



ارائه یک چارچوب توزیع شده برای انتخاب ویژگی چندمتغیره

منا شریف‌نژاد^{*}، محسن رحمانی و حسین غفاریان
گروه کامپیوتر، دانشکده فنی و مهندسی، دانشگاه اراک

چکیده

در بسیاری از مسائل یادگیری ماشین، انتخاب ویژگی‌های مرتبط و اجتناب از ویژگی‌های افزونه، برای بهبود کارایی انتخاب ویژگی ضروری است. در اکثر رویکردهای موجود، از الگوریتم‌های فیلتر چندمتغیره برای این منظور استفاده می‌شود که در آن‌ها تعامل با طبقه‌بند نادیده گرفته می‌شود. این مقاله با ارائه یک چارچوب، ترکیب روش‌های نهفته با روش‌های فیلتر چندمتغیره را پیشنهاد می‌دهد تا با در نظر گرفتن تعامل با طبقه‌بند در انتخاب ویژگی‌ها، این مشکل را برطرف کند. در چارچوب پیشنهاد شده، ارتباط بین هر ویژگی و برچسب‌های رده به وسیله الگوریتم‌های نهفته محاسبه و افزونگی بین ویژگی‌ها از طریق الگوریتم‌های فیلتر چندمتغیره بررسی می‌شود. در فرایند انتخاب ویژگی پیشنهاد شده، به جای استفاده از همه مجموعه داده‌ها، از توزیع افقی آن‌ها استفاده شده است. این خصوصیت برای مجموعه داده‌هایی که دارای نمونه‌های زیادی هستند و نیز در محیط‌هایی که داده‌ها متمرکز نیستند، باعث کاهش زمان اجرای فرایند انتخاب ویژگی شده است. کیفیت روش ما با استفاده از شش مجموعه داده ارزیابی شده است. نتایج ثابت می‌کنند که چارچوب پیشنهاد شده، می‌تواند دقت طبقه‌بندی را در مقایسه با روش‌های صرفاً مبتنی بر فیلتر چندمتغیره بهبود دهد. همچنین سرعت اجرا می‌تواند در مقایسه با روش‌های متمرکز بهبود یابد.

کلمات کلیدی: انتخاب ویژگی، فیلتر چندمتغیره، انتخاب ویژگی نهفته، طبقه‌بندی، توزیع شدگی

New Framework for Distributed Multivariate Feature Selection

Mona Sharifnezhad, Mohsen Rahmani, Hosein Ghaffarian
Department of Computer Engineering, Faculty of Engineering,
Arak University, Arak, Iran

Abstract

Feature selection is a process for discovering relevant features and putting the irrelevant and redundant features aside. In recent years, feature domain studies have increased in machine learning and pattern recognition significantly, in a way that several techniques have been developed to solve the problem of feature reduction. Feature selection is an essential activity in data preprocessing. It describes features that significantly reduce the effects of noise and irrelevant and redundant data and provides desirable prediction results. Feature selection algorithms are divided into three groups: Filter methods, are typically divided into two subsets: univariate evaluation and multivariate evaluation. While in univariate evaluation, features are evaluated based on attributing matching weights with their relevance degree to the classes, multivariate evaluation approaches can manage the feature relevance to the classes together with their redundancy. Wrapper methods, in these methods, the feature set is selected based on the classifier in a way that all possible subsets of features are considered and evaluating all states, the best subset that has more accurate classification is selected. Embedded methods, in these methods, the search for an optimal feature set is located inside the classifier structure. In many current feature selection approaches, redundancy and relevance between features and class vector are considered

^{*} Corresponding author

^{*} نویسنده عهده‌دار مکاتبات

simultaneously. However, these multivariate filter algorithms calculate the suitability of features by only the intrinsic characteristics of the data. In this paper, we propose a new framework to offset the multivariate feature selection problem regarding the interaction with classifiers in feature selection. Our proposed framework calculates the relevance of each feature to class labels by embedded algorithms. Then, redundancy among features is examined through multivariate filter algorithms. In addition, in the proposed framework, instead of using all records at once, we use their horizontal distribution. This approach reduces the runtime of the process in datasets with many instances and in environments without centralized data. The results of the evaluation show that the suggested framework can improve classification accuracy compared with the methods just based on multivariate filters. In addition, the runtime speed is improved compared with the centralized methods.

Key words: Multivariate filter feature selection, Embedded feature selection, Classification, Distribution

۱- مقدمه

انتخاب ویژگی، فرایند کشف ویژگی‌های مرتبط و کنار گذاشتن ویژگی‌های غیرمرتبط و افزونه است [1,2]. در سال‌های گذشته دامنه ویژگی‌ها، از دهه‌ها به صدها یا هزاران متغیر در زمینه‌هایی چون یادگیری ماشین یا شناسایی الگو با افزایش چشم‌گیری مواجه شده‌اند؛ به‌طوری‌که چندین روش برای حل مشکل کاهش ویژگی‌ها توسعه داده شده‌اند. انتخاب ویژگی به‌عنوان یک فعالیت مهم در پیش‌پردازش داده‌ها، بر روی ویژگی‌های ورودی متمرکز شده‌است که با کاهش تأثیرات نوفه، داده‌های غیرمرتبط و افزونه، ویژگی‌ها را به‌طور مؤثر توصیف می‌کند و نتایج پیش‌بینی خوبی را فراهم می‌کند [3,4,5].

به‌طور معمول روش‌های انتخاب ویژگی به‌صورت متمرکز به کار گرفته شده‌اند [6]. اما در سال‌های اخیر، به‌دلیل به کار گرفتن داده‌های بزرگ و توزیع شدن آن‌ها در مکان‌های مختلف، استفاده از انتخاب ویژگی توزیع‌شده ضروری است [4,6,7,8,9]. لذا تعداد زیادی از روش‌های توزیع‌شده به‌جای رویکردهای متمرکز توسعه داده شده‌اند. روش‌های اصلی برای بخش‌بندی و توزیع داده‌ها می‌توانند به‌صورت افقی یا عمودی باشند. در بخش‌بندی افقی، مجموعه داده‌ها به چندین بسته تقسیم می‌شوند که دارای ویژگی‌های مشابه با مجموعه داده‌های اصلی هستند، که هر کدام شامل یک زیرمجموعه از نمونه‌های اصلی هستند. در بخش‌بندی عمودی، مجموعه داده‌های اصلی به چندین بسته تقسیم می‌شود که دارای تعداد مشابهی از نمونه‌های اصلی هستند که هر کدام دارای زیرمجموعه‌ای از مجموعه اصلی ویژگی‌ها هستند [8]. روش‌های انتخاب ویژگی به سه دسته تقسیم می‌شوند:

۱. روش فیلتر: به‌طور معمول رویکردهای فیلتر به دو زیرمجموعه مختلف تقسیم می‌شوند: ارزیابی

تک‌متغیره^۱ و ارزیابی چندمتغیره^۲. در ارزیابی تک‌متغیره، ویژگی‌ها برپایه نسبت‌دادن وزن‌هایی مطابق با درجه ارتباط آن‌ها با رده‌ها، مورد ارزیابی قرار می‌گیرند. درحالی‌که رویکردهای مبتنی بر ارزیابی چندمتغیره، می‌توانند افزونگی ویژگی‌ها را همراه با ارتباط آن‌ها با رده‌ها، مدیریت کنند. ویژگی‌هایی که به‌احتمال دارای رتبه‌های مشابه هستند، بررسی شده و از این طریق موارد افزونه حذف می‌شوند. بررسی افزونگی در بین ویژگی‌ها به زمان محاسباتی بیشتری نیاز دارد و دقت طبقه‌بندی را بهبود می‌دهد [10,11].

۲. روش پوشاننده^۳: در این روش، مجموعه ویژگی‌ها برپایه طبقه‌بند، انتخاب می‌شوند. به‌طوری‌که تمامی زیرمجموعه‌های ممکن از ویژگی‌ها در نظر گرفته شده و با ارزیابی همه حالت‌ها بهترین زیرمجموعه که دارای دقت طبقه‌بندی بالاتری است انتخاب می‌شود [6,12,13].

۳. روش نهفته^۴: از مزایای هر دو روش قبلی با استفاده از معیارهای ارزیابی مختلف آن‌ها در انتخاب ویژگی سود می‌برد. در این روش، جستجو برای یک مجموعه ویژگی بهینه، درون ساختار طبقه‌بند قرار دارد [7,9]. به‌طور کلی تعدادی از زیرمجموعه‌های خوب از ویژگی‌ها با استفاده از روش فیلتر انتخاب می‌شوند. سپس روش پوشاننده، روی این زیرمجموعه‌های خوب به کار برده می‌شود تا یکی از بهترین‌ها به‌دست آید [13].

در نظر گرفتن تنها یک معیار در انتخاب ویژگی، یک استراتژی مناسب برای پیدا کردن زیرمجموعه‌های خوب از ویژگی‌ها نیست. در روش‌های فیلتر مبتنی بر ارزیابی تک‌متغیره، انتخاب ویژگی برپایه بیشینه ارتباط ویژگی با

¹ Univariate Evaluation

² Multivariate Evaluation

³ Wrapper

⁴ Embedded

رده است. به طوری که از سرعت بالایی برخوردارند و مستقل از طبقه‌بند هستند. اما در این روش‌ها وابستگی بین ویژگی‌ها در نظر گرفته نشده که این امر در کاهش کارایی طبقه‌بندی تأثیرگذار است. با استفاده از روش‌های فیلتر مبتنی بر ارزیابی چندمتغیره، می‌توان با حذف ویژگی‌های افزونه در جهت افزایش کارایی طبقه‌بندی گام موثرتری برداشت؛ ولی به دلیل نیاز به انجام پردازش بیشتر، سرعت کمتری دارد. در روش‌های پوشاننده، می‌توان از طریق طبقه‌بند نسبت به روش‌های فیلتر نتایج بهتری تولید کرد؛ ولی اجرای این روش‌ها پرهزینه است و در بعضی از مواقع که تعداد ویژگی‌ها زیاد هستند، منجر به شکست می‌شوند. همچنین وابستگی بین ویژگی‌ها در نظر گرفته نمی‌شود. روش‌های نهفته نیز با افزونگی ویژگی‌ها مواجه هستند، به طوری که ارتباط بین ویژگی‌ها نادیده گرفته می‌شود.

ما یک چارچوب با ترکیبی از انتخاب ویژگی فیلتر چندمتغیره همراه با انتخاب ویژگی نهفته پیشنهاد می‌دهیم. این ترکیب کمک می‌کند که انتخاب ویژگی مبتنی بر چند معیار در جهت افزایش بهبود کارایی طبقه‌بندی انجام شود: انتخاب ویژگی مبتنی بر بیشینه ارتباط ویژگی با رده، کمینه افزونگی و بیشینه دقت طبقه‌بند است. چارچوب پیشنهاد شده شامل ۵ مرحله است: در مرحله نخست، بعد از تعیین داده‌های آموزشی و تست، داده‌های آموزشی به صورت افقی توزیع می‌شوند. به طوری که تمام زیرمجموعه‌ها دارای تعداد ویژگی‌های یکسانی هستند. در مرحله دوم، روی هر زیرمجموعه از ویژگی‌ها، با استفاده از یک روش انتخاب ویژگی نهفته، به ویژگی‌ها امتیاز داده می‌شود، به طوری که از طریق دقت طبقه‌بندی ارتباط هر ویژگی با رده تعیین و ویژگی‌هایی که رتبه کمتر دارند، حذف می‌شوند. بدین ترتیب در این مرحله هر دو معیار بیشینه ارتباط با رده و بیشینه دقت طبقه‌بندی بررسی می‌شود. همچنین از طریق انتخاب ویژگی فیلتر چندمتغیره روی هر زیرمجموعه از ویژگی‌ها، ارتباط بین ویژگی‌ها ارزیابی شده، به طوری که به هر ارتباط امتیاز داده می‌شود و ویژگی‌هایی که امتیاز بیشتری کسب کرده‌اند، حذف می‌شوند؛ بنابراین در این مرحله معیار کمینه افزونگی نیز بررسی می‌شود؛ سپس ویژگی‌های انتخاب شده از طریق انتخاب ویژگی نهفته و فیلتر چندمتغیره با هم ترکیب و سایر ویژگی‌ها حذف می‌شوند. در مرحله سوم، به ویژگی‌هایی که حذف شدند، رأی داده می‌شود تا در پایان، آن ویژگی‌هایی که بیشینه

رأی را دریافت کرده‌اند، از مجموعه ویژگی‌های نهایی حذف شوند. در مرحله چهارم، سایر ویژگی‌های باقی‌مانده با هم ادغام شده تا مجموعه نهایی از ویژگی‌های انتخاب شده، مشخص شود. در مرحله آخر با استفاده از داده‌های آموزشی نهایی و داده‌های تست، دقت طبقه‌بندی ارزیابی می‌شود.

ارزیابی نتایج نشان می‌دهد که چارچوب پیشنهادی در بهبود دقت طبقه‌بندی مؤثر بوده است. همچنین زمان اجرا به دلیل توزیع‌شدگی داده‌ها نسبت به روش‌های متمرکز کاهش یافته است.

ادامه مقاله به صورت زیر تدوین شده است. روش‌های انتخاب ویژگی در بخش ۲، چارچوب پیشنهادی برای انتخاب ویژگی در بخش ۳، نتایج تجربی و تشریح مطالب در بخش ۴ آورده شده است. سرانجام بخش ۵ شامل نتیجه‌گیری مقاله و کارهای پیشنهادی آینده است.

۲. روش‌های انتخاب ویژگی

در این بخش، ابتدا چند روش انتخاب ویژگی به صورت متمرکز بررسی، سپس انواع انتخاب ویژگی توزیع شده مرور و همچنین بعضی از روش‌های ادغام که بعد از اعمال انتخاب ویژگی روی بخش‌های توزیع شده استفاده می‌شوند، بررسی شده است.

۲-۱- انتخاب ویژگی متمرکز

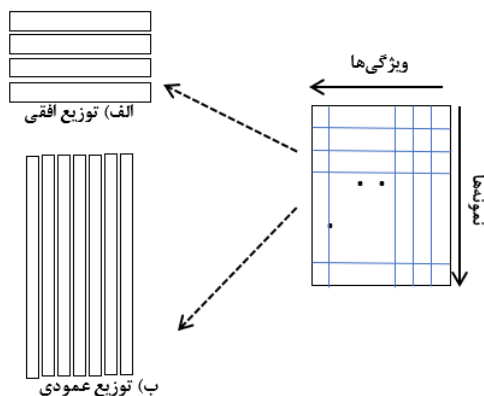
بسیاری از روش‌های انتخاب ویژگی بر روی مجموعه داده‌های متمرکز بررسی شده‌اند. در [15]، یک الگوریتم دو مرحله‌ای برای انتخاب ویژگی معرفی شده است، به طوری که ویژگی‌ها با استفاده از دو روش مختلف فیلتر به نام نسبت فیشر^۵ و روش ارتباط بیشینه و کمینه افزونگی^۶، به طور مجزا رتبه‌بندی و ارزیابی می‌شوند؛ سپس ویژگی‌های انتخاب شده مشترک بین هر دو روش بالا، به عنوان بهترین ویژگی‌ها در نظر گرفته می‌شود. در مرحله دوم، با استفاده از روش حذف بازگشتی ویژگی‌ها به وسیله بردار پشتیبان^۷ ویژگی‌هایی به غیر از آن‌هایی که در مرحله نخست به دست آمده، بررسی می‌شوند. در نهایت، از اجتماع ویژگی‌های به دست آمده در مرحله نخست و دوم، به عنوان ویژگی‌های

⁵ Fisher Ratio

⁶ maximum Relevance Minimum Redundancy (mRMR)

⁷ Support Vector Machine Recursive Feature Elimination (SVM-RFE)

الگوریتم انتخاب ویژگی مورد نظر روی هر زیرمجموعه به طور مجزا می تواند اعمال شود. در مرحله بعد، برای ترکیب نتایج، استفاده از یک روش ادغام ضروری است.



هاهای توزیع داده (شکل ۱). تکنیک
(Figure-1): Techniques for Data Distribution

به طوری که باید ویژگی های انتخاب شده در هر زیرمجموعه ادغام شوند و یک مجموعه ی واحد را تشکیل دهند.

روش های مختلفی برای ادغام نتایج از طریق تعیین یک حد آستانه، ارائه شده است. در [7,14]، برای محاسبه حد آستانه از خطای طبقه بندی و درصد ویژگی های باقیمانده استفاده شده است. به طوری که هر دو مقدار باید تا حد امکان به حداقل برسند. به ازای هر ویژگی که ممکن است در هر زیرمجموعه برچسب حذف را دریافت کرده باشد، مقدار آن ها تعیین می شود. کمینه مقدار به عنوان حد آستانه در نظر گرفته می شود. ویژگی هایی که تعداد برچسب های حذف شان بیشتر از مقدار حد آستانه باشند، حذف می شوند.

در [3,4]، دقت طبقه بندی نخستین زیرمجموعه ویژگی انتخاب شده را محاسبه کرده و آن را به عنوان مبنا^{۱۳} در نظر گرفته اند؛ سپس دقت طبقه بندی سایر زیرمجموعه ویژگی های باقیمانده به طور مجزا با نخستین زیرمجموعه محاسبه شده است. در صورتی که مقدار مبنا را بهبود دهند، آن ها نیز بخشی از انتخاب نهایی خواهند شد. در این روش، نخستین زیرمجموعه ویژگی انتخاب شده همیشه به عنوان بخشی از انتخاب نهایی خواهد بود.

در [4,8,9]، با استفاده از روش شاخص پیچیدگی^{۱۴}، حد آستانه تعیین شده است. ویژگی هایی که به عنوان نامزد خوب در نظر گرفته می شوند، باعث کاهش

انتخاب شده، استفاده، سپس از طریق ماشین بردار پشتیبان، دقت و صحت آن ها بررسی می شود. در [16]، روشی به نام $FWHM^8$ در دو مرحله ارائه شده است: در مرحله نخست، مجموعه داده ها به وسیله شش روش متفاوت از فیلترهای تک متغیره امتیازدهی می شوند. علت استفاده از این تعداد فیلتر این است که روش های فیلتر هر کدام به تنهایی منطبق بر یک معیار هستند، برای کاهش وابستگی مربوط به هر معیار از آن ها استفاده شده است. از طریق آن ها می توان نتایج مختلفی از رتبه بندی یا امتیازدهی به دست آورد؛ در نهایت از میانگین آن ها برای رتبه بندی هر ویژگی استفاده می شود. در مرحله دوم، از استراتژی های جستجوی تصادفی مثل الگوریتم ژنتیک^۹ به عنوان یک روش پوشاننده استفاده شده به طوری که میانگین رتبه بندی به دست آمده برای تولید جمعیت اولیه منظور شده است. ویژگی هایی که دارای رتبه ی بالاتری هستند، فرصت بیشتری برای انتخاب شدن دارند. در [17]، برای انتخاب ژن ها از روش SVM-RFE استفاده شده است. به طوری که دقت طبقه بندی در تشخیص ژن های سرطانی افزایش یافته است. در [18]، از سه روش جستجوی محلی LS^{10} ، SLS^{11} و VNS^{12} برای انتخاب ویژگی های مرتبط در زمینه Credit Scoring برای مسائل مالی و بانکداری استفاده شده، سپس نتایج هر کدام از روش ها به طور جداگانه با طبقه بند SVM ترکیب شده است. در [19]، با معرفی یک روش انتخاب ویژگی سلسله مراتبی در شبکه های عصبی تک لایه ویژگی های افزونه و نوفه دار هرس می شوند.

۲-۲- انتخاب ویژگی توزیع شده

ایده توزیع داده ها به دو صورت افقی و عمودی قابل انجام است. همان طور که در شکل (۱) نشان داده شده است، داده ها در توزیع افقی بر پایه نمونه ها و در توزیع عمودی، بر اساس ویژگی ها تقسیم می شوند. به طور معمول اگر تعداد نمونه ها بیشتر باشند، از توزیع افقی استفاده می شود و اگر تعداد ویژگی ها بیشتر باشند، توزیع عمودی به کار برده می شود [4,6,7,8,9].

با تعیین نوع توزیع، مجموعه داده ها به زیرمجموعه هایی کوچک تر تقسیم می شوند. سپس

⁸ Filter-Wrapper Hybrid Method

⁹ Genetic Algorithm

¹⁰ Local Search Method

¹¹ Stochastic Local Search Method

¹² Variable Neighborhood Search Method

¹³ baseline

¹⁴ Complexity Measure

طریق معیارهای مختلف، پیشنهاد داده‌ایم. این چارچوب پیشنهادی دارای ۵ مرحله به شرح زیر باطشد:

مرحله نخست شامل توزیع افقی مجموعه داده‌های آموزشی در چندین زیرمجموعه است. مرحله دوم، در دو بخش پیشنهاد شده‌است. ابتدا با استفاده از انتخاب ویژگی نهفته روی هر زیرمجموعه، ویژگی‌ها با بیشینه ارتباط بررسی، سپس از طریق انتخاب ویژگی فیلتر چندمتغیره^{۱۷} روی هر زیرمجموعه، ویژگی‌هایی با کمینه افزونگی محاسبه می‌شوند. در این مرحله ویژگی‌ها با بیشینه ارتباط و کمینه افزونگی ترکیب شده و سایر ویژگی‌ها حذف می‌شوند. در مرحله سوم به ویژگی‌های حذف شده در هر زیرمجموعه، رأی داده می‌شود، سپس زیرمجموعه ویژگی‌های انتخاب شده در یک زیرمجموعه نهایی ترکیب می‌شوند؛ در نهایت کارایی طبقه‌بند^{۱۸}، روی زیرمجموعه نهایی از ویژگی‌های انتخاب شده و مجموعه داده‌های تست بررسی می‌شود.

طرح کلی چارچوب پیشنهادی در (شکل ۲) و جزئیات آن در (شکل ۳)، نمایش داده شده‌است. پارامترهای استفاده شده در جدول ۱، شرح داده شده‌است.

(جدول ۱-): پارامترهای روش پیشنهادی
(Table-1): Parameters of Suggested Method

پارامتر	کاربرد
D	مجموعه داده‌های آموزشی
D_i	زیرمجموعه‌ای از D
F	مجموعه‌ی غیرتهی و محدود از ویژگی‌ها
C	مجموعه‌ی غیرتهی و محدود از رده‌ها
X_F	روش انتخاب ویژگی نهفته
W_{fi}	وزن محاسبه شده برای وابستگی ویژگی f_i به مجموعه رده C از طریق روش انتخاب ویژگی نهفته
W_{Fx}	بردار وزن ویژگی‌ها در روش انتخاب ویژگی نهفته
$W^{(f_i, f_j)}$	وزن محاسبه شده برای ارتباط بین ویژگی f_i با ویژگی f_j از طریق روش انتخاب ویژگی فیلتر چندمتغیره
W_{FY}	بردار وزن برای ارتباط بین ویژگی‌ها در روش انتخاب ویژگی فیلتر چندمتغیره
Δ	تابع ترکیب W_{FY} و W_{Fx}
Λ	ویژگی‌های حذف شده
FFS	ویژگی‌های انتخاب شده نهایی

برپایه شکل (۳)، در ابتدا، مجموعه داده‌های آموزشی (D) به صورت افقی به چند زیرمجموعه مجزا (D_i) و بدون جای‌گذاری تقسیم می‌شوند. D_i به شکل زیر تعریف می‌شوند:

¹⁷ Multivariate Filter

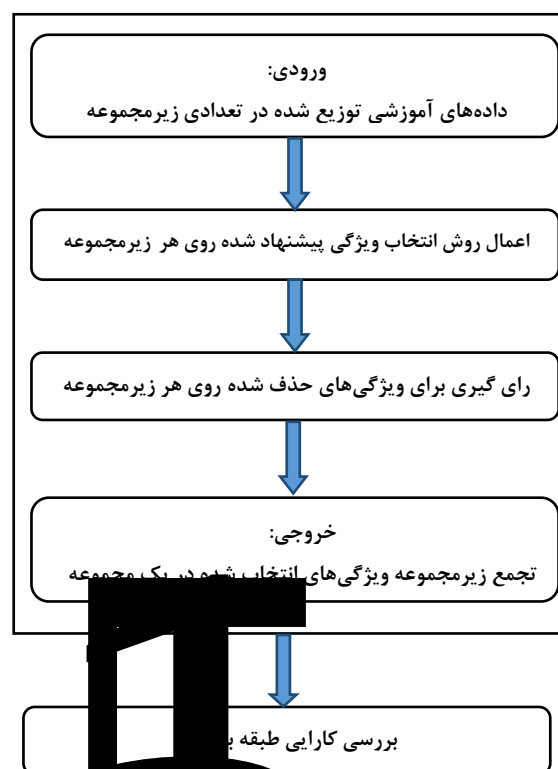
¹⁸ Performance of Classifier

پیچیدگی شده و باید حفظ شوند. درحالی که ویژگی‌هایی که نامزد بدی محسوب می‌شوند، سبب افزایش پیچیدگی می‌شوند و باید از بین بروند. این روش مستقل از طبقه‌بند^{۱۵} است.

در نهایت بعد از اتمام فرایند ادغام زیرمجموعه‌ها توسط یکی از روش‌های بالا، دقت مدل ساخته شده با استفاده از طبقه‌بند بر روی مجموعه داده‌ها محاسبه می‌شود.

۳- چارچوب پیشنهادی

در بسیاری از روش‌های انتخاب ویژگی ارائه شده، بررسی ویژگی‌های مرتبط و حذف ویژگی‌های افزونه مبتنی بر الگوریتم‌های فیلتر بوده است. به طوری که برای انتخاب ویژگی‌ها، مستقل از طبقه‌بند عمل می‌کنند که این امر می‌تواند باعث کاهش کارایی در انتخاب ویژگی‌ها شود. همچنین در بعضی از رویکردها، به دلیل استفاده از روش‌های متمرکز در انتخاب ویژگی‌ها، برای مجموعه داده‌هایی که دارای نمونه‌های زیادی هستند، زمان‌بر بوده و مناسب نیستند.



(شکل ۲-): طرح کلی چارچوب پیشنهادی
(Figure-2): Schematic of the Suggested Method

در این بخش ما یک چارچوب انتخاب ویژگی چندمتغیره به صورت توزیع شده مبتنی بر طبقه‌بند^{۱۶} و ویژگی‌ها از

¹⁶ Ranking

$$W_{F_X} = \{w_{fi}\} \text{ for } f_i \in F$$

از طریق بردار وزن (W_F) ، وزن مجموعه ویژگی‌ها برپایه دقت طبقه‌بند مطابق با رویکرد نهفته X_F محاسبه می‌شود. سپس ویژگی‌ها با استفاده از W_F ، رتبه‌بندی می‌شوند.

Y_F = Multivariate Filter method with tuning parameter $W_{(fi, fj)}$ for F

$$W_{F_Y} = \{w_{(fi, fj)}\} \text{ for } f_i, f_j \in F, i \neq j$$

$$\lambda = \arg \min \{\delta(W_{F_X}, W_{F_Y})\}$$

$$FFS(D_i) = F - \{\lambda\}$$

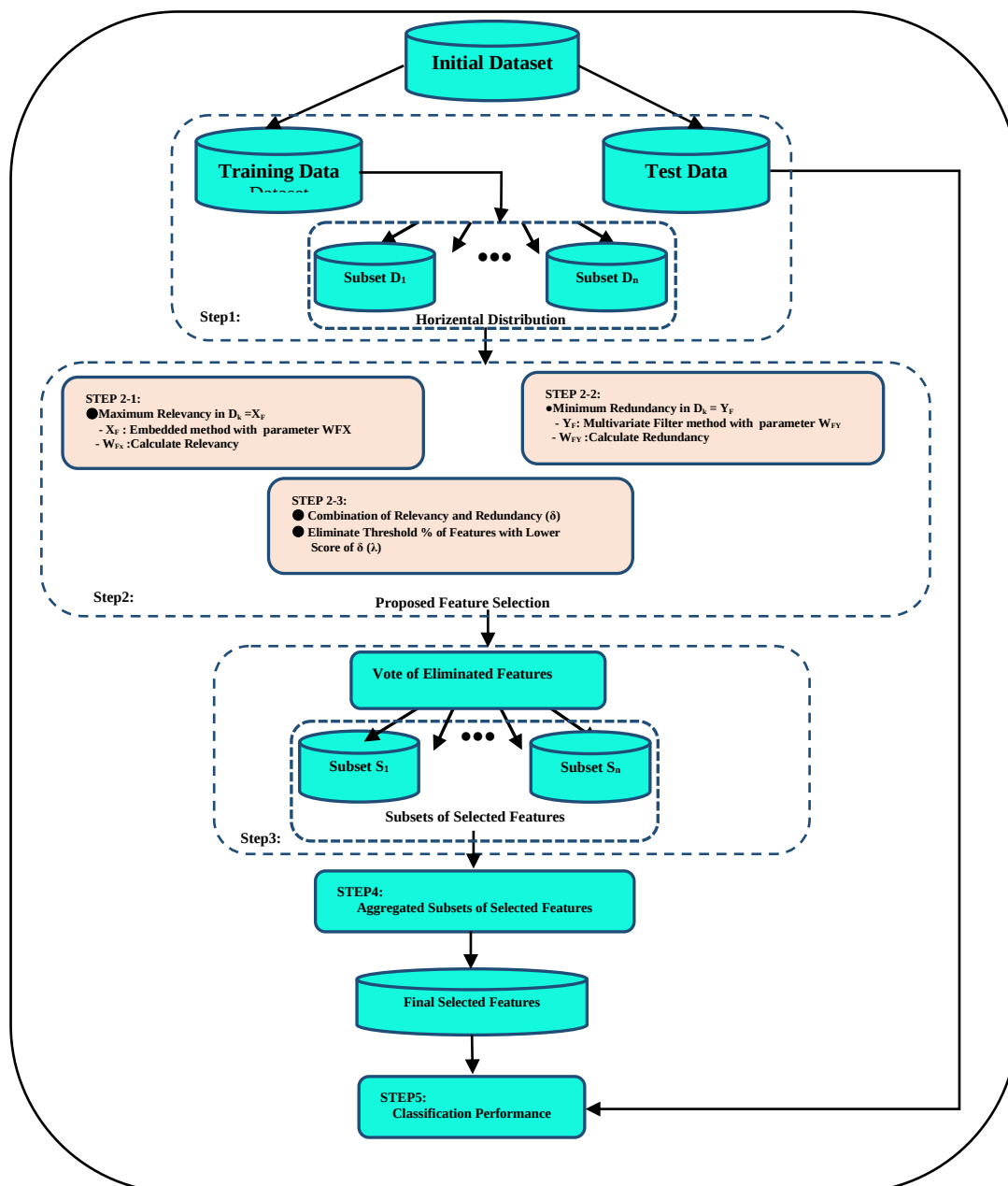
$$D = \{(S, F) \mid S \in \text{samples}, F \in \text{input features}\}$$

$$D_i = \{(s_j, F) \mid s_j \in S, j \in n\}$$

$$S = \{s_1, s_2, \dots, s_n\}, F = \{f_1, f_2, \dots, f_m\}$$

که در این تعاریف F و S به ترتیب مجموعه‌های غیرتهی و محدود از نمونه‌ها و ویژگی‌ها هستند. با محاسبه وابستگی ویژگی f_i به مجموعه رده C با استفاده از روش‌های انتخاب ویژگی نهفته، می‌توان از طریق فاکتور دقت طبقه‌بند آن را به صورت زیر در نظر گرفت:

X_F = Embedded method with tuning parameter W_{fi} for F



(شکل - ۳): جزئیات روش پیشنهادی
(Figure-3): Details of Suggested Me

Compute:

1. Horizontal Partition on Train Data
 D_1 to D_n : Subset of Train Datasets
2. Proposed Framework of Feature selection on each Subset of Train Data
 - For each D_i Subset of Train Data
 - a. Maximum Relevancy by Embedded Feature Selection (W_{Fx})
 - b. Minimum Redundancy by Multivariate Filter Feature Selection (W_{FY})
 - c. δ (W_{Fx} , W_{FY}): Combine of Relevancy and Redundancy
 - d. Remove Threshold Percent of Features with Lower Score (λ)
 - e. Increment one Vote in Vector V for each Remove Feature
3. Aggregated Subsets of Selected Features
 a. Th: A threshold of votes to remove a Feature
- b. FFS: Subset of Final Selected Features
4. Obtaining Accuracy by Classifier

۴- نتایج تجربی

در این بخش، ابتدا درباره تنظیمات اجرا توضیح می‌دهیم؛ سپس نتایج چارچوب پیشنهادی را از مشخصات این مجموعه داده‌ها در (جدول ۲) آورده شده است که شامل تعداد نمونه‌ها، تعداد ویژگی‌ها و تعداد رده‌ها است. نمونه‌ها در مجموعه داده‌ها به دو مجموعه آموزش و آزمون تقسیم می‌شوند. به‌طور معمول این تقسیم‌بندی به صورت $\frac{2}{3}$ مجموعه داده‌ها برای آموزش و $\frac{1}{3}$ برای آزمون در نظر گرفته می‌شود.

همه نتایج روی یک PC با مشخصات Intel (R) Core (TM) i7-2670QM CPU@ 2.2GHz 2.2.GHz و 8 GB RAM، با استفاده از نرم‌افزار متلب (نسخه ۲۰۱۳) در محیط ویندوز ۷ اجرا شده است.

۴-۱- تنظیمات اجرا

در این بخش برای ارزیابی چارچوب پیشنهادی از شش مجموعه داده‌ها [20] استفاده کردیم. چندمتغیره استفاده کرد. در این جا ما از یکی از روش‌های معروف در انتخاب ویژگی نهفته به نام SVM-RFE و نیز برای انتخاب ویژگی فیلتر چندمتغیره از روش mRMR استفاده کردیم. همچنین چارچوب پیشنهادی را با چند روش معروف از انتخاب ویژگی فیلتر مانند بهره اطلاعاتی^{۲۲} و ReliefF و نیز روش‌های غیرترکیبی انتخاب ویژگی SVM-RFE و mRMR مقایسه کردیم.

²² Information Gain (IG)

در مرحله بعد، برای بررسی ویژگی‌های افزونه، از یکی از روش‌های انتخاب ویژگی فیلتر چندمتغیره بر روی هر زیرمجموعه استفاده می‌شود. به‌طوری که به ارتباط هر ویژگی با سایر ویژگی‌ها امتیاز داده می‌شود (W_{FY}). امتیازدهی برای ارتباط بین ویژگی‌ها می‌تواند به‌وسیله هر نوع انتخاب ویژگی مبتنی بر فیلتر انجام شود. سپس دو رابطه ارتباط و افزونگی با هم ترکیب می‌شوند (δ). δ ، تابعی است که برپایه بیشینه کردن ارتباط بین ویژگی‌ها با رده و کمینه کردن افزونگی بین ویژگی‌ها عمل می‌کند. معیارهای ترکیب رابطه ارتباط و افزونگی می‌تواند به روش‌های مختلفی باشد. در [15]، از معیار تفاضل اطلاعات متقابل^{۱۹} بین بیشینه ارتباط و کمینه افزونگی استفاده شده است. در [11,16]، علاوه بر معیار MID، از نسبت اطلاعات متقابل^{۲۰} نیز استفاده شده است؛ سپس ویژگی‌هایی که کمترین رتبه را دارند، برپایه یک حدآستانه معین، از مجموعه ویژگی‌های نهایی حذف می‌شوند (λ).

در مرحله بعد به ویژگی‌هایی که در مرحله قبل برای حذف شدن انتخاب شده‌اند، رأی داده می‌شود. این رأی‌گیری کمک می‌کند تا ویژگی‌هایی که بیشتر از یک حدآستانه معین رأی آوردند، حذف شوند؛ سپس ویژگی‌های انتخاب شده در هر زیرمجموعه، با استفاده از یک تابع ادغام در یک زیرمجموعه نهایی جمع می‌شوند. راه‌کار ادغام به‌کاررفته در این مقاله، مشابه [4,8,9] مبتنی بر استفاده از روش شاخص پیچیدگی^{۲۱} است. در پایان با استفاده از مجموعه داده‌های نهایی آموزش و مجموعه داده‌های تست، کارایی و دقت طبقه‌بندی بررسی می‌شود. مشخصات جزئیات بیشتر در الگوریتم ۱ نشان داده شده است.

Algorithm 1 Pseudo-code for Proposed Approach

Input:

- Load Dataset D ($m \times s$) with m Samples and s Features
 - a. Train Dataset = $D(m_1 \times s)$
 - b. Test Dataset = $D(m_2 \times s)$
 - c. Vector of votes for Removing Features = V , length of Number of Features, Initialize the V to 0

Output:

- FS: Subset of Final Selected Features

¹⁹ Mutual Information Difference criterion (MID)

²⁰ Mutual Information Quotient criterion (MIQ)

²¹ Complexity Measure

• حذف بازگشتی ویژگی‌ها به وسیله بردار

پشتیبان

انتخاب ویژگی SVM-RFE توسط گایون²³ و همکاران به عنوان روشی برای انتخاب ویژگی معرفی شده است [15]. مطابق با این روش ویژگی‌ها به صورت بازگشتی حذف می‌شوند، به طوری که در هر مرحله به وسیله بردار پشتیبان وزن ویژگی‌ها محاسبه و به هر یک از آن‌ها یک امتیاز داده می‌شود و ویژگی که دارای کمترین امتیاز باشد، حذف خواهد شد [21]. این فرآیند تکرار خواهد شد تا زمانی که تمام ویژگی‌ها حذف شوند؛ در نهایت فهرستی از رتبه‌بندی ویژگی‌ها وجود خواهد داشت که انتخاب ویژگی می‌تواند با انتخاب یک گروه از ویژگی‌های برتر به دست آید.

• انتخاب ویژگی ارتباط بیشینه و کمینه افزونگی

انتخاب ویژگی mRMR یک الگوریتم انتخاب ویژگی چندمتغیره مبتنی بر فیلتر است که در میان ویژگی‌های انتخاب شده به طور هم‌زمان ارتباط را بیشینه و افزونگی را کمینه می‌کند. این روش از اطلاعات متقابل²⁴ برای تجزیه و تحلیل مرتبط بودن و افزونگی استفاده می‌کند؛ به طوری که بیشینه ارتباط از طریق محاسبه بیشینه مقدار میانگین از تمام اطلاعات متقابل بین تک‌تک ویژگی‌های موجود و بردار رده به دست می‌آید. همچنین کمینه افزونگی از طریق محاسبه کمینه مقدار میانگین از تمام اطلاعات متقابل بین بردار ویژگی‌ها بررسی می‌شود؛ سپس روابط بالا را می‌توان با یکی از دو روش تفاضل اطلاعات متقابل یا روش نسبت اطلاعات متقابل پیاده‌سازی کرد [10,11,15,22].

• انتخاب ویژگی بهره اطلاعاتی

IG یکی از رایج‌ترین روش‌های تک‌متغیره مبتنی بر فیلتر برای ارزیابی ویژگی‌ها است. این روش براساس بهره اطلاعاتی، به ارتباط هر ویژگی با رده رتبه و امتیاز می‌دهد؛ سپس با استفاده از یک حد‌آستانه، تعداد معینی از ویژگی‌ها که رتبه بالاتری را به دست آوردند را انتخاب می‌کند [23].

• انتخاب ویژگی ReliefF

این روش، توسعه یافته الگوریتم Relief است که

قابلیت کار کردن با مجموعه داده‌های چندرده، نوفه دار و ناقص (غیر کامل) را فراهم می‌کند [24]. به جای یافتن n نمونه نزدیک‌ترین برخورد و نزدیک‌ترین شکست، ReliefF از هر رده n نمونه را انتخاب می‌کند و سهم هر غیر هم‌رده‌ی در وزن‌دهی بر پایه احتمالات پیشین آن‌ها محاسبه می‌شود. نوفه در یک رده و یا مقدار یک ویژگی، تأثیر قابل توجه‌ای در انتخاب نزدیک‌ترین برخورد و نزدیک‌ترین شکست دارد. برای این که انتخاب نزدیک‌ترین برخورد و نزدیک‌ترین شکست با دقت و اطمینان بالاتری صورت گیرد الگوریتم ReliefF، از n نزدیک‌ترین برخورد و n نزدیک‌ترین شکست استفاده می‌کند و در تخمین کیفیت یک ویژگی میانگین سهم هر کدام را در نظر می‌گیرد. زمانی که نمونه‌ها در بعضی از ویژگی‌ها فاقد مقدار باشند، الگوریتم ReliefF تابع مربوطه را تغییر می‌دهد و آن را مدیریت می‌کند [25].

۴-۲- نتایج صحت طبقه‌بندی

در این بخش، دقت طبقه‌بندی به دست آمده به وسیله KNN²⁵، NB²⁶ و SVM²⁷ که برای ارزیابی بسیاری از الگوریتم‌های انتخاب ویژگی استفاده می‌شوند [3,4,6,7-15,17,9] را بر روی چارچوب پیشنهادی و چهار رویکرد مختلف انتخاب ویژگی IG، ReliefF، SVM-RFE و mRMR مقایسه کردیم. جدول‌های ۳، ۴ و ۵ نتایج تجربی حاصل از دقت طبقه‌بندی را به دو شکل: توزیع شده (مشخص شده با پیشوند D در ابتدای نام روش‌ها در جداول) و متمرکز (مشخص شده با پیشوند C در ابتدای نام روش‌ها در جداول) بر روی مجموعه داده‌های یاد شده، نشان می‌دهند. مواردی که به عنوان بهترین نتایج در هر مجموعه داده به دست آمده، به صورت پررنگ نشان داده شده‌اند. نتایج به دست آمده متغیر هستند و به مجموعه داده‌ها و طبقه‌بند وابسته هستند. در ستون آخر از هر جدول، نتایج حاصل از دقت طبقه‌بندی بدون استفاده از الگوریتم انتخاب ویژگی گزارش شده است. در سطر آخر، میانگین دقت طبقه‌بند مورد نظر محاسبه شده است. برای مجموعه داده‌های Ozon، Madelon، Isolet و Connect4 با طبقه‌بند KNN بهترین نتیجه برای چارچوب پیشنهادی به صورت توزیع گزارش شده است. مجموعه داده Ozon، بهترین مقدار صحت را

²⁵ k-Nearest Neighbor

²⁶ Naive Bayes

²⁷ Support Vector Machine

²³ Guyon et. al

²⁴ Mutual Information

آزمایش شده، مجموعه داده Spambase بیشترین تعداد بهبود را در هر سه طبقه بند بدون استفاده از الگوریتم انتخاب ویژگی، نسبت به سایرین داشته است. بیشترین تعداد بهبود دقت طبقه بندی با استفاده از الگوریتم انتخاب ویژگی در مقایسه با بدون استفاده از آن، به ترتیب مربوط به مجموعه داده های Madelon، Isolet و Connect4 است. برای رویکرد پیشنهادی نسبت به سایر روش ها در کمتر از $\frac{1}{20}$ موارد آزمایش شده، دقت رده بندی آن در مقایسه با حالت بدون استفاده از الگوریتم انتخاب ویژگی، کاهش یافته است. SVM در مقایسه با سایر رده بندی های مورد بررسی، دارای تعداد نتایج بهبود یافته بیشتری در استفاده از الگوریتم انتخاب ویژگی نسبت به بدون استفاده از آن، بوده است. نمودار میانگین دقت طبقه بندی برای هر سه طبقه بند KNN، SVM و NB، به ترتیب در شکل های ۴، ۵ و ۶ نشان داده شده است. میانگین دقت رده بندی برای رده بندی های KNN و NB در حالت توزیع شده نسبت به متمرکز در تمام روش های انتخاب ویژگی آزمایش شده، بهبود یافته است. میانگین دقت طبقه بندی برای چارچوب پیشنهادی نسبت به سایر روش های انتخاب ویژگی در حالت توزیع شده با KNN و SVM بیشتر گزارش شده است. اما میانگین دقت رده بندی NB برای چارچوب پیشنهادی نسبت به سایر روش های انتخاب ویژگی توزیع شده کاهش یافته است. به طوری که در سایر روش های انتخاب ویژگی، حالت توزیع شده نسبت به متمرکز بهبود یافته است.

برای حالت متمرکز همراه با انتخاب ویژگی ReliefF داشته است. مجموعه داده Madelon، هم در حالت متمرکز و هم در حالت توزیع شده در چارچوب پیشنهادی با طبقه بند KNN بهترین نتیجه را به خود اختصاص داده است. در طبقه بند SVM، برای تمام مجموعه داده ها به جز Spect و Ozon، بهترین نتیجه مربوط به چارچوب پیشنهادی است. مجموعه داده Ozon، بهترین مقدار صحت را برای حالت متمرکز در تمام انتخاب ویژگی ها به غیر از ReliefF داشته است. مجموعه داده Madelon، هم در حالت توزیع شده و هم در حالت متمرکز در تمام انتخاب ویژگی ها دارای نتایج مشابهی است. اما در طبقه بند NB بهترین میانگین صحت مربوط به نسخه توزیع شده الگوریتم IG است.

چارچوب پیشنهادی در حالت متمرکز، بهترین نتیجه را در میانگین صحت دارد. همچنین برای مجموعه داده های Isolet، Ozon و Connect4 با طبقه بند NB بهترین نتیجه مربوط به چارچوب پیشنهادی به صورت توزیع شده است. البته در مجموعه داده Connect4 در تمام انتخاب ویژگی ها در حالت توزیع شده دارای نتایج مشابهی است. همچنین بهترین نتیجه در حالت توزیع شده در این طبقه بند مربوط به مجموعه داده های Ozon و Spambase است که دارای مقادیر مشابهی هستند. از طرفی با توجه به نتایج حاصل از ستون آخر در هر جدول، تنها در کمتر از $\frac{1}{3}$ موارد آزمایش شده، دقت طبقه بندی بدون استفاده از الگوریتم انتخاب ویژگی، بهتر شده است. در بین مجموعه داده های

(جدول ۲): مشخصات مجموعه داده ها
(Table-2): Characteristics of used Dataset

Dataset	Features	Samples		Classes
		Training	Test	
Ozone	72	1691	845	2
Madelon	500	1600	800	2
Spambase	57	3067	1534	2
Connect4	42	45038	22519	3
Spect	22	178	89	2
Isolet	617	5198	2599	26

(جدول-۳): نتایج تست دقت طبقه‌بندی با KNN
(Table-3): Results of KNN Classifier Accuracy

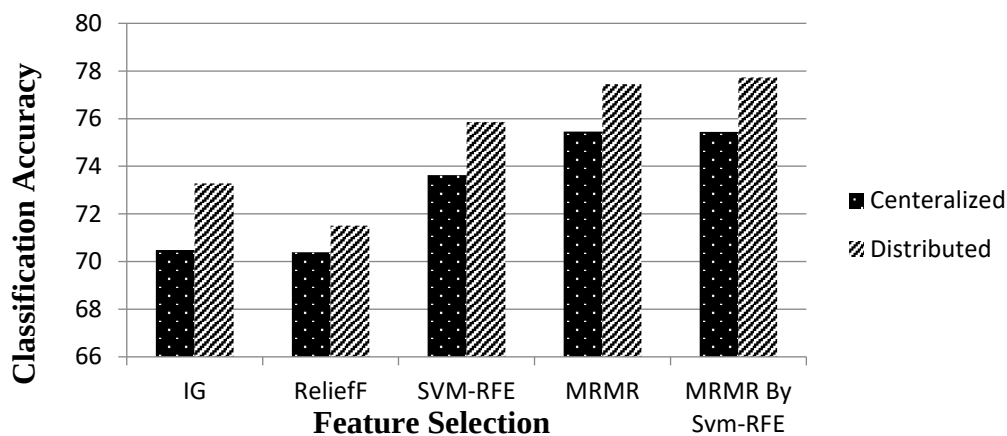
Dataset	C-IG D-IG	C-ReliefF D-ReliefF	C-SVM_RFE D-SVM_RFE	C-mRMR D-mRMR	SVM-RFE_By_C-mRMR SVM-RFE_By_D-mRMR	With out Feature Selection
Ozone	98.70 98.82	98.22 93.79	98.70 90.18	98.70 98.70	98.70 98.58	98.12
Madelon	49.62 49.62	49.62 49.62	49.62 49.62	49.62 49.62	49.62 49.62	48.62
Spambase	95.56 94.98	77.69 88.00	94.65 91.86	92.37 96.41	90.54 99.61	91.65
Spect	68.54 64.04	66.29 64.04	64.04 70.79	65.17 65.17	64.04 62.92	62.92
Isolet	92.90 92.99	95.17 95.19	94.25 95.30	91.02 93.35	92.68 95.60	90.42
Connect4	66.42 66.42	60.42 60.42	71.22 74.72	72.26 73.12	73.12 75.42	65.58
Average	78.62 77.81	74.57 75.18	78.75 78.74	78.19 79.39	78.12 80.29	-----

(جدول-۴): نتایج تست دقت طبقه‌بندی با SVM
(Table-4): Results of SVM Classifier Accuracy

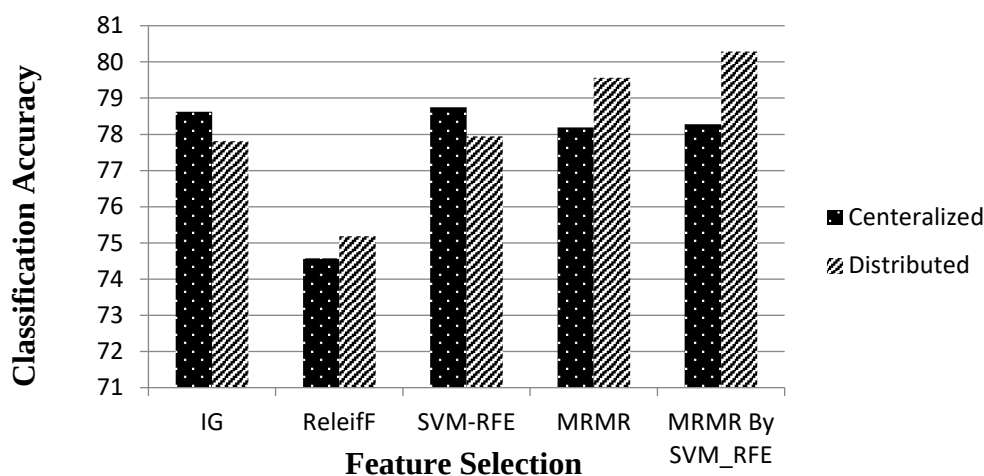
Dataset	C-IG D-IG	C-ReliefF D-ReliefF	C-SVM_RFE D-SVM_RFE	C-MRMR D-MRMR	SVM-RFE_By_C-MRMR SVM-RFE_By_D-MRMR	With out Feature Selection
Ozone	98.75 98.70	98.70 98.50	98.58 98.34	98.58 98.50	98.58 98.00	98.35
Madelon	49.37 49.13	48.25 50.87	49.50 50.75	50.15 53.13	51.26 53.75	48.25
Spambase	77.62 78.80	75.60 79.71	72.73 75.53	74.49 76.02	73.84 77.69	77.10
Spect	73.03 73.03	71.91 76.40	65.17 71.91	74.15 76.50	69.66 72.79	73.03
Isolet	73.64 81.14	90.62 90.72	88.26 89.14	90.06 91.35	91.45 92.82	85.23
Connect4	50.53 58.92	37.26 32.74	67.54 69.51	66.35 70.12	67.87 71.22	62.06
Average	70.49 73.28	70.39 71.50	73.63 75.86	75.63 77.45	75.44 77.84	-----

(جدول-۵): نتایج تست دقت طبقه‌بندی با NB
(Table-5): Results of NB Classifier Accuracy

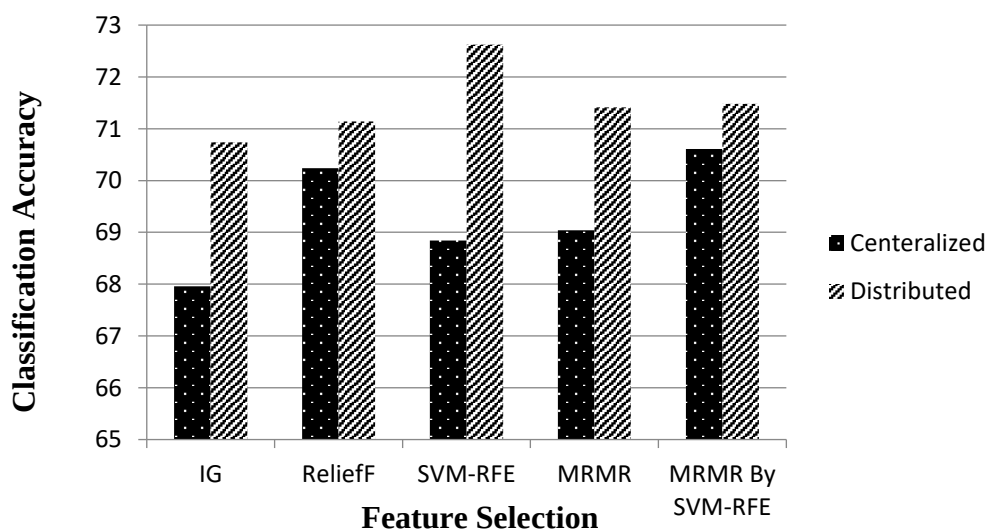
Dataset	C-IG D-IG	C-ReliefF D-ReliefF	C-SVM_RFE D-SVM_RFE	C-mRMR D-mRMR	SVM-RFE_By_C-mRMR SVM-RFE_By_D-mRMR	With out Feature Selection
Ozone	70.41 73.61	71.84 71.71	71.24 79.53	73.67 73.50	72.43 72.20	71.95
Madelon	47.25 47.87	46.87 50.00	48.85 49.5	46.50 47.0	47.42 48.88	46.87
Spambase	76.53 89.70	92.85 91.25	80.12 89.68	78.88 93.22	90.35 93.46	93.20
Spect	73.03 74.16	76.50 75.28	71.91 74.16	74.15 71.91	70.79 74.16	74.16
Isolet	79.98 78.68	73.10 78.18	80.46 82.42	81.25 82.42	82.24 86.45	80.34
Connect4	60.58 60.42	60.30 60.42	60.42 60.42	60.10 60.42	60.42 60.42	58.62
Average	67.96 70.74	70.24 71.14	68.84 72.62	69.09 71.41	70.61 72.59	-----



(شکل-۴): میانگین دقت طبقه‌بندی KNN روی رویکردهای مختلف انتخاب ویژگی
(Figure-4): Mean of KNN classifier accuracy on different feature selection methods



(شکل-۵): میانگین دقت طبقه‌بندی SVM روی رویکردهای مختلف انتخاب ویژگی
(Figure-5): Mean of SVM classifier accuracy on different feature selection methods



(شکل-۶): میانگین دقت طبقه‌بندی NB روی رویکردهای مختلف انتخاب ویژگی
(Figure-6): Mean of NB classifier accuracy on different feature selection methods

(جدول - ۶): زمان اجرای الگوریتم‌های انتخاب ویژگی تست شده بر حسب ثانیه
(Table-6): Runtime (seconds) for all tested feature selection algorithms

		Ozone	Madelon	Spect	Spambase	Connect4	Isolet	Average
mRMR	C	40.450	3631.541	0.646	18.917	1565.512	5700.237	1826.217
	D	31.561	98.21	0.288	13.052	62.820	164.64	61.762
SVM-RFE	C	4.728	6.960	0.680	8.962	12.250	9.146	7.121
	D	0.245	0.712	0.224	0.645	0.828	0.947	0.600
IG	C	2.116	22.008	0.316	3.077	7.847	30.742	11.018
	D	1.429	16.037	0.096	2.737	2.183	20.05	7.089
RELEIFF	C	10.284	61.802	0.711	15.920	1502.104	603.497	365.720
	D	0.805	15.310	0.024	2.712	8.195	123.17	25.036
_By_mRMR SVM-RFE	C	88.703	3925.284	1.015	26.374	1743.731	6028.157	1968.877
	D	41.863	155.683	0.771	7.348	37.387	256.091	83.190

(جدول - ۷): تعداد ویژگی‌های انتخاب شده توسط رویکردهای مختلف انتخاب ویژگی
(Table-7): Number of selected features by tested feature selection algorithm

		Ozone	Madelon	Spect	Spambase	Connect4	Isolet
Full Set		72	500	22	57	42	617
mRMR	C	55	185	21	42	30	226
	D	31	52	16	29	32	124
SVM-RFE	C	41	175	16	42	30	226
	D	16	40	7	30	33	102
IG	C	41	175	16	42	30	226
	D	37	155	10	37	30	168
ReliefF	C	40	174	16	42	30	226
	D	17	68	10	21	34	175
SVM-RFE_By_mRMR	C	59	175	16	42	30	226
	D	37	50	18	34	32	131

نسبت به سایرین دارای کمترین تعداد ویژگی و نمونه است. در بین سایر مجموعه داده‌ها، کمترین زمان اجرای توزیع شده مربوط به انتخاب ویژگی SVM-RFE است. ستون آخر از جدول ۶، میانگین زمان اجرا را در هر دو حالت توزیع شده و متمرکز نشان می‌دهد. با توجه به داده‌های این جدول، بیشترین اختلاف بین زمان اجرای متمرکز و توزیع شده مربوط به چارچوب پیشنهادی است؛ به طوری که به میزان قابل توجهی کاهش یافته است. (از ۱۹۶۸/۸۷۷ ثانیه به ۸۳/۱۹۰ کاهش یافته است).

۴-۴- نتایج تعداد ویژگی‌ها

تعداد ویژگی‌های انتخاب شده هم در حالت توزیع شده (D) و هم در حالت متمرکز (C)، با استفاده از پنج روش مختلف انتخاب ویژگی بر روی مجموعه داده‌های یاد شده در (جدول ۷)، نشان داده شده است. در تمام حالت‌ها از حد آستانه ۲۵ درصد برای حذف ویژگی‌ها استفاده شده است. همان‌طور که دیده می‌شود، تعداد ویژگی‌های

۴-۳- نتایج زمان اجرا

زمان اجرای الگوریتم‌های انتخاب ویژگی به صورت توزیع شده (سطرهای با برچسب D در جدول) و متمرکز (سطرهای با برچسب C در جدول) در (جدول ۶)، بر روی شش مجموعه داده نشان داده شده است. در رویکرد توزیع شده با توجه به این که همه زیرمجموعه‌ها می‌توانند در یک زمان پردازش شوند؛ لذا زمان نمایش داده شده در این جدول، بیشینه زمان مورد نیاز برای الگوریتم انتخاب ویژگی روی هر زیرمجموعه ایجاد شده در مرحله تقسیم‌بندی است. در این آزمایش، همه زیرمجموعه‌ها در یک ماشین پردازش شده‌اند، اما الگوریتم پیشنهاد شده می‌تواند بر روی چندین پردازنده اجرا شود [3,6,8,9].

در مقایسه با روش متمرکز، زمان اجرای مورد نیاز در روش‌های توزیع شده روی تمام مجموعه داده‌ها کاهش یافته است. کمترین زمان اجرا در بین مجموعه داده‌های آزمایش شده به صورت توزیع شده، متعلق به مجموعه داده Spect با انتخاب ویژگی ReliefF است. این مجموعه داده

انتخاب ویژگی به منظور بررسی ویژگی‌های مرتبط و افزونه توانسته دقت طبقه‌بندی را بهبود دهد. از طرفی زمان اجرای فرایند انتخاب ویژگی پیشنهادی به صورت توزیع شده نسبت به حالت متمرکز کاهش یافته است. همچنین کاهش تعداد ویژگی‌ها باعث کاهش دقت طبقه‌بندی نشده و در اکثر موارد نسبت به حالت متمرکز بهبود داشته است.

روش پیشنهاد شده، می‌تواند با هر الگوریتم انتخاب ویژگی از نوع نهفته و از نوع فیلتر چندمتغیره استفاده شود. به عنوان کار و پژوهش آینده، می‌توان توزیع و تقسیم‌بندی عمودی مجموعه داده‌ها را به جای افقی بررسی کرد، همچنین می‌توان استراتژی‌های جدید دیگری را برای ترکیب سایر روش‌های انتخاب ویژگی، توسعه داد.

6- Refrence

۶- منابع

- [1] I. Guyon, A. Elisseeff, (2003), "An introduction to variable and feature selection", Journal of Machine Learning Research, Vol.3, pp.1157-1182.
- [2] I. Guyon, S. Gunn, M. Nikravesh and L. A. Zadeh, (2006), "Feature Extraction: Foundations and Applications", vol. 207, Springer, ISBN-10: 9783540354871.
- [3] V. Bolón-Canedo, N. Sánchez-Marño, and A. Alonso-Betanzos (2013), "A Distributed Wrapper Approach for Feature Selection", ESANN proceedings, Computational Intelligence and Machine Learning, ISBN 978-2-87419-081-0.
- [4] V. Bolón-Canedo, N. Sánchez-Marño, and A. Alonso-Betanzos (2015), "A Distributed Feature Selection Approach Based on a Complexity Measure", Advances in Computational Intelligence, pp. 15-28.
- [5] G. Chandrashekar and F. Sahin (2014), "A survey on feature selection methods", Journal of Computers and Electrical Engineering vol. 40, pp.16-28.
- [6] V. Bolón-Canedo, N. Sánchez-Marño, and A. Alonso-Betanzos (2015), "Distributed feature selection: An application to microarray data classification", Applied Soft Computing, vol. 30, pp. 136-150.
- [7] V. Bolón-Canedo, N. Sánchez-Marño, and J. Cerviño-Rabuñal (2013), "Scaling up feature selection: a distributed filter approach", Advances in Artificial Intelligence, pp. 121-130.
- [8] L. Morán-Fernández, V. Bolón-Canedo, and A. Alonso-Betanzos (2016), "Centralized vs. distributed feature selection methods based on data complexity measures", Journal of Knowledge-Based Systems, vol. 117, pp.27-45.

انتخاب شده به وسیله رویکردهای متمرکز نسبت به توزیع شده در بیشتر مجموعه داده‌ها، بیشتر است. در مواردی که دقت رده‌بندی در حالت متمرکز نسبت به توزیع شده بیشتر است، می‌توان به این نتیجه رسید که از طریق توزیع داده‌ها تنزل قابل توجهی در دقت طبقه‌بندی رخ نداده است. در مجموعه داده Ozon در بیشتر موارد دقت رده‌بندی به وسیله رویکردهای متمرکز بیشتر از توزیع شده به دست آمده است. در بین مجموعه داده‌های مورد ارزیابی، بالاترین دقت در حالت متمرکز نسبت به توزیع شده مربوط به مجموعه داده Spect است. این در حالی است که نتایج به دست آمده به وسیله رویکردهای توزیع شده تنها دو یا سه درصد پایین تر هستند. برای سایر مجموعه داده‌ها، بهترین نتایج برای رویکردهای توزیع شده گزارش شده است. در مجموعه داده Connect4 در اکثر موارد دقت طبقه‌بندی به وسیله رویکردهای توزیع شده بهبود یافته است که این ممکن است، به علت تعداد کم ویژگی‌های انتخاب شده توسط رویکردهای متمرکز باشد. بالاترین بهبود مربوط به رویکرد پیشنهادی (MRMRBY SVM-RFE) با مجموعه داده Spambase همراه با طبقه‌بند SVM است که با وجود کاهش تعداد ویژگی‌های انتخاب شده در حالت توزیع شده، دقت آن کاهش نیافته است.

۵- نتیجه گیری

در این مقاله یک چارچوب جدید برای انتخاب ویژگی مبتنی بر الگوریتم‌های نهفته و فیلتر چندمتغیره پیشنهاد شده است. در این چارچوب پیشنهادی، فرایند انتخاب ویژگی به صورت توزیع افقی برای مجموعه داده‌هایی که تعداد نمونه هایشان بیشتر است، ارائه شده است. مراحل اصلی این چارچوب عبارتند از: توزیع افقی مجموعه داده‌ها، اجرای فرایند انتخاب ویژگی پیشنهاد شده روی هر زیرمجموعه و ترکیب نتایج نهایی از انتخاب ویژگی پیشنهادی در یک زیرمجموعه. در فرایند انتخاب ویژگی پیشنهادی، استفاده از الگوریتم‌های فیلتر چندمتغیره به وسیله الگوریتم‌های نهفته معرفی شده است. به طوری که انتخاب ویژگی‌های مرتبط مبتنی بر طبقه‌بند از طریق الگوریتم‌های نهفته بررسی شده و تشخیص و کاهش ویژگی‌های افزونه از طریق الگوریتم‌های فیلتر چندمتغیره محاسبه شده است. نتایج نشان می‌دهد که چارچوب پیشنهاد شده در بیشتر مجموعه داده‌های آزمایش شده، به دلیل استفاده از دو نوع روش متنوع برای

Trans. Pattern Anal. Mach. Intell., vol. 27, no. 8, pp. 1226–1238, Aug..

- [23] M.A. Hall, L.A. Smith, (1998), "Practical feature subset selection for machine learning", Comput. Sci.98,181–191
- [24] I. Kononenko,(1994)," Estimating attributes: analysis and extensions of RELIEF", Machine Learning: ECML-94, vol. 784, pp 171-182
- [25] M. Robnik-Šikonja and I. Kononenko, (2003), "Theoretical and empirical analysis of ReliefF and RReliefF", Machine learning, vol. 53, Issue:1-2, pp. 23-69.



منا شریف‌نژاد، فارغ‌التحصیل دوره کارشناسی‌ارشد رشته مهندسی کامپیوتر- نرم‌افزار است. در حال حاضر دانشجوی دوره دکترای در دانشگاه اراک است. ایشان هم‌اکنون عضو هیئت علمی دانشگاه آزاد اسلامی اراک است. حوزه تخصصی ایشان هوش مصنوعی است.

نشانی رایانامه ایشان عبارت است از:

m.sharifnezhad98@gmail.com



محسن رحمانی، مدرک دکترای خود را در سال ۱۳۸۷ از دانشگاه علم و صنعت ایران دریافت کردند. ایشان هم اکنون عضو هیئت علمی دانشگاه اراک

هستند. تخصص ایشان پردازش سیگنال‌های صوتی و هوش مصنوعی است.

نشانی رایانامه ایشان عبارت است از:

m-rahmani@araku.ac.ir



حسین غفاریان، مدرک دکترای خود را از دانشگاه علم و صنعت ایران دریافت کردند. ایشان هم‌اکنون عضو هیئت علمی دانشگاه اراک هستند.

تخصص ایشان شبکه‌های رایانه‌ای، کیفیت خدمات، کنترل‌کننده‌های فازی، داده‌کاوی، الگوریتم‌های تکاملی، سامانه‌های حمل و نقل هوشمند است.

نشانی رایانامه ایشان عبارت است از:

h-ghaffarian@araku.ac.ir

- [9] L. Mor'an-Fern'andez, V. Bol'on-Canedo, and A. Alonso-Betanzos(2015), "A Time Efficient Approach for Distributed Feature Selection Partitioning by Features", Lecture Notes in Computer Science book series (LNCS), vol. 9422, pp.245–254.
- [10] L. Yu, H. Liu, (2004)," Efficient feature selection via analysis of relevance and redundancy", J. Mach. Learn. Res. 5 , 1205–1224.
- [11] C. Ding, H. Peng, (2005) "Minimum redundancy feature selection from microarray gene expression data", Journal of Bioinformatics and computational Biology, Vol.03, No.02, pp.185–205.
- [12] R. Kohavi, GH. John (1997), "Wrappers for feature subset selection", Artificial Intelligence, Vol. 97, Issues 1-2, pp.273–324.
- [13] J.Li, K.Cheng, S.Wang, F. Morstatter, and R. P. Trevino(2018)," Feature Selection: A Data Perspective", Journal of ACM Computing Surveys (CSUR), Vol. 50 ,Issue 6.
- [14] A.De Haro Garc'ia, (2011), "Scaling data mining algorithms. Application to instance and feature selection", Ph.D. Thesis, University of Granada.
- [15] H. Djellali, N. Ghoulmi Zine and N. Azizi (2016), "Two Stages Feature Selection Based on Filter Ranking Methods and SVMRFE on Medical Applications", Modelling and Implementation of Complex Systems. Lecture Notes in Networks and Systems, Springer, Cham,, vol. 1, pp. 281-293.
- [16] H. Min and W. Fangfang (2010), "Filter-Wrapper Hybrid Method on Feature Selection ", Second WRI Global Congress on Intelligent Systems (GCIS), pp.98-101.
- [17] I. Guyon, J. Weston, S. Barnhill, and V. Vapnik (2002), "Gene selection for cancer classification using support vector machines", Journal of Machine Learning, vol. 46, Issue 1-3, pp. 389–422.
- [18] D. Boughaci and A.A Alkhalwaldeh (2018), "Three local search-based methods for feature selection in credit scoring", Vietnam Journal of Computer Science, May 2018, Vol. 5, Issue 2, pp. 107–121.
- [19] Q.Wang , J. Wan, F. Nie , B. Liu , C.Yan , and X. Li (2019), "Hierarchical Feature Selection for Random Projection", IEEE Transactions on Neural Networks and Learning Systems, Vol. 30 , Issue 5, pp. 1581 – 1586.
- [20] <http://archive.ics.uci.edu/ml/datasets/>
- [21] I.Guyon, J.Weston, S.Barnhill and V.Vapnik,(2002), "Gene selection for cancer classification using support vector machines," Jorنال of Machine Learning, vol.46, pp.389–422.
- [22] H. Peng, F. Long, and C. Ding, (2005), "Feature selection based on mutual information: Criteria of max-dependency, max-relevance, and minredundancy, "IEEE