

تصحیح خودکار داده‌ها مبتنی بر وابستگی تابعی و سیستم یادگیری مرکب

مهدیه عطایان^۱، کارشناس ارشد؛ نگین دانشپور^۲، استادیار

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

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

چکیده: صحت داده‌ها یکی از مهم‌ترین ابعاد کیفیت داده‌ها به‌شمار می‌رود. با توجه به حجم بالای منابع داده‌ای نیاز به روش‌هایی خودکار وجود دارد. در این مقاله راهکاری خودکار برای تصحیح داده‌هایی با انواع داده‌ای متفاوت ارائه شده است. در این راهکار در ابتدا رکوردهایی که احتمالاً حاوی ویژگی خطا است با استفاده از وابستگی تابعی شناسایی می‌گردد، بدین‌صورت که رکوردی که به ازای یک وابستگی تابعی با بیش از δ از رکوردها در تناقض باشد، مشکوک به خطا است. سپس به ازای هر ویژگی از منبع داده مورد بررسی، سیستم یادگیری مرکب ساخته می‌شود. سیستم یادگیری مرکب از سه طبقه‌بند بیز، درخت تصمیم و شبکه عصبی MLP تشکیل شده است و دارای استراتژی ترکیب رأی اکثریت است. سیستم یادگیری مرکب به‌وسیله رکوردهای صحیح شناسایی شده مورد آموزش قرار داده می‌شود. پس از آموزش طبقه‌بندها، هر ویژگی غلط به‌عنوان کلاس هدف سیستم یادگیری مرکب قرار می‌گیرد و مقداری برای آن پیش‌بینی می‌گردد. روش پیشنهادی قادر است چندین خطا در یک رکورد را شناسایی نماید. آزمایش‌ها نشان می‌دهد که true negative rate الگوریتم پیشنهادی در بخش تشخیص خطا به‌طور متوسط ۹۳/۷٪ و در بخش تصحیح خطا به‌طور متوسط ۹۰/۶٪ است. هم‌چنین آزمایش‌ها نشان می‌دهد که میزان پارامترهای ارزیابی در الگوریتم پیشنهادی در مقایسه با دو الگوریتم مشابه مبتنی بر وابستگی تابعی بهبود داشته است.

واژه‌های کلیدی: تصحیح داده، تشخیص خطا، وابستگی تابعی، سیستم یادگیری مرکب.

Automatic Data Cleaning using Functional Dependency and Ensemble Learning

M. Ataeyan¹, M.Sc; N. Daneshpour², Assistant Professor

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

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

Abstract: Data accuracy is one of the important aspects of data quality. According to high volume of data sources an automatic method is required. In this article an automatic method is proposed for cleaning of data with various data types. Initially, records that may contain incorrect attributes are detected using functional dependency, so that each record that inconsistent more than δ records for one functional dependency, probably is error. Thereafter ensemble learning is done for each attribute of data source. Ensemble learning contains 3 classifier naïve bayes, decision tree and MLP, and voting is combination strategy. It is trained using correct records that identified in the initial steps. After training, each incorrect attribute is placed as target and predict values for it. Proposed method is able to clean data in data sources with different data types. Experiments show the true negative rate in detecting error part of the proposed algorithm is averagely 93.7% and in cleaning error part is averagely 90.6%. Also experiments show that evaluation parameters are improved in proposed method compared with 2 similar methods.

Keywords: Data cleaning, error detection, functional dependency, ensemble learning.

تاریخ ارسال مقاله: ۱۳۹۵/۰۹/۱۰

تاریخ اصلاح مقاله: ۱۳۹۶/۰۱/۲۰

تاریخ پذیرش مقاله: ۱۳۹۶/۰۵/۰۹

نام نویسنده مسئول: نگین دانشپور

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

۱- مقدمه

امروزه سازمان‌ها جهت اتخاذ تصمیمات خود نیاز به جمع‌آوری داده، ذخیره‌سازی و تجزیه و تحلیل آن دارند. داده جمع‌آوری شده اغلب شامل داده‌های ناصحیح و مقادیر ازدست‌رفته^۱ است [۱]. به همین دلیل تصمیم‌گیری مبتنی بر این گونه داده‌ها موجب هزینه‌های بالا و درنهایت شکست آن‌ها می‌گردد. بررسی‌ها نشان داده که ۳۰٪-۸۰٪ زمان توسعه پروژه‌های پایگاه داده تحلیلی^۲ صرف تصحیح داده‌ها می‌گردد [۲]. علاوه بر این گزارش‌ها منتشر شده نشان می‌دهد سازمان‌های آمریکایی سالانه چندین تریلیون دلار به دلیل داده‌های ناصحیح متحمل هزینه می‌گردند [۲]. به همین علت، علی‌رغم زمان و هزینه بالا جهت تصحیح داده‌ها، فرایند تصحیح ضروری و غیرقابل اغماض است [۳].

به‌طور کلی فرایند پیش‌پردازش شامل انتساب مقدار به مقادیر ازدست‌رفته، شناسایی داده‌های پرت^۳، برطرف نمودن داده‌های نویزی و اصلاح داده‌های ناسازگار^۴ است [۴]. این مشکلات ممکن است به دلایلی از قبیل معیوب بودن تجهیزات جمع‌آوری اطلاعات، اشتباه کاربر و یا رایانه در هنگام وارد نمودن اطلاعات و یا هنگام انتقال داده‌ها ایجاد شوند. علاوه بر این به دلیل فرمت‌های ناهمگون و نام‌گذاری‌های ناسازگار، در طی فرآیند جمع‌آوری داده از منابع توزیع شده نیز ممکن است با این مشکلات و یا مشکل رکوردهای تکراری^۵ مواجه شوند. درواقع فرایند پیش‌پردازش نقش اساسی در حصول اطمینان از کیفیت داده‌ها^۶ دارد [۴].

تصحیح داده‌ها^۷ اغلب پردازشی دو مرحله‌ای است، که در مرحله اول داده‌های خطا دار شناسایی و در مرحله دوم مقداری صحیح به داده خطا دار انتساب داده می‌شود [۵]. تاکنون روش‌های متعددی برای تصحیح داده‌ها ارائه شده که در این راهکارها از وابستگی تابعی شرطی^۸ [۲]، طبقه‌بندها [۱، ۴، ۶، ۷]، قوانین^۹ [۸، ۹]، هستان‌شناسی^{۱۰} [۱۰]، وابستگی‌های تابعی^{۱۱} [۱۲، ۱۱] و قوانین فازی [۱۳] استفاده شده است. قابل ذکر است که از این راهکارها در حوزه‌های دیگر از جمله صنایع نفت و گاز [۱۴]، پردازش تصویر [۱۵] و علوم پزشکی [۱۶] استفاده می‌شود. این روش‌ها از دیدگاه‌هایی از جمله خودکار یا نیمه خودکار^{۱۲} بودن، نوع داده‌ای^{۱۳} مورد تصحیح، دارا بودن فاز تشخیص خطا و تفکیک پذیری فاز تشخیص و تصحیح خطا با یکدیگر تمایز دارند. قابل ذکر است که با توجه به این که تصحیح داده‌ها جز فرایند پیش‌پردازش است که زمان در آن فاقد اهمیت است، از این رو هیچ‌یک از راهکارهایی که تاکنون ارائه شده‌اند، از بعد پیچیدگی زمانی با یکدیگر مورد مقایسه قرار نگرفته‌اند.

راهکارهای پیشین دارای معایبی می‌باشند از جمله این که برخی از آن‌ها نیمه خودکار می‌باشند [۱۱، ۱۰]. درواقع این راهکارها حین تصحیح، نیازمند تعامل با کاربر هستند که در حال حاضر با توجه به حجم بالای منابع داده‌ای، استفاده از چنین روش‌هایی برای تصحیح، امکان پذیر نیست. از معایب دیگر برخی از روش‌ها نوع داده‌ای است که آن‌ها قادر به تصحیح می‌باشند. این روش‌ها تنها توانایی تصحیح منابع داده‌ای با نوع داده‌ای خاص را دارند [۴، ۱]، اما اغلب منابع داده‌ای موجود دارای ویژگی‌هایی با انواع داده‌ای گوناگون هستند. از این رو روشی کارآمد

است که توانایی تصحیح تمام انواع داده‌ای را دارا باشد. علاوه بر این برخی از روش‌های موجود تنها توانایی تشخیص و تصحیح یک خطا را دارند، درحالی‌که در دنیای واقعی امکان بروز بیش از یک خطا در یک رکورد وجود دارد. همچنین در برخی از راهکارهای ارائه شده، فاز تشخیص و تصحیح غیرقابل تفکیک است، لذا برای کاربر امکان تعامل در صورت دلخواه پس از فاز تشخیص وجود ندارد [۱۹-۱۷]. از این رو در این مقاله روشی خودکار برای تشخیص و تصحیح خطا ارائه شده است که توانایی تشخیص خطا در مجموعه داده‌ای با انواع داده‌ای متفاوت را دارا است. همچنین این روش توانایی تشخیص بیش از یک خطا را در یک رکورد دارا است. در روش پیشنهادی فاز تشخیص و تصحیح کاملاً تفکیک پذیر بوده و لذا کاربر در صورت تمایل می‌تواند پس از تشخیص خطا هر رویکرد دلخواهی را برای هر خطا اتخاذ نماید.

راهکار پیشنهادی در این مقاله از دو فاز تشخیص خطا و تصحیح آن تشکیل شده است. در فاز تشخیص ابتدا رکوردهای مشکوک به خطا با استفاده از وابستگی‌های تابعی از رکوردهای صحیح تشخیص داده می‌شوند، بدین صورت که رکوردی که به ازای یک وابستگی تابعی با بیش از δ از رکوردها در تناقض باشد، مشکوک به خطا است. سپس به وسیله رکوردهای صحیح و وابستگی‌های تابعی محل وقوع خطاها در هر رکورد مشکوک به خطا تعیین می‌شود. برای انجام این کار تمام مقادیری که به ازای هر وابستگی تابعی در رکوردهای صحیح وجود دارد، استخراج می‌شود. سپس براساس این مقادیر تعداد دفعاتی که هر ویژگی وابستگی‌های تابعی را نقض می‌کند، محاسبه می‌شود. براساس نظر کاربر حداکثر تعداد خطا در هر رکورد (β) مشخص می‌شود. سپس β ویژگی که به دفعات بیشتری مقادیر ذخیره شده برای هر وابستگی تابعی را نقض می‌کنند، به عنوان خطا انتخاب می‌شوند.

در فاز تصحیح راهکار پیشنهادی، خطای شناسایی شده در فاز قبل به وسیله سیستم یادگیری مرکب^{۱۴} اصلاح می‌شود. در راهکار پیشنهادی سیستم یادگیری مرکب شامل ۳ طبقه بند، درخت تصمیم^{۱۵}، طبقه‌بند بیز^{۱۶} و شبکه عصبی MLP است و از رأی اکثریت به عنوان استراتژی ترکیب در آن استفاده شده است. به ازای هر ویژگی یک سیستم یادگیری مرکب ساخته و طبقه‌بندهای آن به وسیله رکوردهای فاقد خطا، آموزش داده می‌شوند. پس از مرحله آموزش، هر ویژگی خطا دار به عنوان کلاس هدف انتخاب و توسط سیستم یادگیری مرکب مربوطه تصحیح می‌شوند.

تاکنون راهکارهای متعددی از وابستگی تابعی برای تصحیح داده‌ها استفاده نموده‌اند؛ ولی در راهکار پیشنهادی این مقاله از ترکیب وابستگی تابعی و سیستم یادگیری مرکب استفاده شده است که به نسبت روش‌هایی که فقط از وابستگی تابعی برای تصحیح استفاده نموده‌اند کارایی بالاتری داشته است.

بخش‌های بعدی مقاله بدین شرح است. در بخش دوم مروری بر راهکارهای پیشین در حوزه تصحیح داده‌ها صورت می‌گیرد. در بخش چهارم راهکار پیشنهادی به‌طور کامل تشریح می‌شود و در بخش پنجم

صحیح است. برای استفاده از درست‌نمایی از دو معیار برای بیشینه کردن سود و کمینه کردن هزینه تغییرات بهره‌گرفته‌شده است. در این روش به ازای هر ویژگی یک طبقه‌بند ایجاد و مقدار هر ویژگی غلط پیش‌بینی می‌شود. در مرحله بعد برای هر رکورد یک گراف که گره‌ها مقادیر پیش‌بینی شده و یال‌ها وابستگی میان ویژگی‌ها را نشان می‌دهد، رسم می‌شود. پس از رسم گراف تا زمانی که از هر ویژگی تنها یک مقدار باقی‌ماند، گره‌ای که یال‌های آن دارای کمترین وزن هستند، حذف می‌شود. این الگوریتم می‌تواند چندین ویژگی ناصحیح در یک رکورد را پیش‌بینی نماید، به علاوه این که در حین پیش‌بینی وابستگی بین داده‌های صحیح و داده‌های ناسازگار را به صورت مجزا در نظر می‌گیرد. از معایب این روش این است که در این روش فرض شده که خطا تشخیص داده‌شده و فقط به تصحیح پرداخته است.

الگوریتم ارائه‌شده در [۸] با استفاده از محدودیت‌ها به تصحیح یک منبع ناسازگار می‌پردازد. این محدودیت‌ها از سه بخش اصلی تشکیل و فرم عمومی آن‌ها به صورت $t_p^+[B] \rightarrow ((X, t_p[X]), (B, t_p^-[B]))$ می‌باشد. X بیانگر ویژگی‌هایی است که با یکدیگر در ارتباط هستند، $t_p[X]$ مقادیر آن ویژگی‌ها، B ویژگی‌هایی است که برای آن الگوهای منفی $t_p^-[B]$ ارائه شده است و $t_p^+[B]$ مقادیر صحیح ارائه‌شده برای مقادیر منفی است. در ابتدا هر یک از ویژگی‌های ارائه‌شده در بخش الگوهای شاهد، به‌تنهایی به‌عنوان یک کلید انتخاب می‌شود و هر یک با مقادیر خود در یک لیست با نام قانون ذخیره می‌شوند. به ازای هر رکورد موجود در منبع بررسی می‌شود که آیا از میان مقادیر کلیدهای ذخیره‌شده در لیست با رکورد مورد نظر تطابق یافت می‌شود. در صورتی که تطابق یافته نشد، نشانگر این است که این رکورد دارای مقادیر صحیح است. اما در صورتی که تطابق یافت گردید باید آن قانون بازیابی و از قسمت الگو صحیح آن رکورد تصحیح گردد. به ازای یک رکورد باید تمامی کلیدها مورد بررسی قرار گیرند. این روند باید برای تمامی رکوردها تکرار شود. هدف از تعیین کلید در این الگوریتم صرفه‌جویی در هزینه مقایسات است، زیرا با این کار لازم نیست تمام قسمت‌های یک قانون برای یک کلید مورد بررسی قرار گیرد.

در [۱۲] مدلی برای تصحیح داده‌ها مبتنی بر وابستگی‌های تابعی معرفی شده است. در این راهکار وابستگی‌های تابعی استخراج شده ثابت نیستند بلکه هم‌زمان با تصحیح داده‌ها ممکن است وابستگی‌های تابعی نیز تغییر نمایند. در واقع زمانی که تعداد زیادی از رکوردها یک وابستگی تابعی را نقض می‌کنند، با اضافه نمودن یک ویژگی به وابستگی‌های تابعی به گونه‌ای که سازگاری رکوردها با آن برقرار شود، تصحیح انجام می‌گیرد. هم‌چنین این روش هنگامی که دو رکورد یک وابستگی تابعی را نقض می‌کنند، تنها رویکرد تصحیح ویژگی سمت راست وابستگی تابعی را به کار نمی‌گیرد. این روش به جستجوی رکوردی که فاصله کمی با یکی از دو رکورد متناقض دارد و هم‌چنین وابستگی تابعی مورد بررسی را نقض نمی‌کند، می‌پردازد. در صورت یافتن این رکورد بر اساس تابع

آزمایش‌ها انجام گرفته بر روی مجموعه داده‌ای با انواع داده‌ای اسمی، ترتیبی و عددی بیان می‌شود. علاوه بر این در بخش پنجم روش پیشنهادی با روشی مشابه مورد مقایسه قرار گرفته است. در بخش پایانی خلاصه و نتیجه‌گیری از مقاله ارائه می‌گردد.

۲- پیشینه تحقیق

پیش از آن که داده‌های دنیای واقعی مورد استفاده قرار گیرند، باید از کیفیت آن‌ها اطمینان حاصل گردد؛ از این رو تأمین کیفیت، پیش‌پردازش محسوب می‌شود. برای تأمین صحت داده‌ها باید مقادیر از دست‌رفته، ناسازگاری‌ها و داده‌های ناقص شناسایی و تصحیح گردد. برای تصحیح داده‌ها رویکردهای متفاوتی ارائه شده است که در این بخش به اختصار به شرح برخی از آن‌ها پرداخته می‌شود.

در [۱۰] راهکاری متعامل با کاربر مبتنی بر هستان‌شناسی برای تصحیح داده‌ها ارائه شده است. در مرحله اول جهت شناسایی خطاها از هستان‌شناسی استفاده می‌شود، سپس مقادیر صحیح از هستان‌شناسی استخراج و جهت انتخاب به کاربر نمایش داده می‌شود. پس از انتخاب کاربر، این رفتار او در سیستم ذخیره می‌شود. زمانی که انتخاب‌های کاربر به یک حد آستانه رسید، از رفتار او قوانین استخراج و از این پس برای تصحیح از قوانین استفاده می‌شود. این روش به دلیل تعامل با کاربر در بخش ابتدایی در حجم بالای داده‌ای امکان‌پذیر نیست.

روش DMI [۴] و SiMI [۱] برای انتساب مقدار صحیح به رکوردهایی که مقادیر خود را از دست داده‌اند، ارائه شده است. این دو الگوریتم بر پایه الگوریتم EM هستند با این تفاوت که SiMI بهبود یافته الگوریتم DMI است. در الگوریتم DMI پس از تفکیک داده‌های صحیح و داده‌های حاوی مقادیر از دست‌رفته، برای هر ویژگی از مجموعه داده مورد بررسی، درخت تصمیمی ایجاد می‌شود. برای تعیین مقدار ویژگی در برگ‌ها چنانچه ویژگی عددی باشد از الگوریتم EM و در غیر این صورت از رأی اکثریت در آن برگ استفاده می‌گردد. سپس هر ویژگی که مقدار خود را از دست داده است به وسیله درخت متناظر آن تصحیح می‌گردد. در الگوریتم SiMI که توسعه یافته DMI است، از جنگل تصمیم استفاده شده است. در این روش نیز به ازای هر ویژگی یک جنگل تصمیم ساخته می‌شود و به جستجوی تقاطع میان برگ‌ها می‌پردازد. سپس به ادغام تقاطع‌ها با هدف بیشینه نمودن شباهت بین رکوردها پرداخته و با انتساب هر رکورد دارای مقدار از دست‌رفته به یک تقاطع آن را تصحیح می‌نماید. این دو الگوریتم تنها قادر به تصحیح ویژگی‌هایی با دو نوع داده‌ای عددی و دسته‌ای می‌باشند.

الگوریتم SCARE با استفاده از الگوریتم‌های یادگیرنده و درست‌نمایی به تصحیح خطا پرداخته است [۷]. هدف در این الگوریتم حفظ ارتباط بین ویژگی در هنگام تصحیح است. هم‌چنین در این الگوریتم برخلاف برخی روش‌ها که هدف تنها کمینه کردن تغییرات است، علاوه بر محدود نمودن تغییرات، بیشینه کردن درست‌نمایی نیز در نظر گرفته شده است. روش تصحیح مقادیر ناصحیح نیز به وسیله مقادیر

الگوریتم یافتن یک مجموعه سازگار است که فاصله آن از مجموعه اولیه، نزدیک به فاصله تصحیح بهینه از مجموعه اولیه باشد. در واقع در این روش هدف به دست آوردن رکوردهایی است که با وابستگی‌های تابعی تعریف شده در تناقض می‌باشند و اصلاح آن‌ها به وسیله وابستگی تابعی به نحوی است که دارای کمترین فاصله از مجموعه داده اولیه باشد. این الگوریتم به جستجوی مقادیری که وابستگی‌های تابعی تعریف شده را نقض می‌کنند، می‌پردازد و آن‌ها را در یک گراف قرار می‌دهد. در این گراف برای مجموعه ناسازگاری‌ها، گره‌ها مقادیر ناسازگار و یال‌ها وابستگی بین مقادیر را تعیین می‌کنند. پس از این که تمامی گراف‌ها محاسبه گردید؛ الگوریتم، کوچک‌ترین گراف همپوشان را پیدای می‌کند. سپس به ازای هر دو رکورد t_1 و t_2 که هر وابستگی تابعی را نقض می‌کند، در صورتی که (t_1, a) تنها مقدار در گراف همپوشان باشد، مقدار t_2 بر اساس t_1 تصحیح می‌شود و t_1 را از گراف همپوشان حذف می‌نماید. پس از این که تصحیح انجام شد، امکان دارد در اثر رفع برخی تناقض‌ها در رکوردها اشکال ایجاد شود و مجدداً این اشکال مورد بررسی قرار می‌گیرد. در ابتدا وابستگی‌های تابعی بر اساس امتیازی که از میانگین دو پارامتر حاصل می‌شود، در یک لیست به صورت نزولی مرتب می‌گردند. پارامتر اول درجه ناسازگاری هر وابستگی تابعی را با منبع داده مشخص می‌کند و پارامتر دوم میزان تناقض وابستگی‌های تابعی دیگر با وابستگی تابعی در نظر گرفته شده را تعیین می‌کند. وابستگی‌های تابعی به ترتیب از لیست برداشته می‌شوند. این الگوریتم به صورت حریصانه عمل می‌کند. میزان هزینه تصحیح تناقض با وابستگی تابعی مورد نظر محاسبه می‌شود؛ سپس هزینه تصحیح ناسازگاری وابستگی تابعی مورد نظر با وابستگی‌های دیگر مورد محاسبه قرار می‌گیرد. تصحیحی که کمترین هزینه را داشته باشد، اعمال می‌شود و این روند برای وابستگی تابعی دیگر اجرا می‌شود.

در [۲۱] الگوریتمی ارائه شده که به طور هم‌زمان به تصحیح وابستگی‌های تابعی و داده‌ها می‌پردازد. در این الگوریتم به ازای هر وابستگی تابعی که توسط رکوردها نقض می‌شود، تصحیح‌های متنوعی برای وابستگی تابعی و داده‌ها ارائه می‌شود، به نحوی که دارای کمترین تغییرات باشند. این تصحیح‌ها به کاربر نمایش داده می‌شود، تا او با توجه به دانشی که نسبت به مجموعه داده‌ها دارد، بهترین تصحیح را انتخاب نماید. هدف اصلی در این روش محدود نمودن حداکثر تعداد فیلدهایی است که در منبع داده دچار تغییر می‌گردند و به دست آوردن مجموعه‌ای از وابستگی‌های تابعی است که به مجموعه اولیه وابستگی تابعی نزدیک و تناقضی با مجموعه داده‌ها نداشته باشد. در این راهکار ابتدا به کمک گراف تناقض، مجموعه داده و وابستگی‌های تابعی اولیه به جستجوی تصحیحاتی برای وابستگی‌های تابعی می‌پردازد. این روش برای یافتن مجموعه وابستگی تابعی بهینه، گراف تناقض را با ترکیب A^* و یک الگوریتم تجربی حرص می‌نماید. پس از اصلاح وابستگی تابعی، به ازای مجموعه جدید وابستگی تابعی و مجموعه داده‌ها گراف تناقض رسم می‌شود و سپس هر تراکنش به صورت تصادفی انتخاب و جهت

هزینه تصمیم‌گیری می‌نماید که به تصحیح ویژگی سمت راست بپردازد و یا این که ویژگی سمت چپ را بر اساس رکورد یافته شده، اصلاح نماید. در [۱۹-۱۷] هدف شناسایی ویژگی‌های مشکوک به خطا و اصلاح آن به وسیله الگوریتم *polishing* است. این الگوریتم از دو فاز پیش‌بینی و تنظیم مقادیر تشکیل شده است. در این الگوریتم از یک رویه به صورت بازگشتی استفاده می‌شود، که ترکیب‌های مختلف از مقادیر مختلف ویژگی‌ها را امتحان می‌نماید. در فاز پیش‌بینی از میان الگوریتم‌های طبقه‌بندی، یک روش انتخاب می‌شود. به ازای هر ویژگی ۱۰ دسته بر روی داده‌ها ایجاد می‌کند و برچسب آن ویژگی به عنوان هدف در نظر گرفته می‌شود. مقدار ویژگی هدف در هر رکورد به ازای ۱۰ طبقه‌بندی پیش‌بینی می‌شود. در صورتی که مقدار اصلی آن ویژگی در رکورد مورد نظر با مقادیر پیش‌بینی شده در تناقض باشد، مقادیر حاصل از پیش‌بینی برای آن رکورد ذخیره می‌گردد. این مقادیر در فاز بعدی جهت تصحیح مورد استفاده قرار می‌گیرد. در فاز تنظیم مجدداً به ازای هر ویژگی ۱۰ دسته بر روی داده‌ها ایجاد می‌شود. سپس به ازای هر ویژگی متناقض شناسایی شده در فاز قبل، مقدار آن بر اساس مقادیر پیش‌بینی شده و توسط ۱۰ طبقه‌بندی تصحیح می‌شود. مقدار جایگزین باید از میان مقادیر پیش‌بینی شده در مرحله قبل انتخاب شود. باید توجه نمود که در این روش تمامی مقادیر تغییر نمی‌کند، بلکه تنها مقادیری که دارای مقدار غیر طبقه‌بندی هستند، دچار تغییر می‌گردند، به علاوه این که سعی می‌شود تغییرات را به حداقل برساند. در این روش نیز فاز تصحیح و تشخیص غیر قابل تفکیک می‌باشند و لذا کاربر نمی‌تواند رویکرد دیگری را برای فاز تصحیح به کار گیرد.

الگوریتم [۱۱] تصحیحی مبتنی بر وابستگی تابعی ارائه نموده و برای انتخاب بهترین تصحیح از دو تابع هزینه و تنوع استفاده کرده است. برای محاسبه میزان تنوع از تابع فاصله استفاده می‌شود که میزان عدم شباهت رکوردها را محاسبه می‌نماید. برای تعیین فاصله از وزن‌های اختصاص داده شده به فیلدها که در واقع میزان اعتماد کاربر به آن‌ها است، استفاده می‌گردد. در این روش برای محاسبه میزان تنوع از دو تابع استفاده شده است. در تابع تنوع اول تفاضل بین مجموع هزینه را با مجموع فاصله برای تمام تصحیح‌ها ارائه شده، محاسبه می‌کند. در این تابع پارامتری برای تعادل بین هزینه و فاصله وجود دارد. زمانی که این پارامتر افزایش می‌یابد، برای بهینه شدن تابع تنوع مقدار تابع هزینه نیز باید افزایش یابد. مقدار بین ۰/۳-۰ برای این پارامتر پیشنهاد می‌شود. تابع تنوع دوم تفاضل بین حداکثر هزینه و حداقل فاصله را بیان می‌کند. سپس این روش با استفاده از الگوریتمی از میان K تصحیح بهینه با هدف کمینه نمودن دو تابع تنوع که در واقع حداقل نمودن بیشینه هزینه و حداکثر نمودن کمینه فاصله است، یک تصحیح را به عنوان تصحیح بهینه معرفی می‌کند. این الگوریتم به دلیل تعامل با کاربر برای تعیین وزن‌های اولیه در حجم بالای داده‌ای مناسب نیست.

در [۲۰] یک الگوریتم تقریبی برای تصحیح مجموعه داده ناسازگار مبتنی بر مجموعه وابستگی‌های تابعی ثابت بیان شده است. هدف این

ظهور پایگاه‌های دانش^{۱۷} موجب ارائه فرصت‌های بیشتر برای دستیابی به داده‌های صحیح در مقیاس بزرگ‌تر شده است. در [۲۴] روشی مبتنی بر پایگاه دانش و سیستم‌های هوشمندسازی محاسبات^{۱۸} ارائه شده است. در این شیوه جدول حاوی خطا، پایگاه دانش و *crowd* به عنوان ورودی دریافت می‌شوند. در ابتدا جدول الگوها برای نگاشت جدول داده‌ها به پایگاه دانش استخراج می‌شود. سپس برای هم‌تراز کردن با پایگاه دانش جدول معانی را تفسیر می‌کند. پس از آن با جدول الگوها داده‌ها را به عنوان صحیح و ناصحیح به وسیله جایگذاری پایگاه دانش مشخص می‌کند. برای داده‌های ناصحیح k بالاترین نگاشت را از پایگاه دانش به عنوان مقدار صحیح جهت تصحیح معرفی می‌کند. در صورتی که اطلاعات لازم برای انتخاب k تصحیح در پایگاه دانش وجود نداشته باشد، این روش از فرد خبره در حوزه پایگاه دانش استفاده می‌کند.

در [۲۵] یک مسئله جدید به نام اثر متقابل بین تصحیح داده‌ها و تطابق آن‌ها بیان شده است. در این راهکار بیان شده که تصحیح داده‌ها به طور مؤثر به تطابق داده‌ها کمک می‌کند. در این شیوه چارچوبی متشکل از عملیات تطابق و تصحیح برای دستیابی به مجموعه داده صحیح مبتنی بر محدودیت‌ها، قوانین تطابق و *master data* ارائه شده است. در این راهکار از وابستگی تابعی شرطی و وابستگی تطابق به عنوان محدودیت‌ها برای کشف ناسازگاری‌ها استفاده شده است. برای اطمینان از دقت رفع خطاها سه پارامتر، میزان اطمینان کاربر به داده‌ها، اندازه‌گیری آنتروپی و داده *master* به کار گرفته شده است. بدین ترتیب سه الگوریتم برای رفع خطاها در این چارچوب ارائه شده است. الگوریتم اول مبتنی بر آنالیز مقدار اطمینان و داده *master* به رفع قطعی خطاها می‌پردازد، که دقت بالاتری نسبت به دو الگوریتم دیگر دارد. زمانی که اطمینان مقدار پایینی داشته باشد یا ناموجود باشد، الگوریتم برای رفع منطقی خطاها از آنتروپی استفاده می‌کند. برای باقی‌مانده خطاها یک روش تجربی ارائه شده است.

در [۲۶] یک سیستم یکپارچگی برای شناسایی مقادیر ناسازگار داده‌ها همراه با مقدار پیشنهادی جهت تصحیح آن، مبتنی بر وابستگی تابعی بیان شده است. در این روش برای رفع ناسازگاری‌ها تصحیحاتی بر روی داده یا وابستگی‌های تابعی انجام می‌گیرد. در این راهکار برای تصحیح وابستگی‌های تابعی ۳ شیوه مورد استفاده قرار گرفته است: (۱) اضافه نمودن ویژگی به وابستگی تابعی، (۲) تبدیل وابستگی تابعی به وابستگی تابعی شرطی، (۳) شناسایی افزونگی در وابستگی تابعی. این سیستم از ۴ جز تشکیل شده است: (۱) مازول تشخیص‌دهنده ناسازگاری، (۲) مازول تصحیح‌کننده داده، (۳) مازول تصحیح‌کننده وابستگی تابعی، (۴) موتور تصحیح. در ابتدا داده و وابستگی‌های تابعی برای تشخیص داده ناسازگار به بخش تشخیص‌دهنده ناسازگاری هدایت می‌شوند، سپس هر رکورد که با هر وابستگی تابعی در تناقض باشد، شناسایی می‌شود. رکورد ناسازگار وارد مازول تصحیح‌کننده داده و وابستگی تابعی نقض شده وارد تصحیح‌کننده وابستگی تابعی می‌شود. مازول تصحیح‌کننده داده، الگوهای رکوردها را با یکدیگر مورد مقایسه

تصحیح مورد بررسی قرار می‌گیرد. هم‌چنین بار دیگر رکوردهای متناقض بدون اصلاح وابستگی تابعی اصلاح می‌گردد. تصحیحات یافته شده برای مجموعه وابستگی‌های تابعی و مجموعه داده‌ها برای مشخص نمودن تصحیح بهینه به کاربر پیشنهاد می‌گردد. این روش به دلیل تعامل با کاربر در حجم داده‌ای بالا مناسب نیست.

در [۲۲] چارچوبی برای تصحیح داده‌هایی که به طور مداوم در حال تغییر می‌باشند، ارائه شده است. این چارچوب مبتنی بر وابستگی تابعی بوده، و در هنگام یافتن تناقض علاوه بر جستجو برای یافتن تصحیح داده‌ها، به جستجوی تصحیح وابستگی تابعی نقض شده می‌پردازد. سپس به وسیله یک طبقه‌بند مشخص می‌کند که تصحیح باید بر روی وابستگی تابعی، داده و یا به صورت ترکیبی بر روی هر دو صورت گیرد. نویسنده این مقاله معتقد است که روش‌های تصحیح داده اغلب ذهنی است، لذا زمانی که طبقه‌بند تشخیص می‌دهد که داده‌ها باید اصلاح گردد، داده‌های متناقض را به همراه مقادیری جهت تصحیح به کاربر نمایش می‌دهد و از او می‌خواهد تا آن‌ها را اصلاح نماید. پس از انتخاب کاربر، داده اصلاح شده جهت یادگیری به طبقه‌بند داده می‌شود. در این روش از پارامترهای آماری نیز برای یافتن تصحیح بهینه استفاده شده است. مرحله اول یادگیری طبقه‌بند است. در این مرحله طبقه‌بند بر اساس تصحیحی که توسط کاربر انتخاب شده است، تحت آموزش قرار می‌گیرد. در مرحله بعد تناقض بعدی شناسایی می‌شود و طبقه‌بند بر اساس داده‌های آماری و آموزشی که صورت گرفته، نوع تصحیح (تصحیح داده، تصحیح وابستگی تابعی و یا تصحیح ترکیبی) را پیش‌بینی می‌نماید. پس از تعیین نوع تصحیح وارد مرحله سوم می‌گردد. در این مرحله به جستجوی تمام مقادیر معتبر برای رفع ناسازگاری‌ها می‌پردازد. پس از یافتن تمام مقادیر معتبر جهت جایگزینی، این مقادیر جهت انتخاب تصحیح بهینه به کاربر نمایش داده می‌شود. پس از انتخاب کاربر این تصحیح منتخب بر روی مجموعه داده‌ها و یا وابستگی تابعی اعمال می‌گردد و هم‌چنین این انتخاب مجدداً جهت آموزش وارد مرحله اول می‌شود. این الگوریتم به دلیل تعامل با کاربر در هر دور از چرخه در حجم بالای داده‌ای مناسب نیست.

در [۲۳] راهکاری برای تصحیح ناسازگاری‌ها با نمونه‌برداری از تصحیحات مبتنی بر وابستگی تابعی شرطی ارائه شده است. این راهکار به ازای هر ناسازگاری در مجموعه داده‌ها تنها یک تصحیح ارائه نمی‌کند، بلکه به ازای هر ناسازگاری چندین تصحیح از فضای وابستگی تابعی شرطی ارائه می‌کند. سپس با استفاده از الگوریتمی که ارائه شده، به طور تصادفی یک تصحیح کارآمد را از فضای تصحیحات انتخاب می‌کند. هم‌چنین این روش به گونه‌ای طراحی شده، که علاوه بر وابستگی تابعی و وابستگی تابعی شرطی، می‌تواند محدودیت‌هایی را که از سوی کاربر تعیین می‌گردد، در تصحیحات به کار گیرد. محدودیت‌های کاربر می‌تواند شامل فیلدهایی باشد که کاربر در نظر دارد در طول تصحیح بدون تغییر باقی‌بمانند. این روش تنها رکوردهای ناسازگار را تصحیح می‌نماید و برای رفع ناسازگاری هیچ رکوردی را حذف نمی‌کند.

۳- پیش‌زمینه

در این بخش تعریفی از انواع داده‌ای و سیستم یادگیری مرکب ارائه می‌گردد.

تعریف ۱: ویژگی‌های یک مجموعه داده در فرمت، دامنه و نوع با یکدیگر متفاوت هستند. ویژگی‌های یک منبع داده ممکن است همگی از یک نوع داده‌ای باشند و یا امکان دارد از انواع داده‌ای مختلف در یک مجموعه داده وجود داشته باشد [۲۹]. انواع داده‌ای شامل نوع داده‌ای اسمی، ترتیبی، عددی و باینری است که در ادامه به شرح هر یک پرداخته می‌شود.

- ویژگی اسمی: نوع داده‌ای است که مربوط به اسمی است. مقدار یک ویژگی اسمی سِمبل‌ها یا نام اشیاء است. هر مقدار نشان‌دهنده نوعی از دسته‌ها، کدها و حالت‌ها است، از این رو به ویژگی‌های اسمی، دسته‌ای^{۱۹} نیز گفته می‌شود. این مقادیر هیچ ترتیب معناداری ندارند.
 - ویژگی باینری: ویژگی باینری، نوعی از ویژگی اسمی است که تنها دو مقدار ۰ و ۱ دارد، که ۰ نشان‌دهنده این است که آن ویژگی حضور ندارد و ۱ به معنای این است که آن ویژگی وجود دارد.
 - ویژگی ترتیبی: نوعی از ویژگی‌ها است که مقادیر آن دارای ترتیب معناداری هستند و یا این‌که بین مقادیر آن نوعی رتبه‌بندی وجود دارد.
 - ویژگی عددی: ویژگی‌های عددی کمی هستند؛ در واقع دارای یک مقدار قابل اندازه‌گیری است که صحیح یا حقیقی است.
- تعریف ۲: مدل‌های یادگیری مرکب، الگوریتم‌های یادگیری‌ای هستند که از مجموعه‌ای از طبقه‌بندها ساخته می‌شوند. سیستم‌های یادگیری مرکب به‌طور عمده شامل دو بخش هستند. بخش اول نحوه تقسیم‌بندی مجموعه داده آموزش، انتخاب نوع طبقه‌بندها، تعداد طبقه‌بندها و ویژگی‌های مناسب برای هر طبقه‌بند است. راهکارهای متنوعی برای تقسیم‌بندی مجموعه داده آموزش وجود دارد که از جمله می‌توان به cross validation، bootstrapping و تقسیم‌بندی بر اساس میزان شباهت اشاره نمود. بخش دوم شامل نحوه ترکیب خروجی طبقه‌بندها به‌منظور حصول بهترین نتیجه برای طبقه‌بندی الگوهاست. روش‌های گوناگونی برای ترکیب طبقه‌بندها وجود دارد که از پرکاربردترین شیوه‌ها می‌توان به میانگین، رأی اکثریت، boosting، borda counts، ECO و mixture of expert اشاره نمود [۳۵].

۴- راهکار پیشنهادی

در این بخش یک راهکار خودکار برای تشخیص و تصحیح خطا ارائه شده‌است. در فاز تشخیص، ابتدا رکوردهای مشکوک به خطا و سپس ویژگی‌های مشکوک به خطا در هر رکورد با استفاده از وابستگی تابعی تشخیص داده می‌شوند. در فاز تصحیح خطاهای شناسایی شده به‌وسیله یادگیری مرکب که شامل الگوریتم‌های درخت تصمیم، طبقه‌بند بیز و

قرار می‌دهد و مجموعه‌ای از تصحیحات را به موتور تصحیح می‌فرستد. در مازول تصحیح‌کننده وابستگی تابعی نیز با یکی از ۳ روش بیان شده تصحیحاتی ارائه و به موتور تصحیح فرستاده می‌شود. در موتور تصحیح، تصحیحی که کم‌ترین هزینه را داشته باشد، انتخاب می‌گردد. هم‌چنین در این سیستم امکان این که بهترین تصحیح توسط کاربر انتخاب گردد، فراهم شده‌است.

در [۲۷] راهکاری برای تصحیح داده‌های ناصحیح مبتنی بر محدودیت‌ها ارائه شده‌است. در این روش تصحیح داده‌ها با تغییرات کوچک شامل درج یا حذف ویژگی بر روی محدودیت‌ها همراه است. این مقاله تصحیحی کمینه برای داده‌ها ارائه می‌کند، به‌نحوی که حداقل یک نوع محدودیت که دارای تنوع بیشتر از θ نسبت به محدودیت‌های دیگر داده شده باشد را برآورده نماید. این روش قابلیت تصحیح ویژگی‌های عددی را دارا است، که این امکان در روش‌های مشابه مبتنی بر محدودیت وجود نداشته‌است.

در [۲۸] یک مدل مبتنی بر هزینه که داده و محدودیت در آن به‌طور تساوی با یکدیگر مورد مقایسه قرار می‌گیرند، ارائه شده‌است. در این راهکار از وابستگی تابعی به‌عنوان محدودیت استفاده شده‌است. در این مدل دو الگوریتم جداگانه برای تعمیر داده و محدودیت بیان شده‌است. در الگوریتم تعمیر داده به ازای هر داده ناسازگار، به جستجوی تغییراتی برای داده می‌پردازد، به‌نحوی که با محدودیت‌ها سازگار و هزینه تعمیرات را حداقل گرداند. برای تصحیح دو رکورد ناسازگار با یک وابستگی تابعی، مقدار ویژگی سمت راست وابستگی تابعی در دو رکورد را یکسان قرار می‌دهد یا این که مقدار ویژگی سمت چپ آن‌ها را متفاوت می‌گذارد. در الگوریتم تصحیح وابستگی تابعی برای رفع ناسازگاری‌ها، دو رویکرد ارائه شده‌است. در رویکرد اول به جستجوی ویژگی‌ای می‌پردازد که با اضافه نمودن آن به وابستگی تابعی، ناسازگاری رفع می‌شود. برای انتخاب ویژگی از مفهوم واریانس اطلاعات بین ویژگی جدید و وابستگی تابعی ناسازگار استفاده می‌کند، به‌صورتی که ویژگی بدون کاهش افزونگی باعث تغییر ناسازگاری گردد. در رویکرد دوم مجموعه‌ای از مقادیر ویژگی‌ها از رکوردهایی که فقط آن‌ها وابستگی تابعی را برآورده می‌کنند، شناسایی می‌گردد. این مقادیر، زیرمجموعه‌ای از رکوردها را تعیین می‌کند، که با وابستگی تابعی سازگار است. سپس مدلی مبتنی بر هزینه برای انتخاب بهترین تصحیح اجرا می‌شود. هدف اصلی در این مدل انتخاب تعمیری با کمترین میزان تغییرات است.

راهکارهای تصحیح معرفی شده از چندین دیدگاه با یکدیگر تمایز دارند که از جمله می‌توان به خودکار یا نیمه خودکار بودن، انواع داده‌ای که قادر به تصحیح آن هستند، تفکیک پذیری فاز تشخیص و تصحیح و تعداد خطایی که می‌توانند در یک رکورد شناسایی کنند، اشاره نمود. هدف اصلی این مقاله معرفی راهکاری جهت تصحیح خودکار داده‌ها برای تمام انواع داده‌ای است. هم‌چنین روش پیشنهادی توانایی تشخیص چندین خطا را در یک رکورد دارد. در بخش پنجم جزئیات راهکار پیشنهادی به‌طور کامل شرح داده می‌شوند.

م تفاوت هستند، اما کیفیت وابستگی‌های تابعی استخراج‌شده در تمام الگوریتم‌ها یکسان است [۳۰]. هم‌چنین تصحیح داده‌ها جز فرایند پیش‌پردازش است و لذا پیچیدگی زمانی، پارامتر فاقد اهمیت است. بنابراین در راهکار پیشنهادی از الگوریتم TANE [۳۱] که پایه روش‌های بعدی بوده است، استفاده شده است. با استفاده از وابستگی‌های تابعی استخراج‌شده، باید رکوردهای مشکوک به خطا شناسایی گردد. رکوردهایی که بیش از یک حد آستانه وابستگی‌های تابعی را نقض نمایند، مشکوک به خطا می‌باشند. این حد آستانه توسط کاربر نیز قابل تنظیم است، در این مقاله هر رکوردی که به ازای یک وابستگی تابعی با δ رکوردها ناسازگار باشد، مشکوک به خطا در نظر گرفته شده است. پس از شناسایی رکوردهای مشکوک به خطا، رکوردهای صحیح از این رکوردها تفکیک می‌گردد. این قسمت از مرحله تشخیص خطا بخش اول نامیده می‌شود و شبه کد آن در شکل ۱ آورده شده است.

شبکه عصبی MLP است، تصحیح می‌گردد. این راهکار قادر است که خطا را در تمام انواع داده‌ای تشخیص و تصحیح نماید. هم‌چنین به دلیل تفکیک‌پذیری فاز تشخیص و تصحیح، کاربر امکان تعامل با الگوریتم را دارد، تا در صورت دلخواه در خطاها تغییر ایجاد نماید. البته در صورت عدم تعامل کاربر هیچ مشکلی برای راهکار تصحیح پیش نمی‌آید و راهکار بدون اشکال به تصحیح می‌پردازد. علاوه بر این، راهکار پیشنهادی توانایی تشخیص چندین خطا را در یک رکورد دارا است. در بخش‌های ۴-۱ و ۴-۲ به ترتیب به شرح فاز تشخیص و تصحیح پرداخته می‌شود.

۴-۱- مرحله اول. تشخیص خطا

در این مرحله پیش از شروع الگوریتم، ویژگی‌های عددی باید گسسته‌سازی^{۲۰} شوند و به ویژگی‌های ترتیبی تبدیل گردند. برای این کار می‌توان از یکی از الگوریتم‌های موجود استفاده نمود [۲۹]. پس از این مرحله وابستگی‌های تابعی منبع داده باید استخراج گردد. قابل ذکر است که الگوریتم‌های استخراج وابستگی در پیچیدگی زمانی با یکدیگر

```

Input
D: a set of records
δ: a threshold for number of inconsistent records
Output
D: a set of correct records
D': a set of erroneous records
Procedure
1 foreach numeric attribute in D
2     D ← discretized numeric attribute
3 end
4 FD ← functional dependencies of D extracted with
    TANE algorithm
5 D' ← ∅
6 foreach record in D
7     foreach fdi in FD
8         if record inconsistent with more than δ records
9             D' ← D' ∪ record
10            D ← remove record from D
11        end
12    end
13 end
14 return(D, D')

```

شکل ۱: شبه کد بخش اول مرحله تشخیص خطا راهکار پیشنهادی

یکسان داشته باشند. به ازای این وابستگی مقادیر "32, master, 100\$" و "26, Bachelor, 70\$" در مجموعه داده حداقل دومرتبه تکرار شده است. چنین مقادیری باید به ازای تمام وابستگی‌های تابعی استخراج گردد. این مقادیر به ازای هر وابستگی تابعی در ماتریس φ_i ذخیره می‌شوند. چنانچه m وابستگی تابعی $\{fd_1, fd_2, \dots, fd_m\}$ از منبع داده استخراج شده باشد، مجموعه این مقادیر به صورت $\varphi_1, \varphi_2, \dots, \varphi_m$ است. پس از استخراج این مقادیر، به کمک آن‌ها محل وقوع خطا در رکوردهای خطا دار شناسایی می‌شود.

حال به وسیله رکوردهای صحیح و وابستگی‌های تابعی باید محل وقوع خطا در هر رکورد مشخص گردد. این قسمت از مرحله تشخیص خطا بخش دوم نامیده می‌شود و شبه کد آن در شکل ۲ آورده شده است. ابتدا باید تمام مقادیری که حداقل دومرتبه، به ازای یک وابستگی تابعی در مجموعه داده صحیح تکرار شده است، استخراج شود. خطوط ۱-۶ در شکل ۲ این قسمت از روش پیشنهادی را بیان می‌کند. به طور مثال فرض کنید در یک پایگاه داده مربوط به اطلاعات کارمندان، وابستگی تابعی $age, educate \rightarrow income$ استخراج شده است، که بدین مفهوم است که کسانی که سن و مدرک تحصیلی یکسان دارند، باید درآمد

Input
 FD : a set of functional dependency $\{fd_1, fd_2, \dots, fd_m\}$
 β : number of user thinks that each record had error
 D : a set of correct records
 D' : a set of erroneous records

Output
 Out : set of erroneous attribute in each record in D'

Procedure

```

1 foreach  $fd_i$  in  $FD$ 
2   foreach record in  $D$ 
3      $\gamma \leftarrow$  value of attribute in  $fd_i$  from record
4   end
5    $\varphi_i \leftarrow$  the value is repeated at least twice in  $\gamma$ 
6 end
7 foreach record in  $D'$ 
8    $left \leftarrow 0$ 
9    $right \leftarrow 0$ 
10  foreach  $fd_i$  in  $FD$ 
11     $attl \leftarrow$  attribute in the left side of  $fd_i$ 
12    if value of  $attl$  attribute in record exist in  $\varphi_i$ 
13       $attr \leftarrow$  attribute in the right side of  $fd_i$ 
14      if value of  $attr$  attribute in record inconsistency with value of  $attr$ 
15        attribute in founded  $\varphi_i$ 
16         $right \leftarrow right + 1$ 
17         $cache_r \leftarrow attr$  attribute
18      end
19    else
20       $left \leftarrow left + 1$ 
21       $cache_l \leftarrow attl$  attribute
22    end
23    if  $right \geq \frac{m}{k} \cdot \frac{1}{2}$ 
24      error  $\leftarrow$  find repeated number of each attribute in  $cache_r$  &
25      select  $\beta$  the most frequent
26    else if  $left \geq \frac{m}{k} \cdot \frac{1}{2}$ 
27      error  $\leftarrow$  find repeated number of each attribute in  $cache_l$  &
28      select  $\beta$  the most frequent
29    end
30  end
31  return error
32 end

```

شکل ۲: شبه کد بخش دوم تشخیص خطا در الگوریتم تصحیح مبتنی بر وابستگی تابعی و یادگیری مرکب

لذا این ویژگی ذخیره و به پارامتر $right$ که بیانگر تعداد دفعاتی است که ویژگی‌های سمت راست مشکوک به خطا هستند، یک واحد اضافه می‌شود. خطوط ۱۰-۲۱ در شکل ۲ این بخش از راهکار پیشنهادی را نشان می‌دهد. این روند به ازای تمام وابستگی‌های تابعی انجام می‌گیرد. پس از اتمام این مرحله، مقدار $left$ و $right$ مورد بررسی قرار می‌گیرد. فرض می‌شود k تعداد ویژگی‌های منبع داده‌ای باشد، در صورتی که مقدار هر کدام از $\frac{m}{k}$ وابستگی‌های تابعی بیشتر باشد، ویژگی آن سمت از وابستگی تابعی مشکوک به خطا هستند. پس از این که مشخص گردید، ویژگی‌های کدام سمت مشکوک به خطا هستند، به بررسی تعداد دفعاتی که هر ویژگی در آن سمت مشکوک به خطا معرفی شده، پرداخته می‌شود. این تعداد بر تعداد دفعاتی که هر ویژگی در وابستگی تابعی تکرار شده، تقسیم می‌شود و β ویژگی که این مقدار را بیشینه نمایند، به عنوان خطای نهایی معرفی می‌شوند. این بخش از راهکار، در خطوط ۲۲-۲۷ بیان شده است.

ابتدا کاربر لازم است تعیین نماید، که در یک رکورد حداکثر چند ویژگی ناصحیح ممکن است موجود باشد. این مقدار β نامیده می‌شود. به ازای هر رکورد r باید تمام وابستگی‌های تابعی مورد بررسی قرار گیرند. فرض می‌شود fd_i در حال بررسی است. ابتدا ویژگی‌های سمت چپ fd_i مورد بررسی قرار می‌گیرد. به ازای ویژگی‌های موجود در fd_i ، به جستجوی مقدار این ویژگی‌ها در رکورد r ، در مقادیر ذخیره شده در ماتریس φ_i ، پرداخته می‌شود. چنانچه ترکیب مقادیر موجود در رکورد r در φ_i یافت نگردید، ویژگی‌های سمت چپ مشکوک به خطا هستند. لذا این ویژگی‌ها ذخیره شده و به پارامتر $left$ که نشان می‌دهد که ویژگی‌های سمت چپ مشکوک به خطا هستند، یک واحد اضافه می‌گردد. اما چنانچه این ترکیب مقادیر در رکورد r در φ_i موجود بود، ویژگی سمت راست مورد بررسی قرار می‌گیرد. چنانچه مقدار ویژگی سمت راست در رکورد r ، با مقدار آن ویژگی در ترکیبی که در φ_i یافت شد، برابر نباشد، ویژگی سمت راست مشکوک به خطا است.

۴-۲- مرحله دوم. تصحیح خطا

در این مرحله به ازای خطاهای یافته‌شده در هر رکورد، از یادگیری مرکب برای تصحیح خطا استفاده می‌شود. روش‌های متنوعی برای ساخت یک سیستم یادگیری مرکب وجود دارد، که یکی از این رویکردها استفاده از الگوریتم‌های متنوع یادگیری است. در راهکار پیشنهادی از سه الگوریتم طبقه‌بند بیز، درخت تصمیم و شبکه عصبی MLP استفاده شده است. هدف از انتخاب این سه الگوریتم قدرتمندی آن نسبت به الگوریتم‌های

طبقه‌بندی دیگر مانند نزدیک‌ترین همسایگی است. هم‌چنین استفاده از درخت تصمیم، طبقه‌بند بیز و شبکه عصبی MLP به‌تنهایی در راهکارهای پیشین کارایی نسبتاً خوبی در پیش‌بینی مقادیر صحیح داشته است [۱، ۴، ۶، ۷، ۱۷-۱۹]. سیستم‌های یادگیری مرکب دارای استراتژی‌های ترکیب متنوعی است، که در راهکار پیشنهادی از رأی اکثریت برای استراتژی ترکیب استفاده شده است.

```

Input
D: a set of correct records with k attributes  $\{t_1, t_2, \dots, t_k\}$ 
D': a set of erroneous records with k attributes  $\{t_1, t_2, \dots, t_k\}$ 
ER: incorrect attributes in erroneous records
Output
Out: set of correct values for incorrect attributes in each record in D'

Procedure
1 foreach  $t_i$  in D
2    $tree_i \leftarrow$  train decision tree with D data set &  $t_i$  as target class
3    $bayes\_model_i \leftarrow$  train Bayes with D data set &  $t_i$  as target class
4    $MLP\_model_i \leftarrow$  train MLP with D data set &  $t_i$  as target class
5 end
6 foreach record in D'
7   foreach incorrect attribute e in ER
8      $tree\_estimate\_class \leftarrow$  select  $tree_e$  & predict value for attribute e with  $tree_e$ 
9      $bayes\_estimate\_class \leftarrow$  select  $bayes\_model_e$  & predict value for attribute e
      with  $bayes\_model_e$ 
10     $MLP\_estimate\_class \leftarrow$  select  $MLP\_model_e$  & predict value for attribute e
      with  $MLP\_model_e$ 
11     $correct\_value \leftarrow$  most repeated value between  $tree\_estimate\_class$ ,
       $bayes\_estimate\_class$ ,  $MLP\_estimate\_class$ 
12    return  $correct\_value$ 
13  end
14 end

```

شکل ۳: شبه کد بخش تصحیح خطا از الگوریتم تصحیح مبتنی بر وابستگی تابعی و سیستم یادگیری مرکب

متفاوت است. در این سیستم یادگیری مرکب به ازای هر ویژگی، ۵ شبکه عصبی با پارامترهای اولیه متفاوت استفاده شده است. پس از آموزش طبقه‌بندها، می‌توان مقادیر صحیح را برای ویژگی‌های ناصحیح پیش‌بینی نمود. هر بار یک ویژگی از یک رکورد وارد سیستم یادگیری مرکب می‌گردد. سپس در درخت تصمیم، از میان درختان ساخته‌شده، درخت مربوط به ویژگی ورودی انتخاب و مقداری برحسب رکورد ورودی برای ویژگی ناصحیح به‌وسیله درخت پیش‌بینی می‌شود. در طبقه‌بند بیز و شبکه عصبی نیز به همین روش مدل آموزش‌داده شده برای ویژگی ناصحیح انتخاب و مقدار صحیح به‌وسیله مدل پیش‌بینی می‌گردد. پس از این‌که پیش‌بینی به‌وسیله تمام طبقه‌بندها صورت گرفت، به‌وسیله استراتژی یادگیری مرکب که رأی اکثریت است، مقداری که بیشترین تکرار را داشته به‌عنوان مقدار صحیح معرفی می‌شود. شبه کد این بخش از الگوریتم در شکل ۳ آمده است. در این شبه کد فرض شده که کاربر پارامترهای اولیه شبکه عصبی را می‌داند و لذا از یک شبکه عصبی ساخته می‌شود که در آن‌ها پارامترهای اولیه شبکه عصبی

ابتدا درخت تصمیم مورد آموزش قرار می‌گیرد. برای آموزش هر بار یک ویژگی به‌عنوان کلاس هدف قرار می‌گیرد و درخت به‌وسیله رکوردهای صحیح آموزش ساخته می‌شود. برای ساخت درختان از روش CART استفاده شده است. سپس هر درخت ساخته‌شده برای دستیابی به درخت بهینه، هرس می‌گردد. پس از هرس نمودن درختان برای فاز پیش‌بینی ذخیره می‌شوند. در طبقه‌بند بیز نیز، به ازای هر ویژگی ناصحیح، آن ویژگی به‌عنوان کلاس هدف قرار می‌گیرد و طبقه‌بند بیز ساخته می‌شود. در طبقه‌بند بیز از توزیع گوسی^{۲۱} و احتمال پیشین^{۲۲} استفاده می‌شود. برای آموزش شبکه عصبی MLP چنان‌چه کاربر درباره انتخاب پارامترهای شبکه عصبی از جمله ضریب یادگیری، تعداد لایه‌های میانی و تعداد تکرار برای داده مورد بررسی دانشی داشته باشد، می‌تواند به ساخت یک شبکه عصبی به ازای هر ویژگی اکتفا نماید. اما در اغلب موارد کاربر چنین دانشی را در اختیار ندارد، لذا به ازای هر ویژگی چندین شبکه عصبی ساخته می‌شود که در آن‌ها پارامترهای اولیه شبکه عصبی

بالای راهکار پی‌شنهادی در تشخیص خطا است. مقدار *recall* نیز در مجموعه داده *adult* به‌طور متوسط ۸۶٪ و در مجموعه داده *protein* به‌طور متوسط ۸۴٫۷٪ است. این معیار در حقیقت نشان‌دهنده فیلدهای صحیحی می‌باشند که به‌درستی به‌عنوان فیلد صحیح تشخیص داده شده‌است. به دلیل این‌که هیچ اطلاعاتی درباره تعداد خطا در هر رکورد وجود ندارد، یک مقداری به‌عنوان تعداد خطا حدس زده می‌شود و این باعث می‌گردد تعدادی فیلد صحیح به‌عنوان خطا تشخیص داده شوند و به همین علت مقدار معیار *recall* کمتر از *true negative rate* است. اما آزمایش‌ها نشان خواهد داد، که به اکثر این فیلدها در فاز تصحیح مقداری یکسان یا نزدیک به مقدار اولیه تخصیص داده می‌شود. هم‌چنین میزان *false alarm rate* در تمام نرخ‌های خطا در هر دو مجموعه داده کمتر از ۷٫۵٪ است و مقدار *error rate* به‌طور متوسط در مجموعه داده *adult* ۱۴٪ و در مجموعه داده *protein* ۱۵٪ است. هم‌چنین مقدار *precision* برای هر دو مجموعه داده بالای ۹۹٪ است. در شکل ۵ عملکرد بخش تشخیص خطای الگوریتم در مجموعه داده *adult* که شامل انواع داده‌ای متفاوت است، نشان داده شده‌است. همان‌طور که در شکل مشخص است بهترین عملکرد در تشخیص خطا در نوع داده‌ای اسمی و بدترین عملکرد در نوع داده‌ای ترتیبی رخ داده است. در بدترین عملکرد که در نوع داده‌ای ترتیبی رخ داده است، مقدار *true negative rate* به‌طور متوسط ۹۰٪ است و در نوع داده‌ای اسمی بالای ۹۷٪ است. هم‌چنین مقدار *error rate* در نوع داده‌ای ترتیبی به‌طور متوسط ۱۱٪ و در نوع داده‌ای اسمی ۶٪ است. مقدار *precision* نیز در تمام انواع داده‌ای و در تمام نرخ‌های خطا بالاتر از ۹۸٪ است.

جدول ۱: بررسی مقادیر مختلف برای پارامتر δ

Error rate	true negative rate	δ
۳۱٪	۹۸٫۵٪	$\frac{n}{4}$
۹٪	۹۶٫۷٪	$\frac{n}{3}$
۱۵٪	۹۲٪	$\frac{n}{2}$
۳۰٪	۸۱٪	$\frac{2n}{3}$

پس از تشخیص خطا، با استفاده از یادگیری مرکب به تصحیح خطا پرداخته می‌شود. همان‌طور که در بخش پیش اشاره شد، از درخت تصمیم، طبقه‌بند بیز و شبکه عصبی MLP برای ساخت سیستم یادگیری مرکب استفاده شده‌است. برای ساخت درخت تصمیم برای هر ویژگی از روش CART استفاده شده، سپس درختان هرس می‌گردد، جدول ۲ و ۳ به ترتیب بیان می‌کند که درختان ساخته شده از هر ویژگی، در مجموعه داده *adult* و *protein* در کدام سطح هرس شده‌اند. برخی از درختان نیز بدون هرس باقی‌مانده که با خط تیره در جدول نشان داده شده‌است. برای ساخت شبکه عصبی MLP به ازای هر ویژگی همان‌طور که در بخش

کاربر اطلاعاتی از تعداد خطاها در یک رکورد ندارد، کاربر برای حداکثر تعداد ویژگی‌هایی که در یک رکورد ناصحیح است مقداری را تخمین می‌زند. لذا ممکن است یک ویژگی صحیح به‌عنوان غلط در نظر گرفته شده باشد، اما آزمایش‌ها نشان داده که چنان‌چه این رویداد پیش آید، سیستم یادگیری مرکب، مقداری بسیار نزدیک به مقدار اولیه برای آن ویژگی پیش‌بینی می‌نماید.

۵- آزمایش‌ها تجربی

در این بخش راهکار پیشنهادی بر روی دو مجموعه داده آزمایش می‌شود که مجموعه اول (*adult*) درباره درآمد افراد براساس اطلاعات سرشماری^{۲۳} و مجموعه دوم (*protein*) درباره ویژگی‌های نوعی پروتئین^{۲۴} است. مجموعه داده *adult* شامل ۴۸۸۴۲ رکورد و ۱۵ ویژگی است، که ۶ ویژگی عددی، ۱ ویژگی ترتیبی و ۸ ویژگی اسمی است. مجموعه داده *protein* شامل ۴۵۷۳۰ رکورد و ۹ ویژگی است که تمام ویژگی‌های آن عددی است. هدف از انتخاب این مجموعه داده‌ها این است که دو مجموعه داده‌ای هستند که از لحاظ انواع داده‌ای با یکدیگر کاملاً متفاوت می‌باشند، مجموعه *adult* شامل تمام انواع داده‌ای است درحالی‌که مجموعه *protein* فقط شامل یک نوع داده‌ای است. هم‌چنین این مجموعه داده‌ها از پرستفاده‌ترین منابع داده‌ای است [۳۲-۳۴]. در این مجموعه داده‌ها که فاقد خطا هستند، با استفاده از توزیع نرمال، ۵٪، ۱۰٪، ۱۵٪ و ۲۰٪ خطا ایجاد شده، و نحوه عملکرد با پنج معیاری که در فرمول ۱ تا ۵ آورده شده مورد ارزیابی قرار گرفته است.

$$\text{false alarm rate} = \frac{FP}{FP + TN} \quad (1)$$

$$\text{error rate} = \frac{FP + FN}{N + P} \quad (2)$$

$$\text{true negative rate} = \frac{TN}{TN + FP} \quad (3)$$

$$\text{recall} = \frac{TP}{TP + FN} \quad (4)$$

$$\text{precision} = \frac{TP}{TP + FP} \quad (5)$$

در ابتدای بخش ۴-۱ بیان شد که رکوردی که به ازای یک وابستگی تابعی با بیش از δ رکورد ناسازگار باشد، آن رکورد به‌عنوان رکورد مشکوک به خطا در نظر گرفته می‌شود. جدول ۱ مقادیر مختلف را برای این پارامتر بر روی دو مجموعه داده مورد بررسی قرار داده است. *true negative rate* و *error rate* برای دو مجموعه داده به‌صورت میانگین بیان شده‌است. با توجه به این دو پارامتر بهترین مقدار برای پارامتر δ برای دو مجموعه داده مورد بررسی در مقدار $\frac{n}{3}$ رخ داده است.

راهکار پیشنهادی از دو فاز تشخیص و تصحیح تشکیل شده، که در ابتدا به بررسی کارایی بخش تشخیص پرداخته می‌شود. همان‌طور که در شکل ۴ نشان داده شده مقدار *true negative rate* در تمام نرخ‌های خطا برای هر دو مجموعه داده بالای ۹۲٪ است، که نشان‌دهنده توانایی

در جدول ۴ و ۵ آورده شده است. در واقع این مقادیر پس از آزمایش‌ها حاصل شده، و پیش از آزمایش‌ها به ۵ شبکه عصبی مقادیر مختلف داده شده بود، که بهترین نتایج از این مقادیر حاصل گردید.

قبل بیان گردید، به دلیل عدم آگاهی از پارامترهای اولیه شبکه عصبی، از ۵ شبکه عصبی با پارامترهای اولیه متفاوت استفاده می‌گردد. بهترین مقادیر براساس آزمایش‌ها انجام شده برای هر ویژگی از مجموعه داده‌ها

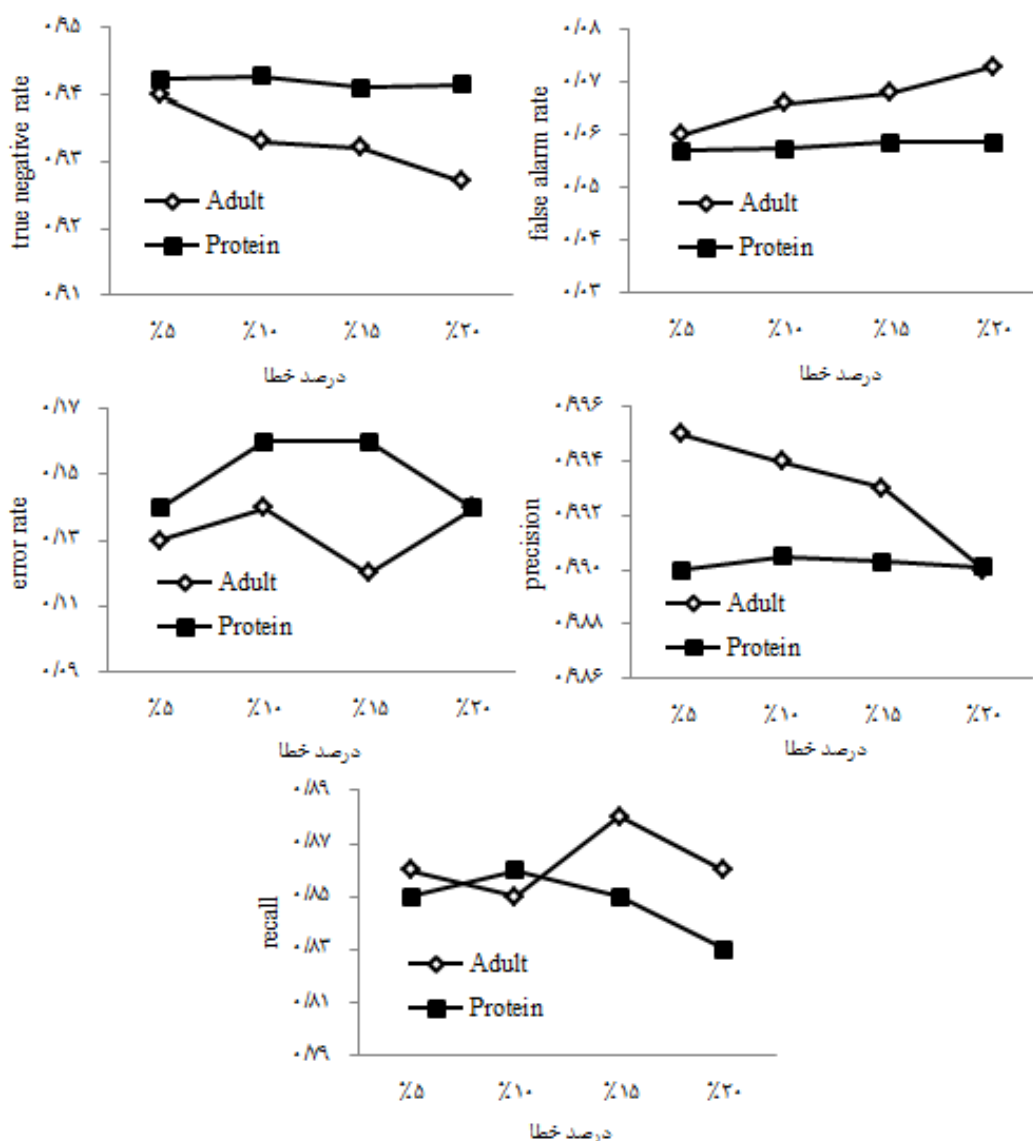
جدول ۲: سطح هرس شده برای هر درخت ساخته شده به ازای

ویژگی	۱	۲	۳	۴	۵	۶	۷	۸
سطح	۱۱۰	۹۵	۱۰۰	-	-	۶۰	۱۰۵	۸۰
ویژگی	۹	۱۰	۱۱	۱۲	۱۳	۱۴	۱۵	
سطح	۶۵	۷۵	-	-	۱۱۵	۲۸	۸۰	

جدول ۳: سطح هرس شده برای هر درخت ساخته شده به ازای

هر ویژگی برای مجموعه داده protein

ویژگی	۱	۲	۳	۴	۵
سطح	۱۴۵	۱۵۶	۱۰۰	۱۴۱	۹۵
ویژگی	۶	۷	۸	۹	
سطح	۱۶۴	۱۵۳	۱۴۶	۸۶	



شکل ۴: نحوه عملکرد بخش تشخیص خطا در الگوریتم پیشنهادی

به‌طور متوسط در مجموعه داده adult ۱۲٪ و در مجموعه داده protein نیز به‌طور متوسط ۶/۲٪ است. مقدار recall و precision نیز در این سیستم به‌طور میانگین در مجموعه داده adult ۹۳٪ و ۹۳/۷٪ و در مجموعه داده protein ۹۱/۵٪ و ۹۶/۸٪ است. مقدار error rate در مجموعه داده adult کمتر از ۱۰٪ و در مجموعه داده protein کمتر از ۹٪ است.

در شکل ۱۰ همانند شکل ۵ کارایی کل الگوریتم در مجموعه داده adult که شامل انواع داده‌ای مختلف است، نشان داده شده است. در این شکل همان ویژگی‌های موجود در شکل ۵ که کارایی بخش تشخیص را نشان داد، انتخاب شده است. عملکرد کل الگوریتم در نوع داده‌ای اسمی بهترین و در نوع داده‌ای عددی بدترین بوده است. در نوع داده‌ای اسمی میزان true negative rate به‌طور متوسط ۹۶٪ و recall ۹۷٪ است. در نوع داده‌ای عددی میزان true negative rate ۸۱٪ و recall ۹۷٪ است. همان‌طور که مشخص است میزان recall در کل الگوریتم نسبت به مرحله تشخیص خطا افزایش یافته است. علت این افزایش در این معیار این است که برخی از ویژگی‌هایی که در فاز تشخیص به‌عنوان خطا در نظر گرفته شدند، در مرحله تصحیح مقداری نزدیک یا مشابه مقدار اصلی به آن‌ها اختصاص داده شد و لذا باعث بهبود معیار recall شده است.

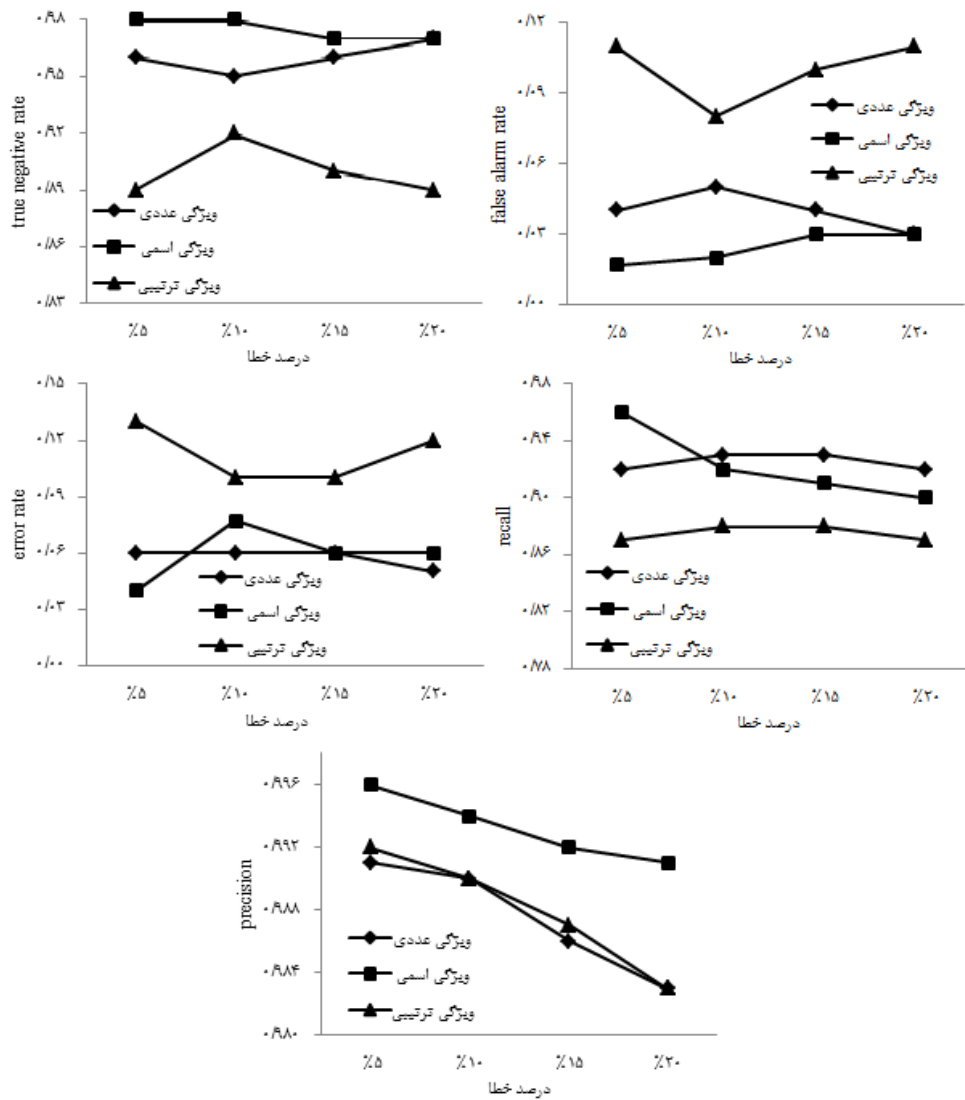
در شکل ۶، ۷ و ۸ نحوه عملکرد بخش تصحیح در طبقه‌بندی‌های مختلف نشان داده شده است. از میان این سه طبقه‌بند، در مجموعه داده adult، شبکه عصبی MLP و در مجموعه داده protein درخت تصمیم بهترین عملکرد را داشته‌اند. میزان true negative rate به ترتیب در مجموعه داده adult، ۸۲/۷٪ و در مجموعه داده protein، ۹۱/۵٪ محاسبه شده است. میزان false alarm rate و error rate نیز در این طبقه‌بندها در تمام نرخ‌های خطا در مجموعه داده adult، کمتر از ۲۰٪ و ۱۱٪ و مجموعه داده protein کمتر از ۱۰٪ و ۱۲٪ است. هم‌چنین مقدار precision و recall نیز در تمام نرخ‌های خطا در تمام طبقه‌بندها، در مجموعه داده adult بالاتر از ۹۱٪ و ۹۲٪ و در مجموعه داده protein بالاتر از ۸۸٪ و ۹۱٪ است. در شکل ۹ نیز عملکرد ترکیب این طبقه‌بندها در سیستم یادگیری مرکب که در واقع استراتژی تصحیح در راهکار پیشنهادی است، بیان شده است. در سیستم یادگیری از رأی اکثریت طبقه‌بندها به‌عنوان استراتژی ترکیب استفاده شده است. همان‌طور که در شکل ۹ مشخص است، میزان true negative rate، که دراقع بیان‌کننده توانایی سیستم یادگیری مرکب در تصحیح خطاها است، در تمام نرخ‌های خطا در مجموعه داده adult بالای ۸۶٪ و در مجموعه داده protein بالای ۹۳٪ است. هم‌چنین میزان false alarm rate در سیستم یادگیری مرکب

جدول ۴: پارامترهای ورودی شبکه عصبی MLP برای مجموعه داده Adult

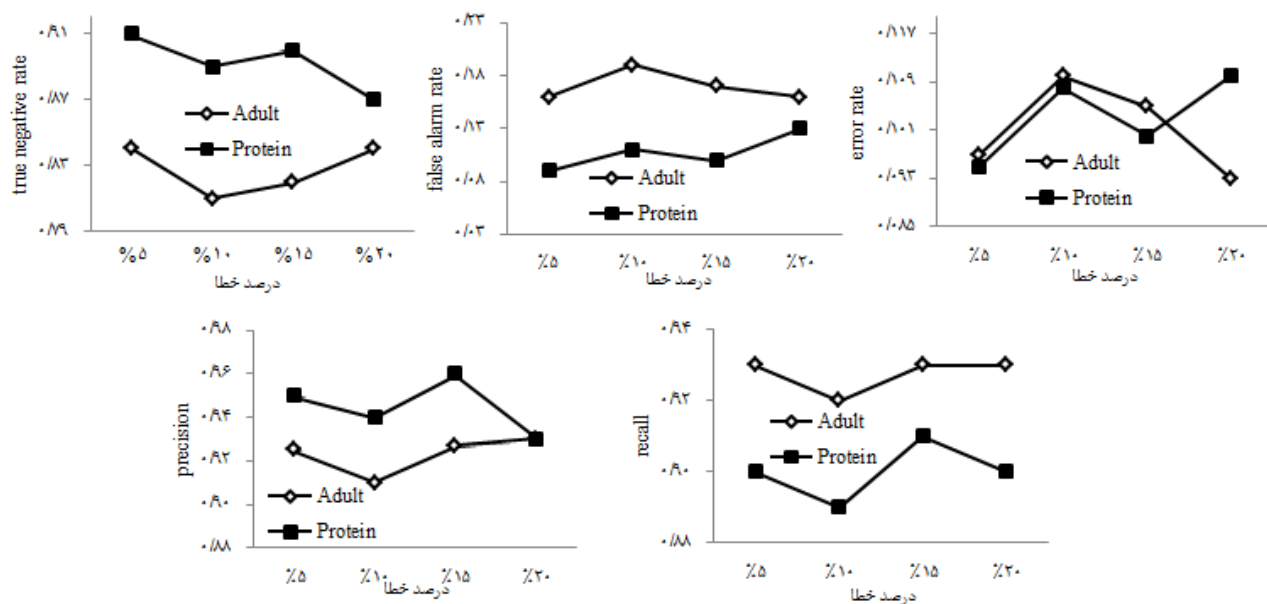
ویژگی	۱	۲	۳	۴	۵	۶	۷	۸	۹	۱۰	۱۱	۱۲	۱۳	۱۴	۱۵
ضریب یادگیری	۰/۱۴	۰/۱	۰/۱۲	۰/۱	۰/۱	۰/۱۲	۰/۱۴	۰/۱۴	۰/۱۲	۰/۱۴	۰/۱	۰/۱	۰/۱	۰/۱۱	۰/۱۲
Momentum	۰/۹	۰/۸	۰/۷	۰/۹	۰/۸	۰/۹	۰/۶	۰/۹	۰/۷	۰/۷	۰/۶	۰/۹	۰/۹	۰/۸	۰/۸
Epoch	۱۰	۱۰	۸	۵	۵	۱۰	۸	۵	۱۰	۵	۱۰	۱۰	۸	۸	۵
تعداد دسته‌ها	۵	۸	۵	۵	۵	۵	۵	۵	۵	۵	۵	۵	۵	۵	۵

جدول ۵: پارامترهای ورودی شبکه عصبی MLP برای مجموعه داده Protein

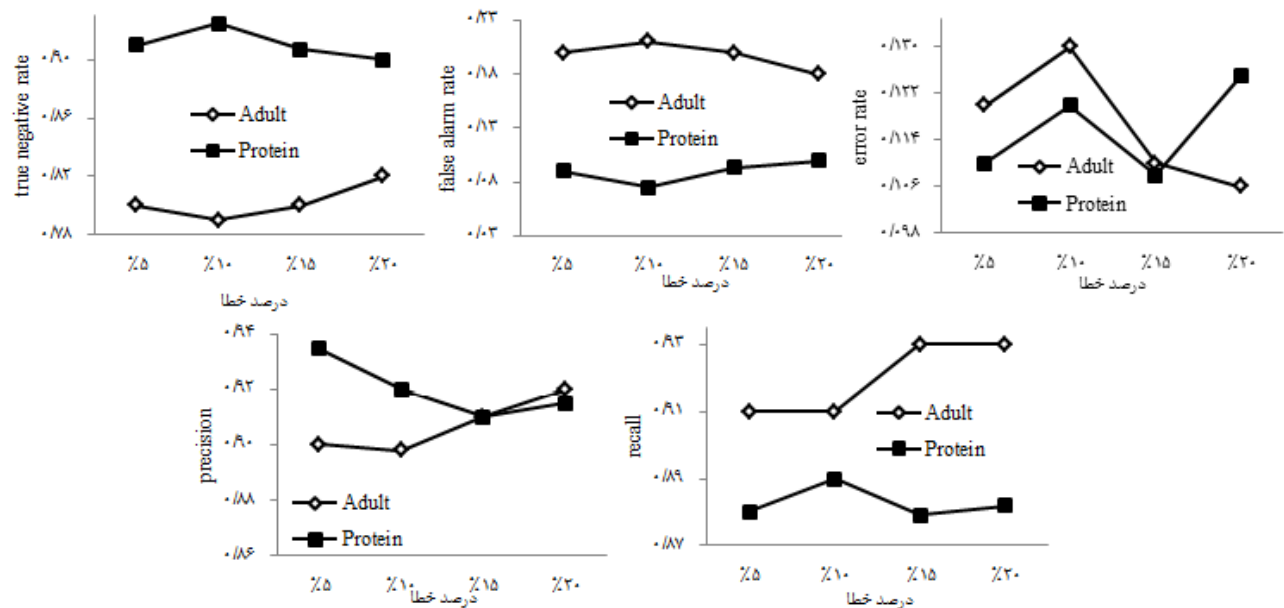
ویژگی	۱	۲	۳	۴	۵	۶	۷	۸	۹
ضریب یادگیری	۰/۲۳	۰/۲۱	۰/۲۴	۰/۱۲	۰/۲	۰/۲۳	۰/۲	۰/۱	۰/۲۱
Momentum	۰/۷	۰/۶	۰/۶	۰/۷	۰/۶	۰/۶	۰/۸	۰/۹	۰/۸
Epoch	۵	۸	۱۲	۲۰	۵۰	۱۰	۸	۲۰	۲۵
تعداد دسته‌ها	۱۰	۱۰	۵	۷	۵	۵	۶	۸	۹



شکل ۵: نحوه عملکرد الگوریتم پیشنهادی در بخش تشخیص خطا در انواع داده‌ای متفاوت در مجموعه داده adult



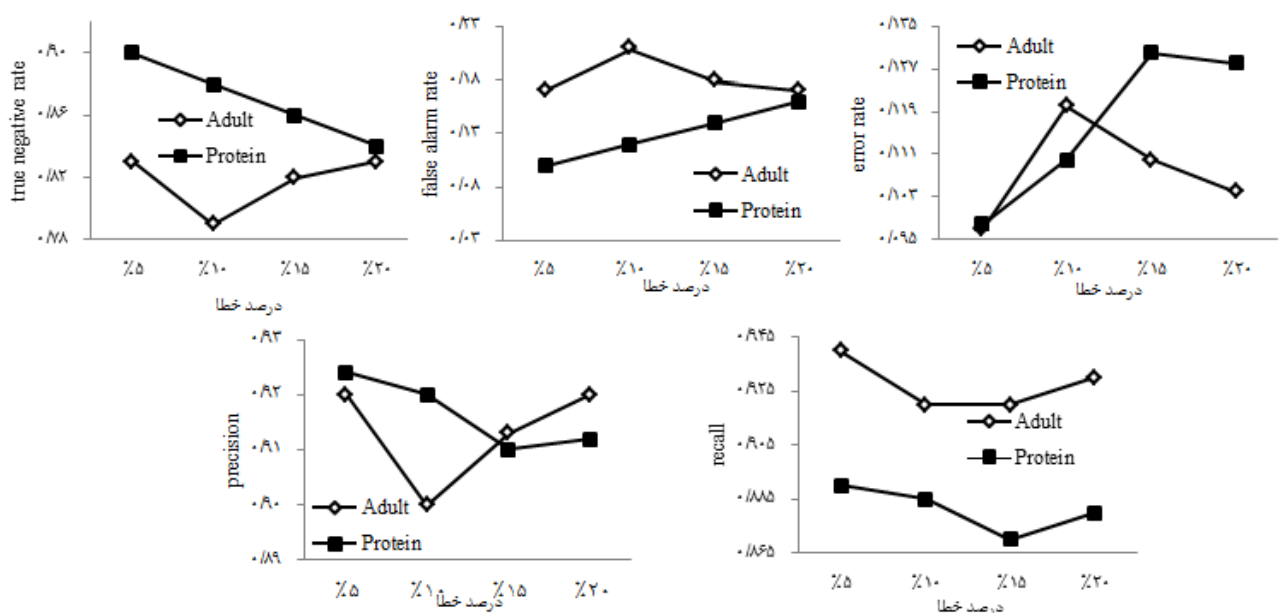
شکل ۶: نحوه عملکرد الگوریتم پیشنهادی در مرحله تصحیح خطا در شبکه عصبی MLP



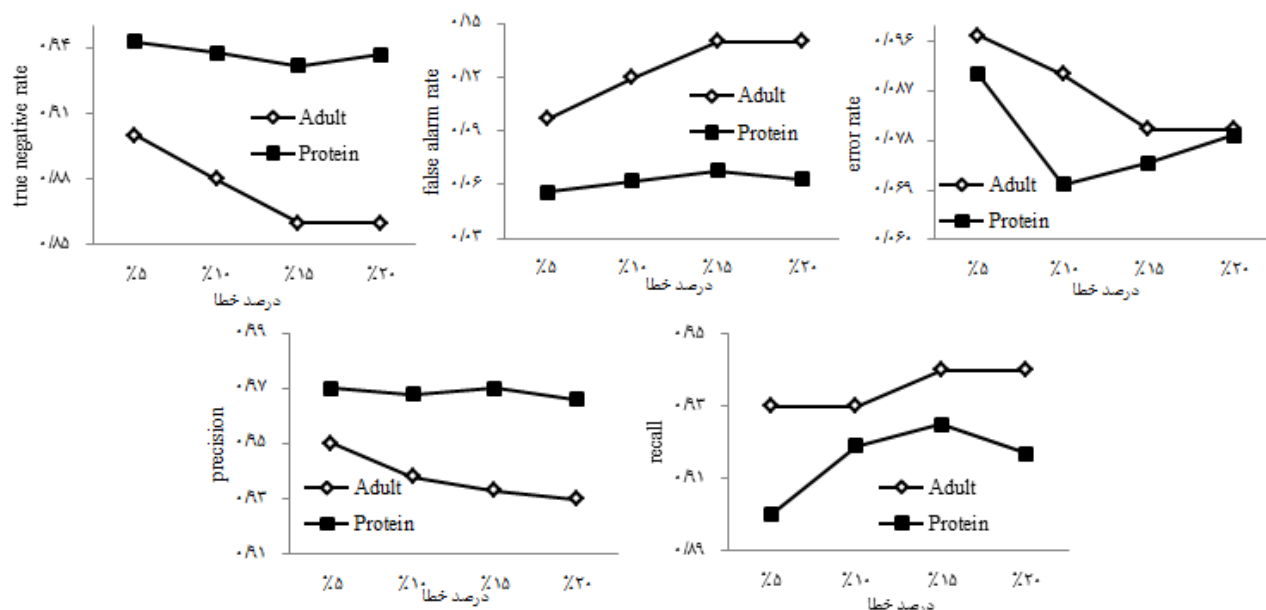
شکل ۷: نحوه عملکرد الگوریتم پیشنهادی در مرحله تصحیح خطا در درخت تصمیم

در شکل ۱۲ راهکار پیشنهادی با روش [۱۲] و [۲۸] که با استفاده از وابستگی تابعی به صورت خودکار به تصحیح داده‌ها پرداخته‌است، مورد مقایسه قرار گرفته‌است. هدف از انتخاب روش [۱۲] برای مقایسه این است که برخی از روش‌های پیشین که مبتنی بر وابستگی بودند [۲۱]، [۲۲] برای مقایسه روش پیشنهادی خود از این روش استفاده نموده‌اند. براساس آزمایش‌ها میزان true negative rate در روش پیشنهادی به طور متوسط ۷٪ بالاتر از روش تصحیح با وابستگی تابعی [۱۲] است. هم‌چنین true negative rate در راهکار پیشنهادی به طور متوسط ۳٪ بالاتر از راهکار تصحیح [۲۸] است.

در شکل ۱۱ کارایی کل الگوریتم بر روی هر دو مجموعه داده مورد بررسی نشان داده شده‌است. میزان recall در کل الگوریتم به طور متوسط در مجموعه داده adult ۹۴٪ و در مجموعه داده protein به طور متوسط ۹۱/۵٪ است. مقدار error rate در تمام نرخ‌های خطا در مجموعه داده adult کمتر از ۱۰٪ و در مجموعه داده protein کمتر از ۹٪ است. مقدار true negative rate به طور متوسط در مجموعه داده adult ۸۳٪ و در مجموعه داده protein ۸۹/۸٪ است. میزان false alarm rate در تمام نرخ‌ها در مجموعه داده adult کمتر از ۱۷٪ و در مجموعه داده protein کمتر از ۱۱٪ است. مقدار precision به طور میانگین در مجموعه داده adult ۹۰/۹٪ و در مجموعه داده protein ۹۴/۴٪ است.



شکل ۸: نحوه عملکرد الگوریتم پیشنهادی در مرحله تصحیح خطا در طبقه بند بیز

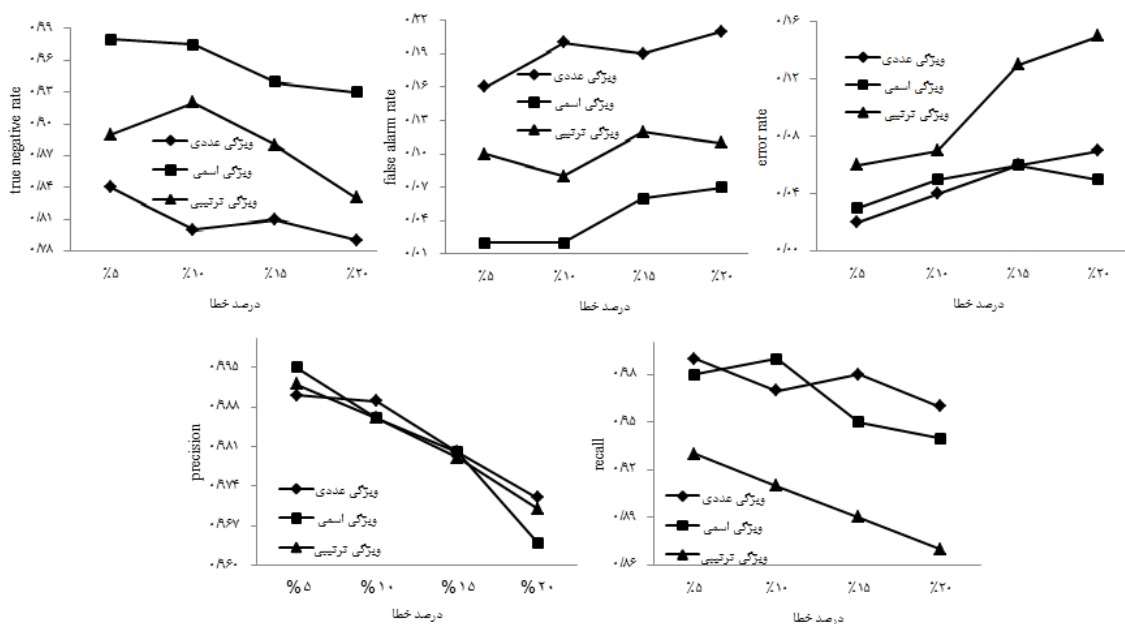


شکل ۹: نحوه عملکرد الگوریتم پیشنهادی در مرحله تصحیح خطا در سیستم یادگیری مرکب

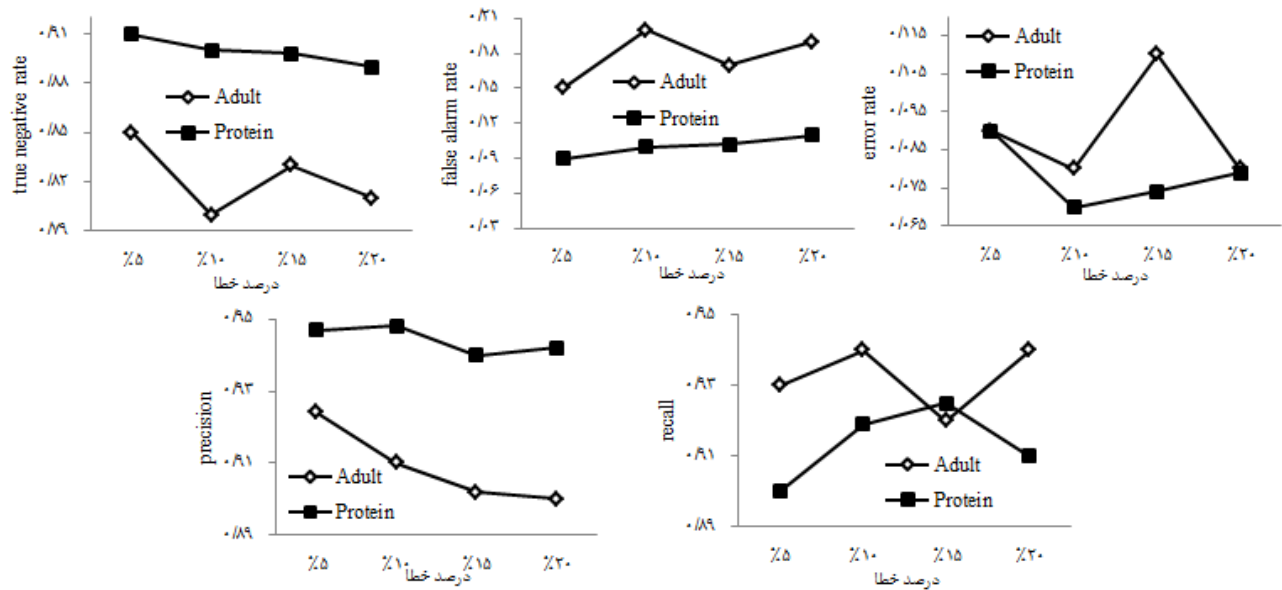
۶- نتیجه‌گیری

مرکب اصلاح گردید. سیستم یادگیری مرکب مورد استفاده در راهکار پیشنهادی از سه طبقه‌بند بیز، درخت تصمیم و شبکه عصبی MLP تشکیل شده و استراتژی ترکیب آرای طبقه‌بندها رأی اکثریت است. بنابر آزمایش‌ها انجام‌شده کارایی روش پیشنهادی در تشخیص خطا، به‌طور متوسط ۹۳/۷٪ و در بخش تصحیح خطا به‌طور متوسط ۹۰/۶٪ است. همچنین روش پیشنهادی با دو روش مشابه مبتنی بر وابستگی تابعی مورد مقایسه قرار گرفت که روش پیشنهادی ۵٪ توانایی بالاتری در تصحیح خطا نسبت به دو روش مورد مقایسه دارا است.

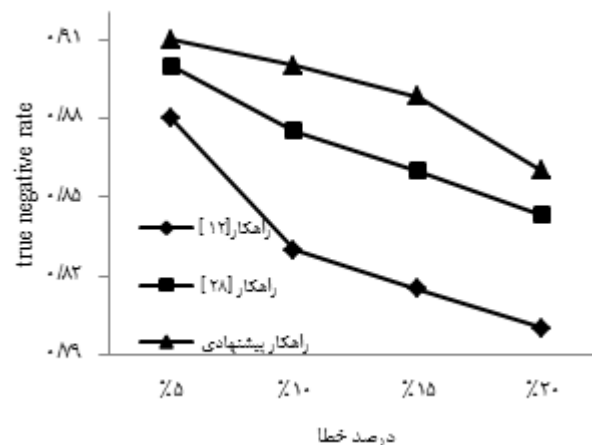
در این مقاله یک راهکار خودکار برای تشخیص و تصحیح فیلد حاوی خطا پیشنهاد شد. راهکار پیشنهادی توانایی تصحیح خطا بر روی مجموعه داده با انواع داده‌های مختلف را دارد. همچنین راهکار پیشنهادی توانایی تصحیح چندین خطا در یک رکورد را دارا است. در روش ترکیبی پیشنهادی، برای تشخیص خطا از وابستگی تابعی استفاده شد. در فاز تصحیح نیز فیلد خطا دار تشخیص داده‌شده، به‌وسیله سیستم یادگیری



شکل ۱۰: نحوه عملکرد کل الگوریتم پیشنهادی بر روی انواع داده‌ای متفاوت در مجموعه داده adult



شکل ۱۱: نحوه عملکرد کل الگوریتم در نرخ‌های متفاوت خطا



شکل ۱۲: مقایسه راهکار پیشنهادی با راهکار [۱۲] و [۲۸] در نرخ‌های متفاوت خطا

مراجع

- [1] G. Rahman and Z. Islam, "Missing value imputation using decision trees and decision forests by splitting and merging records: Two novel techniques," *Knowledge-Based Systems*, vol. 53, pp. 51 – 65, 2013.
- [2] W. Fan, F. Geerts, X. Jia and A. Kementsietsidis, "Conditional functional dependencies for capturing data inconsistencies," *ACM Transactions on Database Systems (TODS)*, vol. 33, no. 2, pp. 1-48, 2008.
- [3] P. H. Williams, C. R. Margules, and D. W. Hilbert, "Data requirements and data sources for biodiversity priority area selection," *Journal of biosciences*, vol. 27, no. 4, pp. 327–338, 2002.
- [4] G. Rahman and Z. Islam, "Decision Tree-based Missing Value Imputation Technique for Data Pre-processing," *Proceedings of the 9th Conference on Australasian Data Mining*, pp. 41-50, Ballarat, Australia, 2011.
- [5] X. Chu and I. F. Ilyas, "Qualitative Data Cleaning," *The VLDB Journal*, vol. 9, no. 13, pp. 1605 - 1608, 2016.
- [6] L. Breiman, "Bagging predictors," *Machine Learning*, vol. 24, no. 2, pp.123-140, 1996.
- [7] M. Yakout, L. Berti-Équill, and A. K. Elmagarmid, "Don't be SCARED: Use SCalable Automatic REpairing with Maximal Likelihood and Bounded Changes", *Proceedings of the 13th International Conference on Management of Data*, pp. 553-564, 2013.
- [8] N. Tang, "Big Data Cleaning", *Proceedings of the 16th International Conference on Web Technologies and Applications*, pp.13-24, 2014.
- [9] J. Hipp, U. Güntzer, and U. Grimmer, "Data quality mining-making a virtue of necessity," *Proceedings of the 6th Workshop on Research Issues in Data Mining and Knowledge Discovery (DMKD)*, pp. 52-57, 2001.
- [10] S. Bruggemann, "Rule Mining for Automatic Ontology Based Data Cleaning," *Proceedings of the 10th International Conference on Asia-Pacific Web Conference Progress*, pp. 522-527, 2008.
- [11] C. He, Z. Tan, Q. Chen, C. Sha, Z. Wang, and W. Wang, "Repair Diversification for Functional Dependency Violations," *Proceedings of the 19th International Conference on Database Systems for Advanced Applications*, pp. 468-482, 2014.

- International Conference on Data Engineering*, pp. 244-255, 2014.
- [23] G. Beskales, I. F. Ilyas, L. Golab and A. Galiullin, "Sampling from Repairs of Conditional Functional Dependency Violations," *The VLDB Journal*, vol. 23, no. 1, pp. 103 - 128, 2014.
- [24] X. Chu, J. Morcos, I. F. Ilyas, M. Ouzzani, P. Papotti, N. Tang and Y. Ye, "KATARA: A Data Cleaning System Powered by Knowledge Bases and Crowdsourcing," *Proceedings of the 15th International Conference on Management of Data*, pp. 1247 - 1261, 2015.
- [25] W. Fan, S. Ma, N. Tang and W. Yu, "Interaction Between Record Matching and Data Repairing," *J. Data and Information Quality*, vol. 4, no. 4, 2014.
- [26] J. Segeren, D. Gairola and F. Chiang, "CONDOR: A System for CONstraint DiscOverY and Repair," in *Proceedings of the 23rd International Conference on Information and Knowledge Management*, pp. 2078 - 2089, 2014.
- [27] S. Song, H. Zhu and J. Wang, "Constraint-Variance Tolerant Data Repairing," *Proceedings of the 16th International Conference on Management of Data*, pp. 877 - 892, 2016.
- [28] F. Chiang and S. Sitaramachandran, "Unifying Data and Constraint Repairs," *J. Data and Information Quality*, vol. 7, no. 3, pp. 1 - 26, 2016.
- [29] J. Han, M. Kamber and J. Pei, *Data Mining Concept And Techniques*, Morgan Kaufmann Publishers, 3 Edition, 2011.
- [30] H. Yao and H. J. Hamilton, "Mining functional dependencies from data," *Data Mining and Knowledge Discovery*, vol. 16, no. 2, pp. 197-219, 2008.
- [31] Y. Huhtala, J. Kärkkäinen, P. Porkka and H. Toivonen, "Tane: An Efficient Algorithm for Discovering Functional and Approximate Dependencies," *The Computer Journal*, vol. 42, no. 2, pp. 100-111, 1999.
- [32] P. Kontkanen, J. Lahtinen, P. Myllymaki, T. Silander and H. Tirri, "Pre- and Post-processing in Machine Learning and Data Mining: Theoretical Aspects and Applications," *The workshop within Machine Learning and Applications Complex Systems Computation Group (CoSCo)*, pp. 38-47, 1999.
- [33] R. Agrawal and R. Ikant and D. Thomas, "Privacy Preserving OLAP," *Proceedings of the 5th International Conference on Management of Data*, pp. 251-262, 2005.
- [34] I. Yoncaci, "Maximum a Posteriori Tree Augmented Naive Bayes Classifiers," *Lecture Notes in Artificial Intelligence*, 2004.
- [35] S. Haykin, *Neural Networks: A Comprehensive Foundation*, 2nd. Englewood Cliffs, NJ: Prentice-Hall, 1999.
- [12] F. Chiang and R. J. Miller, "A Unified Model for Data and Constraint Repair," *Proceedings of the 27th International Conference on Data Engineering*, pp. 46-457, 2011.
- [۱۳] الناز زعفرانی معطر، محمدرضا فیضی درخشی و آزاده روحانی، «تشخیص هوشمند و خودکار غلط‌های تایپی در پایگاه داده‌های بزرگ بدون استفاده از لغت نامه»، مجله علمی پژوهشی مهندسی برق دانشگاه تبریز، جلد ۴۷، شماره ۱، صفحات ۸۱-۹۱، ۱۳۹۶.
- [۱۴] مرتضی خرم کشکولی و مریم دهقانی، «تشخیص، شناسایی و جداسازی عیب توربین گاز پالایشگاه دوم پارس جنوبی با استفاده از روش‌های ترکیبی داده کاوی، 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.
- [۱۵] یوکابد صدری، علی آقاگلزاده و مهدی ازوجی، «ادغام تصویر چندفوکوسه با استفاده از همدوسی فاز و خوشه‌بند K-means»، مجله علمی پژوهشی مهندسی برق دانشگاه تبریز، جلد ۴۵، شماره ۴، صفحات ۱۱۵-۱۲۵، ۱۳۹۴.
- [16] M. Fijałkowska and B. Antoszewski, "Classification of congenital nasal deformities: a proposal to amend the existing classification," *European Archives of Oto-Rhino-Laryngology*, vol. 274, no. 3, pp. 1231-1235.
- [17] C. M. Teng, "Correcting Noisy Data," *Proceedings of the 16th International Conference on Machine Learning*, pp. 239-248, 1999.
- [18] C. M. Teng, "A Comparison of Noise Handling Techniques," *Proceedings of the 14th International Conference on Artificial Intelligence Research Society*, pp. 269-273, 2001.
- [19] C. M. Teng, "Polishing Blemishes: Issues in Data Correction," *Intelligent Systems*, vol. 19, no. 2, pp. 34-39, 2004.
- [20] S. Kolahi and L. Lakshmanan, "On Approximating Optimum Repairs for Functional Dependency Violations," *Proceedings of the 12th International Conference on Database Theory*, pp. 53-62, 2009.
- [21] G. Beskales, I. F. Ilyas, L. Golab and A. Galiullin, "On the Relative Trust between Inconsistent Data and Inaccurate Constraints," *Proceedings of the 29th International Conference on Data Engineering*, pp. 541-552, 2013.
- [22] M. Volkovs, F. Chiang, J. Szlichta and R. J. Miller, "Continuous data cleaning," *Proceedings of the 30th*

زیرنویس‌ها

^{۱۰} Ontology

^{۱۱} Functional dependency

^{۱۲} Semi-automatic

^{۱۳} Data type

^{۱۴} Ensemble learning

^{۱۵} Decision tree

^{۱۶} Bayes theorem

^{۱۷} Knowledge base

^{۱۸} Crowd powered

^۱ Missing value

^۲ Data warehouse

^۳ Outlier

^۴ Inconsistent

^۵ Duplicate

^۶ Data quality

^۷ Data cleaning

^۸ Conditional functional dependency

^۹ Rules

^{۱۹} Categorical

^{۲۰} Discretization

^{۲۱} Gaussian distribution

^{۲۲} Prior probability distribution

^{۲۳} <https://archive.ics.uci.edu/ml/datasets/Adult>

^{۲۴} <http://archive.ics.uci.edu/ml/datasets/Physicochemical+Properties+of+Protein+Tertiary+Structure>