

دریافت مقاله: ۱۴۰۴/۰۸/۲۶

پذیرش مقاله: ۱۴۰۴/۰۹/۲۲

نوع مقاله: پژوهشی

کشف تقلب در کارت‌های اعتباری با استخراج ویژگی از شبکه عصبی و طبقه‌بندی تقویت گرادیان سبک وزن

مائده نصیری

دانشجو، گروه ریاضیات و کاربردها، دانشکده علوم پایه، دانشگاه شاهد، تهران.

پست الکترونیکی: maedenasiri2000@gmail.com

ابوالفضل تاری مرزآباد*

دانشیار، گروه ریاضیات و کاربردها، دانشکده علوم پایه، دانشگاه شاهد، تهران.

پست الکترونیکی: tari@shahed.ac.ir

چکیده

با افزایش استفاده از کارت‌های اعتباری، تشخیص خودکار تراکنش‌های متقلبانه اهمیت بیشتری یافته است. در این پژوهش، دو مدل ترکیبی مبتنی بر شبکه عصبی پیچشی و پرسپترون چندلایه برای استخراج ویژگی به کار گرفته شد و ویژگی‌های استخراج شده به الگوریتم تقویت گرادیان سبک‌وزن منتقل گردید. برای انتخاب بهترین لایه، پنج معیار کمی شامل شاخص‌های کالینسکی-هارابان، دیویس-بولدین، اطلاعات متقابل، میزان افزونگی و ابعاد لایه برای همه لایه‌ها محاسبه و لایه بهینه براساس میانگین نرمال‌شده معیارها تعیین شد. نتایج نشان داد مدل‌های ترکیبی عملکرد بهتری نسبت به مدل منفرد دارند و مدل مبتنی بر شبکه پیچشی بیشترین کارایی را ارائه کرده است. این مدل موجب افزایش امتیاز F1، کاهش قابل‌توجه خطای مثبت کاذب و افزایش اندک خطای منفی کاذب شد؛ اما این افزایش در برابر کاهش محسوس خطای مثبت کاذب قابل پذیرش بوده و تعادل میان دقت و بازیابی را حفظ می‌کند.

آزمون ویلکاکسون تفاوت عملکرد را از نظر آماری معنادار نشان نداد، با این حال بهبود معیارهای عملکردی اهمیت استفاده از مدل ترکیبی مبتنی بر شبکه هم‌آمیختی را در سناریوهای واقعی تشخیص تقلب برجسته می‌سازد. تحلیل SHAP نیز نشان داد ویژگی دوم استخراج شده از شبکه هم‌آمیختی بیشترین نقش را داشته و تحت تأثیر متغیر «مقدار تراکنش» شکل گرفته است.

کلید واژه‌ها: کشف تقلب، مدل‌های ترکیبی، استخراج ویژگی، شبکه عصبی، تقویت گرادیان سبک وزن

۱- مقدمه

مهاجرت کسب و کارها به اینترنت و گسترش تراکنش‌های الکترونیکی در اقتصاد بدون پول نقد که به‌طور مداوم در حال رشد است [۱]، باعث افزایش چشمگیر حجم تراکنش‌ها شده است. این موضوع همچنین پیچیدگی و تنوع روش‌های تقلب را افزایش داده و مقابله با آن را دشوارتر کرده است. شرکت‌های صادرکننده کارت اعتباری معمولاً

* نویسنده مسئول

آمار خسارت‌ها را منتشر نمی‌کنند. بنابراین، دستیابی به برآورد دقیق خسارت‌ها دشوار است. با این حال، برخی داده‌های عمومی درباره زیان‌های مالی ناشی از قلب موجود است.

استفاده از کارت‌های اعتباری بدون امنیت کافی، میلیارد‌ها دلار خسارت مالی به همراه دارد [۲]. در گذشته، سامانه‌های تشخیص قلب با تکیه بر معیارهای از پیش تعریف‌شده، تراکنش‌های مشکوک را شناسایی می‌کردند. با این حال، محدودیت این سامانه‌ها در انطباق با الگوهای نوظهور و نرخ بالای هشدارهای کاذب، عملکرد آنها را تضعیف کرد. در نتیجه، رویکردهای مبتنی بر یادگیری ماشین و یادگیری عمیق با هدف ارتقای دقت و تطبیق‌پذیری^۱ به تدریج جایگزین روش‌های سنتی شدند [۳]. این الگوریتم‌ها قادرند با بهره‌گیری از داده‌های حجیم و تحلیل الگوهای رفتاری پیچیده، ناهنجاری‌هایی را که ممکن است نشان‌دهنده فعالیت متقلبانه باشند، شناسایی کنند [۴].

یکی از رویکردهای مطرح و مؤثر در این زمینه، بهره‌گیری از الگوریتم‌های یادگیری ماشین برای طبقه‌بندی تراکنش‌های مشکوک است. در سال‌های اخیر، تحقیقات زیادی در مورد استخراج ویژگی از شبکه‌های عصبی و ترکیب آن با الگوریتم‌های یادگیری ماشین برای افزایش دقت در کشف قلب انجام گرفته است. با این حال، بررسی مطالعات موجود نشان می‌دهد که برخی چالش‌های مهم همچنان بدون پاسخ باقی مانده‌اند، از جمله:

- عدم توجه به انتخاب لایه بهینه برای استخراج ویژگی از شبکه‌های عصبی که می‌تواند نقش تعیین‌کننده‌ای در عملکرد مدل داشته باشد.

- فقدان تحلیل آماری دقیق برای ارزیابی معناداری تفاوت عملکرد بین مدل‌ها.

- کم‌توجهی به تفسیرپذیری مدل‌ها و تحلیل علل تصمیم‌گیری آنها، که می‌تواند بینش مهمی درباره رفتار مدل ارائه دهد.

پژوهش حاضر با هدف رفع این شکاف‌ها، چارچوبی

ترکیبی برای تشخیص قلب ارائه می‌دهد. در این چارچوب، از شبکه عصبی هم‌آمیختگی^۲ و شبکه عصبی پرسپترون چند لایه^۳ برای استخراج ویژگی‌های عمیق استفاده شده و طبقه‌بندی با الگوریتم تقویت گرادیان سبک وزن^۴ انجام می‌شود. همچنین برای ارزیابی دقیق عملکرد مدل، از تحلیل‌های آماری و ابزارهای تفسیرپذیری بهره گرفته‌ایم. ساختار این مقاله به شرح زیر است:

در بخش دوم، مرور ادبیات پژوهش ارائه می‌شود. بخش سوم به معرفی مجموعه داده و مدل‌های پایه اختصاص دارد. در این بخش، فرایند پیش‌پردازش، معیارهای ارزیابی عملکرد، ساختار شبکه‌های عصبی و الگوریتم تقویت گرادیان سبک وزن تشریح می‌شوند. روش پیشنهادی در بخش چهارم توصیف می‌شود. نتایج تجربی در بخش پنجم گزارش و تحلیل شده و تفسیر نتایج مدل در بخش ششم مورد بررسی قرار می‌گیرد. در نهایت، بخش هفتم به نتیجه‌گیری و ارائه پیشنهادهایی برای تحقیقات آینده اختصاص دارد.

۲- پیشینه تحقیق

در یکی از مطالعات اخیر روشی ترکیبی با ترکیب رمزگذار خودکار^۵ و جنگل تصادفی احتمالاتی^۶ برای تشخیص قلب در

تراکنش‌های کارت اعتباری ارائه شده است. در این روش، ابتدا با استفاده از رمزگذار خودکار ابعاد داده‌ها کاهش یافته و ویژگی‌های مهم استخراج می‌شوند، سپس با استفاده از جنگل تصادفی با طبقه‌بندی احتمالات تراکنش‌ها به صورت عدم قلب یا قلب طبقه‌بندی می‌شود [۵].

در [۶]، با به کارگیری ترکیب شبکه عصبی هم‌آمیختگی به منظور استخراج مؤثر ویژگی‌ها و ماشین بردار پشتیبان^۷ به عنوان طبقه‌بند نهایی، چارچوبی ارائه شده است که در

2-Convolutional neural network

3-Multilayer perceptron

4-Lightweight Gradient Boosting Machine (LightGBM)

5-Autoencoder

6-Probabilistic Random Forest

7-Support vector machines

1-Adaptability

مقایسه با رویکردهای متداول، دقت بالاتری در تشخیص تقلب در تراکنش‌های کارت اعتباری ارائه می‌دهد.

در پژوهشی دیگر، مدلی ترکیبی مبتنی بر شبکه عصبی هم‌آمیختگی، حافظه طولانی کوتاه مدت^۸ و الگوریتم k-نزدیک‌ترین همسایه^۹ برای تشخیص تقلب در تراکنش‌های بانکی ارائه شده است. مدل پیشنهادی با صحت ۹۹/۵ درصد، میزان بازیابی ۹۲/۱ درصد و مساحت زیر منحنی مشخصه عملکرد گیرنده^{۱۰} برابر با ۹۷/۵ درصد، عملکرد بهتری نسبت به الگوریتم‌های سنتی مانند رگرسیون لجستیک^{۱۱}، جنگل تصادفی و درخت تقویت‌شده تدریجی^{۱۲} دارد. ترکیب لایه‌های شبکه عصبی هم‌آمیختگی برای درک روابط پیچیده میان ویژگی‌ها و لایه‌های حافظه بلندمدت-کوتاه‌مدت برای تحلیل الگوهای زمانی، موجب شناسایی دقیق‌تر تراکنش‌های متقلبانه شده است. این مدل، علی‌رغم پیچیدگی و چالش‌هایی در زمینه تفسیرپذیری، راهکاری نویدبخش برای استفاده در سامانه‌های مالی ارائه می‌دهد[۷].

در [۸] با تلفیق شبکه عصبی هم‌آمیختگی به عنوان استخراج‌کننده ویژگی و الگوریتم جنگل تصادفی به عنوان طبقه‌بند، یک مدل ترکیبی برای تشخیص تقلب در تراکنش‌های کارت اعتباری ارائه شده است. نتایج تجربی نشان می‌دهند که این مدل با دستیابی به صحت ۹۹/۹۸ درصد، عملکردی فراتر از سایر الگوریتم‌های کلاسیک مانند ماشین بردار پشتیبان با حداکثر صحت ۹۹/۹۴ درصد داشته است.

در [۹]، روشی خلاقانه برای تشخیص تقلب در کارت‌های اعتباری ارائه شده است که از ترکیب یادگیری عمیق و بینایی ماشین بهره می‌برد. در این رویکرد، داده‌های متنی تراکنش‌ها با استفاده از تکنیکی نو به نام تبدیل متن به تصویر^{۱۳} به تصاویر کوچکی تبدیل شده و سپس به

یک معماری شبکه عصبی هم‌آمیختگی تغذیه می‌شوند. برای مقابله با عدم تعادل شدید کلاس‌ها، از روش وزن‌دهی معکوس فرکانس^{۱۴} استفاده شده است. نتایج تجربی نشان می‌دهند که ویژگی‌های عمیق استخراج‌شده از شبکه عصبی هم‌آمیختگی، زمانی که به الگوریتم‌های کلاسیک یادگیری ماشین مانند نزدیک‌ترین همسایه خشن^{۱۵} داده می‌شوند، منجر به دقتی معادل ۹۹/۸۷ درصد در تشخیص تراکنش‌های متقلبانه شده‌اند. این روش با معرفی بعدی جدید در تبدیل داده‌های متنی به تصاویر، مسیرهای جدیدی برای به‌کارگیری تکنیک‌های بینایی ماشین در حوزه تشخیص تقلب پیشنهاد می‌دهد.

در [۱۰]، مدلی ترکیبی دیگری برای تشخیص تقلب در تراکنش‌های کارت اعتباری معرفی شده است که از ادغام پرسپترون چندلایه به عنوان استخراج‌کننده ویژگی و طبقه‌بند جنگل تصادفی بهره می‌برد. هدف این مدل، بهره‌گیری هم‌زمان از قدرت یادگیری الگوهای پیچیده توسط پرسپترون چندلایه و قابلیت تصمیم‌گیری قوی جنگل تصادفی است. پس از پیش‌پردازش داده‌ها و متعادل‌سازی آنها با استفاده از تکنیک نمونه‌برداری افزایشی مصنوعی برای کلاس اقلیت، عملکرد مدل پیشنهادی با مدل‌های مستقل پرسپترون چندلایه و جنگل تصادفی مقایسه شده است. نتایج نشان می‌دهد که مدل ترکیبی عملکرد بهتری نسبت به مدل‌های مستقل دارد و با صحت ۹۹/۹۴ درصد، دقت ۸۷/۰۹ درصد، بازیابی ۸۲/۶۵ درصد و امتیاز F1 معادل ۸۴/۸۱ درصد، توانسته است بهبود قابل‌توجهی در تشخیص تراکنش‌های متقلبانه ایجاد کند.

در [۱۱]، برای تشخیص تقلب، ویژگی‌های استخراج‌شده از شبکه عصبی هم‌آمیختگی به عنوان ورودی به مدل‌های یادگیری ماشین داده شدند. مدل‌های یادگیری ماشین شامل بیز ساده^{۱۶}، k-نزدیک‌ترین همسایه، درخت تصمیم^{۱۷} و ماشین بردار پشتیبان به صورت تجمعی با ویژگی‌های

14-Inverse Frequency Weighting

15-Coarse-KNN

16-Naïve Bayes

17-Decision Tree

8-Long Short-Term Memory

9-K-Nearest Neighbor

10-Area Under the Curve

11-Logistic regression

12-Gradient boosting

13-Text2IMG

جدول ۱. مروری مقایسه‌ای از پژوهش‌های منتخب

دسته‌بندی مدل ترکیبی	مرجع	مجموعه داده	مدل	نوع پیش‌پرازش	لایه استخراج ویژگی
شبکه عصبی هم‌آمیختگی + الگوریتم‌های طبقه‌بند	[۶]	داده‌های اروپایی	شبکه عصبی هم‌آمیختگی + ماشین بردار پشتیبان	متعادل سازی	لایه تخت شده
	[۸]	داده‌های اروپایی	شبکه عصبی هم‌آمیختگی + جنگل تصادفی	نرمال سازی + حذف ویژگی‌های غیرضروری	لایه تخت شده
	[۱۱]	داده‌های اروپایی	شبکه عصبی هم‌آمیختگی + یادگیری تجمیعی (k_ نزدیک ترین همسایه، درخت تصمیم، بیزساده، ماشین بردار پشتیبان)	متعادل سازی	لایه سوم تمام متصل
	[۹]	داده‌های اروپایی	شبکه عصبی هم‌آمیختگی + k- نزدیک ترین همسایه	تبدیل داده‌ها تصاویر ۵×۶ پیکسلی + وزن دهی معکوس کلاسی‌ها	لایه تمام متصل
پرسپترون چندلایه + الگوریتم‌های طبقه‌بند	[۱۰]	داده‌های اروپایی	پرسپترون چندلایه + جنگل تصادفی	پاک سازی داده‌ها + نرمال سازی + متعادل سازی	لایه آخر
چند شبکه عصبی + الگوریتم‌های طبقه‌بند	[۷]	داده بانکی شبیه سازی شده	شبکه عصبی هم‌آمیختگی + حافظه طولانی کوتاه مدت k + - نزدیک ترین همسایه	متعادل سازی	آخرین لایه تمام متصل
خودرمزگذار + الگوریتم‌های طبقه‌بند	[۱۲]	داده‌های اروپایی	خودرمزگذار + تقویت گرادیان سبک وزن	پاک سازی داده‌ها + متعادل سازی + نرمال سازی + حذف ویژگی	لایه فشرده خودرمزگذار
	[۵]	داده‌های اروپایی	خودرمزگذار + جنگل تصادفی	متعادل سازی	لایه فشرده خودرمزگذار
خودرمزگذار + شبکه عصبی	[۱۳]	داده آلمانی	خودرمزگذار + شبکه عصبی	نرمال سازی + متعادل سازی	لایه فشرده خودرمزگذار

استخراج شده آموزش دیده‌اند. نتایج نشان دادند که در برخی معیارهای ارزیابی، عملکرد مدل‌ها به نزدیک به ۱۰۰ درصد رسید.

در [۱۲]، از خودرمزگذار همراه با الگوریتم تقویت گرادیان سبک وزن استفاده شده است. در این روش، ابتدا از خودرمزگذار برای کاهش ابعاد داده‌ها و استخراج ویژگی‌های کلیدی استفاده می‌شود و سپس با استفاده از الگوریتم تقویت گرادیان سبک وزن طبقه‌بندی انجام می‌شود. نتایج نشان می‌دهد که این روش به‌طور طبیعی با داده‌های نامتوازن مطابقت دارد و در مقایسه با روش‌های سنتی مانند الگوریتم k- نزدیک ترین همسایه، جنگل تصادفی و الگوریتم تقویت گرادیان سبک وزن عملکرد بهتری ارائه می‌دهد.

در [۱۳]، ویژگی‌های معنادار و مهم توسط خودرمزگذار استخراج شده و این ویژگی‌ها به شبکه‌های عصبی دیگر برای طبقه‌بندی و شناسایی تراکنش‌های تقلبی داده شده‌اند. این مدل ترکیبی نه تنها نسبت به مدل‌های منفرد شبکه عصبی عملکرد بهتری داشته، بلکه در مقایسه با روش تحلیل مولفه‌های اصلی نیز نمایش بهتری از داده‌ها و کاهش بعد مؤثرتری ارائه کرده است. این نتایج نشان‌دهنده کارآمدی شبکه‌های عصبی عمیق در بهبود عملکرد روش‌های تشخیص تقلب کارت اعتباری است.

در بسیاری از مطالعات پیشین، روش‌های ترکیبی مختلفی برای شناسایی تقلب ارائه شده است. با این وجود، این پژوهش‌ها به گزارش نتایج عددی بسنده کرده‌اند و کمتر به تحلیل آماری پرداخته‌اند. در جدول ۱ مروری مقایسه‌ای

ارزیابی عملکرد مدل‌ها و ساختار مدل‌های پایه، شامل شبکه‌های عصبی و الگوریتم تقویت گرادیان سبک وزن را شرح می‌دهد.

۳-۱ مجموعه داده و پیش‌پردازش

مجموعه داده مورد استفاده در این مقاله شامل ۲۸۴۸۰۷ تراکنش کارت اعتباری اروپایی در سپتامبر ۲۰۱۳ و دارای دو کلاس «تقلب» و «عدم تقلب» است که از این میان تنها ۴۹۲ تراکنش تقلبی بوده و در نتیجه مجموعه داده به شدت نامتعادل است. این داده‌ها شامل ۳۰ ویژگی هستند، ۲۸ ویژگی ناشناس (که از 1V تا 28V نام‌گذاری شده‌اند) حاصل از تحلیل مؤلفه‌های اصلی و دو ویژگی اصلی «زمان» و «مقدار تراکنش» بدون تغییر هستند.

در مرحله پیش‌پردازش، داده‌های تکراری حذف و ویژگی زمان کنار گذاشته شد. این ویژگی تنها فاصله هر تراکنش از آغاز فرایند ثبت داده‌ها را نشان می‌دهد و زمان واقعی وقوع تراکنش‌ها را در طول روز یا هفته مشخص نمی‌کند. افزون بر این، الگوریتم مورد استفاده وابستگی‌های زمانی میان تراکنش‌ها را در نظر نمی‌گیرد.

برای ساخت مدل، داده‌ها به نسبت ۸۰ درصد برای آموزش (۲۵ درصد از آن برای اعتبارسنجی) و ۲۰ درصد برای آزمون نهایی تقسیم شدند. تقسیم داده‌ها پیش از هرگونه عملیات پیش‌پردازش انجام شد تا از بروز هرگونه نشت داده جلوگیری گردد.

برای نرمال‌سازی داده‌ها از روش Z-Score استفاده شد، به طوری که میانگین و انحراف معیار صرفاً بر اساس داده‌های آموزشی محاسبه گردید و همان پارامترها بدون هیچ‌گونه بازآموزی بر داده‌های اعتبارسنجی و آزمون اعمال شد. این رویکرد تضمین می‌کند که هیچ اطلاعاتی از توزیع داده‌های آزمون وارد فرآیند آموزش نشود.

به منظور رفع عدم تعادل کلاس‌ها، عملیات نمونه‌برداری افزایشی تنها بر روی داده‌های آموزشی نرمال‌شده اجرا شد. در این روش، با درونیابی میان هر نمونه از کلاس

از پژوهش‌های منتخب سال‌های ۲۰۲۰ تا ۲۰۲۵ ارائه شده است. در این جدول، مجموعه داده، مدل، پیش‌پردازش و لایه استخراج ویژگی ذکر شده است.

تحلیل منابع بررسی شده نشان می‌دهد که علی‌رغم عملکرد قابل قبول روش‌های ترکیبی مبتنی بر استخراج ویژگی از شبکه‌های عصبی و استفاده از الگوریتم‌های کلاسیک یادگیری ماشین، اغلب مقایسه‌ها فقط به ارزیابی عددی مدل‌ها بسنده کرده‌اند و آزمون‌های آماری برای بررسی پایداری و قابل اعتماد بودن نتایج، ارائه نشده است. در زمینه استخراج ویژگی از شبکه‌های عصبی نیز، در ادبیات موجود معیار دقیق و مشخصی برای تعیین لایه بهینه ارائه نشده است. به‌ویژه در معماری‌هایی که کمتر مورد استفاده قرار گرفته‌اند، مشخص نیست که کدام لایه می‌تواند بهترین بازنمایی را برای تشخیص تقلب فراهم کند. این خلأ می‌تواند منجر به استخراج ویژگی‌هایی با قدرت بازنمایی ناکافی شود، که در نهایت باعث کاهش دقت و پایداری مدل در مسائل حساس نظیر تشخیص تقلب می‌شود.

موضوع تفسیرپذیری مدل‌ها اغلب نادیده گرفته شده و کمتر مطالعه‌ای به بررسی میزان درک‌پذیری خروجی مدل از نظر کاربران نهایی یا کارشناسان بانکی پرداخته است. این کاستی بیانگر شکاف‌های پژوهشی مهمی هستند که پژوهش حاضر در راستای پر کردن آنها طراحی شده است.

در ادامه، روش پیشنهادی ما که بر انتخاب لایه بهینه از شبکه عصبی هم‌آمیختگی و پرسپترون چندلایه برای استخراج ویژگی، استفاده از الگوریتم تقویت گرادیان سبک وزن برای طبقه‌بندی با ویژگی‌های استخراج شده، انجام آزمون آماری و تقویت تفسیرپذیری مبتنی است، معرفی می‌شود.

۳- مجموعه داده و مدل‌های پایه

این بخش، مجموعه داده و پیش‌پردازش آن، معیارهای

۳-۳-۱ ساختار شبکه‌های عصبی

رویکردهای مبتنی بر یادگیری عمیق، روش‌های پیشرفته‌ای هستند که برای تحلیل و شناسایی الگوهای پیچیده در مجموعه‌های داده بزرگ، مانند تراکنش‌های کارت اعتباری، استفاده می‌شوند [۱۴]. شبکه عصبی پرسپترون چندلایه و شبکه عصبی هم‌آمیختی از زیرشاخه‌های یادگیری عمیق هستند.

شبکه عصبی پرسپترون چندلایه از مجموعه‌ای از نورون‌ها تشکیل شده و این نورون‌ها از طریق وزن‌های پیوندی^{۲۲} به یکدیگر متصل‌اند. عملکرد پرسپترون چندلایه به‌گونه‌ای است که یک مجموعه ورودی را به خروجی دلخواه تبدیل می‌کند [۱۵]. شبکه عصبی پیچشی بر پایه ساختار پرسپترون چندلایه توسعه یافته‌است [۱۶].

شبکه‌های عصبی مورد استفاده در این پژوهش شامل دو مدل مستقل، پرسپترون چندلایه با چهار لایه و شبکه عصبی هم‌آمیختی با ۵ لایه بودند که آموزش آنها با حداکثر ۱۵۰ و ۳۰ دوره با استفاده از توقف زودهنگام برای جلوگیری از بیش‌برازش انجام شد و مدل‌ها پس از ۷۴ دوره و ۱۴ دوره به دلیل عدم پیشرفت معیار صحت در داده‌های اعتبارسنجی طی چهار دوره پیاپی، متوقف شدند برای بهبود فرایند یادگیری مدل‌ها، از لایه نرمال‌سازی دسته‌ای^{۲۳}، برای جلوگیری از بیش‌برازش از لایه حذف تصادفی^{۲۴}، برای کاهش ابعاد ویژگی از لایه‌ی تجمیع بیشینه^{۲۵} و برای ایجاد توانایی یادگیری روابط غیرخطی در شبکه از توابع فعال‌ساز^{۲۶} LeakyReLU، ReLU و Sigmoid استفاده شد. در پایان آموزش، مدل پرسپترون چندلایه به صحتی ۰/۹۹۵۲ در داده‌های آزمون و ۰/۹۹۶۲ در داده‌های اعتبارسنجی دست یافت و مقادیر زیان متناظر به‌ترتیب برابر با ۰/۰۲۲ و ۰/۰۲۰ گزارش شد. همچنین، مدل شبکه عصبی هم‌آمیختی صحتی معادل ۰/۹۹۹۹ در داده‌های آموزشی و ۰/۹۹۹۴ در داده‌های اعتبارسنجی

اقلیت و K همسایه نزدیک آن، داده‌های مصنوعی جدید ایجاد می‌شود تا تعادل نسبی میان کلاس‌ها برقرار گردد. تعداد همسایه‌ها $k=5$ در نظر گرفته شد تا تنوع کافی در داده‌های تولید شده ایجاد شود و از هم‌پوشانی بیش از حد میان کلاس‌ها جلوگیری گردد. لازم به ذکر است که هیچ‌گونه نمونه‌برداری یا تغییر ساختاری بر روی مجموعه‌های اعتبارسنجی و آزمون اعمال نشد تا ساختار واقعی داده‌ها حفظ شود و ارزیابی مدل به صورت بی‌طرفانه انجام گیرد

۳-۲ معیارهای ارزیابی

معیارهای ارزیابی این پژوهش شامل صحت^{۱۸}، دقت^{۱۹}، بازیابی^{۲۰} و امتیاز F1^{۲۱} هستند.

● صحت: نسبت کل پیش‌بینی‌های درست انجام شده به کل داده‌ها.

● دقت: نسبت پیش‌بینی‌های درست مثبت به کل پیش‌بینی‌های مثبت.

● بازیابی: نسبت داده‌های مثبت واقعی که درست شناسایی شده‌اند به کل داده‌های مثبت واقعی.

● امتیاز F1: میانگین هارمونیک دقت و بازیابی.

و همچنین از دو مفهوم خطای مثبت کاذب و خطای منفی کاذب نیز استفاده شده است.

● خطای مثبت کاذب: مدل تراکنش «عدم تقلب» را به اشتباه «تقلب» تشخیص می‌دهد.

● خطای منفی کاذب: مدل تراکنش «تقلب» را به اشتباه «عدم تقلب» تشخیص می‌دهد.

۳-۳ ساختار مدل‌ها

در این بخش، شبکه‌های عصبی و الگوریتم تقویت گرادیان سبک وزن مورد استفاده، همراه با ویژگی‌ها و ساختار آنها ارائه شده‌است.

22-Linking Weights
23-Batch Normalization
24-Dropout
25-MaxPooling
26-Activation Functions

18-Accuracy
19-Precision
20-Recall
21-F1 Score

مناسب برای استخراج ویژگی‌ها از شبکه‌های عصبی به صورت شفاف توضیح داده نشده است. در اغلب موارد، یا اساس انتخاب لایه ذکر نمی‌شود، یا صرفاً بر پایه نتایج تجربی، لایه بهتر انتخاب می‌شود. این روش‌ها، علاوه بر ایجاد ابهام، باعث افزایش پیچیدگی محاسباتی و افزایش زمان اجرای مدل خواهند شد.

در این پژوهش، مجموعه‌ای از شاخص‌های کمی معرفی و ارزیابی شده‌اند تا با تحلیل یکپارچه آنها بتوان بدون نیاز به آزمون‌های مکرر و پرهزینه، لایه بهینه برای استخراج ویژگی‌ها را انتخاب نمود. این شاخص‌ها عبارتند از:

۱) شاخص کالینسکی-هاراباز

ویژگی‌های استخراج‌شده باید بتوانند داده‌ها را به خوبی از یکدیگر متمایز کنند و شرایط مناسبی برای خوشه‌بندی فراهم نمایند بنابراین از شاخص کالینسکی-هاراباز^{۲۹} که به صورت نسبت پراکندگی بین خوشه‌ای به پراکندگی درون خوشه‌ای تعریف می‌شود، برای ارزیابی کیفیت خوشه‌بندی استفاده شده است. مقدار بالاتر این شاخص نشان‌دهنده کیفیت بهتر خوشه‌بندی و در نتیجه، مناسب‌تر بودن ویژگی‌های استخراج‌شده برای تمایز داده‌ها است [۱۸] با این وجود، شاخص کالینسکی-هاراباز به طوری ذاتی برای خوشه‌های دارای ساختار محدب طراحی شده است و در شرایطی که خوشه‌ها دارای ساختارهای پیچیده‌تر یا غیرمحدب هستند، ممکن است عملکرد این شاخص کاهش یابد [۱۸]، بنابراین بهتر است که از شاخص دیویس-بولدین^{۳۰} به عنوان مکمل استفاده شده است.

۲) شاخص دیویس-بولدین

این شاخص میانگین شباهت بین هر خوشه و مشابه‌ترین خوشه دیگر را اندازه‌گیری می‌کند. این شباهت بر اساس نسبت پراکندگی درونی خوشه‌ها به فاصله بین مراکز آنها سنجیده می‌شود. هر چه مقدار این شاخص به صفر نزدیک‌تر باشد، نشان‌دهنده خوشه‌بندی بهتر و تفکیک‌پذیری بالاتر داده‌هاست [۱۸].

با مقادیر زیان ۰/۰۰۸۶ و ۰/۰۰۵۸ به دست آورد. نتایج مذکور بیانگر همگرایی مناسب هر دو مدل و عدم بروز بیش‌برازش محسوس در فرایند یادگیری هستند.

۳-۲-۳ الگوریتم تقویت گرادیان سبک وزن

ماشین تقویت گرادیان سبک‌وزن یک چارچوب قدرتمند برای پیاده‌سازی الگوریتم درخت تصمیم تقویت گرادیانی است که هم کارآمد بوده و هم به طور گسترده در پردازش مجموعه داده‌های بزرگ و ساختاریافته، به ویژه در حوزه‌های مالی مانند تشخیص تقلب کارت اعتباری، مورد استفاده قرار می‌گیرد. کارایی این الگوریتم در مدیریت مجموعه داده‌های بزرگ برای پردازش حجم بالای داده‌های تراکنش بسیار حیاتی است [۱۷].

بر خلاف بسیاری از ابزارهای سنتی درخت تصمیم تقویت گرادیانی که از رشد درخت بر پایه سطح^{۲۷} استفاده می‌کنند، ماشین تقویت گرادیان سبک‌وزن از یک راهبرد رشد مبتنی بر برگ^{۲۸} با محدودیت عمق بهره می‌برد. در این روش، در هر مرحله، برگی انتخاب می‌شود که بیشترین سود تقسیم را دارد، سپس آن برگ تقسیم می‌شود. این فرایند تا رسیدن به شرایط توقف ادامه می‌یابد. مزیت اصلی این رویکرد در مقایسه با روش سطحی آن است که می‌تواند با تعداد تقسیمات کمتر، خطای آموزش را کاهش داده و به دقت بالاتری برسد [۱۲].

۴- روش‌شناسی

در این بخش، معیارهای کمی به کاررفته برای انتخاب لایه بهینه در استخراج ویژگی‌ها تشریح شده است، سپس فرایند استخراج ویژگی، آزمون آماری و تحلیل تفسیرپذیری مبتنی بر چارچوب SHAP توضیح داده می‌شود.

۴-۱ انتخاب لایه بهینه

در بسیاری از پژوهش‌های پیشین، فرایند انتخاب لایه

۳) اطلاعات متقابل

اطلاعات متقابل^{۳۱} مفهومی از نظریه اطلاعات است که میزان وابستگی بین جفت متغیرها را برآورد می‌کند. این معیار، بدون وابستگی به بازپارامترسازی، میزان وابستگی میان متغیرها را اندازه‌گیری می‌کند [۱۹]. هرچه این معیار بیشتر باشد، ارتباط آماری بین این دو متغیر قوی‌تر است. در این پژوهش، اطلاعات متقابل میان هر ویژگی و برچسب با استفاده از تخمین‌گر غیرپارامتریک مبتنی بر k -نزدیک‌ترین همسایه^{۳۲} محاسبه شد. این روش وابستگی مستقیم به بُعد داده‌ها ندارد [۲۰] و می‌تواند روابط خطی و غیرخطی میان ویژگی‌ها و برچسب را شناسایی کند [۲۱].

۴) افزونگی

یک ویژگی بولی^{۳۳}، افزونه^{۳۴} است اگر حذف آن تاثیری بر هیچ یک از جفت نمونه‌هایی که در سایر ویژگی‌ها کاملاً یکسان هستند، نداشته باشد.

برای محاسبه افزونگی، از مفهوم اطلاعات متقابل استفاده می‌شود. در حالی که اطلاعات متقابل معرفی شده در بخش قبل، میزان همبستگی بین ویژگی‌ها و برچسب‌ها را اندازه‌گیری می‌کند، می‌توان از همان روابط ریاضی برای سنجش همبستگی بین ویژگی‌ها با یکدیگر نیز بهره گرفت. اگر مقدار اطلاعات متقابل بین دو ویژگی زیاد باشد، این نشان‌دهنده افزونگی اطلاعاتی بین آن‌هاست؛ به این معنی که یکی از آنها احتمالاً اطلاعات جدیدی نسبت به دیگری ارائه نمی‌دهد و وجود آن ممکن است تنها موجب افزایش حجم داده و اشغال حافظه شود.

۵) ابعاد لایه‌ها

در این پژوهش، ابعاد خروجی لایه‌ها به‌منظور تحلیل پیچیدگی محاسباتی مورد بررسی قرار گرفته‌اند.

در ادبیات مرتبط با این حوزه، هیچ روش استاندارد یا اجماعی برای تعیین وزن نسبی میان این معیارها گزارش نشده است. بنابراین، برای جلوگیری از ایجاد سوگیری و

حفظ شفافیت، وزن برابر برای تمام معیارها در نظر گرفته شد. در گام نهایی، کلیه معیارهای ارزیابی مربوط به هر لایه با استفاده از روش حداقل- حداکثر^{۳۵} نرمال شدند. در این فرایند، به این نکته توجه شده است که در برخی معیارها، مقدار کمتر نشان‌دهنده عملکرد بهتر است، در حالی که در برخی دیگر، مقدار بیشتر مطلوب‌تر است، بنابراین نرمال‌سازی مطابق با ماهیت هر معیار انجام شده است. در نهایت، میانگین نهایی تمام معیارهای استاندارد شده برای هر لایه محاسبه گردیده و لایه‌ای که بالاترین میانگین را به‌دست آورده، به‌عنوان لایه بهینه انتخاب و مورد استفاده قرار گرفته است.

۲-۴ استخراج ویژگی

پس از انتخاب لایه بهینه، ویژگی‌های استخراج شده با تکنیک Z-score، نرمال‌سازی و به الگوریتم تقویت گرادیان سبک‌وزن منتقل شدند.

به‌منظور اطمینان از پایداری عملکرد مدل و کاهش احتمال تصادفی بودن نتایج، فرایند آموزش و ارزیابی پنج مرتبه تکرار گردید. در هر تکرار، تقسیم‌بندی داده‌ها بین مجموعه‌های آموزش و آزمون ثابت بود و تنها حالت تصادفی شروع آموزش الگوریتم طبقه‌بند تغییر می‌کرد. در هر تکرار، مدل با استفاده از شاخص‌های صحت، دقت، بازیابی، امتیاز F1، خطای مثبت کاذب و خطای منفی کاذب مورد ارزیابی قرار گرفت. در پایان، برای هر یک از معیارهای مذکور به‌صورت مستقل، میانگین مقادیر به‌دست آمده در پنج اجرا محاسبه و به‌عنوان نتیجه نهایی گزارش شد.

۳-۴ آزمون آماری

به منظور ارائه مقایسه عددی دقیق‌تر و تقویت تحلیل آماری، مدل منفرد الگوریتم سبک وزن نیز ارزیابی شد. مشابه مدل‌های ترکیبی، این مدل پنج بار اجرا شد و

31-Mutual Information

32-k-Nearest Neighbors non-parametric estimator

33-Boolean Feature

34-Redundancy

35-Min-Max Normalization

۵- نتایج

در این بخش، نتایج حاصل از مراحل مختلف پژوهش شامل انتخاب لایه بهینه برای استخراج ویژگی، ارزیابی عملکرد مدلها، بررسی معناداری تفاوت میان آنها و تحلیل تفسیرپذیری مدل منتخب ارائه می شود. هدف از این بخش، ارزیابی اثربخشی رویکردهای پیشنهادی در ارتقای عملکرد تشخیص تقلب و تحلیل نقش ویژگیهای استخراج شده در فرآیند تصمیم گیری مدل است.

۵-۱ نتیجه انتخاب بهینه لایه

برای انتخاب لایه بهینه جهت استخراج ویژگی، مجموعه ای از معیارهای کمی شامل شاخص کالینسکی-هاراباز، شاخص دیویس-بولدین، اطلاعات متقابل، میزان افزونگی ویژگیها و ابعاد لایه بررسی شد. با توجه به این که برخی معیارها (کالینسکی-هاراباز و اطلاعات متقابل) مقادیر بالاتر و برخی دیگر (ابعاد لایه، دیویس-بولدین و افزونگی) مقادیر کمتر، نشان دهنده کیفیت بهتر هستند، پیش از اعمال تکنیک نرمال سازی حداقل-حداکثر جهت هر شاخص به گونه ای تنظیم شد که قابلیت مقایسه عادلانه فراهم گردد. مقادیر نرمال شده در جدول ۲ و میانگین آنها برای تصمیم گیری نهایی در جدول ۳ ارائه شده است.

بر اساس نتایج به دست آمده، لایه تمام متصل از شبکه عصبی هم آمیختی و لایه پنهان تمام متصل سوم از شبکه عصبی پرسپترون چندلایه به عنوان لایه های بهینه شناسایی شدند. ویژگیهای استخراج شده از این لایه ها بعد از نرمال سازی به الگوریتم تقویت گرادیان سبکوزن تزریق شد.

۵-۲ نتایج مدلها

همان طور که در بخش ۲-۴ و ۳-۴ توضیح داده شد، هر سه مدل شامل مدل ترکیبی مبتنی بر شبکه عصبی هم آمیختی، مدل ترکیبی مبتنی بر پرسپترون چندلایه و مدل منفرد، پنج مرتبه اجرا شدند. مقادیر معیارهای عملکرد

میانگین نتایج برای مقایسه با دو مدل دیگر مورد استفاده قرار گرفت تا اثر تصادفی بودن نتایج به حداقل برسد. برای ارزیابی معناداری آماری روش پیشنهادی، آزمون ویلکاکسون^{۳۶} برای سه گروه مختلف شامل مدل ترکیبی مبتنی بر شبکه عصبی پیچشی، مدل ترکیبی مبتنی بر شبکه عصبی پرسپترون چندلایه و مدل منفرد انجام شد. مقایسه ها به صورت جفتی و با در نظر گرفتن تمام معیارهای ارزیابی صورت گرفت.

۴-۴ تفسیرپذیری

هدف از به کارگیری تفسیرپذیری در این پژوهش، شفاف سازی سازوکار تصمیم گیری مدل و مشخص کردن این موضوع است که کدام ویژگیهای استخراج شده بیشترین نقش را در طبقه بندی تراکنشها داشته اند. آنجا که سامانه های تشخیص تقلب در محیطهای حساس مالی مورد استفاده قرار می گیرند، امکان توضیح و توجیه تصمیم مدل برای تحلیل گران و واحدهای کنترلی اهمیت بالایی دارد. تفسیرپذیری کمک می کند تا در صورت شناسایی یک تراکنش به عنوان تقلب یا عدم تقلب، مشخص شود کدام متغیرهای اولیه بیشترین تأثیر را در فعال شدن ویژگیهای کلیدی مدل داشته اند. این موضوع علاوه بر افزایش شفافیت و اعتمادپذیری مدل، امکان بازبینی هدفمند هشدارها و تنظیم مناسب تر آستانه های تصمیم گیری را فراهم می سازد. با هدف بررسی تفسیرپذیری تصمیم مدل، تحلیل مبتنی بر SHAP^{۳۷} بر روی مدلی که عملکرد بهتری نشان داده بود، انجام شد. در این تحلیل، علاوه بر شناسایی ویژگیهای مؤثر استخراج شده توسط شبکه عصبی که در تصمیم نهایی الگوریتم طبقه بند نقش دارند، میزان تأثیر هر یک از این ویژگیها و همچنین ویژگیهای اولیه داده ها که شبکه عصبی برای ساخت مؤثرترین ویژگی ترکیب کرده بود، مورد بررسی قرار گرفت.

36-Wilcoxon signed-rank test

37-SHapley Additive exPlanations

جدول ۲. مقادیر نرمال شده معیارهای کمی ارزیابی لایه‌های شبکه عصبی

نام شبکه عصبی		لایه‌های شبکه عصبی هم آمیختی						لایه‌های شبکه عصبی پرسپترون چندلایه			
معیار	نام لایه	لایه گسترش یافته تمام متصل	لایه هم آمیختی یک بعدی اول	لایه هم آمیختی یک بعدی دوم	لایه تخت سازی	لایه‌های تمام متصل	لایه پنهان تمام متصل اول	لایه پنهان تمام متصل دوم	لایه پنهان سوم	لایه پنهان چهارم	
		۰/۹۶	۱	۰/۸۲	۰/۵	۰	۱	۰/۵۴	۰	۰/۴۶	
میزان افزونگی بین ویژگی‌ها		۰/۹۷	۰/۰۰۰۲۲	۰/۴۸۱	۰/۷۲۵	۱	۰	۰/۷۷	۱	۰/۷۳	
شاخص دیویس - بولدین		۰/۰۰۱	۰	۰/۰۴	۰/۱۲	۱	۰	۰/۵۹	۱	۰/۴۳	
شاخص کالینسکی - هارباز		۰/۰۵	۰	۰/۴۷	۱	۰/۹	۰	۰/۶۲	۰/۹۹	۱	
اطلاعات متقابل		۰	۱	۱	۰	۱	۰	۰/۵۷	۰/۸۵	۱	
ابعاد ویژگی هر لایه											

جدول ۳. میانگین مقادیر نرمال شده معیارها برای انتخاب لایه بهینه

نام شبکه	نام لایه	نمره نهایی
شبکه عصبی هم‌آمیختگی	لایه گسترش یافته تمام متصل	۰/۲۲
	لایه هم‌آمیختگی یک بعدی اول	۰/۴
	لایه هم‌آمیختگی یک بعدی دوم	۰/۵۶
	لایه تخت سازی	۰/۴۷
	لایه تمام متصل	۰/۷۸
شبکه عصبی پرسپترون چندلایه	لایه پنهان تمام متصل اول	۰/۲
	لایه پنهان تمام متصل دوم	۰/۶۲
	لایه پنهان تمام متصل سوم	۰/۷۷
	لایه پنهان تمام متصل چهارم	۰/۷۶

هر اجرا ثبت و میانگین نتایج برای مقایسه نهایی محاسبه شد و در جدول ۴ گزارش شده است.

بررسی نتایج نشان می‌دهد که استفاده از ویژگی‌های استخراج شده از شبکه عصبی هم‌آمیختگی در مقایسه با ویژگی‌های استخراج شده از شبکه عصبی پرسپترون چندلایه و مدل منفرد با ویژگی‌های اولیه، منجر به بهبود عملکرد مدل شده است. مدل ترکیبی شبکه عصبی هم‌آمیختگی + الگوریتم تقویت گرادیان سبک وزن در اغلب معیارها برتری محسوس دارد. به طور کلی، ترکیب عصبی هم‌آمیختگی با الگوریتم تقویت گرادیان سبک وزن توانسته است تعادل بهتری میان معیارهای اصلی عملکرد برقرار سازد اما هیچ کدام از مدل‌های ترکیبی، با وجود

جدول ۴. نتایج مدل‌های ترکیبی و منفرد بر اساس معیارهای عملکرد

نام مدل	صحت	دقت	بازیابی	امتیاز f1	خطای مثبت کاذب	خطای منفی کاذب
شبکه عصبی هم‌آمیختگی + الگوریتم تقویت گرادیان سبک وزن	۹۹/۹۵	۷۵/۵۵	۹۱/۶۴	۸۲/۸۲	۶	۲۲
شبکه عصبی پرسپترون چندلایه + الگوریتم تقویت گرادیان سبک وزن	۹۹/۹۳	۷۱/۱۱	۸۲/۴۹	۷۶/۳۸	۱۴	۲۶
الگوریتم تقویت گرادیان سبک وزن	۹۹/۸۳	۵۷/۸۳	۷۷/۷۷	۶۵/۷۹	۵۳	۲۱

بهبود عملکرد، در کاهش خطای منفی کاذب موفق نبوده و با افزایش اندک همراه بوده است. این رفتار عمدتاً ناشی از ساختار ترکیبی شبکه‌های عصبی و الگوریتم تقویت گرادیان سبک وزن است. در مرحله استخراج ویژگی، بخش عصبی الگوهای پرتکرار و عمومی داده را در قالب نمایش‌های فشرده‌تر بازنمایی می‌کند که اگرچه موجب افزایش دقت کلی مدل می‌شود، اما می‌تواند موجب کاهش تمایزهای ظریف میان ترائکشن‌های «تقلب» و «عدم تقلب» گردد. از سوی دیگر، الگوریتم تقویت گرادیان سبک وزن،

مثبت کاذب را نسبت به مدل منفرد کاهش دهند. موضوعی که در سامانه‌های واقعی تشخیص تقلب اهمیت حیاتی دارد، زیرا کاهش هشدارهای اشتباه موجب جلوگیری از مسدود شدن غیرضروری کارت‌های بانکی مشتریان و حفظ تجربه کاربری و رضایت مشتریان می‌شود.

۴-۵ تفسیرپذیری

از میان سه مدل مقایسه‌شده، مدل ترکیبی شبکه عصبی هم‌آمیختی + الگوریتم تقویت گرادیان سبک وزن به دلیل عملکرد بهتر نسبت به سایر مدل‌ها به‌عنوان مدل منتخب جهت تحلیل تفسیرپذیری انتخاب شد. هدف از تفسیرپذیری، بررسی نحوه تأثیرگذاری ویژگی‌های استخراج‌شده توسط شبکه عصبی هم‌آمیختی بر تصمیم‌گیری نهایی است. برای این منظور، از تحلیل مقادیر SHAP بر خروجی الگوریتم تقویت گرادیان سبک وزن استفاده گردید. از آنجا که خروجی شبکه شامل ویژگی‌های ترکیبی بدون نام مشخص بود، این ویژگی‌ها به‌ترتیب خروجی شماره‌گذاری شدند. در شکل ۱، بیست ویژگی با بیشترین میانگین قدر مطلق مقادیر SHAP، به‌عنوان اثرگذارترین ویژگی‌ها نمایش داده شده‌اند. مطابق با شکل ۱، ویژگی ۲، یعنی دومین ویژگی استخراج‌شده از شبکه عصبی هم‌آمیختی، بیشترین سهم را در فرآیند تصمیم‌گیری مدل داشته است. برای بررسی دقیق‌تر این بیست ویژگی موثر، مقادیر SHAP برای هر نمونه داده محاسبه و نتایج در شکل ۲ نشان ارائه شده است. همان‌طور که در شکل ۲ مشاهده می‌شود، مقادیر SHAP در محور افقی نشان‌دهنده جهت و میزان تأثیر هر ویژگی بر خروجی مدل هستند؛ مقادیر مثبت بیانگر افزایش احتمال وقوع «تقلب» و مقادیر منفی بیانگر کاهش این احتمال‌اند. هر نقطه یک نمونه تراکنش را نشان می‌دهد و رنگ نقاط مقدار واقعی ویژگی را بازنمایی می‌کند. الگوهای مشاهده شده بیانگر آن است که تنها تعداد محدودی از ویژگی‌ها دارای تأثیر مثبت یا منفی معنادار هستند و بخش عمده ویژگی‌ها نقشی محدود در تصمیم‌گیری مدل دارند. این موضوع حاکی از آن است که مدل برای تمایز بین

به دلیل ماهیت درخت‌گونه و راهبرد تقویت گرادیانی، تقسیم‌های محافظه‌کارانه‌ای را در نواحی با بیشترین اطمینان انجام می‌دهد و بنابراین خطای مثبت کاذب را کاهش می‌دهد. به این ترتیب، میانگین خطای منفی کاذب در چند اجرای مستقل مدل ترکیبی تغییر قابل‌توجهی نداشته و این رفتار می‌تواند به‌عنوان یک ویژگی ساختاری و پایدار مدل تلقی شود. با این وجود، این افزایش در مقایسه با کاهش محسوس خطای نوع اول قابل قبول است و نشان می‌دهد که تعادل میان دقت و بازیابی حفظ شده است.

۵-۳ تحلیل آماری

برای ارزیابی معناداری تفاوت با استفاده از آزمون ویلکاکسون، مقایسه جفتی بین مدل‌ها صورت گرفت، که نتایج آن در جدول ۵ گزارش شده است.

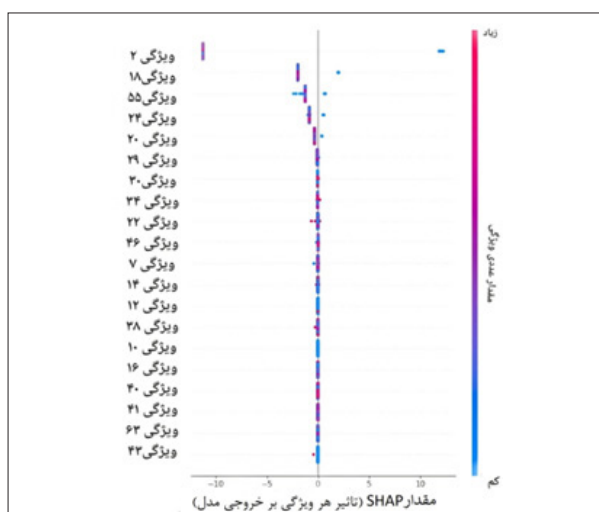
جدول ۵. مقایسه جفتی بین مدل‌ها با آزمون ویلکاکسون

مقایسه مدل‌ها	p-value	آماره آزمون
شبکه عصبی هم‌آمیختی + الگوریتم تقویت گرادیان سبک وزن و شبکه عصبی پرسپترون چندلایه + الگوریتم تقویت گرادیان سبک وزن	۰/۰۶۲۵	۰
شبکه عصبی هم‌آمیختی + الگوریتم تقویت گرادیان سبک وزن و الگوریتم تقویت گرادیان سبک وزن	۰/۰۶۲۵	۰
شبکه عصبی پرسپترون چندلایه + الگوریتم تقویت گرادیان سبک وزن و الگوریتم تقویت گرادیان سبک وزن	۰/۰۶۲۵	۰

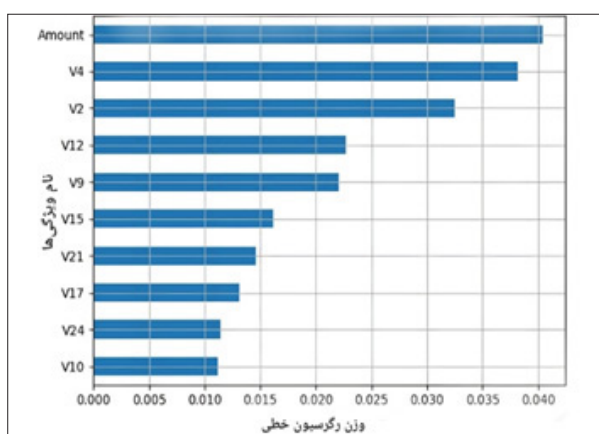
نتایج نشان می‌دهد که مقدار آماره آزمون^{۳۸} در تمامی مقایسه‌ها برابر با صفر است که بدین معناست جهت اختلاف‌ها در تمامی موارد در یک‌راستا بوده و برتری قطعی میان مدل‌ها وجود ندارد. همچنین مقدار p-value برای هر سه مقایسه برابر با ۰/۰۶۲۵ گزارش شد که بزرگتر از سطح خطای متداول (۰/۰۵) است و می‌توان نتیجه گرفت که تفاوت میان مدل‌ها از نظر آماری معنادار نیست. با این حال، عدم معناداری آماری در این پژوهش به معنای فقدان برتری کاربردی نیست. بررسی شاخص‌های عملکرد نشان می‌دهد که مدل‌های پیشنهادی توانسته‌اند میزان خطای



شکل ۱. بیست ویژگی مؤثر در تصمیم‌گیری مدل



شکل ۲. توزیع مقادیر SHAP برای بیست ویژگی مؤثر



شکل ۳. ده ویژگی اولیه مؤثر در شکل‌گیری ویژگی شماره ۲ از لایه تمام‌متصل شبکه عصبی هم‌آمیختی

چندلایه نشان داد. از سوی دیگر، مقدار خطای منفی کاذب در مدل‌های ترکیبی با افزایش اندک همراه بوده است. این

دو کلاس «تقلب» و «عدم تقلب» از مجموعه‌ای مشخص و تکرارشونده از عوامل استفاده می‌کند. به‌ویژه، ویژگی شماره ۲ بیشترین تأثیر را بر افزایش یا کاهش احتمال تقلب داشته و نقش اصلی را در تفکیک نمونه‌های مشکوک از تراکنش‌های عادی ایفا کرده است. از آنجا که ویژگی‌های مورد استفاده‌ی الگوریتم طبقه‌بند، ترکیبی خطی از ویژگی‌های اولیه استخراج‌شده هستند، برای شناسایی مهمترین مؤلفه‌های تشکیل‌دهنده ویژگی ۲، یک مدل رگرسیون خطی آموزش داده شد. در این مدل، ویژگی‌های اولیه به‌عنوان متغیرهای مستقل و ویژگی میانی حاصل از شبکه عصبی به‌عنوان متغیر وابسته در نظر گرفته شدند. نتایج این تحلیل در شکل ۳ نشان داده شده است که بر اساس آن، ویژگی «مقدار تراکنش» بیشترین سهم را در شکل‌گیری ویژگی ۲ دارد.

بنابراین، می‌توان گفت بخش قابل‌توجهی از اطلاعاتی که ویژگی ۲ در تصمیم‌گیری مدل منتقل می‌کند، به «مقدار تراکنش» مربوط است. با توجه به این که ویژگی ۲ مهمترین ویژگی در طبقه‌بندی نهایی شناخته شد، این نتیجه بیان می‌کند که اطلاعات مرتبط با مقدار تراکنش نقش محوری در عملکرد مدل تشخیص تقلب داشته است.

۵-۵ جمع‌بندی

در این پژوهش، سه گروه مدل شامل یک مدل منفرد مبتنی بر ویژگی‌های اولیه و دو مدل ترکیبی مبتنی بر استخراج ویژگی از شبکه‌های عصبی (شبکه عصبی هم‌آمیختی و پرسپترون چندلایه) و ترکیب آنها با الگوریتم تقویت گرادیان سبک‌وزن مورد بررسی قرار گرفتند. نتایج نشان داد که مدل‌های ترکیبی به‌طور کلی عملکرد بهتری نسبت به مدل منفرد دارند. همچنین، در میان مدل‌های ترکیبی، مدل مبتنی بر ویژگی‌های استخراج‌شده از شبکه عصبی هم‌آمیختی در اغلب معیارهای ارزیابی به‌ویژه در کاهش خطای مثبت کاذب، برتری بیشتری نسبت به مدل مبتنی بر ویژگی‌های استخراج‌شده از شبکه پرسپترون

تقلب در تراکنش‌های کارت اعتباری ارائه شد که در آن ویژگی‌های استخراج‌شده از شبکه‌های عصبی هم‌آمیختی و پرسپترون چندلایه به الگوریتم تقویت گرادیان سبک‌وزن منتقل شدند. برای انتخاب لایه مناسب برای استخراج ویژگی، از ترکیب چند معیار کمی استفاده شد و لایه‌ای که بالاترین میانگین نرمال‌شده را کسب کرد، به‌عنوان لایه بهینه تعیین گردید.

نتایج نشان داد که مدل‌های ترکیبی، به‌ویژه ترکیب مبتنی بر شبکه عصبی هم‌آمیختی، در مقایسه با مدل منفرد موجب بهبود امتیاز F1 و کاهش قابل‌توجه خطای مثبت کاذب شدند. در مقابل، مقدار خطای منفی کاذب با افزایش اندک همراه بود، افزایشی که در مقایسه با کاهش قابل توجه خطای مثبت کاذب، توازن عملکرد مدل را تحت تأثیر قرار نمی‌دهد. آزمون ویلکاکسون نیز تفاوت میان مدل‌های ترکیبی و مدل منفرد را از نظر آماری معنادار ندانست. با این حال، کاهش خطای مثبت کاذب از دیدگاه عملی در سامانه‌های بانکی اهمیت دارد، زیرا منجر به کاهش مسدودسازی غیرضروری کارت‌های سالم و بهبود تجربه مشتریان می‌شود.

تحلیل تفسیرپذیری مدل بر پایه‌ی SHAP مشخص کرد که دومین ویژگی استخراج‌شده از شبکه عصبی هم‌آمیختی بیشترین نقش را در طبقه‌بندی داشته و این ویژگی عمدتاً بر پایه‌ی اطلاعات مرتبط با مقدار تراکنش شکل گرفته است. این موضوع نشان می‌دهد که شبکه عصبی توانسته است نمایش ساخت‌یافته‌تری از متغیرهای اولیه ایجاد کند و این نمایش جدید در بهبود عملکرد مدل تأثیرگذار بوده است. در مجموع، یافته‌ها بیان می‌کنند که ترکیب استخراج ویژگی از شبکه‌های عصبی و مدل‌های تقویت گرادیان می‌تواند راهکاری مؤثر برای افزایش کیفیت سیستم‌های تشخیص تقلب بانکی باشد، به‌ویژه در شرایطی که کاهش خطاهای مثبت کاذب اهمیت عملی بالایی دارد.

۶-۲ پیشنهادهایی برای پژوهش‌های آینده

در ادامه، با توجه به یافته‌های این پژوهش، پیشنهادهایی

افزایش در مقایسه با کاهش قابل‌توجه خطای مثبت کاذب، از نظر عملیاتی قابل پذیرش است و نشان می‌دهد که مدل پیشنهادی توانسته است تعادل مؤثری میان دقت و بازیابی برقرار سازد.

همچنین، برای انتخاب لایه‌های مناسب جهت استخراج ویژگی از شبکه‌های عصبی، برای هر لایه مجموعه‌ای از معیارهای کمی شامل شاخص کالینسکی- هاراباز، شاخص دیویس- بولدین، اطلاعات متقابل، میزان افزونگی ویژگی‌ها و ابعاد لایه محاسبه شد. پس از نرمال‌سازی این معیارها و یکسان‌سازی جهت آن‌ها، میانگین نمرات نرمال‌شده به‌عنوان معیار تصمیم‌گیری نهایی مورد استفاده قرار گرفت و لایه‌ای که بالاترین میانگین را کسب کرده بود، به‌عنوان لایه مناسب برای استخراج ویژگی انتخاب شد. سپس، ویژگی‌های حاصل از لایه‌های منتخب پس از نرمال‌سازی مجدد، به الگوریتم تقویت گرادیان سبک‌وزن برای مرحله طبقه‌بندی منتقل شدند.

تحلیل تفسیرپذیری نشان داد که ویژگی شماره ۲، که از شبکه عصبی هم‌آمیختی استخراج شده و بیشترین سهم را در تصمیم‌گیری مدل داشته است، عمدتاً تحت تأثیر متغیر «مقدار تراکنش» شکل گرفته است. به عبارت دیگر، مدل از مقدار تراکنش به‌عنوان ورودی پایه برای ساخت یک ویژگی ترکیبی مؤثر استفاده کرده است و این ویژگی نقش کلیدی در فرایند تشخیص تقلب داشته است. این نتیجه نشان می‌دهد که اطلاعات مربوط به مقدار تراکنش، پس از تبدیل و بازنمایی توسط شبکه عصبی، در بهبود عملکرد مدل طبقه‌بندی نقش قابل‌توجهی ایفا کرده است.

۶- نتیجه‌گیری

این بخش به جمع‌بندی نتایج پژوهش و ارائه راهکارهایی برای مطالعات آتی اختصاص دارد.

۶-۱ یافته‌های کلیدی پژوهش

در این پژوهش، رویکردی ترکیبی برای تشخیص

8. Singh, S., Nimje, H., Fulpatile, A., & Neware, R. 2024. Credit card fraud detection using a hybrid machine learning algorithm, Preprints.
9. Boukhalifa, K., Harous, S., & Bouras, A. 2022. A novel text2IMG mechanism of credit card fraud detection: A deep learning approach, *Electronics*, Vol. 11, No. 4, p. 756.
10. Subagio, A. & Utama, D. N. 2025. Fraud credit card transaction detection using hybrid multilayer perceptron-random forest method, *Edelweiss Applied Science and Technology*, Vol. 9, No. 3, pp. 2482–2494.
11. Ming, R., Abdelrahman, O., Innab, N., & Ibrahim, M. H. K. 2024. Enhancing fraud detection in auto insurance and credit card transactions: A novel approach integrating CNNs and machine learning algorithms, *PeerJ Computer Science*, Vol. 10, p. e2088.
12. Du, H., Lv, L., Guo, A., & Wang, H. 2023. AutoEncoder and LightGBM for credit card fraud detection problems, *Symmetry*, Vol. 15, No. 4, p. 870.
13. Fanai, H. & Abbasimehr, H. 2023. A novel combined approach based on deep Autoencoder and deep classifiers for credit card fraud detection, *Expert Systems with Applications*, Vol. 217, p. 119562.
14. Mienye, I. D. & Jere, N. 2024. Deep learning for credit card fraud detection: A review of algorithms, challenges, and solutions, *IEEE Access*, Vol. 12, pp. 96893–96910.
15. Kasasbeh, B., Aldabaybah, B. & Ahmad, H. 2022. Multilayer perceptron artificial neural networks-based model for credit card fraud detection, *Indonesian Journal of Electrical Engineering and Computer Science*, Vol. 26, No. 1, pp. 362–373.
16. Ke, Q., Liu, J., Bennamoun, M., An, S., Sohel, F., & Boussaid, F. 2018. Computer Vision for Human-Machine Interaction, in **Computer Vision for Assistive Healthcare*, pp. 127–145.
17. Zhao, X., Liu, Y., & Zhao, Q. 2024. Improved LightGBM for extremely imbalanced data and application to credit card fraud detection, *IEEE Access*, Vol. 12, pp. 159316–159335.
18. Gan, M. & Zhang, L. 2021. Iteratively local fisher score for feature selection, *Applied Intelligence*, Vol. 51, No. 8, pp. 6167–6181.
19. Han, W., Chen, H., & Poria, S. 2021. Improving multi-modal fusion with hierarchical mutual information maximization for multimodal sentiment analysis, *arXiv preprint arXiv:2109.00412*.
20. Carrara, N. & Ernst, J. 2020. On the estimation of mutual information, *Proceedings*, vol. 33, no. 1, p. 31-39.
21. Runge, J. 2018. Conditional independence testing based on a nearest-neighbor estimator of conditional mutual information, in *Proceedings of the International Conference on Artificial Intelligence and Statistics*, pp. 938–947.s

برای توسعه مطالعات آینده ارائه می‌شود.

۱. در این پژوهش، تمامی معیارهای کمی برای تعیین لایه بهینه با وزن برابر در نظر گرفته شدند. در مطالعات آینده می‌توان اهمیت نسبی هر معیار را به صورت جداگانه بررسی کرد تا سهم واقعی هر شاخص در تعیین لایه بهینه مشخص شود.
۲. به کارگیری روش‌های مهندسی ویژگی، مانند کاهش ابعاد داده‌ها یا تمرکز بر انتخاب ویژگی‌های موثر، می‌تواند به بهبود مدل و کاهش پیچیدگی محاسباتی کمک کند.
۳. از معماری‌های متنوع شبکه‌های عصبی، مانند شبکه‌های ترنسفورمرها، برای استخراج ویژگی‌ها استفاده شود.
۴. پیشنهاد می‌شود روش ارائه شده بر روی مجموعه داده‌های گوناگون و به روزتر (برای مثال داده‌های تراکنش برخط یا تراکنش‌های مبتنی بر رمز دوم پویا) نیز آزموده شود تا میزان تعمیم‌پذیری و پایداری مدل در سناریوهای واقعی مشخص گردد.

منابع

1. Taha, A. A. & Malebary S. J. 2020 An intelligent approach to credit card fraud detection using an optimized light gradient boosting machine, *IEEE Access*, Vol. 8, pp. 25579–25587.
2. Carneiro, N., Figueira, G., & Costa, M. 2017. A data mining based system for credit-card fraud detection in e-tail, *Decision Support Systems*, Vol. 95, pp. 91–101.
3. Umaru, I. A., Aliyu, A. A., Ibrahim, M., Abdulkadir, S., Ahmed, M. A., Abubakar, M. A., & Tanko, S. A. 2025. An enhanced hybrid model combining LSTM, ResNet, and an attention mechanism for credit card fraud detection, *FUDMA Journal of Science*, Vol. 9, No. 2, pp. 42–48.
4. Carcillo, F., Le Borgne, Y.-A., & Bontempi, G. 2018. Streaming active learning strategies for real-life credit card fraud detection, *International Journal of Data Science and Analytics*, Vol. 7, No. 1, pp. 33–45.
5. Lin, T. H. & Jiang, J. R. 2021. Credit card fraud detection with autoencoder and probabilistic random forest, *Mathematics*, Vol. 9, No. 21, pp. 16-1.
6. Berhane, T., Melese, T., Walelign, A., & Mohammed, A. 2023. A hybrid convolutional neural network and support vector machine-based credit card fraud detection model, *Mathematical Problems in Engineering*, Vol. 2023, No. 1, pp. 1–10.
7. Nancy, A. M., Kumar, G. S., Veena, S., Vinoth, N. S., & Bandyopadhyay, M. 2020. Fraud detection in credit card transaction using hybrid model, *AIP Conference Proceedings*, pp. 1–10.