/۸۸، و امتیاز F1 ۰/۸۹ عملکرد خوبی داشته است. در حالت پنج کلاسه نیز بهترین عملکرد مربوط به مدل جنگل تصادفی با درستی ۰/۹۱، دقت ۰/۹۲، پوشش

ماشین اجرا شده در حالت‌های دوکلاسه، سه کلاسه، و پنج کلاسه پرداخته شده است. در بین کارهای مرتبط بیان شده در بخش‌های قبل، مطالعه‌ی شتی و همکاران (۲۰۲۳)^[۱۲] با عنوان کاوش احساس‌های بازخورد داده‌های تجارت الکترونیک با استفاده از الگوریتم‌های یادگیری ماشین، دارای متغیر هدف یکسان با مطالعه‌ی حاضر بوده است، اما فقط در حالت دوکلاسه صورت گرفته است. در جدول ۴، به منظور مقایسه‌ی معیارهای ارزیابی، دو مطالعه‌ی صورت گرفته طراحی شده‌اند. معیارهای ارزیابی بیان شده در جدول اخیر، هر دو بر روی مجموعه‌ی داده‌ی یکسان و با استفاده از الگوریتم استخراج ویژگی TF-IDF در نظر گرفته شده‌اند. همان‌طور که مشاهده می‌شود، بالاترین درستی در مطالعه‌ی شتی و همکاران، مربوط به مدل ماشین بردار پشتیبان با درستی ۰/۹۴ و بهترین عملکرد در مطالعه‌ی حاضر مربوط به مدل جنگل تصادفی با

درستی ۰/۹۸ در حالت دوکلاسه، ۰/۹۵ در حالت سه کلاسه، و ۰/۹۱ در حالت پنج کلاسه برآورد شده اند.

مطالعه ای حاضر، ترکیبی از متن کاوی و پردازش زبان طبیعی است، اما براساس اهداف و روش شناسی، بیشتر بر روی حوزه ای متن کاوی متمرکز بوده است. هدف اصلی پژوهش، تحلیل و طبقه بندی نظرهای مشتریان با استفاده از الگوریتم های یادگیری ماشین بوده است. روش های پردازش زبان طبیعی، مانند:

TF-IDF و Word2Vec برای برداری سازی و استخراج ویژگی های متنی استفاده شده و نقش پشتیبانی کننده ای در پیش پردازش داده ها و آماده سازی آن ها برای مدل سازی ایفا کرده اند. علاوه بر این، فرآیندهایی مانند دسته بندی مجدد متغیر هدف، متعادل سازی داده ها با SMOTE و استفاده از الگوریتم های مختلف یادگیری ماشین برای پیش بینی و تحلیل احساس ها، نشان دهنده تمرکز پژوهش بر کشف الگوها و پیش بینی ها از داده های متنی هستند. در حالی که روش های پردازش زبان طبیعی در بخش هایی، مانند مدل سازی موضوعی و استخراج ویژگی به کار رفته اند، نقش اصلی آن ها پشتیبانی از تحلیل داده های متنی و مدل سازی بوده است. بنابراین، پژوهش حاضر، بیشتر در حوزه ای متن کاوی قرار می گیرد و روش های پردازش زبان طبیعی به عنوان ابزارهایی برای تسهیل فرآیند متن کاوی استفاده می شوند. به منظور انجام پژوهش در کارهای آینده می توان علاوه بر مدل های یادگیری ماشین، از مدل های یادگیری عمیق نیز به منظور تجزیه و تحلیل احساس ها استفاده کرد. روش های استخراج ویژگی در پردازش زبان طبیعی نیز الگوریتم های متفاوتی دارند، که می توان به منظور بررسی نظرها و رفتار مشتریان آن ها را به کار گرفت.

۰/۹۱، و امتیاز F1 ۰/۹۱ بوده است. پس از آن مدل ماشین بردار پشتیبان با داشتن مقادیر درستی ۰/۸۳، دقت ۰/۸۳، پوشش ۰/۸۳، و امتیاز F1 ۰/۸۳ عملکرد خوبی داشته است.

۵. نتیجه گیری و پیشنهاد های آینده

با توجه به اهمیت شناخت رفتار مشتریان به منظور ادامه ای فعالیت در دنیای امروز، در مطالعه ای حاضر به بررسی احساس های نظرهای تجارت الکترونیک پوشاک زنان با استفاده از روش های پردازش زبان طبیعی و یادگیری ماشین پرداخته شده است. مدل های یادگیری ماشین استفاده شده، شامل ماشین بردار پشتیبان، رگرسیون لجستیک، درخت تصمیم، جنگل تصادفی، بیز ساده ای چند جمله ای، بیز ساده ای مکمل، XGBoost، و LightGBM هستند. به منظور برداری کردن و استخراج ویژگی متن نظرها از الگوریتم های TF-IDF و Word2vec استفاده شده است. مدل سازی موضوعی با استفاده از روش تخصیص دیریکله ی پنهان و خوشه بندی k-means انجام شده است. معیارهای ارزیابی استفاده شده در مطالعه ای حاضر، شامل درستی، دقت، پوشش، و امتیاز F1 بوده اند. از آنجایی که متغیر هدف امتیاز در حالت معمول پنج کلاس (امتیاز ۱ تا ۵) دارد، به منظور انجام مطالعه و اجرای مدل ها در حالت های دوکلاسه و سه کلاسه، امتیازها دسته بندی شده اند. در حالت دوکلاسه، امتیازهای قبل از ۳ به عنوان کلاس ۰ و امتیاز ۳ و بالاتر به عنوان کلاس ۱ تعریف شده اند. در حالت سه کلاسه، امتیازهای قبل از ۳ به عنوان کلاس ۰، امتیاز ۳ به عنوان کلاس ۱، و امتیازهای بالاتر از آن به عنوان کلاس ۲ در نظر گرفته شده اند. پس از اعمال الگوریتم های یادگیری ماشین و پردازش زبان طبیعی در هر سه حالت، بهترین عملکرد مربوط به مدل جنگل تصادفی با

References – منابع

- Mabrouk, A., Redondo, R.P.D. and Kayed, M., 2021. Seopinion: summarization and exploration of opinion from e-commerce websites. *Sensors*, 21(2), pp.636. <https://doi.org/10.3390/s21020636>.
- Huang, H., Asemi, A. and Mustafa, M.B., 2023, July. Sentiment analysis application in e-Commerce: current models and future directions. In *International Conference on Electronic Government and the Information Systems Perspective*, pp.67-72. https://doi.org/10.1007/978-3-031-39841-4_5.
- Kasimu, M., Hellen, N. and Marvin, G., 2023. Explainable sentiment analysis for textile personalized marketing. In *The Fourth Industrial Revolution and Beyond: Select Proceedings of IC4IR+*, pp.473-488. https://doi.org/10.1007/978-981-19-8032-9_33.
- Zainal, S.R.M. and Xiang, C., 2023. Examining women's e-commerce clothing reviews. *International Journal of Business and Technology Management*, 5(1), pp.22-38. Available at: <https://myjms.mohe.gov.my/index.php/ijbtm/article/view/21637>. Date accessed: 04 aug. 2024.
- Lin, X., 2020, April. Sentiment analysis of e-commerce customer reviews based on natural language processing. In *Proceedings of the 2020 2nd international conference on big data and artificial intelligence*, pp.32-36. <https://doi.org/10.1145/3436286.3436293>.
- Kubrusly, J., Neves, A.L. and Marques, T.L., 2022. A statistical analysis of textual e-commerce reviews using tree-based methods. *Open Journal of Statistics*, 12(3), pp.357-372. <https://doi.org/10.4236/ojs.2022.123023>.
- Muniasamy, A. and Bhatnagar, R., 2022. Analyzing online reviews of customers using machine learning techniques. In *Rising Threats in Expert Applications and Solutions: Proceedings of FICR-TEAS 2022*, pp.485-493. https://doi.org/10.1007/978-981-19-1122-4_51.

8. Sunil, N. and Shirazi, F., 2023, July. Customer review classification using machine learning and deep learning techniques. In *International Conference on Human-Computer Interaction*, pp. 581-597. https://doi.org/10.1007/978-3-031-35915-6_42.
9. Loukili, M., Messaoudi, F. and El Ghazi, M., 2023. Sentiment analysis of product reviews for e-commerce recommendation based on machine learning. *International Journal of Advances in Soft Computing & Its Applications*, 15(1). <https://doi.org/10.15849/IJASCA.230320.01>.
10. Sharma, H., 2023. Using Topic Modeling for Extracting Customers' Expectations: a case of women apparel. *Business Perspectives and Research*, p.22785337221150831. <https://doi.org/10.1177/22785337221150831>.
11. Shetty, A.M., Aljunid, M.F., Manjaiah, D.H. and Shaik Afzal, A.M., 2023, July. Hyperparameter optimization of machine learning models using grid search for amazon review sentiment analysis. In *International Conference on Data Science and Applications*, pp. 451-474. https://doi.org/10.1007/978-981-99-7814-4_36.
12. Shetty, A.M., Aljunid, M.F. and Manjaiah, D.H., 2023, February. Sentiment exploring on feedback of e-commerce data using machine learning algorithms. In *International Conference on Emerging Research in Computing, Information, Communication and Applications*, pp.107-129. https://doi.org/10.1007/978-981-99-7622-5_8.
13. Yabes, I., Aryanti, N., Pradana, R.J., Setiawan, K.E. and Hasani, M.F., 2023, October. Classifying and predicting the rating sentiment of women's e-commerce clothing reviews: a comparative study using SVM, ANN, and BERT Models. In *2023 5th International Conference on Cybernetics and Intelligent System (ICORIS)*, pp.1-6. https://doi.org/10.1109/ICORIS60118.2023.1035218_9.
14. Shashank, S. and Behera, R.K., 2024. Factors influencing recommendations for women's clothing satisfaction: a latent dirichlet allocation approach using online reviews. *Journal of Retailing and Consumer Services*, 81, pp.104011. <https://doi.org/10.1016/j.jretconser.2024.104011>.
15. Padurariu, C. and Breaban, M.E., 2019. Dealing with data imbalance in text classification. *Procedia Computer Science*, 159, pp.736-745. <https://doi.org/10.1016/j.procs.2019.09.229>.
16. Mahdikhani, M., 2023. Exploring commonly used terms from online reviews in the fashion field to predict review helpfulness. *International Journal of Information Management Data Insights*, 3(1), pp.100172. <https://doi.org/10.1016/j.jjime.2023.100172>.
17. Li, M., Zhao, L. and Srinivas, S., 2023. It is about inclusion! Mining online reviews to understand the needs of adaptive clothing customers. *International Journal of Consumer Studies*, 47(3), pp.1157-1172. <https://doi.org/10.1111/ijcs.12895>.
18. Shobha Rani, V., Deepthi, K., Ramu, V. and Shirisha, G., 2024, January. Sentiment forecasting in women's fashion e-Commerce: a machine learning perspective. In *International Conference on Multi-Strategy Learning Environment*, pp.597-610. https://doi.org/10.1007/978-981-97-1488-9_44.
19. NICAPOTATO; Women's e-commerce clothing reviews, 2018. <https://www.kaggle.com/datasets/nicapotato/womens-e-commerce-clothing-reviews>.
20. Heng, Y., Gao, Z., Jiang, Y. and Chen, X., 2018. Exploring hidden factors behind online food shopping from amazon reviews: a topic mining approach. *Journal of Retailing and Consumer Services*, 42, pp.161-168. <https://doi.org/10.1016/j.jretconser.2018.02.006>.
21. Kim, S.W. and Gil, J.M., 2019. Research paper classification systems based on TF-IDF and LDA schemes. *Human-centric Computing and Information Sciences*, 9, pp.1-21. <https://doi.org/10.1186/s13673-019-0192-7>.
22. Zhong, K., Jackson, T., West, A. and Cosma, G., 2024. Natural language processing approaches in industrial maintenance: a systematic literature review. *Procedia Computer Science*, 232, pp.2082-2097. <https://doi.org/10.1016/j.procs.2024.02.029>.
23. Vadivel, A., Meena, K., Sumathy, P., Selvaraj, H., Shanmugavadivu, P., & S. G., S., Eds., 2024. *Interactive and Dynamic Dashboard: Design Principles*, 1st Edn., CRC Press. <https://doi.org/10.1201/9781003542735>.
24. Cui, M., 2020. Introduction to the k-means clustering algorithm based on the elbow method. *Accounting, Auditing and Finance*, 1(1), pp.5-8. <https://dx.doi.org/10.23977/accf.2020.010102>.
25. Januzaj, Y., Beqiri, E. and Luma, A., 2023. Determining the optimal number of clusters using silhouette score as a data mining technique. *International Journal of Online & Biomedical Engineering*, 19(4). <https://doi.org/10.3991/ijoe.v19i04.37059>.
26. Sarker, I.H., 2021. Machine learning: Algorithms, real-world applications and research directions. *SN*

- computer science*, 2(3), pp.160.
<https://doi.org/10.1007/s42979-021-00592-x>.
27. Pedregosa, F., 2011. Scikit-learn: Machine learning in python Fabian. *Journal of machine learning research*, 12, pp.2825.
<https://doi.org/10.1002/hbm.25822>.
28. Noor, A. and Islam, M., 2019, July. Sentiment analysis for women's e-commerce reviews using machine learning algorithms. In *2019 10th International conference on computing, communication and networking technologies (ICCCNT)*, pp.1-6.
<https://doi.org/10.1109/ICCCNT45670.2019.89444.36>.
29. Chrismanto, A.R., Sari, A.K. and Suyanto, Y., 2023. Enhancing spam comment detection on social media with emoji feature and post-comment pairs approach using ensemble methods of machine learning.
<https://doi.org/10.1109/ACCESS.2023.3299853>.