

مدل سازی متغیرهای اثرگذار برای کشف تقلب در صورت های مالی با استفاده از تکنیک های داده کاوی

شکراه خواجوی *

مهرداد ابراهیمی **

تاریخ پذیرش: ۱۳۹۶/۰۱/۱۹

تاریخ دریافت: ۱۳۹۵/۰۷/۱۷

چکیده

در این پژوهش به بررسی این مسئله پرداخته می شود که آیا می توان از طریق شناسایی و انتخاب متغیرهای اثرگذار در کشف تقلب در صورت های مالی و با به کارگیری تکنیک های داده کاوی مدلی برای کشف تقلب در صورت های مالی شرکت های پذیرفته شده در بورس اوراق بهادار تهران ارائه کرد؟ برای پاسخگویی به این سؤال از یک نمونه تطبیقی متشکل از ۶ شرکت متقلب و ۶۴ شرکت غیر متقلب (نمونه نامتوازن) به همراه تکنیک های طبقه بندی شبکه عصبی مصنوعی، شبکه بیزین و الگوریتم جنگل تصادفی طی سال های ۱۳۷۹-۱۳۹۳ استفاده گردید. برای حل مشکل نمونه نامتوازن نیز رویکرد معرفی شده توسط هی و همکاران (۲۰۰۸) بکار رفت. متغیرهای پیش بین، عوامل مالی و غیرمالی خطر تقلب مرتبط با گزارشگری مالی متقلبان ه هستند که با مرور پژوهش های انجام شده در این زمینه تعیین خواهند شد. یافته های پژوهش بیان گر وجود شواهدی دال بر عملکرد مناسب مدل های پیشنهادی و برتری الگوریتم جنگل تصادفی و شبکه بیزین برای پیش بینی تقلب در صورت های مالی است. نتایج حاصل از انتخاب ویژگی به روش مبتنی بر همبستگی حاکی از سودمندی متغیرهای نسبت پوشش بهره، نسبت حساب های دریافتنی به کل دارایی ها، نسبت موجودی کالا به فروش خالص، نسبت نقدی، لگاریتم طبیعی فروش، نسبت سود خالص به فروش و نسبت جمع دارایی های جاری به کل دارایی ها برای کشف تقلب بود.

* استاد گروه حسابداری، دانشگاه شیراز، شیراز، ایران (نویسنده مسئول)

Email: shkhajavi@gmail.com

** دانشجوی دکتری حسابداری، دانشگاه شیراز، شیراز، ایران

Email: mehrdadebrahimi66@gmail.com

واژه های کلیدی: کشف تقلب، تقلب در صورت های مالی، داده کاوی، بورس اوراق بهادار تهران.

۱- مقدمه

حسابداری یک فعالیت خدمتاتی است که هدف آن ارائه اطلاعات مفید مرتبط با واحدهای اقتصادی به تصمیم گیرندگان درون سازمانی و برون سازمانی است. از این رو، حسابداری باید فعالیت های یک واحد اقتصادی را به منظور گزارشگری مؤثر به ذینفعان، شناسایی، تجزیه و تحلیل و به شکل مناسبی ثبت کند (سگال^۱، ۲۰۱۶). اطلاعات ارائه شده درباره وضعیت مالی و عملکرد یک شرکت، اهمیت زیادی برای استفاده کنندگان از صورت های مالی دارد زیرا مبنایی برای تصمیم گیری و تخصیص سرمایه است؛ بنابراین، قابلیت اتکا و شفافیت در گزارشگری مالی به ارائه درست و صادقانه دستاوردهای مالی شرکت منجر خواهد شد و نقش مهمی در پایداری سیستم مالی دارد. با این حال، سقوط شرکت های بسیاری (مانند انرون^۲، وردکام^۳ و گلوبال کراسینگ^۴) اثربخشی حاکمیت شرکتی، کیفیت گزارش های مالی و قابلیت اعتماد کارکردهای حسابرسی را با تردید همراه ساخت (رضایی، ۲۰۰۵). صورت های مالی متقلبانه اثرات منفی بر اقتصادهای دنیا داشته و به زیان های مالی قابل توجهی برای افراد و شرکت ها منجر شده است. از این رو، در یک محیط تجاری مبتنی بر فناوری که با تغییرات سریع همراه است، نیاز مبرمی به روش های مؤثر برای پیشگیری و کشف تقلب وجود دارد. روش های داده کاوی^۵ با به کارگیری موارد قبلی تقلب، مدلی برای شناسایی و تشخیص خطر تقلب ارائه می دهد که قادر به پیشگیری و کشف تقلب در صورت های مالی است (سگال، ۲۰۱۶).

به عبارت دیگر، دانشگاهیان و موسسه های حسابرسی به منظور بهبود کیفیت، صداقت و قابلیت اتکای فرآیند گزارشگری مالی به دنبال راه هایی برای کشف تقلب در صورت های مالی هستند. در حالی که پژوهش های دانشگاهی بسیاری برای استفاده از رویکردهای کمی در سایر زمینه های تقلب مانند بیمه^۶، کارت اعتباری^۷ و ارتباطات^۸، انجام شده، تلاش اندکی برای استفاده از تکنیک های داده کاوی جهت کشف تقلب در صورت های مالی صورت گرفته است؛ بنابراین، در

1 Segal

2 Enron

3 WorldCom

4 Global Crossing

5 Data Mining

6 Insurance Fraud

7 Credit Card Fraud

8 Telecommunication Fraud

این پژوهش به بررسی این مسئله پرداخته می‌شود که آیا می‌توان با به‌کارگیری تکنیک‌های داده‌کاوی و از طریق شناسایی متغیرهای مؤثر در کشف تقلب در صورت‌های مالی، مدلی برای کشف تقلب در صورت‌های مالی شرکت‌های پذیرفته‌شده در بورس اوراق بهادار تهران ارائه کرد؟ این پژوهش با معرفی متغیرهای پیشنهادی و تعیین متغیرهایی که شرکت‌های متقلب را به بهترین شکل از شرکت‌های غیر متقلب تفکیک می‌کند، متون حسابداری و حسابرسی را بهبود می‌بخشد.

از این‌رو، پس از تعیین شرکت‌های متقلب، مطابق با مبانی نظری و پیشینه پژوهش از تکنیک‌های انتخاب ویژگی و متغیرهای مرتبط با اهرم مالی، سودآوری، ترکیب دارایی‌ها، نقدینگی، کارایی، اندازه، رشد، وضعیت مالی کلی، اندازه حسابرس، دوره تصدی حسابرس و تغییر حسابرس برای انتخاب متغیرهای اثرگذار جهت پیش‌بینی تقلب در صورت‌های مالی استفاده می‌گردد. در ادامه به‌منظور پیش‌بینی تقلب در صورت‌های مالی از تکنیک‌های داده‌کاوی شبکه عصبی مصنوعی، شبکه بیزین و الگوریتم جنگل تصادفی استفاده خواهد شد.

۲- مبانی نظری

طی ۱۵۰ سال گذشته یا فراتر از آن، سازمان‌های تجاری از مالک - مدیر یا واحدهای تجاری خانوادگی، با تعداد اندک کارکنان به شرکت‌های چندملیتی گسترده‌ای که هزاران نفر را در استخدام دارند، تغییر شکل داده‌اند. این رشد در سایه وجوهی که از بازارهای سرمایه تأمین مالی گردیده امکان‌پذیر و به‌گذر مدیریت از شکل مالکان - سهامدار به گروه کوچکی از مدیران حرفه‌ای منجر شده است. از این‌رو، مدیران می‌بایست نتایج عملکرد خود را به مالکان واحد تجاری و سایر تأمین‌کنندگان وجوه مانند بانک‌ها و اعتباردهندگان گزارش کنند. باین‌حال، گزارشگری مالی شامل طیف وسیعی از اعداد حسابداری برای نشان دادن سود، جریان نقدی و وضعیت مالی واحد تجاری است؛ بنابراین، مدیریت قادر خواهد بود با استفاده نادرست از حسابداری ساختگی یا متهورانه^۱ وضعیت مالی شرکت را بهتر از واقعیت نشان دهد و از اثرات آن در کوتاه‌مدت بهره‌برد. هرچند ممکن است عملکرد مالی یک شرکت در سال مالی جاری به دلیل استفاده از حسابداری ساختگی قابل‌تحسین باشد، اما درنهایت اخبار بد آشکار خواهند شد. از این قبیل می‌توان به برخی رسوایی‌های اخیر مانند انرون، وردکام، آدلفیا^۲، آهولد^۳ و پارمالات^۴ و سایر موارد

1 Creative Or Aggressive Accounting

2 Adelphia

3 Ahold

4 Parmalat

اشاره کرد. ارائه عمدی اطلاعات مالی ناصحیح و انتشار صورت های مالی نادرست در بازارهای سرمایه توسط مدیران این شرکت ها، پیامدهای مالی مهمی مانند زیان سرمایه گذاری ها، از دست دادن شغل و از بین رفتن اعتماد در بازارهای سرمایه را به همراه داشته است (اکیک^۱، ۲۰۱۱). درک کامل از ماهیت، اهمیت و پیامدهای فعالیت های مرتبط با گزارشگری مالی متقلبانه مستلزم تعریف مناسبی از تقلب در صورت های مالی است. عمدتاً به این دلیل که حرفه حسابداری تنها در طول یک دهه گذشته از واژه تقلب در اعلامیه های حرفه ای خود استفاده کرده است، یافتن تعریف دقیقی از تقلب صورت های مالی در اعلامیه ها و بیانیه های الزام آور دشوار است. انجمن حسابداران رسمی آمریکا^۲ در بیانیه شماره ۸۲ استانداردهای حسابرسی، تقلب در صورت های مالی را به صورت حذف یا ارائه نادرست در صورت های مالی به صورت عمد تعریف می کند (رضایی و ریلی، ۲۰۱۰). مطابق با استاندارد حسابرسی ۲۴۰، تقلب عبارت است از هرگونه اقدام عمدی یا فریبکارانه یک یا چند نفر از مدیران، کارکنان یا اشخاص ثالث، برای برخورداری از یک مزیتی ناروا یا غیرقانونی. هرچند تقلب یک مفهوم گسترده ای دارد، اما آنچه به حسابرس مربوط می شود، اقدامات متقلبانه ای است که به تحریف با اهمیت در صورت های مالی می انجامد. هدف برخی تقلبات ممکن است تحریف صورت های مالی نباشد (کمیتة تدوین استانداردهای حسابرسی، بخش ۲۴۰: بند ۴).

طرح های تقلب در صورت های مالی به طور معمول به صورت زیر است:

- بیش نمایی دارایی ها یا درآمدها
- کم نمایی بدهی ها یا هزینه ها

بیش نمایی دارایی ها و درآمدها با احتساب مبالغ ساختگی بهای تمام شده به حساب دارایی ها یا از طریق شناسایی درآمدهای مصنوعی، وضعیت مالی یک شرکت را بهتر نشان می دهد. کم نمایی بدهی ها و هزینه ها با عدم شناسایی هزینه ها یا تعهدات مالی صورت می گیرد. هر دو این روش ها به افزایش حقوق مالکانه و خالص ثروت شرکت منجر می شوند. دست کاری اعداد و ارقام صورت های مالی افزایش سود هر سهم یا منافع در سود تضامنی را به همراه دارد یا تصویر پایداری از وضعیت واقعی شرکت را نشان می دهد (راهنمای میزان تقلب^۳، ۲۰۱۱).

به منظور درک کامل بیش نمایی دارایی ها و درآمدها و کم نمایی بدهی ها و هزینه ها، طرح های تقلب مورد استفاده برای بهبود صورت های مالی به پنج طبقه دسته بندی می شوند. با توجه به

1 Okike

2 American Institute of Certified Public Accountants

3 Fraud Examiners Manual

این که نگهداری سوابق مالی مستلزم استفاده از سیستم حسابداری دوطرفه است، ثبت‌های حسابداری متقالبانه به‌طور معمول حداقل بر دو حساب و بنابراین بر دودسته از صورت‌های مالی اثر خواهد گذاشت؛ بنابراین، یک معامله متقالبانه همیشه طرف دیگری خواهد داشت. طرح‌های تقلب می‌توانند دربرگیرنده ترکیبی از چندین روش باشند. طرح‌های تقلب در صورت‌های مالی به پنج طبقه به شرح زیر دسته‌بندی می‌شوند (راهنمای ممیزان تقلب، ۲۰۱۱).

- درآمدهای ساختگی
- تفاوت‌های زمانی
- ارزشیابی نادرست دارایی‌ها
- بدهی‌ها و هزینه‌های پنهان
- افشای نامناسب

اگرچه مسئولیت کشف تقلب و اشتباه با مدیریت و افراد عهده‌دار راهبری شرکت است اما استانداردهای جدید حسابداری بین‌الملل (استاندارد حسابرسی ۲۴۰) از حساب‌رسان انتظار دارد که تا حدی مسئولیت ثانویه کشف تقلب و اشتباه در صورت‌های مالی را بپذیرند. حساب‌رسان ملزم هستند که حسابرسی خود را با نگرش تردید حرفه‌ای و ذهنی پرسشگر برنامه‌ریزی و انجام دهند، یعنی ممکن است شرايطی وجود داشته باشد که می‌تواند به ارائه نادرست بااهمیت در صورت‌های مالی منجر گردد. ناکامی در کشف گزارشگری متقالبانه در انجام کار حسابرسی می‌تواند حساب‌رسان را در معرض پیامدهای حقوقی و قانونی نامساعدی قرار دهد. این امر به هزینه‌های دعاوی حقوقی قابل توجه برای حساب‌رسان و صدمات جبران‌ناپذیر به شهرت آن‌ها منجر می‌شود.

امروزه با تکامل تکنولوژی و شبکه‌های ارتباطی پرسرعت و جهانی، ارتکاب تقلب آسان‌تر و در نتیجه کشف آن دشوارتر گردیده است. اعمال متقالبانه به زیان‌های مالی بزرگی برای ایالت‌ها و واحدهای تجاری در سطح جهان منجر شده است. بنابراین در زمان کنونی، کشف و مبارزه با تقلب جهت جبران زیان‌های وارده بسیار بااهمیت است. علاوه بر این، با تکامل تکنولوژی شکل‌های جدیدی از تقلب مانند تقلب ارتباطات نفوذ کامپیوتری^۱، بیمه و کارت‌های اعتباری به وجود آمده است. هر واحد تجاری ویژگی‌های منحصر به فرد خود را دارد و در نتیجه روش‌های کشف و مبارزه در واحدهای تجاری مختلف، متفاوت خواهد بود. در حالت پیچیده‌تر، هر ساله

داده‌های بسیاری تولید می‌شوند و کشف اطلاعات متقلبانه مستلزم تکنیک‌های کاراتری جهت کاوش داده‌ها می‌باشد (آلمیدا^۱، ۲۰۰۹).

رویکردهای جدیدی برای کشف تقلب وجود دارد، اما داده‌کاوی که هدف آن استخراج اطلاعات موردعلاقه از مجموعه گسترده‌ای از داده‌ها است به‌طور گسترده‌ای به‌عنوان یک ابزار تصمیم‌گیری فعال مورد استفاده قرار گرفته است. این امر باعث افزایش توجه پژوهشگران در غنی کردن داده‌ها با هدف تسهیل کشف الگوهای موردعلاقه شده است (وایتینگ و همکاران^۲، ۲۰۱۲). داده‌کاوی، استخراج یا کاوش دانش از حجم زیادی اطلاعات است. الگوها یا قواعد قوی کشف‌شده از طریق تکنیک‌های داده‌کاوی را می‌توان برای پیش‌بینی غیر بديهی^۳ داده‌های جدید استفاده کرد. از این‌رو، داده‌کاوی یک حوزه میان‌رشته‌ای است که از ابزارهای تجزیه و تحلیل مدل‌های آماری، الگوریتم‌های ریاضیاتی و روش‌های یادگیری ماشینی^۴ برای کشف الگوها و روابط معتبر قبلاً ناشناخته در مجموعه داده‌های بزرگ استفاده می‌کند (دو و دو^۵، ۲۰۱۱).

بنابراین، در محیطی فعال از تقلب در صورت‌های مالی، مکانیزم‌های کشف تقلب با کمک رایانه بسیار مؤثرتر و کاراتر خواهد بود. تکنولوژی‌های مبتنی بر آمار و یادگیری ماشینی یک راهکار اثربخش برای پیشگیری و کشف تقلب هستند اما مرتکبین تقلب خود را انطباق می‌دهند و به‌طور معمول قادر به کشف راه‌هایی برای دور زدن این تکنولوژی‌ها هستند. تکنیک‌های کنونی کشف تقلب در اکثر موقعیت‌های مرتبط با تقلب از اصول داده‌کاوی مشابهی استفاده می‌کنند اما می‌توانند از لحاظ دانش قلمرو تخصص متفاوت باشند. زمانی که مدیران مالی درگیر در تقلب از نرم‌افزارها و تکنیک‌های کشف تقلب اطلاع کافی دارند، روش‌هایی را برای ارتکاب تقلب بکار می‌گیرند که کشف آن‌ها به‌ویژه با استفاده از تکنیک‌های فعلی دشوار است. از این‌رو، یک نیاز مبرم برای روش‌هایی که نه تنها کارا باشند بلکه برای کشف فریبکاری‌های مالی انطباقی و نوظهور مؤثر هستند، وجود دارد (ژو و کپور^۵، ۲۰۱۱).

1 Almeida

2 Whiting, et al.

3 Nontrivial

4 Machine Learning Methods

5 Zhou & Kapoor

۳- پیشینه پژوهش

۳-۱- پژوهش‌های خارجی

سپاتیس^۱ (۲۰۰۲)، با نمونه‌ای متشکل از ۷۶ شرکت شامل ۳۸ شرکت متقلب و ۳۸ شرکت غیر متقلب، توانایی اطلاعات حسابداری در کشف گزارشگری مالی متقلبان را مورد بررسی قرار داد. در مجموع، ده متغیر مالی به‌عنوان پیش‌بینی‌کننده‌های بالقوه تقلب در صورت‌های مالی انتخاب شدند. نتایج حاصل از رگرسیون لجستیک گام‌به‌گام حاکی از این بود که شرکت‌های با درصد بالای موجودی کالا به فروش، نسبت بالای بدهی به کل دارایی‌ها، نسبت پایین سود خالص به کل دارایی‌ها، نسبت پایین سرمایه در گردش به کل دارایی‌ها و امتیاز پایین Z، با احتمال بیش‌تری صورت‌های مالی خود را دست‌کاری می‌کنند.

هوگس و همکاران^۲ (۲۰۰۷)، با نمونه‌ای متشکل از ۳۹۰ شرکت شامل ۵۱ شرکت متقلب و ۳۳۹ شرکت غیر متقلب طی سال‌های ۲۰۰۴-۱۹۹۸، به معرفی رویکردی تکاملی برای کشف گزارشگری مالی متقلبان پرداختند. در مجموع، ۸۵ متغیر مالی و غیر مالی به‌عنوان پیش‌بینی‌کننده‌های بالقوه خطر تقلب در صورت‌های مالی انتخاب شدند. نتایج حاصل از الگوریتم ژنتیک^۳ بیان‌گر سودمندی این تکنیک و همچنین اهمیت اطلاعات مالی و غیر مالی در کشف گزارشگری مالی متقلبان است.

کرکوز و همکاران^۴ (۲۰۰۷)، با نمونه‌ای متشکل از ۳۸ شرکت متقلب و ۳۸ شرکت غیر متقلب و با استفاده از ۱۰ نسبت مالی به بررسی سودمندی درخت‌های تصمیم^۵، شبکه‌های عصبی مصنوعی و شبکه‌های باور بیزین^۶ در کشف صورت‌های مالی متقلبان پرداخت. نتایج پژوهش حاکی از این بود که صورت‌های مالی منتشره دربرگیرنده اطلاعات سودمندی برای کشف تقلب در گزارشگری مالی می‌باشند. همچنین، مدل شبکه باور بیزین عملکرد بهتری نسبت به مدل شبکه عصبی مصنوعی و درخت تصمیم داشت و خطای نوع اول در هر سه مدل کمینه بود.

آلدن و همکاران^۷ (۲۰۱۲)، با نمونه‌ای متشکل از ۴۵۸ شرکت شامل ۲۲۹ شرکت متقلب و ۲۲۹ شرکت غیر متقلب، سودمندی طبقه‌بندی‌کننده‌های مبتنی بر قواعد فازی^۸ در کشف

1 Spathis

2 Hoogs, et al.

3 Genetic Algorithm

4 Kirkos, et al.

5 Decision Trees

6 Bayesian Belief Networks

7 Alden et al.

8 Fuzzy Rule-Based Classifiers

الگوهای گزارشگری مالی متقلبانه را موردبررسی قرار داد. در مجموع، ۱۸ متغیر مالی به‌عنوان پیش‌بینی‌کننده‌های بالقوه گزارشگری مالی متقلبانه انتخاب شدند. نتایج حاصل از طبقه‌بندی‌کننده‌های مبتنی بر قواعد فازی با الگوریتم ژنتیک و الگوریتم برآورد توزیع یادگیری مارکو^۱ حاکی از سودمندی مدل‌های تکاملی برای کشف گزارشگری مالی متقلبانه بود.

راویسنکار و همکاران^۲ (۲۰۱۱)، با نمونه‌ای متشکل از ۲۰۲ شرکت شامل ۱۰۱ شرکت متقلب و ۱۰۱ شرکت غیر متقلب به بررسی سودمندی تکنیک‌های داده‌کاوی در کشف تقلب در صورت‌های مالی پرداختند. نتایج حاصل از به‌کارگیری تکنیک‌های داده‌کاوی مانند شبکه‌های عصبی چندلایه پیش‌خور^۳، ماشین‌های بردار پشتیبان^۴، برنامه‌ریزی ژنتیک^۵، روش گروهی مدل‌سازی داده‌ها^۶، رگرسیون لجستیک^۷ و شبکه عصبی احتمالی^۸، حاکی از برتری شبکه‌های عصبی احتمالی بدون انتخاب ویژگی نسبت به سایر تکنیک‌ها در کشف صورت‌های مالی متقلبانه بود. در صورت به‌کارگیری انتخاب ویژگی، برنامه‌ریزی ژنتیک و شبکه‌های عصبی احتمالی با صحت تقریباً یکسانی عملکردی بهتر از سایر روش‌ها داشتند.

هوانگ و همکاران^۹ (۲۰۱۴)، با نمونه‌ای متشکل از ۱۴۴ شرکت شامل ۷۲ شرکت متقلب و ۷۲ شرکت غیر متقلب، به ارائه مدلی مبتنی بر نگاشت خودسازمان‌ده سلسله‌مراتبی در حال رشد^{۱۰} برای کشف تقلب در گزارشگری مالی پرداختند. در مجموع ۲۴ متغیر مستقل برای انجام تحلیل ممیزی جهت تعیین ورودی‌های مدل انتخاب شدند و نتایج حاصل از تحلیل ممیزی حاکی از این بود که اثرات ۸ متغیر از نظر آماری معنادار است. این متغیرها وضعیت یک شرکت را از جنبه‌های سودآوری، نقدینگی، ساختار مالی، دسترسی به وجه نقد، درماندگی مالی و حاکمیت شرکتی اندازه‌گیری می‌کنند. نتایج حاصل از پژوهش بیان‌گر عملکرد مناسب مدل پیشنهادی برای کشف گزارشگری مالی متقلبانه بود.

1 Markovian Learning Estimation of Distribution Algorithm

2 Ravisankar et al.

3 Multilayer Feed Forward Neural Network

4 Support Vector Machines

5 Genetic Programming

6 Group Method of Data Handling

7 Logistic Regression

8 Probabilistic Neural Network

9 Huang et al.

10 Growing Hierarchical Self-Organizing Map

دالنیل و همکاران^۱ (۲۰۱۴)، با نمونه‌ای متشکل از ۶۵ شرکت متقلب و ۶۵ شرکت غیر متقلب پذیرفته‌شده در بورس اوراق بهادار مالزی به بررسی سودمندی متغیرهای مالی در کشف شرکت‌های متقلب طی سال‌های ۲۰۱۱-۲۰۰۰ پرداختند. نتایج پژوهش حاکی از این بود که تفاوت معناداری بین میانگین نسبت‌های کل بدهی به کل حقوق صاحبان سهام و حساب‌های دریافتی به فروش در شرکت‌های متقلب و شرکت‌های غیر متقلب وجود دارد. همچنین، امتیاز Z آلتمن که احتمال ورشکستگی را اندازه‌گیری می‌کند از اهمیت زیادی در کشف گزارشگری مالی متقلبانة برخوردار است.

لین و همکاران^۲ (۲۰۱۶)، با نمونه‌ای متشکل از ۱۲۹ شرکت متقلب و ۴۴۷ شرکت غیر متقلب به طبقه‌بندی و رتبه‌بندی عوامل تقلب طی سال‌های ۲۰۱۰ - ۱۹۹۸ پرداختند. در این پژوهش ۳۲ عامل تقلب که مطابق نظر متخصصان برای کشف تقلب مناسب شناخته شدند بکار رفته است. نتایج حاصل از تکنیک‌های داده‌کاوی رگرسیون لجستیک، درخت تصمیم و شبکه‌های عصبی مصنوعی حاکی از دقت بیش‌تر روش‌های درخت تصمیم و شبکه‌های عصبی مصنوعی نسبت به رگرسیون لجستیک است. در پایان نیز به‌منظور بهبود دستاوردهای پژوهش مقایسه‌ای بین قضاوت متخصصان و تکنیک‌های داده‌کاوی صورت گرفته است.

۳-۲- پژوهش‌های داخلی

صفرزاده (۱۳۸۹)، با نمونه‌ای متشکل از ۶۶ شرکت متقلب و ۱۱۲ شرکت غیر متقلب و با استفاده از ۱۰ نسبت مالی به بررسی نقش داده‌های حسابداری در ایجاد یک الگو برای کشف عوامل مرتبط با تقلب در گزارشگری مالی پرداخت. شرکت‌های متقلب بر مبنای ۱. شمول شرکت در فهرست سازمان بورس و اوراق بهادار به دلایلی مرتبط با تحریف داده‌های مالی و ۲. انجام دادن معاملات نهانی و آرای صادرشده توسط دادگاه در خصوص تحریف در گزارشگری مالی انتخاب شدند. نتایج پژوهش حاکی از عملکرد مناسب الگوی پیشنهادی در طبقه‌بندی شرکت‌های نمونه داشت.

اعتمادی و زلقی (۱۳۹۲)، با نمونه‌ای متشکل از ۶۸ شرکت شامل ۳۴ شرکت متقلب و ۳۴ شرکت غیر متقلب و با استفاده از ۹ نسبت مالی به بررسی کاربرد رگرسیون لجستیک در شناسایی گزارشگری مالی متقلبانة پرداختند. در این پژوهش برای تفکیک شرکت‌های متقلب و غیر متقلب از معیارهای؛ ۱. اظهارنظر غیر مقبول حسابرسی، ۲. وجود اختلاف مالیاتی با حوزه مالیاتی و ۳.

1 Dalnial et al.

2 Lin et al.

وجود تعدیلات سنواتی با اهمیت استفاده شده است. نتایج پژوهش بیان گر سودمندی مدل پیشنهادی در شناسایی گزارشگری مالی متقلبانه بود.

فرقاندوست حقیقی و همکاران (۱۳۹۳)، با نمونه ای متشکل از ۱۱۵ شرکت متقلب و ۱۱۵ شرکت غیر متقلب به بررسی رابطه مدیریت سود و امکان تقلب در صورت های مالی شرکت های پذیرفته شده در بورس اوراق بهادار تهران پرداختند. نتایج این پژوهش حاکی از این بود که در شرکت های با سابقه مدیریت سود، امکان ارتکاب تقلب در صورت های مالی وجود دارد. همچنین، در صورت وجود سابقه مدیریت سود، وجود عوامل انگیزشی سبب افزایش احتمال ارتکاب تقلب در صورت های مالی می گردد.

جهانشاد و سرداری زاده (۱۳۹۳)، با نمونه ای متشکل از ۸۰ شرکت شامل ۴۰ شرکت متقلب و ۴۰ شرکت غیر متقلب، به بررسی رابطه معیار مالی (اختلاف رشد درآمد) و معیار غیر مالی (رشد تعداد کارکنان) با گزارشگری مالی متقلبانه پرداختند. شرکت های متقلب بر مبنای شمول شرکت در فهرست سازمان بورس اوراق بهادار تهران به دلایلی مرتبط با ۱. تحریف داده های مالی و ۲. انجام معاملات نهانی در خصوص تحریف در گزارشگری مالی انتخاب شده اند. نتایج پژوهش بیان گر یک رابطه منفی و معنادار میان رشد درآمد، رشد کارکنان و اختلاف این دو متغیر با گزارشگری مالی متقلبانه بود.

مشایخی و حسین پور (۱۳۹۵)، با نمونه ای متشکل از ۱۰۷ شرکت مشکوک به تقلب به بررسی رابطه مدیریت سود واقعی و مدیریت سود تعهدی در شرکت های مشکوک به تقلب بورس اوراق بهادار تهران پرداختند. شرکت های مشکوک به تقلب بر مبنای یکسری عوامل مرتبط با ارائه نادرست و گمراه کننده اطلاعات حسابداری در گزارشگری مالی گزینش شده اند. نتایج حاصل از پژوهش حاکی این است که در شرکت های بورسی مشکوک به تقلب، مدیریت سود واقعی بر مدیریت سود تعهدی تأثیر منفی و معنی داری دارد.

۵- روش پژوهش

این پژوهش کاربردی و طرح آن از نوع شبه تجربی و با استفاده از رویکرد پس رویدادی است. طرح های شبه تجربی مستلزم تخصیص تصادفی واحدهای آزمون به تیمارهای تجربی یا تخصیص تصادفی تیمار تجربی به واحدهای آزمون نیست (سریجیش و همکاران^۱، ۲۰۱۴). در این حالت، طرح های شبه تجربی به کنترل متغیرهای برونزا کمک می کنند اما به دلیل کنترل عوامل مرتبط با کاهش روایی داخلی به اندازه طرح های تجربی واقعی سودمند نیستند (زیکموند و همکاران^۲،

1 Sreejesh et al.

2 Zikmund et al.

۲۰۱۰). برخی مواقع، طرح‌های شبه تجربی تنها راه برای انجام پژوهش هستند. طرح‌های سری زمانی، معروف‌ترین طرح شبه تجربی مورد استفاده به‌وسیله پژوهشگران است.

۴-۱- سؤال پژوهش

با توجه به مبانی نظری و پیشینه پژوهش، سؤال پژوهش به شرح زیر تدوین شد:
آیا می‌توان با استفاده از اطلاعات مالی و غیرمالی سالانه شرکت‌ها که به سازمان بورس و اوراق بهادار ارائه می‌شود، مدلی کمی برای کشف تقلب ارائه داد که بیان‌گر یک روش تحلیلی خودکار برای تعیین شرکت‌های متقلب باشد؟

۴-۲- متغیرهای پژوهش

متغیر پاسخ

متغیر پاسخ در این پژوهش، تقلب در صورت‌های مالی است که مقدار آن برای شرکت‌های متقلب عدد یک و برای شرکت‌های غیر متقلب عدد صفر است. شرکت‌های متقلب با مراجعه به نشریات ویژه سازمان بورس و اوراق بهادار که در وبسایت سازمان بورس و اوراق بهادار نمایه شده‌اند انتخاب گردیدند. از این‌رو بررسی کامل نشریات ویژه نشان می‌دهد نشریات شماره ۳، ۴ و ۷ دربرگیرنده چهارده رأی قطعی صادره در محاکم دادگستری می‌باشند. هر یک از این چهارده رأی در راستای اجرای ماده ۵۲ قانون بازار اوراق بهادار مصوب آذرماه ۱۳۸۴ به دلایل مختلفی از قبیل ایجاد ظاهری گمراه‌کننده از روند معاملات اوراق بهادار، مبادرت به معاملات اوراق بهادار با استفاده از اطلاعات نهانی، خودداری از ارائه اطلاعات، اسناد و مدارک مهم به سازمان بورس و اوراق بهادار، تخلف از مقررات قانون بازار در تهیه اسناد و مدارک، تقسیم منافع موهوم به استناد صورت‌داری و ترازنامه مزور و موارد دیگر صادر شده‌اند. با توجه به این‌که موضوع پژوهش حاضر به‌طور مشخص تقلب در صورت‌های مالی است، شرکت‌هایی که فقط به دلایل مرتبط با تقلب در صورت‌های مالی علیه مدیران آن‌ها در محاکم دادگستری آرای قطعی صادر شده است و با بررسی اطلاعات ارائه‌شده در نشریات ویژه سازمان بورس و اوراق بهادار قابل تشخیص بودند، به‌عنوان شرکت‌های متقلب در نظر گرفته شدند.

متغیرهای پیش‌بینی‌کننده

در این پژوهش، عوامل مالی و غیرمالی خطر تقلب مرتبط با گزارشگری مالی متقلبانه متغیرهای پیش‌بین هستند که با مرور پژوهش‌های انجام‌شده در این زمینه تعیین خواهند شد. در یک دسته‌بندی کلی، این عوامل عبارت‌اند از: اهرم مالی، سودآوری، ترکیب دارایی‌ها، نقدینگی، کارایی،

اندازه، رشد، وضعیت مالی کلی و کیفیت حسابرسی. جدول (۱) خلاصه متغیرهای مرتبط با هر یک از این عوامل را نشان می دهد.

جدول (۱). خلاصه متغیرهای پیش بینی کننده

عامل	متغیرها	نام پژوهش
اهرم مالی	لگاریتم کل بدهی و نسبت های کل بدهی ها به کل دارایی ها، بدهی های بلندمدت به کل دارایی ها، کل بدهی ها به حقوق صاحبان سهام، بدهی های بلندمدت به حقوق صاحبان سهام و نسبت پوشش بهره.	پرسونس ^۱ (۱۹۹۵)، سپاتیس (۲۰۰۲)، سپاتیس و همکاران ^۲ (۲۰۰۲)، کامینسکی و همکاران ^۳ (۲۰۰۴)، کرکوز و همکاران (۲۰۰۷)، برازل و همکاران ^۴ (۲۰۰۹)، آلدن و همکاران (۲۰۱۲)، چن و همکاران ^۵ (۲۰۱۴).
سودآوری	بازده دارایی ها، بازده حقوق صاحبان سهام، سود قبل از بهره و مالیات، سود هر سهم و نسبت های سود خالص به فروش، سود ناخالص به فروش، سود عملیاتی به فروش، سود ناخالص به کل دارایی ها، سود انباشته به کل دارایی ها و نسبت سود خالص به دارایی های ثابت.	پرسونس (۱۹۹۵)، سپاتیس (۲۰۰۲)، سپاتیس و همکاران (۲۰۰۲)، کامینسکی و همکاران (۲۰۰۴)، کرکوز و همکاران (۲۰۰۷)، برازل و همکاران (۲۰۰۹)، آلدن و همکاران (۲۰۱۲)، چن و همکاران (۲۰۱۴).
ترکیب دارایی ها	تغییر در حساب های دریافتی، تغییر در موجودی کالا و نسبت های جمع دارایی های جاری به کل دارایی ها، حساب های دریافتی به کل دارایی ها، موجودی کالا به فروش خالص، حساب های دریافتی به فروش خالص و نسبت موجودی کالا به جمع دارایی ها.	پرسونس (۱۹۹۵)، سپاتیس (۲۰۰۲)، سپاتیس و همکاران (۲۰۰۲)، کرکوز و همکاران (۲۰۰۷)، برازل و همکاران (۲۰۰۹)، آلدن و همکاران (۲۰۱۴)، راعی و همکاران (۱۳۹۳).
نقدینگی	نسبت های آبی، نقدی، جاری، وجه نقد به کل دارایی ها و نسبت سرمایه در گردش به کل دارایی ها.	پرسونس (۱۹۹۵)، سپاتیس (۲۰۰۲)، سپاتیس و همکاران (۲۰۰۲)، کامینسکی و همکاران (۲۰۰۴)، کرکوز و همکاران (۲۰۰۷)، برازل و همکاران (۲۰۰۹)، آلدن و همکاران (۲۰۱۲)، چن و همکاران (۲۰۱۴).
کارایی	نسبت های فروش به کل دارایی ها، فروش به کل دارایی های ثابت، فروش به حساب های دریافتی و بهای تمام شده کالای فروش رفته به موجودی کالا.	پرسونس (۱۹۹۵)، سپاتیس (۲۰۰۲)، سپاتیس و همکاران (۲۰۰۲)، کامینسکی و همکاران (۲۰۰۴)، کرکوز و همکاران (۲۰۰۷)، برازل و همکاران (۲۰۰۹)، آلدن و همکاران (۲۰۱۲)، چن و همکاران (۲۰۱۴).

1 Persons

2 Spathis et al.

3 Kaminski et al.

4 Brazel et al.

5 Chen et al.

ادامه جدول (۱). خلاصه متغیرهای پیش‌بینی کننده

عامل	متغیرها	نام پژوهش
اندازه	لگاریتم طبیعی کل دارایی‌ها، لگاریتم طبیعی کل فروش و لگاریتم طبیعی ارزش بازار شرکت.	پرسونس (۱۹۹۵)، سیاتیس (۲۰۰۲)، سیاتیس و همکاران (۲۰۰۲)، کامینسکی و همکاران (۲۰۰۴)، کرکوز و همکاران (۲۰۰۷)، برازل و همکاران (۲۰۰۹)، آلدن و همکاران (۲۰۱۲)، چن و همکاران (۲۰۱۴).
وضعیت مالی کلی	از Z آلتمن برای اندازه‌گیری سلامت مالی یک شرکت استفاده می‌شود.	ستایس ^۱ (۱۹۹۱)، پرسونس (۱۹۹۵)، سامرز و سویینی ^۲ (۱۹۹۸)، سیاتیس (۲۰۰۲)، سیاتیس و همکاران (۲۰۰۲)، کرکوز و همکاران (۲۰۰۷).
رشد	برای اندازه‌گیری رشد شرکت از متغیر رشد فروش استفاده می‌شود.	ستایس (۱۹۹۱)، سامرز و سویینی (۱۹۹۸)، کرکوز و همکاران (۲۰۰۷).
کیفیت حسابداری	اندازه حسابداری یک متغیر موهومی است که اگر حسابداری شرکت سازمان حسابداری باشد، مقدار آن یک و در غیر این صورت صفر می‌باشد. تعداد سال‌های پیاپی که حسابداری در استخدام یک شرکت است، دوره تصدی حسابداری نامیده می‌شود. متغیر تغییر حسابداری، به صورت یک متغیر دوتایی تعریف می‌شود که در آن عدد یک نشان‌دهنده صاحبکار جدید و عدد صفر نشان‌دهنده صاحبکار ثابت است.	پیر و اندرسن ^۳ (۱۹۸۴)، تیو و وانگ (۱۹۹۳)، ستایس (۱۹۹۱)، کارسلو و ناجی (۲۰۰۴)، پرسونس (۱۹۹۵)، سامرز و سویینی (۱۹۹۸)، برازل و همکاران (۲۰۰۹)، پورحیدری و همکاران (۱۳۹۴).

۴-۳- مدل‌های پژوهش

داده‌کاوی در رشته‌های گوناگونی مانند مالی، مهندسی، زیست پزشکی^۴ و امنیت سایبری^۵ استفاده می‌شود. روش‌های داده‌کاوی به دودسته تقسیم می‌شوند: نظارت‌شده^۶ و غیر نظارت‌شده^۷. تکنیک‌های داده‌کاوی نظارت‌شده با داده‌های آموزشی یک تابع پنهان را پیش‌بینی می‌کنند. داده‌های آموزشی چندین جفت متغیر ورودی و برجسب‌ها یا کلاس‌های خروجی دارند. خروجی روش قادر خواهد بود برجسب کلاس متغیرهای ورودی را پیش‌بینی کند. طبقه‌بندی و پیش‌بینی مثال‌هایی از کاوش نظارت‌شده هستند. داده‌کاوی غیر نظارتی، تلاش برای شناسایی الگوهای

1 Stice

2 Summers & Sweeny

3 Pierre & Anderson

4 Biomedicine

5 Cybersecurity

6 Supervised

7 Unsupervised

مخفی با استفاده از داده های معین و بدون معرفی داده های آموزشی (مانند جفت های ورودی و برچسب کلاس) هستند. خوشه بندی^۱ و کاوش قواعد وابستگی^۲ نمونه هایی از کاوش غیر نظارتی هستند.

داده کاوی همچنین بخشی از کشف دانش در پایگاه های داده، یک فرآیند تکراری برای استخراج اطلاعات غیر بدیهی از داده ها، است. کشف شناخت در پایگاه داده دربرگیرنده چندین مرحله از جمع آوری داده های خام تا ایجاد شناخت جدید می گردد. این فرآیند تکراری از گام های زیر تشکیل می گردد: پاک سازی داده ها^۳، تجمیع داده ها^۴، انتخاب داده ها^۵، انتقال داده ها^۶، داده کاوی، ارزیابی الگو^۷ و نمایش شناخت^۸ (دوا و دو، ۲۰۱۱).

انتخاب ویژگی

داده های نامتوازن به طور معمول با فضای ویژگی با ابعاد بالا همراه هستند. در میان ویژگی های با ابعاد بالا، وجود تعداد زیادی ویژگی های پر اختلال^{۱۰} می تواند عملکرد طبقه بندی کننده را مختل و کاهش دهد. در سال های اخیر، روش های انتخاب ویژگی و زیر فضا^{۱۱} برای حل این مسئله معرفی و ارزیابی شدند. روش های انتخاب زیرمجموعه ویژگی برای انتخاب زیرمجموعه ویژگی کوچکی از میان ویژگی های با ابعاد بالا مطابق با معیارهای انتخاب ویژگی به کار می رود (دوا و دو، ۲۰۱۱).

روش های انتخاب ویژگی را می توان به دودسته تقسیم کرد؛ انتخاب عددی ویژگی که ویژگی ها را به صورت واحد انتخاب می کنند و انتخاب برداری ویژگی که ویژگی ها را بر مبنای همبستگی متقابل بین ویژگی ها برمی گزینند. انتخاب عددی ویژگی دارای مزیت سادگی محاسباتی است و احتمالاً برای مجموعه داده ای که ویژگی ها دارای همبستگی متقابل هستند مؤثر نیست. روش های انتخاب برداری ویژگی بهترین ترکیب برداری ویژگی را انتخاب می کند. روش های

1 Clustering

2 Associative Rule Mining

3 Data Cleaning

4 Data Integration

5 Data Selection

6 Data Transformation

7 Pattern Evaluation

8 Knowledge Representation

9 Dua & Du

10 Noisy

11 Subspace

انتخاب برداری ویژگی را می‌توان به دودسته روش‌های دسته‌بندی^۱ و روش‌های مبتنی بر فیلتر^۲ تقسیم می‌شوند. روش‌های دسته‌بندی با استفاده از روش‌های یادگیری ماشینی، مانند جعبه سیاه، ویژگی‌های مرتبط‌تر را انتخاب می‌کنند به گونه‌ای که روش یادگیری دارای عملکرد بهینه است. استراتژی‌های جستجو دربرگیرنده جستجوی فراگیر، جستجوی پرتو، شاخه و حد، الگوریتم ژنتیک، روش‌های جستجوی حریصانه و موارد دیگر هستند. هنگام استفاده از روش دسته‌بندی، ویژگی‌های انتخاب‌شده مستعد بیش‌برازش داده‌ها هستند. در روش انتخاب ویژگی مبتنی بر فیلتر، یک ویژگی با طبقه‌ای از ویژگی‌ها و زیرمجموعه ویژگی متناظر آن همبستگی دارد (دوا و دو، ۲۰۱۱).

در این پژوهش به منظور انتخاب متغیرهایی که بیشترین تأثیر را در کشف تقلب در صورت‌های مالی دارند، از روش انتخاب ویژگی مبتنی بر همبستگی ارائه‌شده توسط هال^۳ (۱۹۹۹) استفاده خواهد شد. استفاده از معیار همبستگی به راه‌حل بهینه در انتخاب ویژگی منجر می‌شود؛ بنابراین، این روش بر دو مسئله تمرکز دارد: معیار اندازه‌گیری همبستگی و الگوریتم انتخاب ویژگی. معیار همبستگی می‌تواند ضریب همبستگی پیرسون^۴ (PCC)، اطلاعات متقابل^۵ (MI) و سایر معیارهای مرتبط باشد. ضریب همبستگی پیرسون معیاری برای وابستگی خطی بین متغیرها و ویژگی‌ها می‌باشد و با متغیرهای پیوسته و دوتایی قابلیت کاربرد دارد. اطلاعات متقابل توانایی اندازه‌گیری وابستگی غیرخطی را دارد که بی‌ربطی هریک از متغیرها و ویژگی‌ها را با استفاده از واگرایی کولبک-لیبلر^۶ اندازه‌گیری می‌کند. باین‌حال، برآورد اطلاعات متقابل نسبت به ضریب همبستگی پیرسون، به‌ویژه برای داده‌های پیوسته، سخت‌تر است (دوا و دو، ۲۰۱۱).

طبقه‌بندی

طبقه‌بندی شکلی از تحلیل داده است که به استخراج مدل‌های توضیح‌دهنده طبقات مهم داده می‌پردازد. در این پژوهش برای دسته‌بندی شرکت‌ها به متقلب و غیر متقلب از روش‌های شبکه عصبی مصنوعی، شبکه بیزین^۷ و الگوریتم جنگل تصادفی^۸ استفاده خواهد شد که در ادامه به شرح مختصر آن‌ها پرداخته می‌شود.

1 Wrapper

2 Filter

3 Hall

4 Pearson's Correlation Coefficient

5 Mutual Information

6 Kullback-Leibler

7 Bayesian Network

8 Random Forest

شبکه عصبی مصنوعی

شبکه عصبی مصنوعی یک مدل یادگیری ماشینی است که ورودی ها را از طریق پردازش اطلاعات غیرخطی به خروجی هایی که با اهداف مطابقت داده می شوند در یک گروه متصل از نرون های مصنوعی که لایه های واحدهای پنهان را تشکیل می دهند، متصل می کند. فعالیت هر واحد پنهان و خروجی $\hat{Y}$ از ترکیب ورودی X و یک مجموعه از وزن های نرون W تعیین می شود:

$$\hat{Y} = f(X, W) \quad (1)$$

که در این رابطه به ماتریس بردارهای وزن لایه های پنهان اشاره دارد (دوا و دو، ۲۰۱۱). زمانی که شبکه های عصبی مصنوعی به عنوان یک روش یادگیری ماشینی مورد استفاده قرار می گیرد، تلاش هایی برای تعیین مجموعه وزن هایی جهت کمینه کردن خطای طبقه بندی صورت می گیرد. همگرایی حداقل میانگین مربعات یک روش شناخته شده متداول برای بسیاری از پارادایم های یادگیری است. هدف شبکه های عصبی مصنوعی کمینه کردن خطاهای بین واقعیت زمینی Y و خروجی مورد انتظار $f(X, W)$ یک شبکه عصبی مصنوعی است به طوری که:

$$E(X) = ((f(X, W) - Y)^2 \quad (2)$$

روش های شبکه عصبی مصنوعی برای طبقه بندی یا پیش بینی متغیرهای پنهان که اندازه گیری آن ها دشوار است و حل مسائل طبقه بندی غیرخطی به خوبی عمل می کنند و نسبت به داده های پرت حساسیت ندارند. مدل های شبکه عصبی مصنوعی رابطه بین ورودی و خروجی را به صورت ضمنی تعریف می کنند و از این رو راه حل هایی برای مسائل ملال آور شناخت الگو به ویژه زمانی که استفاده کنندگان هیچ ایده ای درباره رابطه بین متغیرها ندارند، ارائه می کند (دوا و دو، ۲۰۱۱).

شبکه بیزین

شبکه بیزین که شبکه باور نیز نامیده می شود از توزیع احتمال مشترک عاملی در یک مدل گرافیکی برای تصمیم گیری در رابطه با متغیرهای نامطمئن استفاده می کند. طبقه بندی کننده شبکه بیزین بر مبنای قاعده بیز که ارائه دهنده فرضیه H از طبقات و داده x است، قرار دارد، بنابراین؛

$$P(H|x) = \frac{P(x|H)P(H)}{P(x)} \quad (3)$$

که در این رابطه:

$P(H)$ نشان دهنده احتمالات پیشین هر طبقه بدون اطلاعاتی درباره متغیر x است.

$P(H|x)$ بیان گر احتمالات پسین متغیر x بر طبقات ممکن است.

$P(x|H)$ بیان گر احتمالات مشروط متغیر x با احتمال H است.

نایو بیز شکل ساده‌ای از مدل شبکه بیزین است که فرض می‌کند همه متغیرها از یکدیگر مستقل هستند. جهت به کارگیری قاعده بیز برای طبقه‌بندی نایو بیز یافتن فرضیه حداکثر احتمال که برچسب طبقه را برای هر داده آزمون x تعیین می‌کند ضروری است. با مشاهدات x و یک گروه از برچسب طبقات $C = \{c_j\}$ ، طبقه‌بندی کننده نایو بیز از طریق بیشینه کردن فرضیه احتمال پسین (MAP) برای داده‌ها به شرح زیر حل شود.

$$\underset{c_j \in C}{\operatorname{argmax}} P(x|c_j)P(c_j) \quad (۴)$$

نایو بیز به منظور انجام وظایف استنباطی کارا است. با این حال، به شدت بر مبنای فرض مستقل بودن متغیرهای درگیر قرار دارد. با کمال شگفتی، این روش حتی در صورت نقض شرط مستقل بودن متغیرها نیز نتایج خوبی به همراه دارد (دوا و دو، ۲۰۱۱).

جنگل تصادفی

الگوریتم جنگل تصادفی، معروف‌ترین طبقه‌بندی کننده Bagging Ensemble می‌باشد. جنگل تصادفی دربرگیرنده تعداد بسیاری درخت تصمیم می‌باشد. خروجی جنگل تصادفی بر مبنای آرای مشخص هر یک از درخت‌ها تعیین می‌شود. هر درخت تصمیم با طبقه‌بندی نمونه‌های Bootstrap داده‌های ورودی از طریق الگوریتم درخت ساخته می‌شود. سپس، هر درخت برای طبقه‌بندی داده‌های آزمون استفاده خواهد شد. هر درخت یک تصمیم برای برچسب‌گذاری داده‌های آزمون دارد. این برچسب، یک رأی نام دارد. در نهایت جنگل با جمع‌آوری بیش‌ترین آرای درخت‌ها نتیجه طبقه‌بندی را قطعی خواهد کرد. در ادامه با برخی تعاریف ارائه شده توسط بریمن^۱ (۲۰۰۱) به بررسی جنگل‌های تصادفی پرداخته می‌شود. با فرض این که یک جنگل متشکل از K درخت $\{T_1, \dots, T_K\}$ ، بردار تصادفی θ_k برای درخت k ام ایجاد می‌شود، $k = 1, \dots, K$ ، بردارهای $\{\theta_k\}$ مستقل و دارای توزیع مشابه برای مدل‌سازی درخت هستند. این بردارها در ساخت درخت تعریف شده‌اند. برای نمونه در انتخاب تصادفی، این بردارها شامل اعداد صحیح تصادفی هستند که به‌طور تصادفی از $\{1, \dots, N\}$ انتخاب شده‌اند که N عدد جداکننده است. با استفاده از مجموعه داده آموزشی و بردارهای $\{\theta_k\}$ ، درخت رشد می‌کند و یک رأی واحد برای معروف‌ترین طبقه در ورودی x ارائه می‌نماید. طبقه‌بندی کننده درخت k ام به‌صورت

1 Breiman

$f(x, \theta_k)$ نشان داده می‌شود و حاصل یک جنگل تصادفی متشکل از مجموعه آن درخت‌ها خواهد بود.

$$\{f\{x, \theta_k\}\}, k = 1, \dots, K \quad (5)$$

صحت جنگل تصادفی به قوت هر یک از درخت‌ها و معیار وابستگی بین درخت‌ها بستگی دارد. علاوه بر این، الگوریتم جنگل تصادفی برای اجتناب از سوپیه در مدل‌سازی درخت از بوت‌استرپ استفاده می‌کند و در نتیجه اعتبارسنجی متقابل (CV) در آموزش و آزمون موردنیاز نیست. باین‌حال، به دلیل حداکثر سازی صحت پیش‌بینی در الگوریتم جنگل تصادفی، این طبقه‌بندی کننده با مشکل طبقات نامتوازن روبرو خواهد بود. روش‌های مبتنی بر درخت پراکندگی بالایی دارند. ساختار سلسله مراتبی درخت‌ها می‌تواند به نتیجه ناپایدار منجر شود. میانگین تعداد بسیاری درخت، مانند استفاده از بگینگ، می‌تواند پایداری الگوریتم‌های یادگیری **انسیمبل** را بهبود دهد (دوا و دو، ۲۰۱۱).

عملکرد طبقه‌بندی کننده^۱

به منظور ارزیابی عملکرد هر یک از طبقه‌بندی کننده‌های شبکه عصبی مصنوعی، شبکه بیزین و الگوریتم جنگل تصادفی، مطابق با ماتریس درهم‌ریختگی^۲ زیر از معیارهای صحت کلی^۳، دقت^۴، فراخوانی^۵، $F - Measure$ و AUC استفاده خواهد شد. با یک طبقه‌بندی کننده و یک نمونه مشخص، چهار خروجی احتمالی وجود دارد. چنانچه نمونه مثبت باشد و به‌عنوان مثبت نیز طبقه‌بندی شده باشد، مثبت حقیقی به حساب خواهد آمد و چنانچه این نمونه به‌صورت منفی طبقه‌بندی شود، منفی کاذب است. در صورتی که نمونه منفی بوده و به‌عنوان منفی طبقه‌بندی شده باشد، منفی حقیقی به حساب می‌آید و چنانچه این نمونه به‌صورت مثبت طبقه‌بندی شود، مثبت کاذب خواهد بود. نمودار شماره یک نشان‌دهنده ماتریس درهم‌ریختگی و معادلات مربوط به معیارهای متفاوتی است که با استفاده از آن محاسبه می‌شوند.

1 Classifier Performance

2 Confusion Matrix

3 Overall Accuracy

4 Precision

5 Recall

جدول (۲). ماتریس درهم‌ریختگی و محاسبه برخی معیارهای عملکرد متداول
طبقه‌های واقعی

		طبقه‌های پیش‌بینی شده	
		Y	N
جمع ستون‌ها	P	مثبت‌های واقعی (True Positives)	مثبت‌های کاذب (False Positives)
	N	منفی‌های کاذب (False Negatives)	منفی‌های واقعی (True Negatives)
		P	N

$$\text{درصد مثبت کاذب: } fp\ rate = \frac{FP}{P} \quad (۶)$$

$$\text{درصد مثبت واقعی: } tp\ rate = \frac{TP}{N} \quad (۷)$$

$$\text{دقت: } precision = \frac{TP}{TP + FP} \quad (۸)$$

$$\text{فراخوانی: } Recall = \frac{TP}{P} \quad (۹)$$

$$\text{صحت کلی: } Overall\ Accuracy = \frac{TP + TN}{TP + FP + FN + TN} \quad (۱۰)$$

$$F - Measure = \frac{2}{1/precision + 1/Recall} \quad (۱۱)$$

فضای زیر منحنی ROC (AUC)

منحنی ROC ترسیم دوبعدی عملکرد طبقه بندی کننده است که در آن درصد مثبت حقیقی بر روی منحنی Y و درصد مثبت کاذب بر روی منحنی X ترسیم می شود. به منظور مقایسه طبقه بندی کننده ها ممکن است به دنبال کاهش عملکرد ROC به یک مقدار برداری واحد که بیان گر عملکرد مورد انتظار است باشیم. یک معیار متداول، محاسبه مساحت زیر نمودار ROC است که به اختصار AUC نامیده می شود (فاوست^۱، ۲۰۰۶).

که در این روابط؛ Y ، شرکت هایی که در واقعیت متقلب هستند؛ N ، شرکت هایی که در واقعیت غیر متقلب هستند؛ p ، شرکت هایی که مدل به عنوان متقلب پیش بینی می کند؛ n ، شرکت هایی که توسط مدل به عنوان غیر متقلب پیش بینی می شود؛ TP ، مثبت های واقعی؛ FP ، مثبت های کاذب؛ FN ، منفی های کاذب؛ TN ، منفی های واقعی و $F - Measure$ ، معیار عملکرد F است.

۴-۴- جامعه آماری، روش نمونه گیری و محدوده زمانی

جامعه آماری این پژوهش، کلیه شرکت های پذیرفته شده در بورس اوراق بهادار تهران است. در انتخاب نمونه مواردی به این شرح در نظر گرفته خواهد شد که صورت های مالی و یادداشت های همراه شرکت ها در دوره زمانی ۱۳۷۹ الی ۱۳۹۳ به گونه کامل در سایت بورس اوراق بهادار موجود باشد و شرکت انتخابی جزء شرکت های بیمه، بانک ها و موسسه های اعتباری، واسطه گری های مالی و شرکت های چند رشته ای صنعتی نباشد. همچنین، بررسی کامل نشریات ویژه سازمان بورس و اوراق بهادار حاکی از چهارده رأی قطعی صادره در محاکم دادگستری است که با توجه به تعریف تقلب در صورت های مالی، تنها ۶ شرکت از ۱۴ شرکت مزبور به عنوان شرکت هایی که صورت های مالی متقلبانه داشته اند انتخاب گردیدند. علاوه بر این ۶ شرکت، ۶۴ شرکت غیر متقلب نیز با رعایت دقت لازم و به گونه ای انتخاب گردیدند که جزء شرکت های بیمه، بانک ها و موسسه های اعتباری، واسطه گری های مالی و شرکت های چند رشته ای صنعتی نباشد، اندازه مشابهی با شرکت های متقلب داشته باشند و در دوره زمانی یکسانی با شرکت های متقلب قرار گیرند و هیچ کدام از علائم خطر تقلب در گزارش حسابرسی آن ها وجود نداشته باشد. از آنجایی که تعداد شرکت های متقلب در مقایسه با تعداد شرکت های غیر متقلب بسیار کم تر است، یک مجموعه داده نامتوازن وجود دارد. از این رو برای حل این مسئله از رویکرد معرفی شده

توسط هی و همکاران^۱ (۲۰۰۸) استفاده خواهد شد. در این رویکرد که آدیسین^۲ نام دارد، ایده اصلی به کارگیری توزیع موزون برای نمونه‌های طبقات اقلیت مطابق با سطح دشواری آن‌ها در یادگیری است؛ بنابراین، داده‌های ترکیبی بیش‌تری برای نمونه‌های طبقات اقلیتی که یادگیری آن‌ها سخت‌تر است، در مقایسه با نمونه‌های طبقات اقلیتی که یادگیری آن‌ها آسان‌تر می‌باشد، ایجاد می‌شود.

۴-۵- روش گردآوری داده‌ها

در این پژوهش برای جمع‌آوری داده‌ها و اطلاعات، از روش آرشیوی استفاده می‌شود. برای نگارش و جمع‌آوری اطلاعات موردنیاز بخش مبانی نظری از مجلات تخصصی لاتین و برای گردآوری سایر داده‌ها و اطلاعات موردنیاز عمدتاً از طریق بانک‌های اطلاعاتی سازمان بورس اوراق بهادار تهران، گزارش‌های روزانه و هفتگی سازمان بورس و نرم‌افزارهای رهاورد نوین و تدبیر پرداز استفاده خواهد شد.

۵- یافته‌های پژوهش

نتایج پژوهش‌های پیشین حاکی از این است که میانگین دوره زمانی وقوع تقلب قبل از کشف آن ۱۸ ماه می‌باشد (جامعه ممیزان رسمی تقلب^۳، ۲۰۱۲). از این‌رو در ادامه مطابق با مبانی نظری پژوهش از ۴۰ متغیر مالی و غیرمالی مرتبط با اهرم مالی، سودآوری، ترکیب دارایی‌ها، نقدینگی، کارایی، اندازه، رشد، وضعیت مالی کلی، اندازه حسابر، دوره تصدی حسابر و تغییر حسابر به همراه انتخاب ویژگی مبتنی بر همبستگی و روش‌های شبکه عصبی مصنوعی، شبکه بیزین و الگوریتم جنگل تصادفی برای پیش‌بینی تقلب در صورت‌های مالی برای به ترتیب برای هر یک از سال‌های یک سال قبل از وقوع تقلب و دو سال قبل از وقوع آن ارائه شده است. به کارگیری داده‌های آموزشی جهت برآورد عملکرد مدل ممکن است با سویه همراه باشد. در بیش‌تر موارد مدل‌ها به جای **یادگیر** تمایل به حفظ کردن نمونه دارند (بیش برآزش^۴ داده‌ها). در این پژوهش برای اجتناب از این مشکل از اعتبارسنجی متقابل با ۱۰ لایه^۵ استفاده شده است؛ بنابراین برای هر لایه، مدل با استفاده از ۹ لایه باقیمانده آموزش داده می‌شود و با لایه جدا شده آزمون می‌گردد. در نهایت میانگین عملکرد محاسبه خواهد گردید.

1 He et al.

2 ADASYN

3 Association of Certified Fraud Examiners

4 Overfitting

5 10-Fold Cross Validation

جدول (۳). نتایج حاصل از مدل های پیشنهادی یک سال قبل از تقلب ($t - 1$)

صحت کلی	شبکه عصبی مصنوعی	شبکه بیزین	جنگل تصادفی
۹۲/۷۴٪	۹۱/۱۳٪	۹۴/۳۵٪	
۸۶/۹۶٪	۸۴/۵۱٪	۹۰/۷۷٪	
۱۰۰/۰۰٪	۱۰۰/۰۰٪	۹۸/۳۳٪	
۹۳/۰۲٪	۹۱/۶۰٪	۹۴/۴۰٪	
۰/۹۱۳	۰/۹۷۱	۰/۹۹۶	
<i>AUC</i>			
<i>F - Measure</i>			
فراخوانی			
دقت			

همان طور که در جدول بالا مشاهده می شود نتایج معیارهای عملکرد مدل های پیشنهادی یک سال قبل از وقوع تقلب حاکی از مناسب بودن مدل های شبکه عصبی مصنوعی، شبکه بیزین و الگوریتم جنگل تصادفی برای پیش بینی تقلب است. به منظور مقایسه عملکرد مدل های پیشنهادی در کشف تقلب از آزمون های آماری استفاده خواهد شد. همان طور که پیش از این اشاره شد در این پژوهش برای برآورد معیارهای ارزیابی عملکرد از رویکرد اعتبار سنجی متقابل با ۱۰ لایه و در ۱۰ تکرار استفاده شده است؛ بنابراین برای بررسی فرض نرمال بودن توزیع معیارهای عملکرد محاسبه شده از آزمون کولموگوروف - اسمیرنوف (KS) استفاده گردید. نتایج حاصل از این آزمون در سطح معناداری ۵ درصد حاکی از رد فرض H_0 مبنی بر نرمال بودن توزیع معیارهای ارزیابی عملکرد می باشد. بنابراین از آزمون رتبه علامت دار ویلکاکسون (Wilcoxon) برای مقایسه معیارهای ارزیابی عملکرد استفاده خواهد شد. جدول شماره (۴) نتایج حاصل از آزمون رتبه علامت دار ویلکاکسون (Wilcoxon) را برای معیار ارزیابی عملکرد AUC یک سال قبل از وقوع تقلب نشان می دهد. در حالت کلی تفاوت معناداری بین عملکرد الگوریتم جنگل تصادفی با شبکه عصبی مصنوعی و شبکه بیزین وجود دارد. از این رو، الگوریتم جنگل تصادفی روش مناسب تری برای طبقه بندی شرکت های متقلب و غیر متقلب است.

جدول (۴). نتایج حاصل از آزمون رتبه علامت دار ویلکاکسون (Wilcoxon)

طبقه بندی کننده ها	آماره آزمون	آماره Z	احتمال آماره Z
شبکه عصبی مصنوعی و شبکه بیزین	۴۹۹/۵۰۰	۳/۰۰۰	۰/۰۰۳
شبکه عصبی مصنوعی و جنگل تصادفی	۱۱۰۲/۵۰۰	۶/۰۶۲	۰/۰۰۰
شبکه بیزین و جنگل تصادفی	۴۱۸/۵۰۰	۴/۳۸۲	۰/۰۰۰

جدول شماره (۵) نتایج حاصل از برآورد مدل های پیشنهادی برای دو سال قبل از وقوع تقلب ($t - 2$) را نشان می دهد.

جدول (۵). نتایج حاصل از مدل‌های پیشنهادی دو سال قبل از تقلب ($t - 2$)

شبکه عصبی مصنوعی	شبکه بیزین	جنگل تصادفی	
صحت کلی	٪۹۵/۹۷	٪۹۷/۵۸	٪۹۸/۳۹
دقت	٪۹۲/۳۱	٪۹۵/۲۴	٪۹۶/۷۷
فراخوانی	٪۱۰۰/۰۰	٪۱۰۰/۰۰	٪۱۰۰/۰۰
$F - Measure$	٪۹۶/۰۰	٪۹۷/۵۶	٪۹۸/۳۶
AUC	۰/۹۵۹	۱/۰۰۰	۱/۰۰۰

چنانکه در جدول شماره (۵) مشاهده می‌شود، عملکرد مدل‌های پیشنهادی در دو سال قبل از وقوع تقلب نیز با یکدیگر متفاوت است. بررسی مقایسه‌ای عملکرد مدل‌های پیشنهادی حاکی از بهتر بودن معیارهای عملکرد الگوریتم جنگل تصادفی و شبکه بیزین نسبت به شبکه عصبی مصنوعی برای پیش‌بینی تقلب است.

جدول (۶). نتایج حاصل از آزمون رتبه علامت‌دار ویلکاکسون (Wilcoxon)

طبقه‌بندی کننده‌ها	آماره آزمون	آماره Z	احتمال آماره Z
شبکه عصبی مصنوعی و شبکه بیزین	۸۲۰/۰۰۰	۵/۵۶۴	۰/۰۰۰
شبکه عصبی مصنوعی و جنگل تصادفی	۸۲۰/۰۰۰	۵/۵۶۴	۰/۰۰۰
شبکه بیزین و جنگل تصادفی	۰/۰۰۰	NAN	۱/۰۰۰

جدول شماره (۶) نتایج حاصل از آزمون رتبه علامت‌دار ویلکاکسون (Wilcoxon) برای معیار ارزیابی عملکرد AUC را برای دو سال قبل از وقوع تقلب نشان می‌دهد. از این‌رو، تفاوت معناداری بین عملکرد الگوریتم جنگل تصادفی و شبکه بیزین با شبکه عصبی مصنوعی وجود دارد. اما بین عملکرد الگوریتم جنگل تصادفی و شبکه بیزین، تفاوت معناداری مشاهده نشد. بنابراین، الگوریتم جنگل تصادفی و شبکه بیزین روش مناسب‌تری برای طبقه‌بندی شرکت‌های متقلب و غیر متقلب است.

همچنین در این پژوهش از آزمون من‌ویتنی (U آزمون) برای مقایسه معیارهای ارزیابی عملکرد یک سال قبل از وقوع تقلب ($t - 1$) با دو سال قبل از وقوع آن ($t - 2$) استفاده خواهد شد. جدول شماره (۷) نتایج حاصل این آزمون را برای هر یک از مدل‌های پیشنهادی نشان می‌دهد.

جدول (۷). نتایج حاصل از آزمون من ویتنی (آزمون U)

طبقه بندی کننده ها	آماره آزمون	آماره Z	احتمال آماره Z
شبکه عصبی مصنوعی	۶۶۴۰/۰۰۰	۴/۴۱۱	۰/۰۰۰
شبکه بیزین	۶۷۰۰/۰۰۰	۶/۳۸۴	۰/۰۰۰
جنگل تصادفی	۵۵۰۰/۰۰۰	۳/۲۳۶	۰/۰۰۱

همان طور که در جدول شماره (۷) مشاهده می شود، تفاوت معناداری بین عملکرد مدل های پیشنهادی در یک سال قبل از وقوع تقلب ($t - 1$) و دو سال قبل از وقوع آن ($t - 2$) وجود دارد. از این رو می توان نتیجه گرفت استفاده از داده های دو سال قبل برای پیش بینی تقلب در صورت های مالی عملکرد بهتری خواهد داشت.

به طور کلی، نتایج این پژوهش با نتایج پژوهش های پرسونس (۱۹۹۵)، فروز و همکاران (۲۰۰۰)، سپاتیس و همکاران (۲۰۰۲)، کامینسکی و همکاران (۲۰۰۴)، کرکوز و همکاران (۲۰۰۷)، آلدن و همکاران (۲۰۱۲)، هوانگ و همکاران (۲۰۱۴)، چن و همکاران (۲۰۱۴) و لین و همکاران (۲۰۱۶) مبنی بر سودمندی شیوه های داده کاوی برای پیش بینی تقلب در صورت های مالی مطابقت دارد. همچنین، نتایج حاصل از انتخاب ویژگی به روش مبتنی بر همبستگی برای دو سال قبل از وقوع تقلب حاکی از سودمندی متغیرهای نسبت پوشش بهره، نسبت حساب های دریافتنی به کل دارایی ها، نسبت موجودی کالا به فروش خالص، نسبت نقدی، لگاریتم طبیعی فروش، نسبت سود خالص به فروش و نسبت جمع دارایی های جاری به کل دارایی ها برای کشف تقلب داشت. بنابراین می توان نتیجه گرفت که عوامل خطر تقلب مرتبط با اهرم مالی، سودآوری، ترکیب دارایی ها، نقدینگی و اندازه شرکت از اهمیت قابل توجهی در پیش بینی تقلب در صورت های مالی شرکت های پذیرفته شده در بورس اوراق بهادار تهران برخوردار هستند.

۶- نتیجه گیری و پیشنهادها

دو جریان متفاوت در متون پژوهشی، پیشینه و مبنایی برای این پژوهش فراهم می کند؛ یکی پژوهش های انجام شده در رابطه با گزارشگری مالی متقلبانه و دیگری متون مربوط به داده کاوی است. این دو جریان، قلمروهای پژوهشی گسترده ای هستند که پژوهش های متقابل بسیار اندکی در داخل کشور در این زمینه انجام شده است. از این رو، این پژوهش در تلاش است تا با معرفی رویکردهای داده کاوی کارا تر برای طبقه بندی صورت های مالی متقلبانه، متون پژوهشی مربوطه را بهبود بخشد. بنابراین، این سؤال مطرح شد که آیا می توان با استفاده از اطلاعات مالی و غیرمالی گزارش های سالانه شرکت ها مدلی کمی برای کشف تقلب ارائه داد که بیان گر یک روش تحلیلی خودکار برای کشف تقلب بالقوه باشد؟ برای پاسخگویی به این سؤال نیز مطابق با مبانی نظری و

پیشینه پژوهش از ۴۰ متغیر مالی و غیرمالی به همراه تکنیک‌های شبکه عصبی مصنوعی، شبکه بیزین و الگوریتم جنگل تصادفی استفاده گردید. یافته‌های پژوهش بیان‌گر وجود شواهدی دال بر عملکرد مناسب مدل‌های پیشنهادی برای پیش‌بینی تقلب در صورت‌های مالی است. نتایج حاصل از آزمون رتبه علامت‌دار ویلکاکسون (Wilcoxon) نیز حاکی از برتری الگوریتم جنگل تصادفی و شبکه بیزین در مقایسه با شبکه عصبی مصنوعی برای کشف تقلب در صورت‌های مالی بود. همچنین، نتایج آزمون من‌ویتنی (آزمون U) نشان می‌دهد داده‌های دو سال قبل برای پیش‌بینی تقلب در صورت‌های مالی مناسب‌تر از داده‌های یک سال قبل است. با توجه به ماهیت متفاوت فعالیت شرکت‌های بیمه، بانک‌ها و موسسه‌های اعتباری و واسطه‌گری‌های مالی این‌گونه شرکت‌ها از جامعه مورد مطالعه حذف شده‌اند و بنابراین نتایج پژوهش قابل‌تعمیم به تمامی شرکت‌ها نیست. بنابراین، بررسی علائم خطر مرتبط با تقلب در صورت‌های مالی و همچنین ارائه مدلی برای کشف صورت‌های مالی متقلبانه این شرکت‌ها و مؤسسات می‌تواند یک موضوع پژوهشی بالقوه باشد. همچنین، به پژوهشگران آینده پیشنهاد می‌شود با استفاده از سایر تکنیک‌های داده‌کاوی برای خوشه‌بندی، انتخاب ویژگی و طبقه‌بندی، به بررسی و تأیید الگو و چارچوب پیشنهادی اقدام کنند. در انجام این پژوهش محدودیت‌هایی به شرح زیر وجود داشت: ارائه مجدد اکثر صورت‌های مالی در سال بعد که خود نیازمند پژوهشی در مورد علل آن و همچنین معنی‌دار بودن یا نبودن تفاوت بین اطلاعات تجدید ارائه‌شده و تجدید ارائه نشده می‌باشد.

عدم دسترسی به داده‌ها و اطلاعات مالی شرکت‌ها در بلندمدت سبب گردید تا نتایج پژوهش تنها محدود به دوره زمانی کوتاهی باشد و بنابراین نتایج آن قابل‌تعمیم به دوره‌های زمانی بلندمدت نخواهد بود.

۷- منابع

- اعتمادی، حسین و حسن زلفی. (۱۳۹۲). کاربرد رگرسیون لجستیک در شناسایی گزارشگری مالی متقلبانه، *دانش حسابرسی* ۱۳(۵۱)، ۱۶۳-۱۴۵.
- پورحیدری، امید؛ مجتبی صفی‌پور افشار؛ علی گودرز تله جردی و معصومه صفی‌پور افشار. (۱۳۹۴). بررسی تاثیر کیفیت حسابرسی بر هزینه حسابرسی و قیمت‌گذاری کم‌تر از واقع در عرضه‌های اولیه، *فصلنامه حسابداری مالی* ۷(۲۶)، ۵۱-۳۱.
- جهان‌شاد، آریتا و سپیده سرداری‌زاده. (۱۳۹۳). رابطه معیار مالی (اختلاف رشد درآمد) و معیار غیر مالی (رشد تعداد کارکنان) با گزارشگری مالی متقلبانه، *پژوهش حسابداری* ۴(۱۳)، ۱۹۸-۱۸۱.

- راعی، رضا؛ پرویز حسن‌زاده و محسن حمشی. (۱۳۹۳). بررسی تاثیر نسبت‌های مالی بر چسبندگی هزینه‌ها، فصلنامه حسابداری مالی ۶(۲۲)، ۴۴-۲۶.
- صفرزاده، محمدحسین. (۱۳۸۹). توانایی نسبت‌های مالی در کشف تقلب در گزار شگری مالی: تحلیل لاجیت، دانش حسابداری ۱۱(۱)، ۱۶۳-۱۳۷.
- فرقاندوست حقیقی، کامییز؛ سید عباس هاشمی و امین فروغی دهکردی. (۱۳۹۳). مطالعه رابطه بین مدیریت سود و امکان تقلب در صورت‌های مالی شرکت‌های پذیرفته شده در بورس اوراق بهادار تهران، دانش حسابرسی ۱۴(۵۶)، ۴۷-۶۸.
- کمیته تدوین استانداردهای حسابرسی. (۱۳۹۴). اصول و ضوابط حسابداری و حسابرسی: استانداردهای حسابرسی. چاپ بیست و پنجم، تهران: سازمان حسابرسی.
- مشایخی، بیتا و امیرحسین حسین‌پور. (۱۳۹۵). بررسی رابطه بین مدیریت سود واقعی و مدیریت سود تعهدی در شرکت‌های مشکوک به تقلب بورس اوراق بهادار تهران، مطالعات تجربی حسابداری مالی ۱۲(۴۹)، ۵۲-۲۹.
- Alden, M.E., D.M. Bryan, B.J. Lessley, and A. Tripathy. (2012). Detection of financial statement fraud using evolutionary algorithms. **Journal of Emerging Technologies in Accounting**, (9): 71-94.
- Almeida, M.P.S.B. (2010). **Classification for fraud detection with social network analysis**. Dissertation, Engenharia Informática e de Computadores.
- Altman, E.I. (1968). Financial ratios, discriminant analysis, and the prediction of corporate bankruptcy. **Journal of Finance**, (23): 589-609.
- Association of Certified Fraud Examiners (ACFE) (2011). **Fraud Examiners Manual**. Austin, Texas: ACFE.
- Association of Certified Fraud Examiners (ACFE) (2012). **Report to the nations on occupational fraud and abuse**. Austin, Texas: ACFE.
- Brazel, J.F., K.L. Jones, and M.F. Zimbelman. (2009). Using nonfinancial measures to assess fraud risk. **Journal of Accounting Research**, (47): 1135-1166.
- Carcello, J.V., and A.L. Nagy. (2004). Audit firm tenure and fraudulent financial reporting. **Auditing: A Journal of Practice & Theory**, (23): 55-69.

- Chen, F.H., D.J. Chi, and J.Y. Zhu. (2014). Application of random forest, rough set theory, decision tree and neural network to detect financial statement fraud – taking corporate governance into consideration. **Intelligent Computing Theory**, (8588): 221-234.
- Dalnial, H., A. Kamaluddin, Z.M. Sanusi, and K.S. Khairuddin, (2014). Accountability in financial reporting: Detecting fraudulent firms. **Procedia- Social and Behavioral Sciences**, (145): pp. 61-69.
- Dua, S., and X. Du. (2011). **Data mining and machine learning in cybersecurity**. 1st edition, Taylor and Francis Group.
- Fawcett, T., (2006). An introduction to ROC analysis. **Pattern Recognition Letters**, (27): 861-874.
- He, H., Y. Bai, E.A. Garcia, and S. Li. (2008). ADASYN: Adaptive synthetic sampling approach for imbalanced learning. **IEE Neural Networks**, pp. 1322-1328.
- Hoogs, B., T. Kiehl, C. Lacombe, and D. Senturk. (2007). A genetic algorithm approach to detecting temporal patterns indicative of financial statement fraud: Research articles. **International Journal of Intelligent Systems in Accounting, Finance and Management**, (15): 41-56.
- Huang, S.H., R.H. Tsaih, and F. Yu. (2014). Topological pattern discovery and feature extraction for fraudulent financial Reporting. **Expert Systems with Applications**, (41): pp. 4360-4372.
- Kaminski, K.A., T.S. Wetzel, and L. Guan. (2004). Can financial ratios detect fraudulent financial reporting? **Managerial Auditing Journal**, (19): 15-28.
- Kirkos, E., C. Spathis, and Y. Manolopoulos. (2007). Data mining techniques for the detection of fraudulent financial statements. **Expert Systems with Applications**, (32): 995-1003.
- Lin, C-H., A-A. Chiu, S.Y. Huang, and D.C. Yen. (2016). Detecting the financial statement fraud: The analysis of the differences between data mining techniques and experts' judgments. **Knowledge-Based Systems**, Article in Press.
- Okike, E., (2011). **Theory and practice of corporate social responsibility**. 1st edition, Springer Berlin Heidelberg.
- Persons, O., (1995). Using financial statement data to identify factors associated with fraudulent financial reporting. **Journal of Applied Business Research**, (11): 38-46.
- Pierre, K. S., and J.A. Anderson, (1984). An analysis of the factors

- associated with lawsuits against public accountants. **The Accounting Review**, (59): 242-263.
- Ravisankar, P., V. Ravi, G.R. Rao, and I. Bose. (2011). Detection of financial statement fraud and feature selection using data mining techniques. **Decision Support Systems**, (50): 491-500.
- Rezaee, Z., (2005). Causes, consequences, and deterrence of financial statement fraud. **Critical Perspective on Accounting**, (16): 277-298.
- Rezaee, Z., and R. Riley, (2010). **Financial statement fraud-prevention and detection**. 2nd edition, John Wiley & Sons, Inc.
- Segal, S.Y., (2016). Accounting frauds – review of advanced technologies to detect and prevent frauds. **Economics and Business Review**, 16 (4): 45–64.
- Spathis, C.T., (2002). Detecting false financial statements using published data: Some evidence from Greece. **Managerial Auditing Journal**, (17): 179–191.
- Spathis, C., M. Doumpos, and C. Zopounidis. (2002). Detecting falsified financial statements: a comparative study using multicriteria analysis and multivariate statistical techniques. **European Accounting Review**, 11(3): 509-535.
- Sreejesh, S., M.R. Anusree, and S. Mohapatra. (2014). **Business research methods: An applied orientation**. 1st edition, Springer International Publishing.
- Stice, D.J. (1991) Using financial and market information to identify pre-engagement market factors associated with lawsuits against auditors. **Accounting Review**, (66): 516–33.
- Summers, S.L., and J.T. Sweeney. (1998) Fraudulently misstated financial statements and insider trading: An empirical analysis. **Accounting Review**, (73): 131–46.
- Whiting, D.G., J.V. Hansen, J.B. McDonald, C. Albrecht, and W. S. Albrecht. (2012). Machine learning methods for detecting patterns of management fraud. **Computational Intelligence**, (28): 505-527.
- Zhou, W., and G. Kapoor. (2011). Detecting evolutionary financial statement fraud. **Decision Support Systems**, (50): 570-575.
- Zikmund, W.G., B.J. Babin, J.C. Car, and M. Griffin. (2010). **Business research methods**. 8th edition, Cengage Learning.