

اعمال مدل‌های رگرسیون بر زیرمجموعه‌های با همبستگی بالا برای بهبود جایگذاری مقادیر جافتاده عددی

امیرمسعود سفیدیان^۱، کارشناسی ارشد؛ نگین دانشپور^۲، استادیار

۱- دانشکده مهندسی کامپیوتر - دانشگاه تربیت دبیر شهید رجایی - تهران - ایران - amirmasoud.sefidian@sru.ac.ir

۲- دانشکده مهندسی کامپیوتر - دانشگاه تربیت دبیر شهید رجایی - تهران - ایران - ndaneshpour@sru.ac.ir

چکیده: حضور مقادیر جافتاده در داده‌های دنیای واقعی مشکلی بسیار رایج و غیرقابل اجتناب است. بنابراین لازم است تا پیش از عملیات اکتشاف دانش، این مقادیر جافتاده به‌طور دقیق پُر شوند. در این مقاله، سه رویکرد جدید برای تخمین مقادیر جافتاده عددی پیشنهاد می‌شود. در تمامی روش‌های پیشنهادی، مدل‌های رگرسیون بر زیرمجموعه‌هایی با همبستگی بالا اعمال می‌شوند. در انتخاب زیرمجموعه‌های مطلوب سعی می‌شود تا همبستگی بین صفت جافتاده و دیگر صفات حداکثر شود. انتخاب این زیرمجموعه‌ها با استفاده از رویکردهایی مبتنی بر انتخاب روبه‌جلو انجام می‌شود. از معیار ضریب همبستگی برای اندازه‌گیری میزان ارتباط بین صفات استفاده شده است. همچنین در روش‌های پیشنهادی، ترتیب صفات جافتاده برای انجام عمل جایگذاری اولویت‌دهی می‌شوند. عملکرد رویکردهای پیشنهاد شده بر روی پنج مجموعه داده از دنیای واقعی با مقادیر مختلف جافتادگی ارزیابی شده است. عملکرد رویکردهای ارائه شده با پنج رویکرد جایگذاری با مقدار میانگین، جایگذاری با استفاده از نزدیک‌ترین همسایگان، روش جایگذاری با خوشه‌بندی c-means فازی، روش جایگذاری با درخت تصمیم و روشی مبتنی بر رگرسیون به نام «الگوریتم جایگذاری با رگرسیون افزایشی صفات» (IARI) مقایسه شده است. از دو معیار شناخته شده ریشه میانگین مربعات خطا و ضریب تعیین برای مقایسه عملکرد رویکردهای پیشنهادی با دیگر روش‌های جایگذاری استفاده شده است. نتایج آزمایش‌ها نشان می‌دهد که رویکردهای ارائه شده، حتی زمانی که درصد جافتادگی بالا است، بهتر از دیگر روش‌های مقایسه شده عمل می‌کنند.

واژه‌های کلیدی: جایگذاری مقادیر جافتاده، همبستگی، رگرسیون.

Applying Regression Models on Subsets with High Correlations for a Better Numeric Missing Values Imputation

Amir Masoud Sefidian¹, MSc; Negin Daneshpour², Assistant Professor

1- Faculty of Computer Engineering, Shahid Rajaee Teacher Training University, Tehran, Iran, Email: amirmasoud.sefidian@sru.ac.ir

2- Faculty of Computer Engineering, Shahid Rajaee Teacher Training University, Tehran, Iran, Email: ndaneshpour@sru.ac.ir

Abstract: The presence of missing values in the real world data is a very prevalent and inevitable problem. So, it's necessary to fill up these missing values accurately, before they are used for knowledge discovery process. This paper proposes three novel methods to fill numeric missing values. All of the proposed methods apply regression models on subsets of data which there are strong correlations among them. These subsets are selected using forward selection based approaches. In the selection of the desired subsets, it is tried to maximize the correlation between missing attribute and other attributes. The correlation coefficient is used to measure the relationships between attributes. The priority of each missing attribute for imputation purpose is also considered in the proposed methods. The performance of proposed methods is evaluated on five real world datasets with different missing ratios. The efficiency of the proposed methods is compared with five different estimation methods, namely, the mean imputation, the k nearest neighbours imputation, a fuzzy c-means based imputation, a decision tree based imputation, and a regression based imputation algorithm, called "Incremental Attribute Regression Imputation" (IARI) method. Two well-known evaluation criteria, namely, Root Mean Squared Error (RMSE) and Coefficient of Determination (CoD) are used to compare the performance of proposed methods with other imputation methods. Experimental results show that the proposed methods perform better than other compared methods, even when the missing ratio is high.

Keywords: Missing values imputation, Correlation, Regression.

تاریخ ارسال مقاله: ۱۳۹۶/۰۵/۰۸

تاریخ اصلاح مقاله: ۱۳۹۶/۰۸/۰۶

تاریخ پذیرش مقاله: ۱۳۹۶/۰۹/۲۰

نام نویسنده مسئول: نگین دانشپور

نشانی نویسنده مسئول: ایران - تهران - لویزان - دانشگاه تربیت دبیر شهید رجایی - دانشکده مهندسی کامپیوتر.

۱- مقدمه

اولین گام در فرآیند کشف دانش، جمع‌آوری داده‌ها است. در این گام، نمونه‌ها از تک منبع یا منابع چندگانه جمع‌آوری می‌شوند. داده‌کاوی باکیفیت بالا فقط هنگامی حاصل می‌شود که داده‌های جمع‌آوری شده کیفیت بالایی داشته باشند. نواقص مختلفی مثل ناسازگاری‌ها، خطاها، داده‌های پرت و مقادیر جافتاده در داده‌های جمع‌آوری شده وجود دارد. بنابراین برای حاصل شدن نتایج مطلوب از هدف کشف دانش، این داده‌ها نیاز به پیش‌پردازش دارند. پیش‌پردازش داده‌ها اهمیت فراوانی دارد، به نحوی که حدود ۸۰٪ از کل تلاش برای مهندسی داده صرف آماده‌سازی داده‌ها برای مراحل بعد می‌شود [۱]. از این رو، تلاش‌های متعددی برای بهبود کیفیت داده‌ها در حوزه‌های مختلف علمی و صنعتی صورت گرفته است. برای مثال از روش‌های ترکیبی داده‌کاوی برای بهبود کیفیت داده‌های جمع‌آوری شده در صنایع نفت و گاز [۲] و از شبکه‌های عصبی برای پیش‌پردازش در حوزه‌ی پردازش تصویر [۳] استفاده شده است.

یکی از مشکلات اساسی در دنیای واقعی، وجود مقادیر جافتاده است. در این حالت مقدار یک یا چند صفت از یک یا چند نمونه در یک مجموعه داده به دلایل مختلف، نامعلوم یا تهی^۱ است. حضور مقادیر جافتاده در داده‌ها نه تنها مشکلی رایج در بسیاری از تحقیقات تجربی مثل تحقیقات صنعتی، تجاری و علمی است، بلکه مشکلی غیرقابل اجتناب است [۴]. برای مثال، حدود ۴۵٪ از مجموعه‌های موجود در مخزن داده‌ی UCI [۵] (که یکی از مشهورترین مخازن داده برای اهداف یادگیری ماشین است) شامل مقادیر جافتاده هستند [۶]. بنابراین یک وظیفه‌ی مهم در پیش‌پردازش داده‌ها، تخمین تا حد ممکن دقیق مقادیر جافتاده با استدلال از داده‌های کامل مشاهده شده برای افزایش کیفیت داده‌ها است. از این عملیات با نام‌های مختلفی از جمله جایگذاری^۲، تخمین و یا پر کردن مقادیر جافتاده در تحقیقات علمی یاد می‌شود. جایگذاری مقادیر جافتاده همانند دیگر وظایف پیش‌پردازش امری زمان‌بر است [۷].

مقادیر جافتاده به دلایل متفاوتی از جمله خرابی سیستم قدرت، عوامل محیطی، اندازه‌گیری‌های غلط، روبه‌های دستی برای ثبت داده و اختلال در عملکرد تجهیزات ذخیره‌سازی داده رخ می‌دهند [۸-۱۰].

مقادیر جافتاده به‌ویژه هنگامی که نرخ جافتادگی بالا باشد، موجب ایجاد مشکلات متعددی می‌شود. حضور مقادیر جافتاده به‌طور مستقیم کیفیت نتایج نهایی به‌دست‌آمده از فرآیند یادگیری یا کاوش از داده‌ها را تحت تأثیر قرار می‌دهد [۱۱]. بسیاری از الگوریتم‌های تحلیل داده‌ها قابلیت اعمال به مجموعه داده‌هایی با مقادیر جافتاده را ندارند، چراکه اصولاً این الگوریتم‌ها به‌گونه‌ای طراحی شده‌اند که فقط با داده‌های کامل می‌توانند کار کنند [۱۲]. وجود مقادیر جافتاده می‌تواند منجر به نتیجه‌گیری‌های گمراه‌کننده، غیر صحیح و غیرمعقول از یک تحقیق علمی شود و ممکن است کل فرآیند جمع‌آوری و تحلیل داده‌ها را برای کاربران بی‌فایده کند [۱۳]. علاوه بر این، حضور مقادیر جافتاده می‌تواند قابلیت تعمیم یافته‌های یک تحقیق را محدود کند [۱۰].

همبستگی^۳ بین صفات در یک مجموعه داده از ویژگی‌های طبیعی آن مجموعه داده است. اغلب ممکن است زیرمجموعه‌هایی درون یک مجموعه داده یافت که روابط قوی‌تری نسبت به روابط حاکم بر کل مجموعه داده در آن‌ها وجود داشته باشد. در کاربردهای یادگیری ماشین یا داده‌کاوی معمولاً صفات با همبستگی بالا به‌عنوان صفات افزونه^۴ در نظر گرفته می‌شوند و فقط از یکی از این صفات برای انجام تحلیل‌ها استفاده می‌شود. اگرچه، این روابط می‌توانند شامل اطلاعات بسیار مفیدی برای هدف تخمین مقادیر جافتاده باشند و مقادیر نامعلوم را می‌توان با استفاده از این روابط تخمین زد.

روش‌های بسیار زیادی مبتنی بر رگرسیون^۵ برای هدف تخمین مقادیر جافتاده وجود دارند. اکثر این روش‌ها مانند روش جایگذاری با استفاده از رگرسیون افزایشی صفات (IARI^۶) [۱۴]، مدل‌های رگرسیون را بر روی کل مجموعه داده اعمال می‌کنند. با این حال، نکته‌ی مهم در مورد این روش‌ها این است که دقت این روش‌ها هنگامی که روابط در کل مجموعه داده ضعیف باشند، به شدت افت می‌کند. اعمال مدل‌های رگرسیون بر افزایشی از داده‌ها با همبستگی بالا می‌تواند منجر به تولید نتایج بهتری در تخمین داده‌ها شود. بنابراین، عملکرد این روش‌ها با توجه به این نکته می‌تواند بهبود یابد.

با مطالعات انجام‌شده توسط نویسندگان، مشاهده شد که تنها روش DMI^۷ [۱۵] بحث همبستگی در داده‌ها را برای بهبود تخمین داده‌های جافتاده موردتوجه قرار داده است. اگرچه این روش بخش‌هایی با همبستگی بالا را توسط درخت تصمیم پیدا می‌کند، با این حال از معیار دقیقی برای سنجش میزان همبستگی بین صفات استفاده نمی‌کند و همواره فرض می‌کند که استفاده از درخت تصمیم منجر به افزایش همبستگی در داده‌ها می‌شود. از این رو، ضروری به نظر می‌رسد که مقدار همبستگی را توسط معیاری دقیق برای اهداف ذکر شده اندازه‌گیری کرد. با توجه به این نکات، در این پژوهش روش‌هایی معرفی می‌شوند که قادر به یافتن بخش‌هایی از داده‌ها با همبستگی بالا توسط معیاری کمی برای اندازه‌گیری همبستگی قوی هستند. بنابراین با استفاده از روش‌های پیشنهاد شده ابتدا سعی می‌شود تا زیرمجموعه‌هایی با همبستگی بالا را نسبت به هر نمونه‌ی جافتاده یافت. سپس روش رگرسیون را بر آن زیرمجموعه‌ها اعمال کرد تا دقت مدل تخمین‌زننده را تا حد امکان بالا برد. با توجه به مطالعات صورت گرفته، تاکنون چنین رویکردی برای استفاده در تخمین مقادیر جافتاده برای مدل رگرسیون ارائه نشده است. به‌علاوه، همان‌طور که گفته شد، تاکنون از معیار دقیقی برای سنجش میزان همبستگی در یک مجموعه داده برای هدف بهبود تخمین مقادیر جافتاده عددی استفاده نشده است.

با توجه به نکات گفته‌شده، در این مقاله سه روش جدید پیشنهاد می‌شود تا مشکلات روش‌های گذشته را تا حد امکان رفع کنند. به بیان دقیق‌تر، از معیاری کمی به نام ضریب همبستگی، برای یافتن زیرمجموعه‌هایی از داده‌ها با همبستگی بالا نسبت به هر نمونه‌ی جافتاده استفاده می‌شود. ترتیب انتخاب صفات جافتاده برای عمل جایگذاری می‌

قوانین انجمنی را استخراج نمود. پس از یافتن قوانین یافت شده از مجموعه داده‌ی تراکنشی، برای هر مقدار جافتاده، مجموعه قوانینی که در سمت راست آن‌ها همان صفت جافتاده حضور دارند، انتخاب می‌شوند. برای هر یک از قوانین انتخاب شده امتیازی محاسبه می‌گردد. قانونی با بیش‌ترین امتیاز به عنوان قانون برگزیده انتخاب شده و از مقدار سمت راست آن قانون به عنوان مقدار جایگزین صفت مقدار صفت جافتاده استفاده می‌شود. از نقاط قوت این روش می‌توان به دخیل کردن معیار وابستگی دو طرف قوانین انجمنی و طول قوانین و حمایت از مقادیر جافتاده‌ی چندگانه اشاره کرد.

از شبکه‌های عصبی نیز برای پُر کردن مقادیر جافتاده استفاده شده است. رایج‌ترین روش مبتنی بر شبکه‌های عصبی برای تخمین مقادیر جافتاده، استفاده از شبکه‌ی عصبی پرسپترون چندلایه (MLP^{10}) است. در یکی از روش‌ها، داده‌های جافتاده برای هجده مجموعه داده تخمین زده شده‌اند [۷]. همچنین تأثیر روش‌های مختلف یادگیری و پارامترهای مختلف شبکه عصبی بر روی دقت نهایی نیز بررسی شده‌اند. با طراحی و آموزش یک شبکه عصبی نگاشت خودسازمان‌ده 11 (SOM) نیز می‌توان مقادیر جافتاده را تخمین زد [۱۸]. مشکل اصلی شبکه‌های عصبی زمان بر بودن فرآیند یادگیری و ساخت مدل بهینه برای یک مجموعه داده نسبت به دیگر روش‌ها و همچنین قابلیت تفسیر 12 پایین است [۱۹]. به علاوه، اصولاً آموزش مناسب شبکه‌های عصبی نیاز به تعداد نمونه‌های ورودی بالایی دارد. بنابراین هنگامی که تعداد نمونه‌های جافتاده بالا است، دقت تخمین این روش‌ها کاهش می‌یابد.

روش k نزدیک‌ترین همسایگان 13 (kNN) یکی از معروف‌ترین و پایه‌ترین روش‌ها برای تخمین مقادیر جافتاده است. سادگی و نتایج قابل قبول آن نسبت به دیگر روش‌ها از ویژگی‌های آن است. از دیگر مزایای این روش قابلیت استفاده برای هم صفات عددی و هم صفات غیر عددی با تغییر تابع سنجش فاصله است. همچنین در این روش زمانی صرف ساخت مدل برای تخمین داده‌های جافتاده نمی‌شود [۲۰]. یک روش مبتنی بر رویکرد k نزدیک‌ترین همسایگان برای پُر کردن مقادیر جافتاده، استفاده از درجه‌ی جافتادگی (تعداد صفات جافتاده) هر نمونه است [۲۱]. در این رویکرد ابتدا درجه‌ی جافتادگی هر نمونه محاسبه می‌شود. نمونه‌هایی که درجه‌ی جافتادگی آن‌ها بزرگ‌تر از دو است از ادامه‌ی پردازش خارج می‌شوند و سپس مقادیر جافتاده با مقادیر نزدیک‌ترین نمونه‌های کامل پُر می‌شوند. از نقاط قوت این روش می‌توان به سادگی رویکرد ارائه شده و قدرت تعمیم بالا برای کاربردهای مختلف با تغییر معیار شباهت اشاره کرد. عدم توانایی در پُر کردن مقادیر جافتاده‌ی چندگانه، حساسیت به داده‌های پرت و حساسیت بالا به تابع شباهت از نقاط ضعف این روش است. همچنین لزوماً بررسی فقط نزدیک‌ترین نمونه، همواره بهترین گزینه نیست.

یک رویکرد ارائه شده برای تخمین مقادیر جافتاده، روش k نزدیک‌ترین همسایه‌ی تکراری مبتنی بر گری 14 به نام GBKII 15 است [۲۲]. رویکرد مطرح شده در یک فرآیند تکراری بر روی مقادیر جافتاده،

تواند کیفیت جایگذاری را تحت تأثیر قرار دهد. بنابراین در روش‌های پیشنهادی از رویکردی مبتنی بر جنگل تصادفی 16 [۱۶] برای رتبه‌بندی صفات جافتاده جهت بهبود دقت مقادیر تخمین زده شده، استفاده شده است [۱۴].

به طور خلاصه، در روش‌های ارائه شده پس از تعیین اولویت صفات جافتاده برای پُر شدن، دو گام اصلی در رویکردهای پیشنهادی وجود دارد که عبارت‌اند از: گام حداکثرسازی و گام تخمین. در گام حداکثرسازی زیرمجموعه‌هایی با همبستگی بالا نسبت به نمونه جافتاده، پیدا می‌شوند. در گام تخمین، مقادیر جافتاده با اعمال یک یا چند مدل رگرسیون خطی بر زیرمجموعه‌های یافت شده تخمین زده خواهند شد. با ترکیب انواع مختلف گام‌های حداکثرسازی و تخمین، سه رویکرد متفاوت در این مقاله پیشنهاد می‌شود.

بخش‌های بعدی این مقاله بدین شرح است. در بخش دوم مروری بر پیشینه تحقیق مرور می‌شود. در بخش سوم جزئیات رویکرد پیشنهادی ارائه می‌شود. در بخش چهارم جزئیات آزمایش‌های انجام شده و نتایج تحلیل‌ها بیان می‌شوند. در نهایت در بخش پنجم نتیجه‌گیری بیان می‌شود.

۲- پیشینه تحقیق

در این بخش، مروری بر روش‌های مختلف موجود برای هدف جایگذاری مقادیر جافتاده انجام می‌شود.

رویکردهای ساده‌ی متفاوتی برای رسیدگی به مقادیر جافتاده وجود دارد:

- حذف نمونه‌هایی با مقادیر جافتاده
- پُر کردن دستی مقادیر جافتاده
- پُر کردن با مقدار ثابت یا مقادیر محتمل (مانند مقدار میانگین)

در رویکرد اول اگر تعداد مقادیر جافتاده بالا باشد، بسیاری از داده‌های موجود را بهیوده هدر می‌دهد. رویکرد دوم فقط وقتی عملی است که در داده‌ها تعداد نسبتاً کمی مقدار جافتاده وجود دارد. مشکل رویکرد سوم نادیده گرفتن روابط بین صفات است. به علاوه در این رویکرد بحث توزیع داده‌ها نادیده گرفته خواهد شد؛ چراکه تمام مقادیر جافتاده با یک مقدار یکسان پُر خواهند شد. بنابراین این رویکردهای ساده نمی‌توانند دقت کافی برای پُر کردن مقادیر جافتاده را فراهم کنند. به همین دلیل، محققان زیادی به مسئله‌ی جایگذاری مقادیر جافتاده توجه کرده‌اند. این تحقیقات شامل طیف وسیعی از روش‌ها می‌شود؛ از روش‌های ساده‌ی آماری تا روش‌های پیچیده‌ی مبتنی بر داده‌کاوی و یادگیری ماشین.

دسته‌ای از روش‌ها با استفاده از قوانین انجمنی 17 داده‌های جافتاده را پُر می‌کنند. محدودیت اصلی این روش‌ها این است که تنها قابل استفاده برای داده‌های غیر عددی هستند. یکی از روش‌های موجود در این دسته از رویکردها، امتیازدهی به قوانین انجمنی است [۱۷]. ابتدا مجموعه داده‌ی رابطه‌ای به مجموعه داده‌ی تراکنشی تبدیل می‌شود تا بتوان از آن

استفاده از رویکردی ترکیبی مبتنی قوانین انجمنی و k نزدیک‌ترین همسایه، از دیگر روش‌های ارائه‌شده برای پُر کردن مقادیر جافتاده است [۲۹]. از نقاط قوت این روش می‌توان به بهره‌گیری از قدرت ترکیب دو روش متفاوت اشاره کرد. درواقع در این روش سعی در استفاده از ارتباطات و فواصل موجود برای پُر کردن مقادیر جافتاده شده است. این روش قابلیت استفاده برای داده‌های عددی را به‌طور مستقیم ندارد (به دلیل استفاده از قوانین انجمنی).

یک روش ترکیبی برای پُر کردن مقادیر جافتاده، روش DMI^{19} است [۱۵]. این روش مبتنی بر الگوریتم EM^{20} و درخت تصمیم است. ایده‌ی اصلی این روش این است که وابستگی بین افزایندهای مجموعه‌داده قوی تر از وابستگی بر کل مجموعه‌داده است و بنابراین الگوریتم EM بر روی بخش‌های متفاوت مجموعه‌داده بهتر عمل خواهد کرد و درنهایت نتیجه به‌دست‌آمده بهتر خواهد شد. افزاز بندی مجموعه‌داده توسط درخت تصمیم صورت می‌پذیرد. از نقاط قوت این روش توجه به این است که وابستگی بین صفات لزوماً در کل مجموعه‌داده برقرار نیست. از آنجایی که دقت الگوریتم EM هنگامی بالا است که همبستگی بین صفات قوی باشد، این روش برای مجموعه‌داده‌هایی خوب عمل می‌کند که روابط قوی در بین صفات آن‌ها برقرار باشد. به بیان دیگر، عدم وجود همبستگی بالا بین صفات منجر به کاهش دقت تخمین خواهد شد.

رگرسیون نیز به‌عنوان یک روش پایه برای تخمین مقادیر جافتاده می‌تواند مورد استفاده قرار گیرد. زمانی که همبستگی بین صفات یک مجموعه‌داده قوی‌تر باشد، دقت رگرسیون نیز بیشتر خواهد بود. با استفاده از محاسبه‌ی معادله‌ی رگرسیون خطی می‌توان داده‌های جافتاده را از طریق داده‌های موجود تخمین زد. در یکی از روش‌های جایگزینی مقادیر جافتاده به نام روش رگرسیون افزایشی صفات ($IARI$) که اخیراً مطرح شده است، تخمین مقادیر جافتاده با استفاده از یک رویکرد افزایشی و تولید دنباله‌ای از رگرسیون‌ها انجام می‌شود [۱۴]. در این روش پس از رتبه‌بندی صفات، یک تخمین اولیه از مقادیر جافتاده انجام می‌شود. پس از آن صفات جافتاده تک‌به‌تک و به‌صورت افزایشی جایگزینی می‌شوند. عمل جایگزینی در هر مرحله توسط مدل‌های رگرسیون انجام می‌شود. این کار آن‌قدر ادامه می‌یابد تا تمام صفات جافتاده جایگزینی شوند. مشکل اصلی روش‌های مبتنی بر رگرسیون ضعف عملکرد آن‌ها در صورت عدم وجود روابط قوی بین صفات پیش‌بینی کننده و صفت پیش‌بینی شده است.

در سال ۲۰۱۶، دو راهبرد برای بهبود دقت پُر کردن مقادیر جافتاده ارائه شد [۱۱]. ایده‌ی اصلی در هر دو راهبرد، انجام عمل انتخاب نمونه^{۲۱} است. در این عمل سعی می‌شود تا نمونه‌های شامل داده‌های پُر و داده‌های مغشوش^{۲۲} از مجموعه‌داده حذف شوند تا کیفیت داده‌ها پیش از عملیات پُر کردن مقادیر جافتاده افزایش یابد. درواقع مجموعه‌داده از یک فیلتر حذف داده‌های پُر و مغشوش عبور داده می‌شود. در هر دو راهبرد، عملیات انتخاب نمونه، از طریق الگوریتم $DROP^{23}$ [۳۰] انجام شده است. تفاوت دو راهبرد پیشنهادی در ترتیب انجام عملیات

تک‌به‌تک مقادیر جافتاده را پُر می‌نماید. از مزیت‌های این روش امکان پُر کردن مقادیر جافتاده چندگانه و کاراتر بودن پیچیدگی زمانی به دلیل اینکه تنها بین نمونه‌های هم برچسب مقایسه شباهت انجام می‌گردد، است. نیاز به برچسب کلاس داشتن نمونه‌ها، ضعف اصلی این روش است. یک نسخه‌ی بهبودیافته مبتنی روش نزدیک‌ترین همسایگان برای تخمین مقادیر جافتاده، روش k نزدیک‌ترین همسایه و استفاده از خوشه‌بندی داده‌های قرارگرفته در یک دسته ($CKNNI^{16}$) است [۲۰]. مشکل اصلی این روش نیاز به دانستن برچسب کلاس تمام نمونه‌ها است. درواقع اگر کلاس نمونه‌ها مشخص نباشند، نمی‌توان از روش ذکرشده استفاده نمود.

مشکل اصلی روش‌های مبتنی بر نزدیک‌ترین همسایگان این است که برای مجموعه‌داده‌های بزرگ زمان زیادی برای جایگذاری صرف می‌شود، چراکه در این روش‌ها برای هر نمونه‌ی جافتاده، روی کل مجموعه‌داده جستجوی کامل انجام می‌شود. به‌علاوه، تعیین دقیق تابع فاصله و مقدار مناسب برای k ممکن است چالش‌برانگیز باشد. همچنین در این روش‌ها از روابط موجود بین صفات در تخمین صفات استفاده نمی‌شود.

در یک روش مبتنی بر خوشه‌بندی، هر مقدار جافتاده با میانگین دو مقدار جایگزین می‌شود [۲۳]. مقدار اول فاصله‌ی نمونه‌ی جافتاده تا مرکز خوشه‌ای است که نمونه‌ی جافتاده به آن تعلق دارد. مقدار دوم برابر با مقدار صفتی از مرکز خوشه است که متناظر در نمونه‌ی جافتاده، نامعلوم است. چشم‌پوشی از روابط احتمالی بین نمونه‌های موجود در یک مجموعه‌داده مشکل اصلی این روش است. یک نوع خوشه‌بندی برای تخمین مقادیر جافتاده، خوشه‌بندی همسایگان مشترک است [۲۴]. از نقاط قوت این روش می‌توان به استفاده از معیار شباهت برای صفات ترکیبی و دخیل کردن معیار جدیدی برای سنجش شباهت نمونه‌ها اشاره کرد. نحوه توزیع داده‌ها بر تعداد همسایگان مشترک در این روش تأثیرگذار است و ممکن است ماتریس همسایگان مشترک تنک^{۱۷} شود. متداول‌ترین روش‌های جایگزینی مبتنی بر خوشه‌بندی، خوشه‌بندی‌های k -means و c -means فازی هستند [۲۵، ۲۶]. یک روش ترکیبی ارائه‌شده مبتنی بر خوشه‌بندی و شبکه‌های عصبی، شامل دو مرحله است [۲۷]. این رویکرد شامل یک تخمین اولیه از مقادیر جافتاده توسط نزدیک‌ترین مرکز خوشه به روش خوشه‌بندی k -means و یک گام پالایش و بهبود مقدار تخمین زده‌شده از گام قبلی، توسط شبکه‌ی MLP است. یک نسخه‌ی بهبودیافته از روش ذکرشده، استفاده از خوشه‌بندی c -means فازی به جای خوشه‌بندی متداول k -means است [۲۸]. در این روش هر نمونه با یک احتمال متعلق به هر خوشه است. این روش نتایج بهتری را به همراه داشته است. از نقاط قوت این روش‌ها می‌توان به بهبود k -means برای پُر کردن مقادیر جافتاده اشاره کرد. نیاز به آموزش یک شبکه‌ی MLP برای هر مقدار جافتاده و آموزش مجدد روی کل مجموعه آموزشی و تعیین ساختار (به‌هم‌بندی^{۱۸} و وزن‌ها) بهینه‌ی شبکه‌ی MLP برای هر بار آموزش برای پُر کردن هر مقدار جافتاده از نقاط ضعف این روش است.

فرض کنید مجموعه‌داده‌ی رابطه‌ای D شامل n نمونه است و هر نمونه مثل x شامل m صفت عددی $A_1, A_2, \dots, A_m$ است. مقدار صفت x_i^j در نمونه i ام را در D نشان می‌دهد. شکل (۱) رویه‌ی کلی را برای روش‌های ارائه‌شده نمایش می‌دهد.

Input:
 D : An incomplete data set containing n instances, defined on schema $(A_1, A_2, \dots, A_m)$
 p : Percentage of complete records which will be selected as the base set
 θ : Threshold for maximization algorithms

Output:
 D : Imputed data set

```

1  $\Lambda \leftarrow$  Find a set of attributes containing missing values in  $D$ 
2 Calculate priority of each missing attribute for imputation using a Random Forest model
3 while  $\Lambda \neq \emptyset$  do
4    $D_{comp} \leftarrow \text{selectCompleteRecords}(D)$ 
5    $D_{miss} \leftarrow \text{selectMissingRecords}(D)$ 
6    $j \leftarrow$  Index of the most important attribute in  $\Lambda$ 
7   foreach  $x \in D_{miss}$  where  $x^j = \text{NULL}$  do
8      $\mathcal{B} \leftarrow \text{findBaseSet}(D_{comp}, x, p)$ 
9      $D' \leftarrow D_{comp} \cup \mathcal{B}$ 
10     $\text{maximizedDatasets} \leftarrow \text{maximizeCorrelation}(x, j, \mathcal{B}, D', \theta)$ 
11     $x^j \leftarrow \text{estimateByRegression}(x, j, \text{maximizedDatasets})$ 
12  end
13   $\Lambda \leftarrow \text{Delete } A_j \text{ from } \Lambda$ 
14 end
15 return  $D$ 
    
```

شکل ۱: رویه کلی رویکردهای پیشنهادی برای تخمین مقادیر جافتاده

در ابتدا، اولویت و اهمیت هر صفت جافتاده برای پُر شدن، توسط یک مدل جنگل تصادفی محاسبه می‌شود (خطوط ۱ و ۲). پس از آن، مجموعه‌داده‌ی ورودی به دو مجموعه‌داده‌ی کامل (D_{comp}) و مجموعه‌داده‌ی با مقادیر جافتاده (D_{miss}) تقسیم می‌شود (خطوط ۳ و ۴). سپس صفتی با بالاترین اولویت جهت پُر شدن انتخاب می‌شود (خط ۵). آنگاه، نمونه‌هایی که مقدار جافتاده‌ای در صفت j ام (اندیس مهم‌ترین صفت جافتاده) خود دارند، تک به تک برای فرآیند پُر شدن انتخاب می‌شوند (خط ۷). بعد از این کار، برای هر نمونه انتخاب‌شده یک مجموعه پایه از نمونه‌های کامل انتخاب می‌شود. با اعمال یک روش حداکثرسازی، بسته به نوع روش، یک یا چند زیرمجموعه از نمونه‌های کامل با همبستگی بالا ساخته خواهند شد (خطوط ۸ تا ۱۰). حداکثرسازی با اعمال رویکردهایی مبتنی بر انتخاب روبه‌جلو^{۲۴} بر مجموعه‌ی پایه ($\mathcal{B}$) و با کمک نمونه‌های کامل غیرپایه (D') انجام می‌شود. در نهایت، مقدار جافتاده در یک نمونه با اعمال یک یا چند مدل خطی بر روی زیرمجموعه یا زیرمجموعه‌های ساخته‌شده تخمین زده می‌شود. این کار آن قدر ادامه می‌یابد تا تمام صفات جافتاده پُر شوند.

همان‌طور که الگوریتم پیشنهادی نشان می‌دهد، پس از پُر کردن هر صفت جافتاده، مقادیر آن صفت برای تخمین دیگر صفات جافتاده در تکرارهای بعدی استفاده خواهد شد. درواقع، مجموعه‌داده‌های کامل (D_{comp}) پس از هر تکرار حلقه **while** به‌روزرسانی می‌شود. الگوریتم کلی پیشنهادی، دو پارامتر اصلی، یعنی p و θ دارد که نیاز به مقداردهی مناسب دارند. در زیر بخش‌های بعدی این بخش، پارامترها، نحوه‌ی انتخاب مجموعه پایه و گام‌های حداکثرسازی و تخمین و دیگر جزئیات روش‌های پیشنهادی شرح داده خواهند شد. در نهایت ترکیبات مختلف از گام‌های حداکثرسازی و تخمین برای تولید روش‌های پیشنهادی مختلف ذکر خواهند شد.

انتخاب نمونه و پُر کردن داده‌ها است. انجام عملیات انتخاب نمونه منجر به بهبود دقت حاصل از دسته‌بندی داده‌های تخمین زده‌شده است. با بررسی روش‌های موجود، می‌توان پی برد که اصولاً روش‌هایی که روابطی بین نمونه‌ها را مدل می‌کنند و سپس از روابط یافت شده برای تخمین مقادیر جافتاده استفاده می‌کنند، منجر به دقت بالایی در عملیات تخمین می‌شوند. از جمله‌ی این روش‌ها می‌توان به رگرسیون اشاره کرد. اگرچه در این روش‌ها لازم است که روابط قوی بین صفات در مجموعه‌داده برقرار باشد؛ از طرفی ممکن است این روابط به‌طور صریح از مجموعه‌داده اصلی قابل استخراج نباشند. بنابراین در این مقاله سعی می‌شود تا با ارائه‌ی روش‌های مختلف، روابطی که به‌طور ضمنی بین نمونه‌ها وجود دارد را یافته و سپس از آن‌ها برای هدف جایگذاری مقادیر جافتاده استفاده نمود. به‌علاوه، اگرچه تعدادی از روش‌های گذشته، مانند روش‌های مبتنی بر رگرسیون، با فرض قوی بودن روابط بر روی کل مجموعه‌داده به نتایج مناسبی دست یافتند، با این حال، این روش‌ها از معیاری دقیق برای سنجش میزان روابط موجود استفاده نکرده‌اند. بنابراین، لازم است که ابتدا فرض قوی بودن روابط به‌طور دقیق ارزیابی شود و سپس این روش‌ها را اعمال نمود.

۳- رویکرد پیشنهادی

در این بخش و زیر بخش‌های آن مراحل مختلف رویکردهای پیشنهادی برای جایگذاری مقادیر جافتاده شرح داده می‌شوند. همان‌طور که در دو بخش گذشته نیز ذکر شد، دو نکته‌ی مهم در روش‌های پیشنهادی وجود دارد. اول این که ترتیب پُر کردن صفات جافتاده می‌تواند در نتیجه نهایی پُر کردن مقادیر جافتاده تأثیرگذار باشد. دوم این که شناسایی بخش‌هایی از مجموعه‌داده‌ی که روابط خطی قوی بر آن‌ها حاکم است و سپس اعمال مدل‌های خطی داخل این بخش‌ها منجر به تخمین‌های دقیق‌تری از مقادیر جافتاده می‌شود.

به‌طور خلاصه، در روش‌های ارائه‌شده پس از تعیین اولویت صفات جافتاده برای پُر شدن، دو گام اصلی حداکثرسازی و تخمین اجرا می‌شوند. ابتدا در گام حداکثرسازی تلاش می‌شود تا زیرمجموعه‌های با همبستگی بالا نسبت به نمونه جافتاده، پیدا شوند. برای این کار ابتدا یک مجموعه پایه از نمونه‌های کامل انتخاب می‌شود. پس از آن، زیرمجموعه‌هایی با روابط قوی بین صفات، با بسط دادن مجموعه پایه ساخته می‌شوند. در گام تخمین مقادیر جافتاده، با اعمال یک یا چند مدل رگرسیون خطی بر زیرمجموعه‌های یافت شده تخمین زده خواهند شد. رویکردهای پیشنهادی ماهیتی افزایشی دارند؛ یعنی پس از جایگذاری مقادیر جافتاده برای هر صفت جافتاده، از آن صفت به‌عنوان متغیرهای پیش‌بینی کننده برای تخمین‌های دیگر صفات استفاده خواهد شد. ماهیت افزایشی امکان استفاده از اطلاعات موجود در مقادیر جافتاده را فراهم می‌کند.

$$d(x_i, x_j) = \sqrt{\sum_{k=1}^m [(x_i^k - x_j^k)^2 I(k)]}, \quad (1)$$

$$I(k) = \begin{cases} 1, & \text{if } x_i^k \text{ and } x_j^k \neq \text{NULL} \\ 0, & \text{otherwise} \end{cases}$$

همان‌طور که پیداست، در محاسبه‌ی فاصله‌ی دو نمونه، فقط صفات با مقادیر معلوم در محاسبات دخیل می‌شوند.

۳-۳- حداکثر کردن همبستگی بین صفات با بسط دادن مجموعه پایه

پس از انتخاب یک مجموعه پایه مناسب، الگوریتم تلاش می‌کند تا رابطه‌ی خطی بین صفت جافتاده و دیگر صفات نمونه جافتاده را در مجموعه پایه افزایش دهد. این افزایش با استفاده از بسط دادن مجموعه پایه توسط نمونه‌های غیرپایه به نحوی انجام می‌گیرد که همبستگی بین صفت جافتاده و دیگر صفات حداکثر شود. این کار با به‌کارگیری یک رویکرد انتخاب روبه‌جلو انجام می‌شود. ایده‌ی اصلی برای انجام چنین کاری این است که در یک مدل رگرسیون خطی، وجود روابط خطی قوی‌تر بین متغیر (های) پیش‌بینی‌کننده^{۲۶} و متغیر پاسخ^{۲۷} منجر به تخمین بهتر متغیر پاسخ می‌شود.

رابطه خطی بین دو متغیر تصادفی را می‌توان با استفاده از ضریب همبستگی پیرسون^{۲۸} (PCC) اندازه‌گیری کرد. مقدار PCC میزان همبستگی دو متغیر تصادفی را نشان می‌دهد. PCC می‌تواند مقداری بین -۱ و +۱ داشته باشد. علامت ضریب همبستگی پیرسون، جهت رابطه‌ی خطی را نشان می‌دهد. با محاسبه مقدار PCC برای هر جفت از m صفات یک مجموعه داده مثل D ، یک ماتریس مقارن همبستگی با اندازه‌ی $m \times m$ حاصل می‌شود. مدخل i, j در این ماتریس، مقدار PCC بین صفات A_i و A_j را در D نشان می‌دهد. در ادامه مقاله، نماد $\rho(D)$ ماتریس همبستگی برای مجموعه داده‌ی D را نشان می‌دهد. نماد $\rho(D)_{i,j}$ اشاره به مدخل ردیف i ام و ستون j ام در ماتریس همبستگی برای مجموعه داده‌ی D خواهد داشت. معادلات (۲) و (۳) نحوه محاسبه ماتریس همبستگی و همبستگی بین دو متغیر تصادفی A_i و A_j را نشان می‌دهد.

$$\rho(D) = \begin{bmatrix} 1 & \rho_{1,2} & \cdots & \rho_{1,m} \\ \rho_{2,1} & 1 & \cdots & \rho_{2,m} \\ \vdots & \vdots & \ddots & \vdots \\ \rho_{m,1} & \rho_{m,2} & \cdots & 1 \end{bmatrix}_{m \times m} \quad (2)$$

که در آن

$$\rho(D)_{i,j} = \frac{Cov(A_i, A_j)}{\sigma_{A_i} \sigma_{A_j}} \quad (3)$$

که در آن، $Cov(A_i, A_j)$ مقدار کوواریانس^{۲۹} بین مقادیر دو صفت A_i و A_j را نشان می‌دهد. σ_{A_i} و σ_{A_j} به ترتیب بیانگر مقادیر انحراف از معیار برای صفات A_i و A_j هستند.

۳-۱- محاسبه اولویت هر صفت جافتاده برای هدف تخمین داده‌های جافتاده

صفات جافتاده در روش‌های پیشنهادی تک‌به‌تک به ترتیب اهمیتشان طبق رویکردی مبتنی بر جنگل تصادفی پر خواهند شد. صفتی با بیش‌ترین اهمیت، ابتدا و صفتی با کم‌ترین اهمیت در انتها پر خواهد شد. اولویت هر صفت جافتاده برای پُر شدن در سه گام محاسبه می‌شود. ابتدا مقادیر جافتاده در مجموعه داده‌ی ورودی (شامل مقادیر جافتاده) با استفاده از یک رویکرد تخمین ساده، مثل تخمین با مقدار میانگین، پر می‌شوند. سپس، یک مدل جنگل تصادفی بر روی مجموعه داده‌ی پرشده ساخته می‌شود. درنهایت، مدل ساخته‌شده به نمونه‌های خارج از کیسه^{۳۰} اعمال می‌شود تا اهمیت هر صفت اندازه‌گیری شود. پس از پیدا کردن ترتیب هر صفت برای پُر شدن، صفتی با بیش‌ترین اهمیت برای ادامه فرآیند پُر کردن انتخاب خواهد شد. بنابراین با فرض این‌که z اندیس مهم‌ترین صفت جافتاده باشد، آنگاه نمونه‌هایی که در صفت z ام خود مقدار جافتاده دارند، تک‌به‌تک برای ادامه در مراحل بعدی فرآیند جایگذاری انتخاب خواهند شد.

۳-۲- انتخاب مجموعه پایه

تا این مرحله، ترتیب تخمین هریک از صفات جافتاده تعیین شد. از این مرحله به بعد، نحوه پُر کردن مقادیر جافتاده برای هر نمونه‌ی جافتاده شرح داده خواهد شد. در رویکردهای ارائه‌شده، هر صفت جافتاده از یک نمونه‌ی جافتاده با اعمال یک یا چند مدل رگرسیون خطی بر زیرمجموعه (هایی) از نمونه‌های کامل، تخمین زده خواهند شد. ویژگی این زیرمجموعه‌های انتخاب‌شده این است که همبستگی قوی بین صفت جافتاده و دیگر صفات وجود دارد.

در گام‌های بعدی، روش‌های حداکثرسازی همبستگی از رویکردهای مبتنی بر انتخاب روبه‌جلو برای پیدا کردن زیرمجموعه‌های مطلوب استفاده خواهند کرد. از این‌رو، لازم است که یک مجموعه پایه مناسب از نمونه‌ها انتخاب شود. این مجموعه‌ی پایه باید دید کلی خوبی از روابط موجود در نمونه‌هایی که در رویکرد انتخاب روبه‌جلو استفاده خواهند شد، بدهد. برای این منظور از روش معروف k نزدیک‌ترین همسایگان برای تشکیل مجموعه پایه برای هر نمونه‌ی جافتاده استفاده می‌شود. بدین منظور p درصد از نزدیک‌ترین نمونه‌های کامل به نمونه‌ی جافتاده، به‌عنوان k نزدیک‌ترین همسایگان برای آن نمونه جافتاده انتخاب می‌شوند. این همسایگان به‌عنوان مجموعه پایه برای آن نمونه جافتاده انتخاب می‌شوند. بنابراین خروجی این گام برای هر نمونه‌ی جافتاده یک زیرمجموعه از نمونه‌های کامل خواهد بود.

فاصله‌ی بین دو نمونه‌ی x_i و x_j با استفاده از رابطه‌ی فاصله‌ی اقلیدسی مطابق با معادله (۱) محاسبه می‌شود.

صفات A_i و A_j در D_i افزایش یافته است یا خیر (خط ۷). اگر مقدار PCC پس از اضافه کردن نمونه‌ی جدید بیش از مقدار حد آستانه، یعنی θ ، افزایش یافته بود، آن نمونه به‌طور دائم به D_i اضافه خواهد شد (خط ۸). در غیر این صورت از آن نمونه چشم‌پوشی می‌شود و الگوریتم به سراغ نمونه‌ی بعدی می‌رود. این کار آن‌قدر ادامه می‌یابد تا تمام نمونه‌های غیرپایه بررسی شوند.

Input:
 $\mathcal{B}$: A base set of complete records for maximization
 D' : A set of records
 j : Index of attribute which the output data sets will be maximized according to it, i.e., index of missing value ($1 \leq j \leq m$)
 θ : Threshold value for adding a record from D' to $\mathcal{B}$ ($0 \leq \theta \leq 1$)

Output:
 $D_1, D_2, \dots, D_m$: Set of data sets such that correlation between A_i and A_j (missing attribute) is maximized in data set D_i

```

1 for  $i \leftarrow 1$  to  $m$  do
2    $D_i \leftarrow \mathcal{B}$ 
3 end
4 foreach  $x \in D'$  do
5   for  $i \leftarrow 1$  to  $m$  do
6     if  $i \neq j$  then
7       if  $(\rho(D_i \cup x)_{i,j} - \rho(D_i)_{i,j}) > \theta$  then
8          $D_i \leftarrow (D_i \cup x)$ 
9       end
10    end
11  end
12 end
13 return  $D_1, D_2, \dots, D_m$ 

```

شکل ۲: الگوریتم حداکثرسازی همبستگی هریک از صفات نسبت به صفات جافتاده

۳-۳-۲- حداکثر کردن میانگین مقادیر همبستگی بین صفات معلوم و صفت جافتاده

الگوریتم دوم پیشنهادی سعی می‌کند تا میانگین مقادیر همبستگی بین صفات غیر جافتاده و صفت جافتاده را افزایش دهد. ورودی‌های این الگوریتم همانند الگوریتم قبلی هستند، اما خروجی تولیدشده متفاوت خواهد بود. خروجی یک تک مجموعه‌داده‌ی (D_{max}) خواهد بود که در آن میانگین مقادیر همبستگی بین صفات غیر جافتاده و صفت جافتاده حداکثر شده است. مراحل این روش در شکل (۳) نمایش داده شده‌اند.

Input:
 $\mathcal{B}$: A base set of complete records for maximization
 D' : A set of records
 j : Index of attribute which the output data set will be maximized according to it, i.e., index of missing value ($1 \leq j \leq m$)
 θ : Threshold value for adding a record from D' to $\mathcal{B}$ ($0 \leq \theta \leq 1$)

Output:
 D_{max} : A data set in which mean of correlation values between non-missing attributes and A_j (missing attribute) is maximized.

```

1  $D_{max} \leftarrow \mathcal{B}$ 
2 foreach  $x \in D'$  do
3    $\mu_{current} \leftarrow \frac{\sum_{i=1}^m [\rho(D_{max})_{i,j}]}{m-1}$ 
4    $\mu_{new} \leftarrow \frac{\sum_{i=1}^m [\rho(D_{max} \cup x)_{i,j}]}{m-1}$ 
5   if  $(\mu_{new} - \mu_{current}) > \theta$  then
6      $D_{max} \leftarrow (D_{max} \cup x)$ 
7   end
8 end
9 return  $D_{max}$ 

```

شکل ۳: الگوریتم حداکثرسازی میانگین مقادیر همبستگی بین صفات غیر جافتاده نسبت به صفت جافتاده

در ابتدا، مقدار اولیه D_{max} برابر با مجموعه پایه قرار می‌گیرد (خط ۱). سپس الگوریتم به‌طور موقت نمونه‌های غیرپایه (D') را تک‌تک به D_{max} اضافه می‌کند. پس از هر اضافه کردن، الگوریتم بررسی می‌کند

از آنجایی که میزان همبستگی بین صفات (قوی یا ضعیف بودن) در روش‌های ارائه‌شده مهم است، نه جهت همبستگی (همبستگی مثبت یا منفی)، از قدر مطلق مقدار ضریب همبستگی در محاسبات ماتریس همبستگی استفاده می‌شود. بنابراین برای هر مقدار i و j خواهیم داشت:

$$0 \leq \rho(D)_{i,j} \leq 1$$

با داشتن یک نمونه جافتاده مثل x که در آن مقدار x^j جافتاده است (صفت j ام آن نمونه مقدار تهی دارد)، الگوریتم‌های پیشنهادی تلاش خواهند کرد تا همبستگی بین صفت جافتاده و دیگر صفات را با اعمال یک روش انتخاب روبه‌جلو بر روی مجموعه پایه افزایش دهند. یک رویکرد کلی برای انتخاب یک ترکیب k تایی از D که مقدار ضریب همبستگی در آن زیرمجموعه بین صفت جافتاده و دیگر صفات حداکثر است، انجام جستجوی کامل^{۲۰} است؛ یعنی مقدار ضریب همبستگی را برای تمام زیرمجموعه‌های k تایی از D محاسبه کرده و زیرمجموعه‌های با بیشترین مقدار PCC را انتخاب کنیم. با فرض این که اندازه‌ی D برابر با n باشد، برای این کار $n!/k!(n-k)!$ انتخاب مختلف وجود خواهد داشت. بنابراین این رویکرد، به‌ویژه زمانی که اندازه‌ی D بسیار بزرگ است، بسیار زمان‌بر است. بنابراین نحوه‌ی انتخاب زیرمجموعه‌های مطلوب باید هوشمندانه‌تر باشد. پیشنهاد ما استفاده از رویکردی مبتنی بر انتخاب روبه‌جلو برای یافتن زیرمجموعه‌های مطلوب است. یعنی مجموعه‌ی پایه را با استفاده از رویکرد انتخاب روبه‌جلو بسط دهیم تا زیرمجموعه‌ی مطلوب حاصل شود. بدیهی است، این رویکرد یافتن بهترین زیرمجموعه‌ی ممکن را ضمانت نمی‌کند؛ با این حال از رویکرد جستجوی کامل، منطقی‌تر به نظر می‌رسد. دو نوع الگوریتم حداکثرسازی همبستگی در این مقاله ارائه شده است. در زیر بخش‌های بعدی این الگوریتم‌ها شرح داده خواهند شد.

۳-۳-۱- حداکثر کردن همبستگی هر یک از صفات معلوم نسبت به صفت جافتاده

الگوریتم اول حداکثرسازی سعی می‌کند تا مجموعه پایه را با در نظر گرفتن هر صفت غیر جافتاده (معلوم) به‌طور جداگانه بسط دهد. ورودی‌های این الگوریتم عبارت‌اند از: اندیس یک صفت مثل j که صفات معلوم در مجموعه‌داده‌های خروجی نسبت به آن صفت حداکثر می‌شوند، یک مجموعه از نمونه‌های پایه، یک مجموعه دیگر از نمونه‌ها که نامزدهای اضافه شدن به مجموعه پایه هستند و یک پارامتر حد آستانه (θ) که توسط کاربر تعریف می‌شود. خروجی الگوریتم m مجموعه‌داده‌ی $D_1, D_2, \dots, D_m$ خواهد بود، به‌نحوی که همبستگی بین صفت A_j و صفت جافتاده‌ی A_i در D_i حداکثر شده است. جزئیات این الگوریتم در شکل (۲) نمایش داده شده است.

در ابتدا، مقادیر تمام مجموعه داده‌های D_i ($i=1, 2, \dots, m$) برابر با مجموعه پایه قرار داده می‌شوند (خطوط ۱ تا ۳). سپس الگوریتم به‌طور موقت نمونه‌های غیرپایه (D') را تک‌تک به هر D_i اضافه می‌کند. پس از هر اضافه کردن، الگوریتم بررسی می‌کند که آیا مقدار همبستگی بین

ستون اول، شناسه‌ی روش، ستون دوم، نحوه‌ی حداکثرسازی همبستگی و ستون سوم نحوه‌ی تخمین نهایی مقدار جافتاده را نمایش می‌دهد.

جدول ۱: ترکیب‌های مختلف گام‌های حداکثرسازی و تخمین

شناسه	روش حداکثرسازی	روش تخمین
۱	حداکثر کردن همبستگی هریک از صفات	تک مدل خطی چند صفته
۲	حداکثر کردن همبستگی هریک از صفات	نظرخواهی وزن‌دار از چندین مدل خطی تک صفته
۳	حداکثر کردن میانگین مقادیر همبستگی تمام صفات	تک مدل خطی چند صفته

روش‌های ۱ و ۲ از الگوریتم ارائه‌شده در شکل (۲) برای گام حداکثرسازی استفاده می‌کنند. روش ۳ از الگوریتم ارائه‌شده در شکل (۳) برای گام حداکثرسازی استفاده می‌نماید.

در روش ۱ پس از تولید شدن خروجی گام حداکثرسازی ($D_1, D_2, \dots, D_m$)، مجموعه‌داده‌ی D_i که در آن i^* از رابطه‌ی (۶) به دست می‌آید، انتخاب‌شده و سپس برای گام تخمین، به تک مدل خطی چند صفته (رویکرد دوم تخمین) داده می‌شود.

$$i^* = \underset{i}{\operatorname{argmax}} (\rho(D_i)_{i,j}), \quad i^* \neq j \quad (6)$$

که در آن j اندیس صفت جافتاده است.

خروجی گام حداکثرسازی در روش ۲ برای گام تخمین به‌طور مستقیم به تک مدل خطی چند صفته داده می‌شود. خروجی گام حداکثرسازی در روش ۳، یعنی $D_{\max}$ ، برای گام تخمین، به m مدل خطی تک صفته داده می‌شود و نتیجه‌ی نهایی از طریق روش میانگین وزن‌دار (نظرخواهی وزن‌دار) حاصل می‌شود (رویکرد اول تخمین).

۴- تحلیل پیچیدگی زمانی رویکردهای پیشنهادی

در این بخش، پیچیدگی زمانی رویکردهای پیشنهادی تحلیل می‌شود. فرض کنید مجموعه‌داده ورودی شامل n نمونه است. از این n نمونه p تای آن‌ها کامل و q تای آن‌ها شامل مقادیر جافتاده است ($p+q=n$). همچنین فرض کنید هر نمونه شامل m ویژگی است. روش‌های مطرح‌شده شامل سه گام اصلی هستند. در گام اول، با استفاده از یک مدل جنگل تصادفی، صفات رتبه‌بندی می‌شوند. مرتبه زمانی این گام در بدترین حالت برابر با $O(T(mn \log n))$ است؛ که در آن T بیانگر تعداد درخت‌ها در مدل جنگل تصادفی است. در گام دوم (حداکثرسازی)، ابتدا زمانی برابر با $O(m.p)$ صرف عملیات انتخاب مجموعه پایه می‌شود. سپس، زمانی برابر با $O(m.p.p \log p)$ صرف اجرای الگوریتم‌های حداکثرسازی (الگوریتم‌های شکل ۲ و ۳) خواهد شد. در این محاسبات $O(p \log p)$ ، زمان محاسبه‌ی مقدار PCC است. سومین و آخرین گام، گام تخمین مقادیر جافتاده است. در این مرحله روش رگرسیون خطی با پیچیدگی زمانی $O(m^2.p)$ اجرا می‌شود. با در نظر گرفتن موارد گفته‌شده، در بدترین حالت، پیچیدگی الگوریتم‌های ارائه‌شده از مرتبه‌ی

که آیا میانگین مقدار همبستگی بین صفات با مقادیر معلوم و صفت A_j افزایش‌یافته است یا خیر (خط ۷). اگر مقدار PCC پس از اضافه کردن نمونه‌ی جدید بیش از مقدار حد آستانه یعنی θ افزایش‌یافته بود، آن نمونه به‌طور دائم به $D_{\max}$ اضافه خواهد شد (خط ۸). در غیر این صورت از آن نمونه چشم‌پوشی می‌شود و الگوریتم به سراغ نمونه‌ی بعدی می‌رود. این کار آن‌قدر ادامه می‌یابد تا تمام نمونه‌های غیرپایه بررسی شوند.

۳-۴ تخمین مقادیر جافتاده با استفاده از زیرمجموعه‌های یافته شده

پس از پیدا کردن زیرمجموعه‌های مطلوب از داده‌های کامل، آخرین گام، تخمین صفت جافتاده در نمونه‌ی جافتاده است. بسته به نوع روش حداکثرسازی همبستگی در گام قبلی، رویکرد تخمین مقدار جافتاده متفاوت خواهد بود. در رویکرد اول، تعداد m مدل رگرسیون خطی ساده (تک صفتی) بر روی m مجموعه‌داده‌ی ساخته‌شده اعمال خواهد شد؛ به‌نحوی که در مجموعه‌داده‌ی D_i صفت A_i به‌عنوان متغیر مستقل و A_j به‌عنوان متغیر وابسته در ساخت مدل رگرسیون استفاده خواهد شد. هر مدل خطی رگرسیون در این حالت مطابق با معادله‌ی (۴) خواهد بود.

$$A_j = \alpha_0 + \alpha A_i \quad (4)$$

که در آن A_j نشان‌دهنده متغیر وابسته، A_i بیانگر متغیر مستقل و α_0 و α ضرایب مدل رگرسیون ساخته‌شده هستند. رویکرد پیشنهادی، هر مقدار جافتاده را با محاسبه میانگین وزن‌دار m مقدار تخمین زده‌شده برای آن صفت جافتاده جایگزین خواهد کرد. وزن مقدار تخمین زده‌شده‌ی i ام برابر با این است که تا چه میزان دو صفت A_i و A_j توانسته‌اند در D_i به یکدیگر به‌صورت خطی مرتبط شوند. یعنی وزن مقدار تخمین زده‌شده‌ی i ام برابر با $\rho(D_i)_{i,j}$ خواهد بود.

در دومین رویکرد تخمین، فقط یک مدل رگرسیون خطی چند صفتی بر روی خروجی گام حداکثرسازی اعمال خواهد شد. از صفت جافتاده به‌عنوان متغیر وابسته و از باقی صفات (غیر جافتاده) به‌عنوان متغیرهای مستقل برای ساخت مدل خطی استفاده خواهد شد. معادله‌ی یک مدل رگرسیون خطی چند صفتی مطابق با (۵) خواهد بود.

$$A_j = \alpha_0 + \alpha_1 A_1 + \alpha_2 A_2 + \dots + \alpha_m A_m = \alpha_0 + \sum_{i=1}^m \alpha_i A_i \quad (5)$$

که در آن A_j نشان‌دهنده متغیر وابسته، A_i ($i=1, 2, \dots, m$) بیانگر متغیرهای مستقل و α_i ($i=0, 1, 2, \dots, m$) ضرایب مدل رگرسیون ساخته‌شده هستند.

۳-۵ ترکیب گام‌های حداکثرسازی و تخمین

با ترکیب روش‌های مختلف حداکثرسازی و تخمین ارائه‌شده، روش‌های متفاوتی را برای تخمین مقادیر جافتاده می‌توان ارائه داد. جدول (۱) حالات مختلف ممکن برای تولید روش‌های مختلف را نشان می‌دهد.

روش‌های مختلف جایگذاری پُر شدند. از آنجایی که مقادیر اصلی برای مقادیر جافتاده‌ی ساخته‌شده معلوم است، بنابراین عملکرد رویکردهای مختلف با دانستن مقادیر واقعی قابل ارزیابی است.

برای این که امکان مقایسه نتایج بین رویکردهای مختلف وجود داشته باشد و مقایسه‌ها بامعنا باشند، مقادیر هر صفت از مجموعه داده‌ها با استفاده از نرمال‌سازی به روش نمره‌ی z^{19} نرمال شدند. این کار با محاسبه و جایگذاری هر مقدار x_i^j در معادله‌ی (۷) انجام می‌شود.

$$(normalized) x_i^j = \frac{x_i^j - \bar{A}_j}{\sigma_{A_j}} \quad (7)$$

که در آن $\bar{A}_j$ و σ_{A_j} به ترتیب بیانگر مقادیر میانگین و انحراف معیار برای صفت A_j هستند.

به‌منظور ایجاد مقادیر جافتاده در مجموعه داده‌ها بعضی صفات انتخاب می‌شوند و بعضی از مقادیر این صفات به مقدار تهی تغییر می‌یابند. برای این منظور ابتدا نیمی از صفات به‌طور تصادفی انتخاب می‌شوند. سپس در ۳۰٪، ۴۰٪، ۵۰٪ و ۶۰٪ از نمونه‌ها، مقادیر جافتاده به‌طور تصادفی ایجاد می‌شود. بنابراین هر نمونه‌ی جافتاده ممکن است به تعداد ۱ تا $m/2$ صفت جافتاده داشته باشد. همه‌ی مقادیر جافتاده مطابق با الگوی جافتاده به‌طور تصادفی (MAR) ایجاد شده‌اند. در تمامی روش‌ها، از مدل رگرسیون خطی ساخته‌شده از طریق برازش z^{22} بهترین خط بر روی داده‌ها به کمک روش حداقل مربعات z^{23} استفاده شده است. از دو معیار شناخته‌شده به نام‌های ریشه میانگین مربعات خطی z^{24} (RMSE) و ضریب تعیین z^{25} (CoD) یا R^2 برای ارزیابی کیفیت روش‌های مختلف پُر کردن مقادیر جافتاده در این مقاله استفاده شده است. شاخص RMSE نشان‌دهنده‌ی میانگین اختلافات مقادیر واقعی از مقادیر تخمین زده‌شده توسط مدل‌های جایگذاری مقادیر جافتاده است. ضریب تعیین در تحلیل مدل‌های آماری برای ارزیابی این که یک مدل چقدر خوب خروجی‌های آینده را پیش‌بینی می‌کند، استفاده می‌شود. ضریب تعیین مقداری بین ۰ و ۱ دارد. مقادیر کمتر برای میانگین خطای مطلق و مقادیر بیش‌تر برای ضریب تعیین، بیانگر کیفیت بالاتر از تخمین مقادیر است.

معادلات (۸) و (۹) به ترتیب نحوه‌ی محاسبه‌ی معیارهای ریشه میانگین مربعات خطی و ضریب تعیین را نشان می‌دهند.

$$RMSE = \sqrt{\frac{1}{n} \sum_{k=1}^n (y_k - y_k')^2} \quad (8)$$

$$CoD = 1 - \frac{\sum_{k=1}^n (y_k' - y_k)^2}{\sum_{k=1}^n (y_k - \bar{y})^2} \quad (9)$$

که در آن‌ها y_k نشان‌دهنده مقدار واقعی، y_k' مقدار پیش‌بینی‌شده، $\bar{y}$ میانگین مقادیر مشاهده‌شده و n برابر با تعداد مقادیر تخمین زده‌شده است.

$O(m.p^2 \cdot \log p + m^2 \cdot p)$ خواهد بود. با فرض این که $p \gg m$ ، آنگاه پیچیدگی زمانی الگوریتم‌های پیشنهادی از مرتبه‌ی $O(m.p^2 \cdot \log p)$ خواهد بود. اگر فرض کنیم که در بدترین حالت $p \approx n$ ، آنگاه پیچیدگی الگوریتم‌های پیشنهادی از مرتبه‌ی $O(m.n^2 \log n)$ خواهند بود. در ادامه، پیچیدگی زمانی برای دیگر روش‌های مقایسه شده ذکر می‌شوند. پیچیدگی زمانی رویکرد جایگذاری با روش k نزدیک‌ترین همسایگان [۳۱] از مرتبه‌ی $O(nk + nm)$ است؛ که در آن k برابر با تعداد نزدیک‌ترین همسایگان است. پیچیدگی زمانی رویکرد جایگذاری با روش FCM [۲۶] از مرتبه‌ی $O(nmc^2)$ است؛ که در آن c برابر با تعداد خوشه‌ها در الگوریتم خوشه‌بندی c-means فازی است. برای روش‌های DMI [۱۵] و IARI [۱۴]، پیچیدگی زمانی از مرتبه‌ی $O(mn^2)$ است. عملیات پیش‌پردازش تنها یک‌بار و پیش از عملیات اصلی داده‌کاوی انجام می‌شود. بنابراین پیچیدگی زمانی عملیات پیش‌پردازش در اولویت پایین‌تری نسبت به کیفیت این امر قرار می‌گیرد. باین حال برای توصیف دقیق‌تر، پیچیدگی زمانی رویکردهای پیشنهادی در این بخش تحلیل شد.

۵- آزمایش‌ها و نتایج

در این قسمت، روش‌های مورد مقایسه با روش‌های پیشنهادی، مجموعه داده‌های استفاده‌شده برای آزمایش‌ها و معیارهای ارزیابی کارایی رویکردهای متفاوت بررسی می‌شوند. علاوه بر این، آزمایش‌هایی برای بررسی تأثیر پارامترهای p و θ بر روی نتایج تخمین نهایی انجام می‌شود. عملکرد رویکردهای پیشنهادی با پنج روش تخمین با استفاده از مقدار میانگین، تخمین با استفاده از k نزدیک‌ترین همسایگان [۳۱]، روش پُر کردن با خوشه‌بندی c-means فازی (FCM) [۲۶]، روش پُر کردن با درخت تصمیم (DMI) [۱۵] و جایگذاری با استفاده از رگرسیون افزایشی صفات (IARI) [۱۴] مقایسه شده است.

مهم‌ترین عامل تصادفی که ممکن است در کیفیت الگوریتم‌های جایگذاری مقادیر جافتاده تأثیرگذار باشد، تعداد مقادیر جافتاده در هریک از نمونه‌های جافتاده است. به‌منظور حذف عامل تصادفی بودن در آزمایش‌ها، هر آزمایش ۱۰ بار تکرار شده است. میانگین این ۱۰ بار تکرار به‌عنوان نتیجه نهایی در این پژوهش گزارش شده است. کدهای مورد استفاده برای رویکردهای مختلف همگی با زبان برنامه‌نویسی پایتون و با کمک بعضی از بسته‌های یادگیری ماشین از مخزن Scikit-Learn [۳۲] پیاده‌سازی شده‌اند.

۵-۱- مشخصات آزمایش‌ها

روش‌های مختلف جایگذاری مقادیر جافتاده بر روی پنج مجموعه داده از دنیای واقعی که در مخزن داده‌ی یادگیری ماشین UCI به‌طور عمومی در دسترس هستند، اعمال شدند. جدول (۲) مشخصات این مجموعه داده‌ها را نمایش می‌دهد. این مجموعه داده‌ها بدون مقادیر جافتاده می‌باشند. در این مجموعه داده‌های کامل، ابتدا مقادیر جافتاده به‌طور مصنوعی ایجاد شدند. سپس مقادیر جافتاده با استفاده از

۵-۲- مقایسه عملکرد رویکردهای ارائه‌شده با دیگر روش‌ها

در این بخش آزمایش‌هایی به منظور مقایسه عملکرد رویکردهای پیشنهادی با دیگر روش‌های موجود انجام می‌شود. عملکرد روش‌های ارائه‌شده با پنج روش تخمین با استفاده از مقدار میانگین، تخمین با استفاده از k نزدیک‌ترین همسایگان (kNN) [۳۱]، روش پُر کردن با خوشه‌بندی c-means فازی (FCM) [۲۶]، روش پُر کردن با درخت تصمیم (DMI) [۱۵] و جایگذاری با استفاده از رگرسیون افزایشی صفات (IARI) [۱۴] مقایسه شده‌اند. جداول (۳) تا (۷) عملکرد هریک از روش‌های جایگذاری مقادیر جافتاده مبتنی بر همبستگی پیشنهادی را در مقایسه با پنج روش دیگر با توجه به معیارهای RMSE و CoD برای درصد‌های مختلف جافتادگی در مجموعه‌داده‌های مختلف نشان می‌دهد. مواردی که رویکرد پیشنهادی بهتر از تمام دیگر روش‌ها عمل کرده است، به صورت کج و پررنگ نوشته شده‌اند. از جداول ذکر شده پیداست که با افزایش درصد جافتادگی ایجادشده در مجموعه‌داده‌ها، خطای کلی تخمین در همه‌ی روش‌ها افزایش می‌یابد. به عبارت دیگر، تعداد بیش‌تر داده‌های کامل منجر به افزایش دقت در تخمین‌ها می‌شود. این پدیده منطقی است، چراکه با افزایش درصد جافتادگی اطلاعات مفید بیش‌تری از دست می‌رود.

در همه‌ی مجموعه‌داده‌ها به غیر از مجموعه‌داده Haberman، روش IARI تقریباً در همه‌ی موارد از چهار روش DMI، روش میانگین، روش FCM و روش kNN بهتر عمل کرده است. از طرفی در اکثر موارد، عملکرد کلی روش‌های ارائه‌شده برای هر دو معیار مقایسه شده بهتر از IARI بوده است. بنابراین می‌توان گفت که به طور کلی روش‌های ارائه‌شده بهتر از دیگر روش‌های مقایسه شده عمل می‌کنند. قابل ذکر است که روش IARI به طور پیش فرض از برچسب دسته‌ی نمونه‌ها به عنوان متغیر پیش‌بینی کننده استفاده می‌کند.

در مجموعه‌داده Haberman اگرچه روش‌های FCM، DMI، IARI و kNN ضعیف‌تر از روش میانگین عمل کرده‌اند، با این حال مشاهده می‌شود که روش‌های پیشنهادی کمی بهتر از روش میانگین عمل کرده‌اند. ضعف نسبی روش‌های دیگر نسبت به روش میانگین احتمالاً به دلیل تعداد کم متغیرهای پیش‌بینی کننده در این مجموعه‌داده است. برای مجموعه‌داده‌های Iris، Wholesale Customers و Wine، تمام روش‌های پیشنهادی در تمام حالات از دیگر روش‌ها با توجه به دو شاخص ارزیابی بهتر عمل کرده‌اند. برای مجموعه‌داده Glass در شرایطی که ۳۰٪ یا ۶۰٪ جافتادگی در نمونه‌ها وجود دارد، در همه‌ی حالات، روش‌های پیشنهادی بهتر از روش IARI عمل کرده‌اند. روش‌های ۱، ۲ و ۳ هنگامی که ۴۰٪ و فقط روش ۲ هنگامی که ۵۰٪ جافتادگی وجود دارد، ضعیف‌تر از روش IARI عمل کرده‌اند. با این حال، در این حالات، روش‌های پیشنهادی بهتر از چهار روش DMI، FCM، kNN و روش میانگین عمل کرده‌اند.

به طور کلی مهم‌ترین عاملی که در کیفیت هریک از روش‌های پیشنهادی تأثیر می‌گذارد، نوع و میزان همبستگی موجود بین صفات در یک مجموعه‌داده است. برای روش‌های ۱ و ۲ که از روش حداکثر کردن همبستگی هریک از صفات استفاده می‌کنند، هنگامی که چند صفت با همبستگی قوی در یک مجموعه‌داده حضور داشته باشد، این روش‌ها بهتر از روش ۳ عمل می‌کنند. از طرف دیگر، اگر در یک مجموعه‌داده میانگین همبستگی بین صفات بالا باشد (حتی اگر چند صفت بسیار همبسته هم وجود نداشته باشند)، روش ۳ که از روش حداکثر کردن میانگین مقادیر همبستگی تمام صفات استفاده می‌کند، بهتر از روش‌های ۱ و ۲ عمل خواهد کرد. برای مثال در مجموعه‌داده‌های Glass، Wine و Haberman که می‌توان چند صفت با همبستگی بالا را یافت، در اکثر موارد دقت روش‌های ۱ و ۲ بیش‌تر از روش ۳ است. از طرف دیگر چون در مجموعه‌داده Wholesale Customers، چندین صفت با همبستگی بالا یافت نمی‌شود، دقت روش ۳ از روش‌های ۱ و ۲ بیش‌تر است. برای مجموعه‌داده Iris، چون همبستگی بین صفات هم به صورت تک‌به‌تک و هم به طور میانگین بالاست، گام تخمین نیز در عملکرد روش‌ها تأثیرگذار است. نکته‌ی مهم دیگر قابل مشاهده این است که به طور کلی کیفیت جایگذاری دو روش ۱ و ۲ بسیار نزدیک به هم است. یعنی گام حداکثرسازی در دقت جایگذاری تأثیر بیشتری نسبت به گام تخمین دارد.

همان‌طور که جداول نشان می‌دهند، در همه‌ی مجموعه‌داده‌ها به غیر از حالت ۴۰٪ جافتادگی در مجموعه‌داده Glass، حداقل یکی از روش‌های پیشنهادی از باقی روش‌ها بهتر عمل کرده است. برای مجموعه‌داده Haberman، به دلیل کم بودن تعداد صفات پیش‌بینی کننده در این مجموعه‌داده (سه صفت)، عملکرد روش‌های پیشنهادی با اختلاف اندکی بهتر از دیگر روش‌ها بوده است. برای مجموعه‌داده Glass، اگرچه در ۴۰٪ جافتادگی، روش‌های ارائه‌شده ضعیف‌تر از بهترین روش (IARI) عمل کرده‌اند، با این حال برای درصد‌های بالای جافتادگی (۵۰٪ و ۶۰٪) روش‌های پیشنهادی بهتر از بهترین روش عمل کرده‌اند. بنابراین به نظر می‌رسد که عامل تصادفی بودن باعث شده است که فقط در حالت ۴۰٪ جافتادگی از مجموعه‌داده Glass، روش‌های پیشنهادی ضعیف‌تر از روش IARI عمل کنند؛ چراکه در باقی حالات و مجموعه‌داده‌ها، روش‌های پیشنهادی بهتر از باقی روش‌ها عمل کرده‌اند. در مجموع روش‌های پیشنهادی بر روی مجموعه‌داده‌هایی با دو ویژگی، بهتر عمل خواهند کرد. اول این است که تعداد صفات یک مجموعه‌داده بسیار کم نباشد؛ چراکه اصولاً دقت روش‌های رگرسیون با افزایش تعداد صفات پیش‌بینی کننده افزایش می‌یابد. دوم این که اگر همبستگی بین صفات در یک مجموعه‌داده به صورت تک‌به‌تک و یا به صورت میانگین بالا باشد، دقت روش‌های پیشنهادی نیز بالاتر خواهد بود.

جدول ۳: مقایسه عملکرد رویکردهای ارائه‌شده با پنج روش دیگر برای مجموعه‌داده‌ی Wholesale Customers

RMSE						CoD						شناسه روش	درصد جالفتادگی
روش DMI	روش FCM	روش IARI	روش kNN	روش میانگین	روش پیشنهادی	روش DMI	روش FCM	روش IARI	روش kNN	روش میانگین	روش پیشنهادی		
۰/۱۷۱۰	۰/۲۰۹۸	۰/۱۶۱۰	۰/۱۶۶۲	۰/۲۰۰۳	۰/۱۴۳۷	۰/۹۶۴۷	۰/۹۴۷۱	۰/۹۷۱۸	۰/۹۶۹۴	۰/۹۵۶۱	۰/۹۷۷۴	۱	۳۰٪
					۰/۱۴۴۳						۰/۹۷۷۰	۲	
					۰/۱۳۶۸						۰/۹۷۹۵	۳	
۰/۲۰۱۱	۰/۲۵۹۲	۰/۱۹۱۸	۰/۱۹۷۹	۰/۲۳۵۳	۰/۱۷۰۷	۰/۹۵۴۷	۰/۹۲۵۷	۰/۹۶۰۴	۰/۹۵۷۵	۰/۹۴۰۸	۰/۹۶۸۳	۱	۴۰٪
					۰/۱۷۳۳						۰/۹۶۷۲	۲	
					۰/۱۶۰۳						۰/۹۷۲۳	۳	
۰/۲۴۸۰	۰/۲۸۰۷	۰/۲۲۵۸	۰/۲۲۸۲	۰/۲۶۴۱	۰/۱۹۸۶	۰/۹۳۷۵	۰/۹۱۵۷	۰/۹۴۵۱	۰/۹۴۴۵	۰/۹۲۶۳	۰/۹۵۸۲	۱	۵۰٪
					۰/۱۹۸۱						۰/۹۵۸۱	۲	
					۰/۱۹۱۸						۰/۹۶۱۰	۳	
۰/۲۶۴۵	۰/۳۲۱۱	۰/۲۵۲۸	۰/۲۵۴۵	۰/۲۹۹۳	۰/۲۲۴۰	۰/۹۲۴۹	۰/۸۹۲۳	۰/۹۳۱۸	۰/۹۳۱۵	۰/۹۰۵۷	۰/۹۴۷۲	۱	۶۰٪
					۰/۲۲۳۴						۰/۹۴۷۹	۲	
					۰/۲۱۵۷						۰/۹۵۱۴	۳	

جدول ۴: مقایسه عملکرد رویکردهای ارائه‌شده با پنج روش دیگر برای مجموعه‌داده‌ی Wine

RMSE						CoD						شناسه روش	درصد جالفتادگی
روش DMI	روش FCM	روش IARI	روش kNN	روش میانگین	روش پیشنهادی	روش DMI	روش FCM	روش IARI	روش kNN	روش میانگین	روش پیشنهادی		
۰/۱۲۷۸	۰/۱۴۰۴	۰/۱۲۴۲	۰/۱۲۷۹	۰/۱۶۴۵	۰/۱۱۰۱	۰/۹۸۵۷	۰/۹۸۰۱	۰/۹۸۵۲	۰/۹۸۳۱	۰/۹۷۲۶	۰/۹۸۷۴	۱	۳۰٪
					۰/۱۰۰۴						۰/۹۸۹۷	۲	
					۰/۱۱۱۱						۰/۹۸۷۲	۳	
۰/۱۵۱۷	۰/۱۶۳۲	۰/۱۴۴۶	۰/۱۴۶۸	۰/۱۸۶۶	۰/۱۲۶۶	۰/۹۷۸۶	۰/۹۷۳۰	۰/۹۸۱۵	۰/۹۷۷۹	۰/۹۶۴۸	۰/۹۸۳۶	۱	۴۰٪
					۰/۱۱۶۹						۰/۹۸۶۰	۲	
					۰/۱۲۷۲						۰/۹۸۳۳	۳	
۰/۱۷۳۴	۰/۱۸۴۰	۰/۱۷۲۳	۰/۱۶۴۴	۰/۲۱۸۰	۰/۱۴۵۶	۰/۹۷۱۷	۰/۹۶۵۸	۰/۹۷۴۴	۰/۹۷۲۲	۰/۹۵۲۱	۰/۹۷۸۰	۱	۵۰٪
					۰/۱۳۸۴						۰/۹۸۰۴	۲	
					۰/۱۴۴۷						۰/۹۷۸۵	۳	
۰/۱۸۹۳	۰/۲۰۶۴	۰/۱۸۸۸	۰/۱۷۹۶	۰/۲۴۰۶	۰/۱۶۱۴	۰/۹۶۵۹	۰/۹۵۷۲	۰/۹۶۸۵	۰/۹۶۶۹	۰/۹۴۱۷	۰/۹۷۳۷	۱	۶۰٪
					۰/۱۵۲۹						۰/۹۷۶۰	۲	
					۰/۱۵۵۸						۰/۹۷۵۷	۳	

جدول ۵: مقایسه عملکرد رویکردهای ارائه‌شده با پنج روش دیگر برای مجموعه‌داده‌ی Iris

RMSE						CoD						شناسه روش	درصد جالفتادگی
روش DMI	روش FCM	روش IARI	روش kNN	روش میانگین	روش پیشنهادی	روش DMI	روش FCM	روش IARI	روش kNN	روش میانگین	روش پیشنهادی		
۰/۱۴۴۰	۰/۱۹۱۲	۰/۱۴۲۴	۰/۱۶۲۱	۰/۲۹۹۴	۰/۱۲۸۸	۰/۹۷۶۲	۰/۹۶۳۰	۰/۹۷۷۷	۰/۹۷۱۳	۰/۹۰۹۴	۰/۹۸۱۷	۱	۳۰٪
					۰/۱۳۶۰						۰/۹۷۹۷	۲	
					۰/۱۲۹۳						۰/۹۸۱۶	۳	
۰/۱۶۵۴	۰/۲۱۹۸	۰/۱۶۳۸	۰/۱۸۵۴	۰/۳۴۵۷	۰/۱۴۹۷	۰/۹۶۵۴	۰/۹۵۱۴	۰/۹۷۰۵	۰/۹۶۲۹	۰/۸۷۹۳	۰/۹۷۵۴	۱	۴۰٪
					۰/۱۵۷۶						۰/۹۷۲۸	۲	
					۰/۱۴۸۰						۰/۹۷۵۹	۳	
۰/۱۸۷۲	۰/۲۵۵۴	۰/۱۸۲۱	۰/۲۰۹۸	۰/۳۸۰۱	۰/۱۶۸۸	۰/۹۶۱۷	۰/۹۳۴۵	۰/۹۶۴۰	۰/۹۵۳۱	۰/۸۵۴۲	۰/۹۶۹۱	۱	۵۰٪
					۰/۱۷۴۷						۰/۹۶۷۳	۲	
					۰/۱۶۵۵						۰/۹۷۰۳	۳	
۰/۲۰۳۴	۰/۲۹۲۲	۰/۲۰۲۹	۰/۲۳۵۵	۰/۴۱۱۹	۰/۱۹۱۵	۰/۹۵۰۰	۰/۹۱۴۴	۰/۹۵۴۳	۰/۹۴۰۱	۰/۸۲۸۶	۰/۹۶۰۱	۱	۶۰٪
					۰/۱۹۵۶						۰/۹۵۸۴	۲	
					۰/۱۸۸۷						۰/۹۶۱۰	۳	

جدول ۶: مقایسه عملکرد رویکردهای ارائه‌شده با پنج روش دیگر برای مجموعه‌داده‌ی Glass

RMSE						CoD						شناسه روش	درصد جافتادگی
روش DMI	روش FCM	روش IARI	روش kNN	روش میانگین	روش پیشنهادی	روش DMI	روش FCM	روش IARI	روش kNN	روش میانگین	روش پیشنهادی		
۰/۱۴۷۳	۰/۱۹۲۵	۰/۱۲۴۹	۰/۱۳۶۳	۰/۲۰۷۲	۰/۱۰۸۵	۰/۹۷۴۶	۰/۹۵۷۸	۰/۹۸۳۰	۰/۹۷۹۶	۰/۹۵۳۷	۰/۹۸۶۷	۱	۳۰٪
					۰/۱۳۳۰						۰/۹۸۳۱	۲	
					۰/۱۰۳۳						۰/۹۸۸۳	۳	
۰/۱۵۶۱	۰/۲۲۸۴	۰/۱۴۱۵	۰/۱۸۱۸	۰/۲۳۴۳	۰/۱۴۷۳	۰/۹۵۷۴	۰/۹۴۱۹	۰/۹۷۸۴	۰/۹۶۱۳	۰/۹۴۱۸	۰/۹۷۳۹	۱	۴۰٪
					۰/۱۵۷۸						۰/۹۷۱۸	۲	
					۰/۱۴۷۳						۰/۹۷۴۵	۳	
۰/۱۷۲۸	۰/۲۴۶۸	۰/۱۵۹۴	۰/۲۰۵۵	۰/۲۶۳۰	۰/۱۵۸۸	۰/۹۵۰۱	۰/۹۳۲۱	۰/۹۷۲۶	۰/۹۵۳۳	۰/۹۲۷۶	۰/۹۷۳۰	۱	۵۰٪
					۰/۱۷۱۹						۰/۹۶۸۵	۲	
					۰/۱۵۸۶						۰/۹۷۳۳	۳	
۰/۱۹۶۴	۰/۲۶۹۰	۰/۱۷۵۶	۰/۲۱۹۵	۰/۲۸۵۱	۰/۱۴۵۰	۰/۹۴۳۶	۰/۹۲۰۰	۰/۹۶۷۶	۰/۹۴۷۸	۰/۹۱۵۵	۰/۹۷۵۹	۱	۶۰٪
					۰/۱۵۵۱						۰/۹۶۳۰	۲	
					۰/۱۴۶۸						۰/۹۷۵۳	۳	

به‌طور کلی، نتایج نشان داده‌اند که مقادیر خیلی پایین و خیلی بالا برای p مناسب نیستند. مقادیر بسیار بالا، مفهوم انتخاب روبه‌جلو را بی‌معنا می‌کند. در مقابل، مقادیر بسیار پایین دید خوبی از روابط موجود در مجموعه پایه نخواهد داد.

برای بررسی تأثیر پارامتر θ ، مقدار θ برای یک مقدار ثابت p بین ۰/۰ تا ۱/۰ با فواصل ۰/۱ تغییر یافته و مقدار RMSE برای هر مقدار θ ثبت شد. شکل (۷) مقدار RMSE را برای مقادیر مختلف θ بر روی مجموعه‌داده‌ی Iris با ۳۰٪ جافتادگی نشان می‌دهد. مجدداً به دلیل کمبود فضا، نتایج برای باقی مجموعه‌داده‌ها و دیگر معیارها نشان داده نشده‌اند. نمودارها ثابت ماندن مقدار RMSE برای مقادیر بالای θ (وقتی θ به ۱ میل می‌کند) را نشان می‌دهند. این روند منطقی است چراکه در این شرایط الگوریتم‌ها به‌طور سخت‌گیرانه اجازت اضافه شدن رکوردهای غیرپایه را به مجموعه پایه خواهند داد. مقدار کمینه سراسری برای روش‌های ۱ و ۳ در حدود آستانه نزدیک به صفر رخ داده‌اند. برای روش ۲، مقدار بهینه‌ی RMSE در $\theta \approx 1$ ظاهر شده است.

۵-۳- تأثیر پارامترهای p و θ

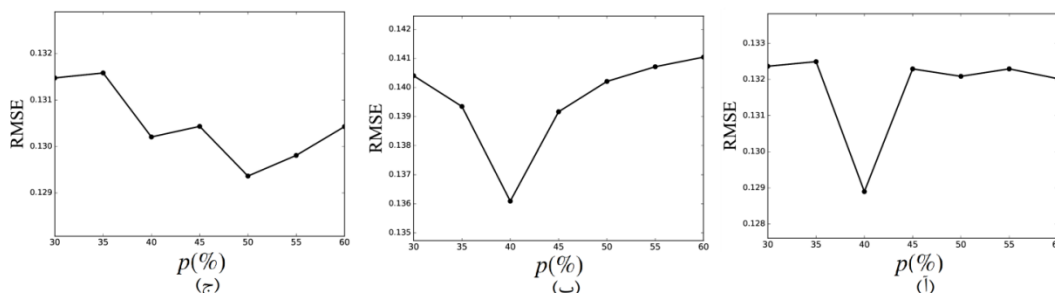
در این بخش تأثیر پارامترهای ورودی بر روی کیفیت عملکرد رویکردهای ارائه‌شده بررسی می‌شود. همان‌طور که در بخش‌های پیش توضیح داده شد، در روش‌های ارائه‌شده، دو پارامتر اصلی p و θ نیاز به تنظیم دارند. پارامتر θ تعیین می‌کند که آیا یک رکورد به مجموعه‌ی پایه اضافه شود یا خیر. پارامتر p درصد انتخاب از نزدیک‌ترین همسایگان کامل نسبت به هر نمونه‌ی جافتاده را (به‌عنوان مجموعه‌ی پایه) تعیین می‌کند.

برای بررسی تأثیر پارامتر p ، مقدار p برای مقدار ثابت θ بین ۳۰ تا ۶۰ با فواصل ۵ تغییر یافته و مقدار RMSE برای هر مقدار p ثبت شد. شکل (۶) مقدار RMSE را برای مقادیر مختلف p بر روی مجموعه‌داده‌ی Iris با ۳۰٪ جافتادگی نشان می‌دهد. نتایج برای باقی مجموعه‌داده‌ها و دیگر معیارها به دلیل کمبود فضا نشان داده نشده‌اند. برای روش ۱، عملکرد بهینه در $p=40$ حاصل شده است. برای روش‌های ۲ و ۳ مقدار RMSE پس از رسیدن به مقدار بهینه روندی صعودی داشته است.

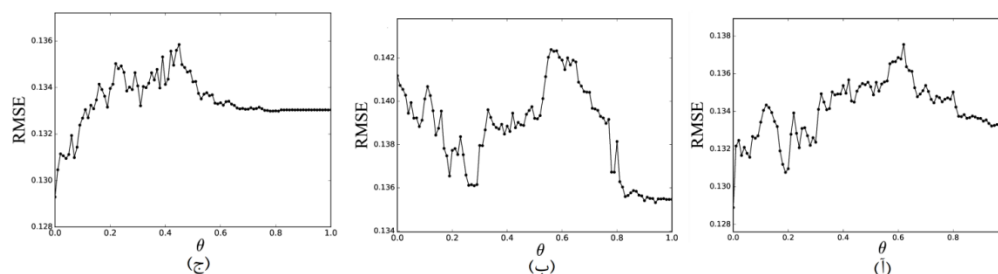
جدول ۷: مقایسه عملکرد رویکردهای ارائه‌شده با پنج روش دیگر برای مجموعه‌داده‌ی Haberman

RMSE						CoD						شناسه روش	درصد جافتادگی
روش DMI	روش FCM	روش IARI	روش kNN	روش میانگین	روش پیشنهادی	روش DMI	روش FCM	روش IARI	روش kNN	روش میانگین	روش پیشنهادی		
۰/۳۷۳۷	۰/۳۴۱۳	۰/۳۶۳۴	۰/۴۱۱۶	۰/۳۳۵۹	۰/۳۲۶۹	۰/۸۵۲۵	۰/۸۸۲۲	۰/۸۶۸۲	۰/۸۳۱۴	۰/۸۸۴۲	۰/۸۹۳۵	۱	۳۰٪
					۰/۳۲۳۳						۰/۸۹۵۹	۲	
					۰/۳۲۶۹						۰/۸۹۳۵	۳	
۰/۴۲۶۷	۰/۴۰۴۴	۰/۴۲۳۷	۰/۴۶۲۲	۰/۳۹۵۹	۰/۳۸۰۱	۰/۸۰۳۴	۰/۸۳۵۷	۰/۸۲۰۸	۰/۷۸۷۹	۰/۸۴۹۷	۰/۸۵۶۴	۱	۴۰٪
					۰/۳۷۴۰						۰/۸۶۱۱	۲	
					۰/۳۸۰۱						۰/۸۵۶۴	۳	
۰/۴۸۵۲	۰/۴۴۱۶	۰/۴۷۶۶	۰/۵۱۶۴	۰/۴۴۷۳	۰/۴۴۱۲	۰/۷۶۳۷	۰/۸۰۴۳	۰/۷۷۳۰	۰/۷۳۶۱	۰/۷۹۸۶	۰/۸۰۶۲	۱	۵۰٪
					۰/۴۳۰۶						۰/۸۱۵۵	۲	

					۰/۴۴۱۵						۰/۱۰۵۹	۳	
					۰/۴۹۳۳						۰/۷۵۸۳	۱	
					۰/۴۷۷۴	۰/۷۱۵۷	۰/۷۶۸۱	۰/۷۳۱۱	۰/۶۶۹۳	۰/۷۵۱۸	۰/۷۷۳۹	۲	۶۰٪
					۰/۴۹۳۳						۰/۷۵۸۳	۳	
۰/۵۲۱۳	۰/۴۸۰۷	۰/۵۱۹۴	۰/۵۷۷۷	۰/۴۹۹۲									



شکل ۶: تأثیر پارامتر p بر مقدار RMSE در مجموعه داده Iris برای (آ) روش ۱، (ب) روش ۲ و (ج) روش ۳



شکل ۷: تأثیر پارامتر θ بر مقدار RMSE در مجموعه داده Iris برای (آ) روش ۱، (ب) روش ۲ و (ج) روش ۳

۶- نتیجه‌گیری

در این مقاله سه روش جدید برای جایگذاری داده‌های جافتاده ارائه شد. ویژگی اصلی و مشترک این روش‌ها تلاش برای حداکثر کردن همبستگی خطی بین صفت جافتاده و دیگر صفات معلوم و سپس اعمال مدل رگرسیون است. نتایج بررسی‌شده بر روی پنج مجموعه داده‌ی مختلف نشان داد که رویکردهای پیشنهادی در اکثر موارد در مقایسه با پنج روش مختلف دیگر، خطای کمتری در تخمین دارند. تأثیر پارامترهای رویکرد پیشنهادی در میزان خطا نیز بررسی شدند. از مشکلات اصلی رویکرد پیشنهادی می‌توان به چالش انتخاب بهینه برای پارامترهای رویکرد پیشنهادی اشاره کرد. اگرچه این چالش نیز با عملیات زمان‌بر جستجوی کامل برای انتخاب بهترین ترکیب پارامترها قابل برطرف شدن است. بررسی تأثیر دیگر معیارهای سنجش همبستگی، مانند معیار اطلاعات متقابل^{۳۶}، در کیفیت پُر کردن مقادیر جافتاده و اعمال دیگر مدل‌های تخمین برای پژوهش‌های آینده در دست بررسی است.

مراجع

- [۲] مرتضی خرم کشکولی و مریم دهقانی، «تشخیص، شناسایی و جداسازی عیب توربین گاز پالایشگاه دوم پارس جنوبی با استفاده از روش‌های ترکیبی داده‌کاوی، k-means، تحلیل مؤلفه‌های اصلی (PCA) و ماشین بردار پشتیبان (SVM)»، مجله علمی پژوهشی مهندسی برق دانشگاه تبریز، جلد ۴۷، شماره ۲، صفحات ۵۰۱-۵۱۵، ۱۳۹۶.
- [۳] علیرضا سردار و رمضان هاوونگی، «بهبود عملکرد الگوریتم خوشه‌یابی خودکار تصاویر رنگی به کمک پیش‌پردازش با شبکه عصبی خودسازمانده (SOM)»، مجله علمی پژوهشی مهندسی برق دانشگاه تبریز، جلد ۴۷، شماره ۳، صفحات ۱۰۸۲-۱۰۷۳، ۱۳۹۶.
- [4] G. Sun, J. Shao, H. Han, and X. Ding, "Missing value imputation for wireless sensory soil data: A comparative study," in 2nd International Conference on Big Data Computing and Communications, pp. 172-184, Springer, Shenyang, China, 2016.
- [5] M. Lichman, UCI Machine Learning Repository, Available online at: <http://archive.ics.uci.edu/ml>, Accessed June 2017.
- [6] P. J. Garcia-Laencina, J. L. Sancho-Gomez, and A. R. Figueiras-Vidal, "Pattern classification with missing data: a review," Neural Computing and Applications, vol. 19, no. 2, pp. 263-282, 2010.
- [7] E. L. Silva-Ramrez, R. Pino-Mejas, and M. Lpez-Coello, "Single imputation with multilayer perceptron and multiple imputation combining multi-layer perceptron and k-nearest neighbours for monotone patterns," Applied Soft Computing, vol. 29, no. 1, pp. 65-74, 2015.
- [8] M. G. Rahman and M. Z. Islam, "Missing value imputation using a fuzzy clustering-based EM approach," Knowledge and Information Systems, vol. 46, no. 2, pp. 389-422, 2016.
- [9] M. Amiri and R. Jensen, "Missing data imputation using fuzzy-rough methods," Neurocomputing, vol. 205, no. 1, pp. 152-164, 2016.

- [1] Y. Qin, S. Zhang, X. Zhu, J. Zhang, and C. Zhang, "Pop algorithm: Kernel-based imputation to treat missing values in knowledge discovery from databases," Expert Systems with Applications, vol. 36, no. 2, pp. 2794-2804, 2009.

- Innovations in Engineering (ICRAIE), pp. 1-5, IEEE, Jaipur, India, 2014.
- [22] C. Zhang, X. Zhu, J. Zhang, Y. Qin, and S. Zhang, "GBKII: An imputation method for missing values," in *Advances in Knowledge Discovery and Data Mining: 11th Pacific-Asia Conference*, pp. 1080-1087, Springer, Nanjing, China, 2007.
- [23] B. M. Patil, R. C. Joshi, and D. Toshniwal, "Missing value imputation based on k-mean clustering with weighted distance," in *3rd International Conference on Contemporary Computing*, pp. 600-609, Springer, Noida, India, 2010.
- [24] V. Ayuyev, J. Jupin, P. Harris, and Z. Obradovic, "Dynamic clustering-based estimation of missing values in mixed type data," in *11th International Conference on Data Warehousing and Knowledge Discovery*, pp. 366-37, Springer, Linz, Austria, 2009.
- [25] D. Li, J. Deogun, W. Spaulding, and B. Shuart, "Towards missing data imputation: A study of fuzzy k-means clustering Method," *Rough Sets and Current Trends in Computing*, vol. 3066, no. 1, pp. 573-579, Springer, 2004.
- [26] P. Raja and K. Thangavel, "Soft clustering based missing value imputation," in *Digital Connectivity-Social Impact: 51st Annual Convention of the Computer Society of India*, pp. 119-133, Springer, Coimbatore, India, 2016.
- [27] N. Ankaiah and V. Ravi, "A novel soft computing hybrid for data imputation," in *7th International Conference on Data Mining (DMIN)*, Las Vegas, USA, 2011.
- [28] S. Azim, S. Aggarwal, "Hybrid model for data imputation: using fuzzy c-means and multi-layer perceptron," in *Advance Computing Conference (IACC)*, 2014 IEEE International, pp. 1281-1285, Gurgaon, India, 2014.
- [29] S. Bashir, S. Razzaq, U. Maqbool, S. Tahir, and A. R. Baig, "Using association rules for better treatment of missing values," in *10th WSEAS International Conference on Computers*, Wisconsin, USA, pp. 1133-1138, 2009.
- [30] D. R. Wilson and T. R. Martinez, "Reduction techniques for instance-based learning algorithms," *Machine learning*, vol. 38, no. 3, pp. 257-286, 2000.
- [31] Batista, G. E., & Monard, M. C. (2002). A study of k-nearest neighbour as an imputation method. *Hybrid Intell Syst (HIS)*, vol. 87, no. 1, pp. 251-260, 2002.
- [32] F. Pedregosa, G. Varoquaux, A. Gramfort, V. Michel, B. Thirion, O. Grisel, M. Blondel, P. Prettenhofer, R. Weiss, V. Dubourg, J. Vanderplas, A. Passos, D. Coumapeau, M. Brucher, M. Perrot, and E. Duchesnay, "Scikit-learn: Machine learning in Python," *Journal of Machine Learning Research*, vol. 12, no. 1, pp. 2825-2830, 2011.
- [10] H. Wang and S. Wang, "Mining incomplete survey data through classification," *Knowledge and information systems*, vol. 24, no. 2, pp. 221-233, 2010.
- [11] C.F. Tsai and F.Y. Chang, "Combining instance selection for better missing value imputation," *Journal of Systems and Software*, vol. 122, no. 1, pp. 63- 71, 2016.
- [12] C. T. Tran, M. Zhang, P. Andraee, and B. Xue, "Improving performance for classification with incomplete data using wrapper-based feature selection," *Evolutionary Intelligence*, vol. 9, no. 3, pp. 81-94, 2016.
- [13] M. G. Rahman and M. Z. Islam, "Data quality improvement by imputation of missing values," in *5th International Conference on Computer Science and Information Technology (CSIT-2013)*, pp. 82- 88, Yogyakarta, Indonesia, 2013.
- [14] B. van Stein and W. Kowalczyk, "An incremental algorithm for repairing training sets with missing values," in *International Conference on Information Processing and Management of Uncertainty in Knowledge-Based Systems*, vol. 611, no. 1, pp. 175-186. Springer International Publishing, Eindhoven, Netherlands, 2016.
- [15] G. Rahman and Z. Islam, "A decision tree-based missing value imputation technique for data pre-processing," in *Proceedings of the Ninth Australasian Data Mining Conference*, vol. 121, no. 1, pp. 41-50. Australian Computer Society, Inc., 2011.
- [16] L. Breiman, "Random Forests," *Machine learning*, vol. 45, no. 1, pp. 5-32, 2001.
- [17] C.H. Wu, C.H. Wun, and H.J. Chou, "Using association rules for completing missing data," in *4th International Conference on Hybrid Intelligent Systems*, pp. 236-241, Kitakyushu, Japan, IEEE, 2004.
- [18] N. Singh, A. Javeed, S. Chhabra, and P. Kumar, "Missing value imputation with unsupervised kohonen self organizing map," *Emerging Research in Computing, Information, Communication and Applications*, vol. 1, no. 1, pp. 61-76. Springer, New Delhi, India, 2015.
- [19] J. Han, M. Kamber, and J. Pei, *Data Mining: Concepts and Techniques*. San Francisco, CA, USA: Morgan Kaufmann Publishers Inc., 3rd Edition, 2011.
- [20] C. Jiang and Z. Yang, "CKNNI: an improved knn-based missing value handling technique," in *11th International Conference on Intelligent Computing*, pp. 441-452, Springer, Fuzhou, China, 2015.
- [21] R. Krishnamoorthy, S. S. Kumar, and B. Neelagund, "A new approach for data cleaning process," in *Recent Advances and*

زیر نویس ها

- ¹⁹ Decision tree based Missing value Imputation technique
- ²⁰ Expectation Maximization
- ²¹ Instance Selection
- ²² Noise Data
- ²³ Decremental Reduction Optimization Procedure
- ²⁴ Forward Selection
- ²⁵ Out of bag
- ²⁶ Predictor
- ²⁷ Response
- ²⁸ Pearson Correlation Coefficient
- ²⁹ Covariance
- ³⁰ Exhaustive Search
- ³¹ z-score
- ³² Fitting
- ³³ Least Square Method
- ³⁴ Root Mean Squared Error
- ³⁵ Coefficient of Determination
- ³⁶ Mutual Information

- ¹ Null
- ² Imputation
- ³ Correlation
- ⁴ Redundant
- ⁵ Regression
- ⁶ Incremental Attribute Regression Imputation
- ⁷ Decision tree based Missing value Imputation technique
- ⁸ Random Forest
- ⁹ Association Rules
- ¹⁰ Multi-Layer Perceptron
- ¹¹ Self Organizing Map
- ¹² Interpretability
- ¹³ k Nearest Neighbors
- ¹⁴ Grey
- ¹⁵ Grey-Based K-NN Iteration Imputation
- ¹⁶ Class-based K-clusters Nearest Neighbor Imputation
- ¹⁷ Sparse
- ¹⁸ Topology