



بررسی تأثیرگذاری اخبار و منابع مختلف خبری بر نرخ ارز با استفاده از یادگیری ماشین

المیرا اصل روستا^۱

علیرضا عرفانی^۲

عبدالمحمد کاشیان^۳

تاریخ پذیرش: ۱۴۰۴/۱۲/۰۹

تاریخ دریافت: ۱۴۰۴/۱۰/۰۹

چکیده

مشاهدات نشان می‌دهد که در سال‌های اخیر اغلب نوسانات شدید نرخ ارز در ایران بعد از انتشار اخبار تأثیرگذار سیاسی-اقتصادی داخلی و خارج رخ داده‌اند. سؤال مهمی که شاید تاکنون پاسخ دقیقی به آن داده نشده این است که آیا می‌توان به صورت روش‌مند تأثیر این اخبار را اندازه‌گیری کرد؟ در این پژوهش با استفاده از رویکرد یادگیری عمیق ابزاری کاربردی به منظور پیش‌بینی نرخ ارز ارائه شده و تلاش شده تا با استفاده از این ابزار میزان تأثیر اخبار و منابع خبری تحلیل شود. برای این منظور اخبار مؤثر بر نرخ ارز از ۱۰ پایگاه اصلی داخلی و خارجی از سال ۱۳۹۳ تا ۱۴۰۲ جمع‌آوری شده و به همراه نرخ ارز و سایر شاخص‌های اقتصادی مؤثر به مدل داده شده است. نتایج بدست آمده حاکی از آن است که وجود اخبار در کنار سایر شاخص‌های اقتصادی تأثیر کاملاً مثبتی در پیش‌بینی مدل‌ها هم در حالت رگرسیون و هم در حالت کلاس‌بندی دارد. همچنین با استفاده از مدل تولید شده و روش‌های تحلیلی، وزن تأثیرگذاری هریک از منابع خبری مشخص گردید. نتایج نشان می‌دهد که میزان تأثیرگذاری اخبار نه تنها به محتوای آن، بلکه به منبع منتشرکننده‌ی خبر نیز به طور چشمگیری وابسته است، به طوری که اگر خبری کاملاً مشابه توسط دو منبع متفاوت اعلام شود، تأثیر متفاوتی خواهد داشت.

واژه‌های کلیدی: پیش‌بینی نرخ ارز، تأثیر اخبار، منابع خبری، یادگیری ماشین، هوش مصنوعی.

طبقه‌بندی: JEL: C45, C51, C53, C55, G17

۱ گروه اقتصاد، دانشکده اقتصاد، مدیریت و علوم اداری، دانشگاه سمنان، سمنان. ایران. e.asleroosta@semnan.ac.ir

۲ گروه اقتصاد، دانشکده اقتصاد، مدیریت و علوم اداری، دانشگاه سمنان، سمنان. ایران. (نویسنده مسئول) a erfani@semnan.ac.ir

۳ گروه اقتصاد، دانشکده اقتصاد، مدیریت و علوم اداری، دانشگاه سمنان، سمنان. ایران. a.m.kashian@profs.semnan.ac.ir



۱- مقدمه

نرخ ارز یکی از مهم‌ترین شاخص‌های اقتصادی است که به‌طور مستقیم در تصمیم‌گیری‌های اقتصادی داخلی و خارجی تأثیرگذار است. این نرخ علاوه بر عوامل اقتصادی مانند تولید ناخالص داخلی، نرخ تورم، بیکاری و قیمت نفت، تحت تأثیر اخبار سیاسی و اقتصادی داخلی و خارجی نیز قرار دارد. در سال‌های اخیر، با افزایش سرعت و گستردگی انتشار اخبار، تأثیرات کوتاه‌مدت این اطلاعات بر نوسانات نرخ ارز به وضوح مشاهده شده است. در اینجا دو سؤال مطرح است؛ نخست آنکه چطور می‌توان اخبار منتشره مستقیماً در مدل‌سازی و پیش‌بینی نرخ ارز دخیل نمود؟ و دوم آنکه چطور می‌توان میزان تأثیرگذاری اخبار و منابع منتشر کننده‌ی خبر را مشخص کرد؟

در گذشته پژوهش‌هایی به منظور بررسی تأثیرات اخبار منتشره بر نرخ ارز در داخل انجام شده که عموماً دو رویکرد داشته‌اند. در رویکرد اول، تأثیرات کاهشی یا افزایشی در بازار را بعد از انتشار اخبار بررسی و تحلیل کرده‌اند. در این رویکرد، اصولاً مدلی برای پیش‌بینی ارائه نشده و پژوهش حول گزارشی از مشاهدات است. در رویکرد دوم، تلاش مختصری شده تا با تحلیل احساسات^۱ برچسب مثبت یا منفی به اخبار منتسب شده و سپس براساس این برچسب و به صورت غیرمستقیم پیش‌بینی از نرخ ارز انجام شود. رویکرد دوم قابل توجه است، اما به دلیل استفاده از روش‌های غیرمستقیم و تبدیل خبر به یک برچسب (مثبت یا منفی)، عموماً قادر به مدل‌سازی معنا و عمق اطلاعات موجود در یک خبر نیست.

در این پژوهش مدلی ارائه شده که در آن اخبار منتشره به همراه سایر عوامل اقتصادی و متغیرهای عددی، مستقیماً در پیش‌بینی نرخ ارز مورد استفاده قرار می‌گیرند. استفاده از روش‌های مدرن در بازنمایی عددی اخبار منتشره، منجر به آن شده که لایه‌های بالاتری از معنا و فحوای موجود در اخبار توسط مدل درک شده و در پیش‌بینی نرخ ارز لحاظ شود.

در ادامه ادبیات موضوع دو حوزه‌ی نرخ ارز و همچنین یادگیری ماشین بیان شده است. پس از مشخص شدن حدود مسئله، ارجاعاتی به پژوهش‌های مرتبط داخلی و خارجی شده است. به دلیل آنکه شاکله‌ی مدل پیشنهادی بر یادگیری ماشینی استوار است، بخش قابل توجهی از مقاله به تشریح فرآیند حل مسئله به صورت یک مسئله‌ی یادگیری ماشین اختصاص داشته و جزئیات خط لوله‌ی^۲ طراحی شده به تفصیل ذکر شده است. به منظور ارزیابی فرضیه‌ی این مقاله مبنی بر اثرگذاری اخبار بر تغییرات نرخ ارز و اندازه‌گیری کمی آن، آزمایش‌های کاملی انجام شده و نتایج و جمع‌بندی آنها در بخش پایانی مقاله منعکس شده است. در انتها پیشنهاداتی برای ادامه‌ی این پژوهش ارائه شده است.

1 Sentiment Analysis

2 Pipeline

۲- مروری پیشینه‌ها

۲-۱- نرخ ارز

به منظور مدل‌سازی نرخ ارز در یکصد سال اخیر تئوری‌های مختلفی مطرح شده که از جمله مهم‌ترین آن‌ها می‌توان به تئوری برابری قدرت خرید^۱، تئوری بازار دارائی^۲، مدل ماندل-فلمنگ^۳ (Mundell، ۱۹۶۳) و مدل پولی با قیمت انعطاف‌پذیر^۴ اشاره کرد. محققین داخلی نیز تلاش‌هایی برای مدل‌سازی نرخ ارز در ایران با توجه به شرایط و اقتضائات اقتصادی کشور انجام داده‌اند که از آن جمله می‌توان به پژوهش (شیرازی و همکاران، ۲۰۱۴) اشاره کرد. در این پژوهش نویسندگان تلاش کرده‌اند تا مدل‌های استاندارد جهانی را برای دادگان اقتصاد ایران پیاده‌سازی و آزمایش نمایند.

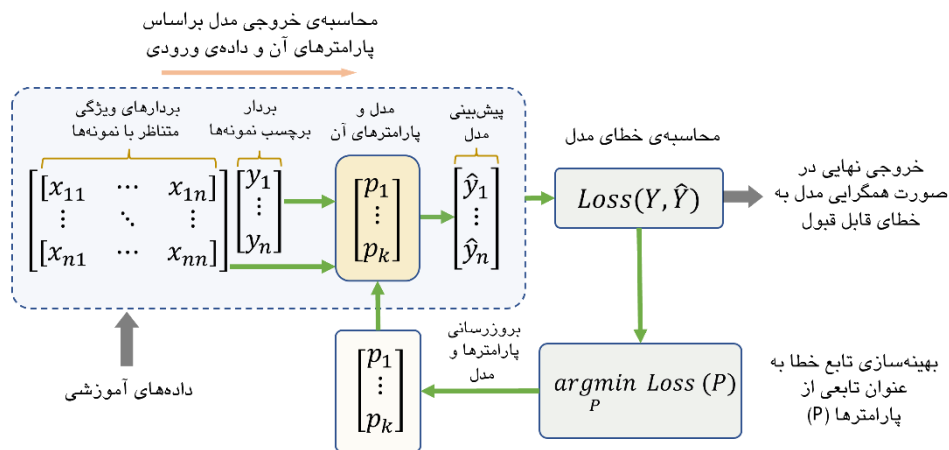
رویکرد دیگر در پیش‌بینی نرخ ارز، استفاده از روش‌های آماری است. این روش‌ها بر اساس مشاهدات بازار و داده‌های تاریخی به عنوان مسائل سری زمانی^۵ به تحلیل می‌پردازند. مهم‌ترین روش در این دسته، مدل‌های خودرگرسیون^۶ مانند ARIMA^۷ هستند که با در نظر گرفتن روند^۸ها و الگوهای فصلی^۹، به پیش‌بینی نرخ ارز در آینده نزدیک می‌پردازند. این مدل‌ها بیشتر بر اساس تحلیل داده‌های گذشته و روندهای موجود کار می‌کنند. در سال‌های اخیر، استفاده از الگوریتم‌های یادگیری ماشین در مدل‌سازی و پیش‌بینی نرخ ارز رایج‌تر شده است. این الگوریتم‌ها به دلیل توانایی در کار با بردارهای بزرگ و داده‌های پیچیده، مانند داده‌های حجیم خبری، بسیار مؤثر هستند. از آنجا که در فرآیند مدل‌سازی نرخ ارز، اخبار به شکل بردارهای عددی بزرگ (حدود ۸۰۰ مؤلفه) وارد مدل می‌شوند، استفاده از روش‌های یادگیری ماشین برای تحلیل این داده‌ها به‌ویژه در پیش‌بینی نرخ ارز مناسب‌ترین گزینه است.

۲-۲- یادگیری ماشین

به صورت غیررسمی می‌توان گفت که یادگیری ماشین زیرمجموعه‌ای از هوش مصنوعی است که در آن سیستم‌ها دانش خود را با استخراج الگوها^{۱۰} از داده‌های فراهم شده بدست می‌آورند. به عبارت دیگر در این سیستم‌ها دانش از قبل برنامه‌ریزی^{۱۱} نشده است (Kapoor و همکاران، ۲۰۲۲). انواع مختلفی از روش‌های یادگیری را می‌توان

-
- 1 Purchasing Power Parity
 - 2 Asset Market Model
 - 3 Mundell-Fleming Model
 - 4 Flexible-Price Monetary Model
 - 5 Time Series
 - 6 Auto-Regressive
 - 7 Auto-Regressive Integrated Moving Average
 - 8 Trend
 - 9 Seasonality Patterns
 - 10 Patterns
 - 11 Programmed

برشمرده، اما عمده‌ی آنها را می‌توان در سه دسته‌ی یادگیری بانظارت^۱، یادگیری بدون نظارت^۲ و یادگیری تقویتی^۳ طبقه‌بندی کرد. یادگیری بانظارت فرآیندی است که در آن به ازای هر نمونه داده یک برچسب^۴ نیز وجود دارد. در این حالت باید مدلی بسازیم که بتواند با داده‌های ورودی این برچسب را به درستی پیش‌بینی کند. روش‌های بسیار زیادی براساس کاربرد و نوع داده‌ها در دسته‌ی یادگیری‌های بانظارت قرار می‌گیرند که از جمله‌ی آنها می‌توان به رگرسیون لجستیک^۵، رگرسیون خطی و غیرخطی، ماشین بردار پشتیبان^۶، درخت تصمیم^۷، جنگل تصادفی^۸، K-نزدیک‌ترین همسایه^۹، کلاس‌بند نایو بیز^{۱۰}، الگوریتم XGBoost، الگوریتم CatBoost و دسته‌ی بزرگی از شبکه‌های عصبی مانند MLP ها اشاره کرد (Morales و Escalante، ۲۰۲۲). الگوریتم پیشنهادی در این پژوهش نیز در زمزه‌ی شبکه‌های عصبی و روش‌های بانظارت جای می‌گیرد. در

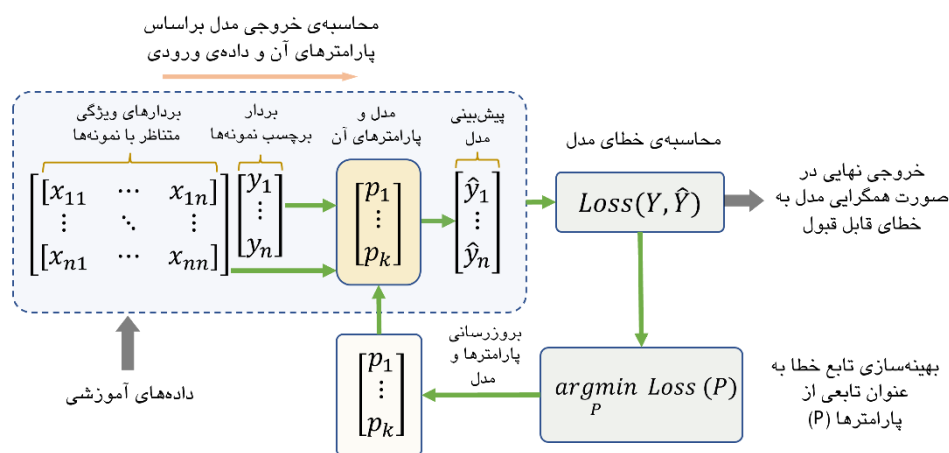


شکل ۱ اجزاء اصلی و فرآیند کلی آموزش در یک سیستم یادگیری نظارت شده نمایش داده شده است. در این فرآیند به ازای داده‌های ورودی به مدل اولیه، یک خروجی محاسبه می‌شود. تفاوت این خروجی با برچسب واقعی متناظر با داده‌ی ورودی، تابع خطا را تشکیل می‌دهد. هدف تنظیم پارامترهای مدل به نحوی است که تابع خطا کمینه شود. انتخاب نحوه‌ی بهینه‌سازی تابع خطا، نحوه‌ی بروزسانی پارامترهای مدل و نحوه‌ی تعریف تابع خطا، بسته به هر الگوریتم یادگیری متفاوت است.

- 1 Supervised Learning
- 2 Unsupervised Learning
- 3 Reinforcement Learning
- 4 Label
- 5 Logistic Regression
- 6 Support Vector Machine
- 7 Decision Tree
- 8 Random Forest
- 9 K-Nearest Neighbors
- 10 Naïve Bayes Classifier

در یادگیری بدون نظارت، برچسبی برای داده‌ها وجود ندارد. در این روش عموماً می‌توانیم داده‌های شبیه به هم را در یک دسته قرار دهیم اما دانشی در مورد برچسب هر داده یا مجموع داده‌ها وجود ندارد. انواع الگوریتم‌های خوشه‌بندی^۱ مانند K-میانگین^۲، خوشه‌بندی سلسله‌مراتبی تجمعی^۳، خوشه‌بندی سلسله‌مراتبی تقسیمی^۴، الگوریتم PCA و بخشی از شبکه‌های عصبی مانند کدگذار^۵ها از جمله روش‌های بدون نظارت هستند (Kapoort و همکاران، ۲۰۲۲).

یادگیری تقویتی را عموماً دانش تصمیم‌گیری تعریف می‌کنند. در این رویکرد، هدف یادگیری بهترین رفتار در یک محیط برای رسیدن به بیشترین بهره یا جایزه است. اساس این روش از یادگیری بر نتیجه‌ی عمل و بازخورد^۶ استوار است. تفاوت این روش با یادگیری با نظارت در آن است که در اینجا برچسب‌های از پیش تعریف شده‌ای وجود نداشته و صرفاً تابعی برای بازخورد دادن به تصمیمات تعریف شده است. به عبارت دیگر در یادگیری تقویتی، داده‌های تصمیمات تا گرفتن تصمیم فعلی و بروزرسانی‌های لازم، مشخص نیست. این در حالیست که در روش با نظارت، همان ابتدا و برای آموزش مدل، ورودی هر نمونه و برچسب خروجی مشخص و آماده است (Escalante و Morales، ۲۰۲۲). این نوع یادگیری در حوزه‌ی معاملات الگوریتمی کاربرد داشته و در پژوهش فعلی مورد استفاده قرار نگرفته است.



شکل ۱ - معماری کلی یک روش یادگیری با نظارت

منبع: یافته‌های پژوهشگر

- 1 Clustering
- 2 K-Means
- 3 Agglomerative Hierarchical Clustering
- 4 Divisive Hierarchical Clustering
- 5 Encoder
- 6 Feedback

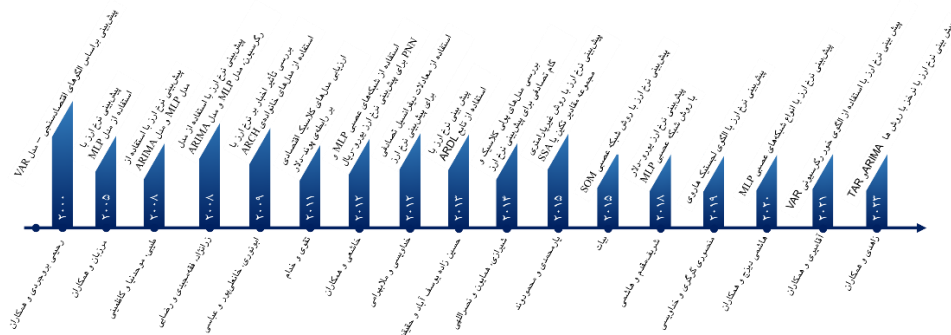
۳- پیشینه پژوهش

به منظور حفظ اختصار مهم‌ترین و مرتبط‌ترین پژوهش‌ها در این بخش ذکر شده‌اند. پژوهش‌های مرتبط با این مقاله، شامل تمامی پژوهش‌هایی هستند که یا مستقیماً درگیر پیش‌بینی نرخ ارز در بازار داخلی یا خارجی بوده‌اند، یا روش پیشنهادی آنها حاوی ایده و نکته‌ی ارزشمندی از جهت به کار بردن اخبار در پیش‌بینی است.

۳-۱- پیشینه داخلی

تا زمان نگارش این مقاله، در هیچ یک از پژوهش‌های داخلی مدلی که اخبار منتشره را در کنار سایر عوامل اقتصادی در پیش‌بینی نرخ ارز لحاظ کند، ارائه نشده است. تنها در پژوهش (ابونوری، خانعلی‌پور و عباسی، ۲۰۰۹) به بررسی آثار اخبار منتشره در بازار پرداخته شده و خود متن اخبار مستقیم یا غیرمستقیم در ساخت مدل پیش‌بینی نرخ ارز دخیل نشده است.

لذا در این بخش صرفاً می‌توان به بیان تمامی پژوهش‌های مرتبط با پیش‌بینی نرخ ارز اکتفا کرد. در شکل ۲، مهم‌ترین پژوهش‌های داخلی انجام شده در مدل‌سازی و پیش‌بینی نرخ ارز نمایش داده شده است. در این پژوهش‌ها تلاش شده تا با استفاده از متغیرهای اقتصادی و به کارگیری روش‌های اقتصادسنجی، آماری، غیرپارامتری و مدل‌های یادگیری ماشین کلاسیک، پیش‌بینی‌ای از نرخ ارز ارائه شود.

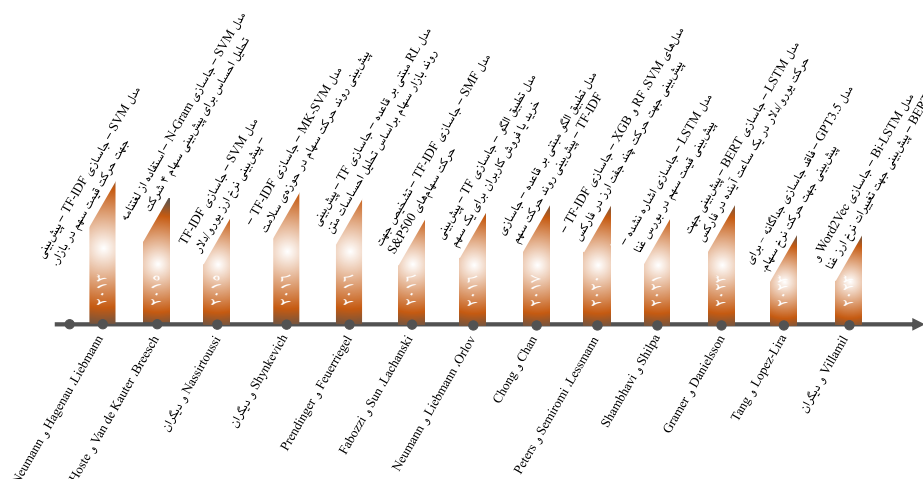


شکل ۲ - مهم‌ترین پژوهش‌های داخلی پیش‌بینی نرخ ارز

منبع: یافته‌های پژوهشگر

۳-۲- پیشینه خارجی

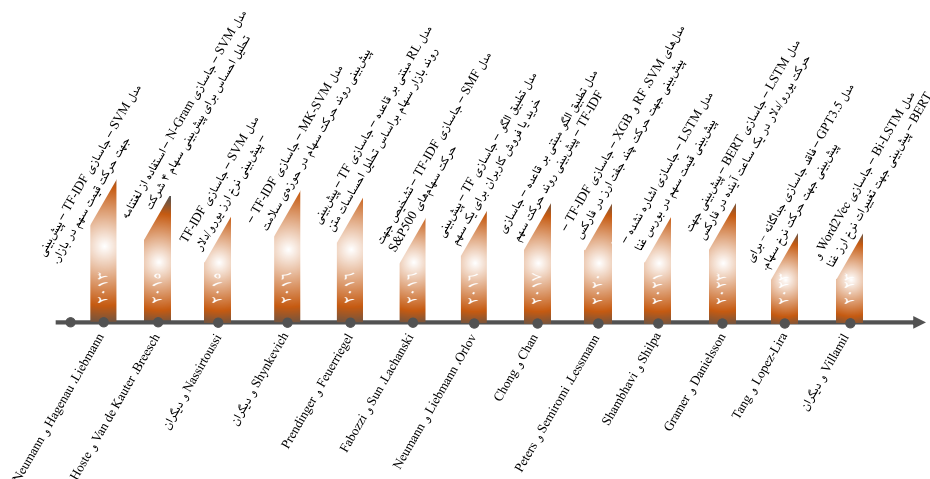
این پژوهش‌ها تلاش شده تا به نحوی داده‌های متنی نیز در کنار سایر انواع داده در پیش‌بینی‌ها لحاظ گردد. هرچند در اغلب این روش‌ها، رویکردی سیستماتیک و مدلی جامع دیده نمی‌شود. در



شکل ۳، مهم‌ترین و مرتبط‌ترین پژوهش‌های خارجی مرتبط نمایش داده شده است.

پژوهش‌های انجام شده در موارد زیر با یکدیگر متفاوت هستند:

- نحوه‌ی استفاده از داده‌های متنی و بازنمایی آنها به بردارهای عددی (جاسازی مورد استفاده)
- بازار و مجموعه‌ی داده‌ی مورد استفاده
- مدل یادگیری ماشین مورد استفاده
- نوع مسئله (رگرسیون یا کلاس‌بندی)



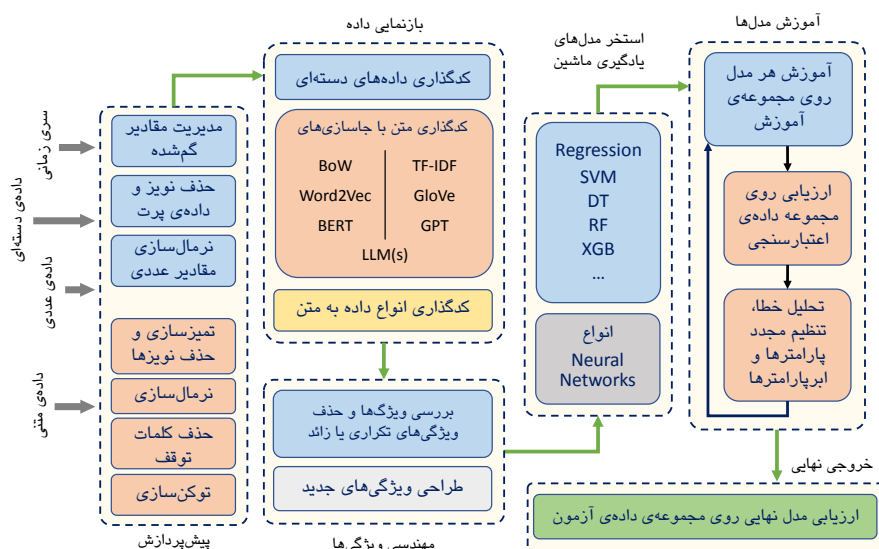
شکل ۳ - مهم‌ترین پژوهش‌های خارجی پیش‌بینی نرخ ارز

منبع: یافته‌های پژوهشگر

در ادامه و در بخش تشریح مدل پیشنهادی، به نکات مربوط به انواع جاسازی و مدل‌های یادگیری ماشین اشاره شده است. ذکر این نکته در اینجا حائز اهمیت است که در پژوهش‌های خارجی تمرکز بر پیش‌بینی نرخ بوده و روند یا روشی به منظور تحلیل تأثیر خبر و منبع منتشر کننده‌ی آن دیده نمی‌شود. نکته‌ی دیگر آنکه به جز پژوهش انجام شده توسط (Lopez-Lira و Tang، ۲۰۲۳) و (Villamil و دیگران، ۲۰۲۳) سایر روش‌ها از جاسازی‌های قدیمی‌تر که توانایی تفهیم سطوح بالاتر معنا را ندارند استفاده کرده‌اند که منجر به کارایی پائین‌تر آنها شده است. همچنین در دو پژوهش مذکور نیز، عملیات تنظیم دقیق جاسازی و مدل گزارش نشده است. این مسئله برای بهبود عملکرد مدل‌ها در این مسئله ضروری به نظر می‌رسد.

۴- مواد و روش‌ها

به‌عنوان یک رویکرد استاندارد برای حل مسئله به‌عنوان یک مسئله یادگیری ماشین، یک خط لوله پردازش داده طراحی شده است. همان‌طور که در شکل ۴ نشان داده شده، این خط لوله شامل پنج مرحله اصلی است که هر کدام دارای چندین زیرمرحله می‌باشند. این ساختار امکان بهبود و گسترش هر یک از مراحل را فراهم کرده و توسعه‌پذیری و سفارشی‌سازی مدل برای کاربردهای مختلف را تضمین می‌کند. گزینه‌های مطرح شده در هر مرحله، نمونه‌هایی از انتخاب‌های ممکن برای هر مرحله هستند. به‌عنوان مثال، در مرحله مدل‌سازی، بسته به نوع داده و مسئله می‌توان انتخاب‌های متفاوتی داشت.



شکل ۴ - خط لوله‌ی پردازش داده، آموزش و تنظیم مدل ترکیبی نهایی

منبع: یافته‌های پژوهشگر.

۴-۱- دادگان مورد استفاده

دادگان مورد استفاده شامل اطلاعات مربوط به نرخ واقعی ارز (ارز معامله شده در بازار آزاد و نه ارز دولتی و سهمیه‌ای) در بازه‌ی زمانی فروردین ۱۳۹۳ تا اسفند ۱۴۰۲ است. این داده‌ها از سه منبع مختلف وبسایت صرافی ملی^۱، وبسایت شبکه‌ی اطلاع‌رسانی نرخ طلا و ارز^۲، و وبسایت بن‌بست^۳ جمع‌آوری و مطابقت داده شده است. هر سطر از این دادگان، شامل اطلاعات مربوط به نرخ باز شدن، نرخ بسته شدن، کمترین نرخ معامله، و بیشترین نرخ معامله در هر روز است.

علاوه بر داده‌های عددی فوق، تمامی اخبار منتشره توسط بانک مرکزی^۴، روزنامه‌ها، پایگاه‌ها و خبرگزاری‌های اقتصادی فارسی زبان داخلی شامل روزنامه‌ی دنیای اقتصاد^۵، روزنامه‌ی صنعت-معدن-تجارت^۶، روزنامه‌ی آسیا^۷،

1 <https://www.mex.co.ir/>

2 <https://www.tgju.org/>

3 <https://www.bonbast.com/>

4 <https://www.cbi.ir/>

5 <https://donya-e-eqtasad.com/>

6 <https://www.smtnews.ir/>

7 <https://www.asianews.ir/>

روزنامه‌ی ابرار اقتصادی^۱، بخش سیاسی-اقتصادی خبرگزاری ایسنا^۲، بخش سیاسی-اقتصادی پایگاه خبرآنلاین^۳، بخشی اقتصادی پایگاه خبری تابناک^۴، و پایگاه‌های خارجی شامل بخش سیاسی-اقتصادی بی‌بی‌سی فارسی^۵ و همچنین بخش سیاسی-اقتصادی صدای امریکای فارسی^۶ در این بازه‌ی زمانی جمع‌آوری شده است. از آنجا که روش‌های یادگیری ماشین مورد استفاده در اینجا از نوع روش‌های یادگیری بانظارت است، هر خبر به برچسبی براساس مشاهدات بازار ارز در زمان انتشار خود نیاز دارد. این برچسب‌ها بر اساس تغییرات نرخ ارز معاصر با زمان انتشار هر خبر اعمال می‌شوند برای این برچسب‌گذاری، نکات زیر در نظر گرفته شده است:

- کوانتوم زمانی یا کوچکترین واحد زمانی که اطلاعات نرخ ارز برای آن موجود است، یک روز است.
 - در هر روز نرخ باز شدن، نرخ بسته شدن، کمترین، بیشترین و میانگین نرخ معاملات را ثبت شده است.
 - تأثیر اخبار منتشره در ساعات غیر کاری و یا ایام تعطیل، در اولین ساعت کاری بعد لحاظ شده است.
 - اگر خبری در زمان باز بودن بازار منتشر شود، تأثیر آن در نرخ‌های همان روز اعمال می‌شود.
 - از آنجا که مدل پیشنهادی یک مدلی مستقل از محتوای و براساس داده‌ی ورودی آموزش می‌بیند. علاوه بر دادگان متنی اخبار و نرخ ارز، تعدادی از شاخص‌های اقتصادی مرتبط با ارز نیز به عنوان ورودی در نظر گرفته شده‌اند. این ایده به دلایل زیر است:
 - مدل با استفاده از تمامی انواع دادگان، پیش‌بینی نزدیک‌تری به واقعیت داشته باشد.
 - در بخش تحلیل اثر اخبار، بتوانیم تأثیر سایر متغیرها و تأثیر اخبار را از یکدیگر تفکیک کرده و اطمینان حاصل کنیم که تغییرات حاصل شده به دلیل تأثیر اخبار بوده است.
 - یک تحلیل روش‌مند در ساختار پردازشی برای استفاده از انواع متغیرهای مؤثر بر نرخ داشته باشیم.
- به طور کلی ساختار مجموعه داده‌ی ساخته شده به صورت شکل ۵ است. در مورد بازنمایی‌های عددی، نرمال‌سازی‌ها و برچسب‌ها در ادامه توضیحات لازم داده شده است.

1 <https://www.abrarnews.com/>

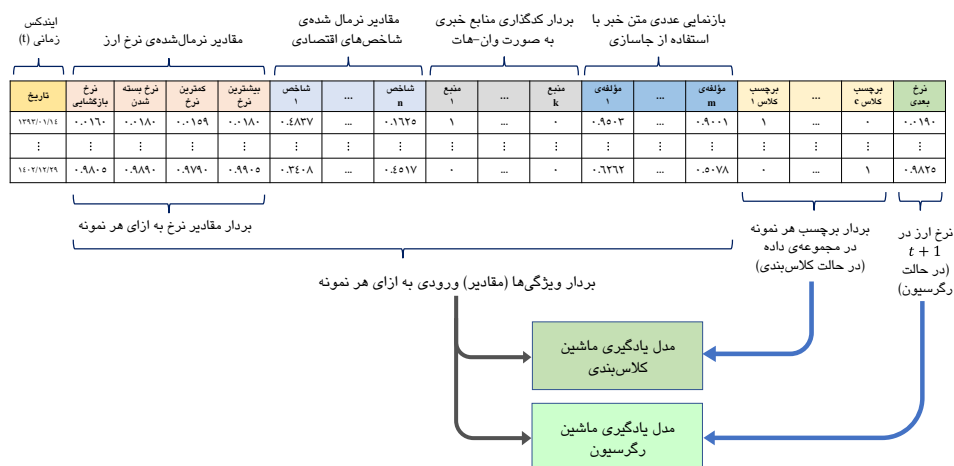
2 <https://www.isna.ir/service/Economy>

3 <https://www.khabaronline.ir/service/Economy>

4 <https://www.tabnak.ir/fa/economic>

5 <https://www.bbc.com/persian/topics/cl819mvlllqt>

6 <https://ir.voanews.com/z/1066>



شکل ۵ - ساخت بردارهای عددی به ازای هر نمونه داده

منبع: یافته‌های پژوهشگر

۴-۲- پیش‌پردازش داده‌ها

در بخش پیش‌پردازش، عملیات آماده‌سازی مربوط به هر نوع داده انجام می‌شود. این عملیات برای داده‌های سری زمانی و عددی مشابه است، اما برای داده‌های متنی کاملاً متفاوت است.

۴-۲-۱- پیش‌پردازش داده‌های عددی و سری‌های زمانی

پردازش داده‌گان عددی در سه دسته‌ی کلی جای می‌گیرد: ۱- نرمال‌سازی بازه‌ی مقادیر^۱ ۲- مدیریت مقادیر ناموجود^۳ ۳- حذف نویز و مقادیر پرت^۴.

از آنجا که مقادیر عددی وارد شده به مدل‌ها ممکن است محدودیت یا کران^۴ مشخصی نداشته باشند، لازم است این مقادیر را به یک بازه معین تبدیل کنیم. برای این کار می‌توان از توابعی مانند لگاریتم استفاده کرد. روش دیگری نیز برای نرمال‌سازی وجود دارد که شامل استفاده از نرمال‌ساز توزیع نرمال می‌شود (Singh و Singh, ۲۰۲۰). در این مقاله از روش استانداردسازی با استفاده از رابطه‌ی زیر، مقادیر هر ویژگی به نحوی نرمال شده‌اند که هر ویژگی میانگین صفر و انحراف معیار استاندارد ۱ دارد.

1 Value Normalization

2 Missing Values

3 Outlier Values

4 Boundary

$$Z = \frac{x - \mu}{\sigma} \quad (1)$$

در این رابطه x همان مشاهدات یا مقدار ویژگی، μ میانگین، و σ انحراف معیار استاندارد است. با توجه به اینکه بسیاری از الگوریتم‌های یادگیری ماشین به صورت توابع احتمال توسعه داده شده و انتظار آنها از ورودی‌ها در بازه ۰ تا ۱ است، این رویکرد مؤثر بوده و در این مقاله نیز از آن استفاده شده است.

در بسیاری از موارد، ممکن است برای یک تاریخ یا محدوده‌ای از تاریخ، داده‌ها در دسترس نباشند. برای مدیریت این داده‌ها، ایده‌های مختلفی وجود دارد (Josse و Husson، ۲۰۱۲)؛ ایده‌ی اول حذف این داده‌ها از مجموعه‌ی داده در صورت عدم آسیب به کلیت مسئله است. ایده‌ی دوم تقریب این داده‌ها براساس داده‌های مجاور آن با استفاده از روش‌هایی چون میانگین‌گیری و یا درون‌یابی است. در این مقاله در مجموع از ۳۶۵۳ روز، ۱۸۴ داده‌ی ناموجود داشتیم که حول آن‌ها نوسانات خاصی در نرخ ارز نبود. لذا این دادگان را از مجموعه‌ی داده حذف کردیم.

در بسیاری از مجموعه‌های داده، ممکن است به دلایل مختلف نویز و مقادیر نامربوط وجود داشته باشد. در گام نخست، این داده‌ها اعتبارسنجی می‌شوند. به عنوان مثال، اگر بیشترین قیمت معامله در یک روز ۱۰ تومان باشد، ولی قیمت میانگین ۲۰ تومان درج شده باشد، این داده نویزی و پرت محسوب می‌شود و باید از مجموعه حذف شود. با اینکه روش‌های متعددی برای تشخیص مقادیر نامربوط و تمایز آنها از مقادیر نادر ولی معتبر وجود دارد، در این مقاله از تحلیل‌های ساده سازگاری داده‌ها استفاده شده است. نتایج نشان می‌دهد که این روش به اندازه کافی مؤثر بوده است. در نهایت، ۵ داده نویزی از مجموعه حذف شدند.

۱-۲-۴- پیش‌پردازش داده‌های دسته‌ای

دادگان دسته‌ای^۱ به داده‌هایی اطلاق می‌شود که مقادیر آنها از یک مجموعه محدود و گسسته انتخاب می‌شوند. به عنوان مثال، برچسب‌ها در یک مسئله کلاس‌بندی نمونه‌ای از دادگان دسته‌ای هستند. برای این نوع دادگان، روش‌های کدگذاری مختلفی پیشنهاد شده است. ایده‌هایی مانند استفاده از اعداد ترتیبی نرمال‌شده حول صفر، کدینگ وان-هات، کدینگ باینری و کدینگ همینگ، هر کدام مزایا و معایب خاص خود را برای مسائل مختلف دارند (Potdar و همکاران، ۲۰۱۷). در این مقاله دو نوع داده‌ی دسته‌ای وجود دارد. نخست برچسب‌های خروجی مدل‌ها در حالت کلاس‌بندی است که شامل ۷ حالت است. دوم متغیر مشخص‌کننده‌ی منبع یک خبر است که ۱۰ مقدار مختلف می‌تواند به خود بگیرد. در اینجا برای هر دو مورد از کدینگ وان-هات استفاده شده که به خوبی با نیازهای این مسئله تطابق داشته و کافی بوده است. در این نوع کدگذاری، به تعداد مقادیر مختلف یک متغیر دسته‌ای، ستون‌های جداگانه در نظر گرفته می‌شود و برای هر نمونه، ستون مربوط به مقدار آن متغیر برابر با یک و سایر ستون‌ها برابر با صفر خواهند بود. به ازای هر نمونه، مقدار آن ویژگی هرچه باشد، ستون متناظر با آن یک شده و مابقی ستون‌ها صفر می‌شوند. در شکل ۵ این کدینگ قابل مشاهده است.

1 Categorical Data

۱-۲-۴- پیش‌پردازش داده‌های متنی

پیش‌پردازش دادگان متنی شامل مجموعه اقداماتی است که هدف آن آماده‌سازی داده به منظور ساخت امبدینگ یا جاسازی^۱ برداری بهتر از متن است. در ادبیات یادگیری ماشین، جاسازی‌ها بازنمایی‌هایی به شکل بردارهای عددی از متون، تصویر و انواع فضاهای ورودی هستند. عملیات پیش‌پردازش متنی عموماً عبارتند از: تمیزسازی و حذف داده‌های نویزی، نرمال‌سازی، توکن‌سازی^۲، حذف کلمات توقف. این مراحل در شکل ۶ نمایش داده شده‌اند. به دلیل آنکه متون اخبار عموماً از صفحات اینترنتی جمع‌آوری شده‌اند، دارای تگ‌های اچ‌تی‌ام‌ال^۳ و کدهای جاوااسکریپت^۴ هستند. این داده‌ها فاقد ارزش بوده و معنایی به متن اصلی اضافه نمی‌کنند؛ لذا در محله اول و به منظور تمیزسازی داده‌ها، آنها را حذف می‌کنیم. در گام دوم، انواع مختلف کاراکترهای موجود در متون مختلف طی فرآیند نرمال‌سازی یکدست و متحدالشکل می‌شوند. در اینجا منظور از نرمال‌سازی، یکدست‌سازی فونت‌ها نیست؛ بلکه هدف یکدست‌سازی کاراکترهایی است که معنی یکسان داشته اما با کد (اسکی^۵ یا یونیکد^۶) متفاوتی در متن ظاهر شده‌اند (Gupta و Singh, ۲۰۱۶).



شکل ۶ - مراحل پیش‌پردازش داده‌های متنی

منبع: یافته‌های پژوهشگر

- 1 Embedding
- 2 Tokenization
- 3 HTML
- 4 JavaScript
- 5 ASCII
- 6 Unicode

برای آنکه بردارهای عددی تولید شده بیش از اندازه بزرگ نشوند، می‌توان برخی از کلمات که بار معنایی خاصی نداشته و اصطلاحاً کلمات توقف^۱ نامیده می‌شوند را از توکن‌ها حذف کرد (Nayak و همکاران، ۲۰۱۶). به عنوان مثال حروف ربط مانند «از، به، در، با، را، بر» از جمله کلمات توقف هستند که آنها را حذف می‌کنیم. در این مقاله در آزمایش‌های اولیه مربوط به جاسازی‌های کوچک، ۱۳۲ کلمه‌ی توقف از توکن‌ها حذف شده‌اند. در ادامه به دلیل استفاده از جاسازی‌های بزرگ و پوشش آن‌ها از انواع حالت‌ها، نیازی به حذف کلمات توقف وجود نداشت. مرحله‌ی بعد، فرآیند توکن‌سازی است که در آن متن نرمال‌شده به دنباله‌ای از کلمات یا واحدهای کوچک‌تر شکسته می‌شود. در ساده‌ترین شکل، پس از تمیزسازی و نرمال‌سازی متن، با جایگزین کردن علائم نگارشی با فاصله، هر بخش از متن که با فاصله از دیگری جدا شده باشد به عنوان یک توکن در نظر گرفته می‌شود. البته روش‌های پیچیده‌تری برای توکن‌سازی وجود دارد (Dridan و Oepen، ۲۰۱۲) اما برای اهداف این مقاله، همین رویکرد ساده کافی است.

۳-۴- بازنمایی داده‌های متنی

برای تبدیل دادگان متنی به بردارهای عددی، می‌توان از انواع مختلف روش‌های جاسازی استفاده کرد. در سال‌های اخیر، مدل‌های متعددی برای جاسازی متنی معرفی شده‌اند که هر یک تلاش کرده‌اند تا نقاط ضعف نسخه‌های قبلی را بهبود بخشند یا مسائل جدیدی را پوشش دهند. (Yang و Li، ۲۰۱۸). قدیمی‌ترین و شاید بدیهی‌ترین روش، روش خورجین کلمات^۲ یا BoW است. در این روش یک ماتریس سند-توکن ساخته می‌شود. سطرهاى این ماتریس بیانگر اسناد متنی و ستون‌های آن بیانگر توکن‌های موجود در این اسناد هستند. اگر توکنی در سندی وجود داشته باشد، درایه‌ی متناظر با آن در این ماتریس برابر با ۱، و در غیر اینصورت برابر با ۰ خواهد بود. ماتریس خروجی این روش، در مسائلی که تعداد توکن‌ها و اسناد زیاد است، بسیار بزرگ خواهد شد. استفاده از روش‌هایی چون ریشه‌یابی توکن‌ها و حذف کلمات توقف می‌تواند به کاهش ابعاد این ماتریس کمک کند، اما همچنان با ماتریس نسبتاً بزرگی مواجه هستیم.

فارغ از ابعاد ماتریس خروجی، روش BoW اهمیت کلمات را در نظر نمی‌گیرد و صرف وجود یا عدم وجود یک کلمه در سند مطرح است. روش TF-IDF ایده‌ی BoW را گسترش داده به نحوی که نرخ رخداد کلمات در اسناد را محاسبه می‌کند. فرکانس عبارت^۳ و معکوس فرکانس سند^۴ براساس روابط زیر محاسبه می‌شوند:

$$(۲) \quad TF = \frac{\text{تعداد دفعاتی که کلمه در سند وجود دارد}}{\text{تعداد کل کلمات در سند}}, \quad IDF = \frac{\text{تعداد کل اسناد}}{\text{تعداد اسنادی که این کلمه در آن وجود دارد}}$$

1 Stop Words

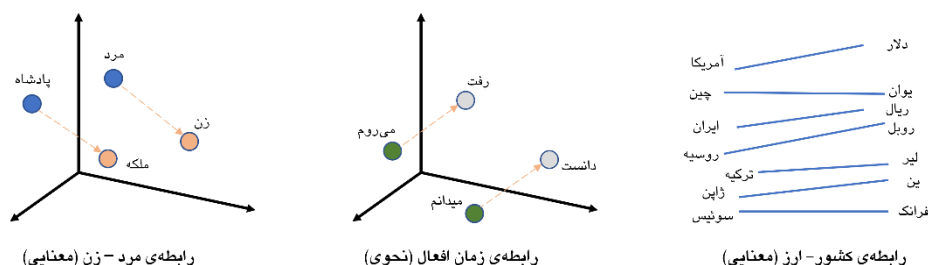
2 Bag of Words

3 Term Frequency

4 Inverse Document Frequency

ماتریس تولید شده در این روش از نظر ابعاد مشابه روش BoW است، با این تفاوت که هر دریای آن از حاصلضرب TF و IDF به ازای هر توکن-سند حاصل می‌شود. این روش با وجود آنکه اهمیت کلمات را در نظر می‌گیرد، اما ترتیب و زمینه‌ی آن‌ها را نادیده می‌گیرد. در این شرایط جایگشت‌های مختلف کلمات تفاوتی باهم نداشته و معنا به خوبی مدل نمی‌شود. هرچند می‌توان در بخش توکن‌ها به جای کلمات تکی، از N-گرم آنها استفاده کرد، اما این فرآیند ابعاد ماتریس خروجی را به صورت نمایی افزایش می‌دهد.

به منظور مدل‌سازی و انعکاس روابط معنایی^۱ و نحوی^۲ زبان‌ها، بازنمایی‌ها و جاسازی‌های متنی پیشرفته‌تری تولید شدند. نمونه‌هایی از روابط معنایی و نحوی توکن‌ها در شکل ۷ نمایش داده شده است.



شکل ۷ - نمونه‌ای از روابط معنایی و نحوی در زبان فارسی

باز طراحی شده توسط نویسندگان مقاله براساس مرجع (Mikolov و همکاران، ۲۰۱۳).

روش Word2Vec که در سال ۲۰۱۳ توسط گوگل معرفی شد، یک شبکه عصبی کم‌عمق است که با تبدیل کلمات به بردارهای عددی، روابط معنایی و نحوی بین کلمات را مدل‌سازی می‌کند (Mikolov و همکاران، ۲۰۱۳). در سال ۲۰۱۴، GloVe با عملکرد مشابه معرفی شد، اما تفاوت‌هایی در روش تولید جاسازی دارد (Pennington و همکاران، ۲۰۱۴). سپس در سال ۲۰۱۷، FastText توسط فیسبوک عرضه شد که با حل مشکلاتی مثل مدیریت کلمات خارج از دامنه (OOV)، نسخه بهبودیافته‌ای از Word2Vec بود (Bojanowski و همکاران، ۲۰۱۷).

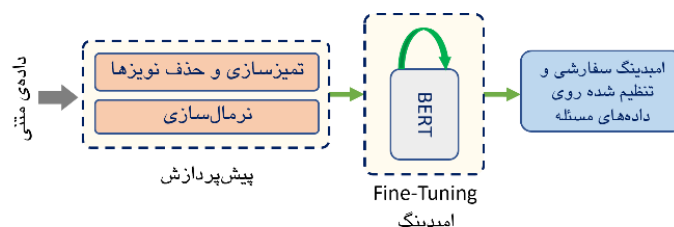
این روش‌ها همگی یک جاسازی ثابت برای هر کلمه تولید می‌کنند، اما در کاربردهایی که معنای کلمه بسته به متن تغییر می‌کند، روش‌های پیشرفته‌تری نیاز است. ELMo، معرفی‌شده در ۲۰۱۸، برای تولید جاسازی یک کلمه، کل جمله را در نظر می‌گیرد و معنای کلمه را براساس زمینه آن تعیین می‌کند (Sarzynska-Wawer و همکاران، ۲۰۲۱). همچنین، در سال ۲۰۱۸ BERT توسط گوگل معرفی شد که با استفاده از معماری ترنسفورمر، عملکرد سریع‌تری داشته و توانایی تولید جاسازی برای کلمات و جملات را دارد (Devlin و همکاران، ۲۰۱۸).

در این مقاله از جاسازی BERT فارسی (Farahani و همکاران، ۲۰۲۱) برای بازنمایی دادگان متنی استفاده شده است. به دلیل آنکه جاسازی‌های بزرگ و مبتنی بر معماری ترنسفورمر، روی حجم عظیم دادگان آموزش داده شده

1 Semantic

2 Syntactic

و برای مدل‌سازی زبان به صورت عمومی^۱ به کار می‌روند، در کاربردهایی که روی بخش خاصی از ادبیات زبانی تمرکز دارند، عموماً کیفیت متوسطی دارند. در این شرایط باید این جاسازی‌ها را با دادگان مرتبط با پژوهش، تنظیم دقیق^۲ یا سفارشی‌سازی کنیم. این فرآیند در واقع بخشی از فرآیند آموزش و تولید جاسازی است که با تمرکز بر داده‌های خاص و سفارشی شده انجام می‌شود. جزئیات این فرآیند در (Sousa و همکاران، ۲۰۱۹) موجود است. با این عملیات کیفیت بازنمایی متون مورد استفاده در پژوهش توسط جاسازی، بالاتر خواهد رفت. خروجی جاسازی تنظیم شده به ازای هر سند ورودی، یک بردار با ابعاد ۷۶۸ است. در این مقاله، عملیات تنظیم دقیق به اندازه‌ی ۸ ایپاک^۳ (سیکل آموزش مدل) انجام شده است. این مقادیر با در نظر گرفتن جلوگیری از بیش‌برازش جاسازی BERT لحاظ شده‌اند. همچنین نرخ یادگیری در بازه‌ی $[1e-5, 5e-5]$ بوده و اندازه‌ی Batch براساس GPU ی در دسترس ۱۶ یا ۳۲ انتخاب شده است.



شکل ۸ - تنظیم دقیق جاسازی روی داده‌های متنی پژوهش

منبع: یافته‌های پژوهشگر.

در سال‌های اخیر، تلاش‌های بسیاری برای توسعه مدل‌های زبانی بزرگ صورت گرفته که برجسته‌ترین آن‌ها مدل‌های زبانی GPT از شرکت OpenAI هستند (Brown و همکاران، ۲۰۲۰). تا زمان نگارش این سند، ۷ نسخه رسمی از این مدل‌ها ارائه شده که آخرین نسخه، GPT-4-O است که به‌ویژه از طریق ChatGPT تأثیر زیادی بر جهان داشته است (Mohamadi و همکاران، ۲۰۲۳). در این پژوهش، مدل‌های زبانی بزرگ به‌عنوان یکی از ابزارهای ارزیابی مسئله مورد استفاده قرار گرفته‌اند. در کار با این مدل‌ها، دیگر نیازی به استفاده از روش‌های جاسازی نیست و داده‌ها مستقیماً به صورت متنی به مدل ارسال می‌شوند. با توجه به اینکه این مدل‌ها به‌صورت سرویس ارائه می‌شوند و کد آنها در دسترس نیست، استفاده از آن‌ها نگرانی‌هایی را ایجاد می‌کند که در بخش پایانی به آن پرداخته خواهد شد.

1 General
2 Fine-Tuning
3 Epoch

۴-۴- مهندسی ویژگی‌ها

در هر مدل یادگیری ماشین، می‌توان از ویژگی‌های متعددی استفاده یا آنها را تعریف کرد. مهندسی ویژگی‌ها فرآیندی است که طی آن ویژگی‌های موجود از نظر سودمندی یا زائد بودن ارزیابی می‌شوند و در صورت نیاز، ویژگی‌های جدیدی تعریف می‌گردند (Nargesian و همکاران، ۲۰۱۷). د

در این پژوهش نیز ویژگی‌های مرتبط با نرخ ارز پیش از استفاده در آموزش مدل بررسی شدند تا ویژگی‌های زائد شناسایی شوند. ویژگی زائد به آن دسته از ویژگی‌ها اطلاق می‌شود که در فرآیند کلاس‌بندی یا رگرسیون نقش مؤثری ایفا نمی‌کنند. این ویژگی‌ها معمولاً یا ارتباط کمی با مسئله دارند یا توسط سایر ویژگی‌ها پوشش داده می‌شوند. روش‌های مختلفی برای یافتن ویژگی‌های زائد و انتخاب ویژگی‌های مفید وجود دارد (Heaton، ۲۰۱۶). در ساده‌ترین حالت می‌توان همبستگی هریک از ویژگی‌ها با یکدیگر و همچنین با برجسب خروجی را محاسبه کرده و ویژگی‌هایی که با خروجی همبستگی پائینی داشته یا توسط ویژگی دیگری توصیف می‌شوند، حذف کرد. اگر برداری از ویژگی‌های به شکل $X = [X_1, \dots, X_n]$ داشته باشیم، همبستگی آن به صورت زیر است:

$$Corr(X) = \begin{bmatrix} \rho_{11} & \cdots & \rho_{1n} \\ \vdots & \ddots & \vdots \\ \rho_{n1} & \cdots & \rho_{nn} \end{bmatrix}, \quad \rho_{nm} = \frac{COV(X_i, X_j)}{X_{\sigma_i} X_{\sigma_j}} \quad (4)$$

اگر در ماتریس همبستگی X عدد، اندازه‌ی بزرگ‌تر مقادیر ρ_{ij} ، نشان‌دهنده هم‌خطی^۱ بودن جدی در دو متغیر i و j است (Nargesian و همکاران، ۲۰۱۷). نتایج این تحلیل در شکل ۹ نمایش داده شده است. براساس این نمودار، به نظر می‌رسد که نرخ تورم، میزان بدهی عمومی و رشد نقدینگی بیشترین همبستگی را با نرخ ارز بازار آزاد دارند. همچنین سایر شاخص‌های اقتصادی دیگر نیز با نرخ ارز دارای همبستگی هستند به طوری هیچ ویژگی کم‌اهمیتی (زائد) در این مجموعه دیده نمی‌شود. این نتیجه مؤید صحت داده‌های جمع‌آوری شده و مطابقت آن با تئوری‌ها و مشاهدات اقتصادی است.

مسئله دیگر بررسی متغیرهایی است که با یکدیگر همبستگی بالایی دارند، نه لزوماً با نرخ ارز. بر اساس شکل ۹، دو متغیر "بدهی عمومی" و "رشد نقدینگی" ارتباط نزدیکی با یکدیگر دارند. از دیدگاه تئوری اطلاعات، این موضوع به این صورت تفسیر می‌شود که این متغیرها اطلاعات جدیدی به مسئله اضافه نمی‌کنند و وجود یکی از آنها برای آموزش و بهبود عملکرد مدل کافی است. بنابراین، از منظر تئوری و به منظور کاهش فضای حالت مسئله، می‌توان یکی از این دو متغیر را حذف کرد. در این میان، "رشد نقدینگی" گزینه مناسبی برای حذف است، زیرا "بدهی عمومی" همبستگی بیشتری با نرخ ارز دارد.

1 Co-Linearity



شکل ۹ - نمودار همبستگی شاخص‌های اقتصادی ایران (۲۰۰۰ - ۲۰۲۳)

* منبع: یافته‌های پژوهشگر.

۴-۵- مدل پیشنهادی

به منظور بررسی اینکه کدام مدل برای مدل‌سازی و پیش‌بینی نرخ ارز مناسب‌تر است، مدل‌های متعددی پیاده‌سازی و مورد ارزیابی قرار گرفته‌اند. این مدل‌ها در دو حالت مختلف آموزش داده شده‌اند؛ در حالت اول، مسئله به عنوان یک مسئله کلاس‌بندی در نظر گرفته شده و در حالت دوم، به عنوان یک مسئله رگرسیون تحلیل شده است. علاوه بر این، برای هر یک از این حالت‌ها، مدل‌ها با دو نوع داده متفاوت آموزش دیده‌اند: نوع اول شامل داده‌های صرفاً عددی و نوع دوم شامل داده‌های عددی همراه با داده‌های متنی بوده است. این رویکرد اتخاذ شده است تا در مرحله تحلیل، بتوان تأثیر واقعی استفاده از اخبار را به طور دقیق ارزیابی کرد.

برای انتخاب مدل، موارد مختلفی را می‌توان مدنظر قرار داد. مدل‌های کلاسیک یادگیری ماشین مانند SVM، RF و امثالهم در مسائل پیچیده و بزرگ کارایی مناسبی ندارند (Khan و همکاران، ۲۰۲۳)؛ به دلیل اندازه بزرگ بردارهای ورودی، پیچیدگی مسئله و نیاز به مدل‌سازی روابط غیرخطی، از مدل‌های یادگیری عمیق استفاده شده است. عملکرد این مدل‌ها به داده‌های ورودی، معماری، ظرفیت و کیفیت آموزش بستگی دارد. تنها برخی از مدل‌های یادگیری ماشین به طور خاص برای داده‌های زمانی و سری‌ها طراحی شده‌اند. این مدل‌ها مانند RNN، LSTM و GRU، روابط موجود در توالی داده‌های زمانی را در نظر می‌گیرند و همگی از خانواده شبکه‌های عصبی هستند.

علاوه بر مدل‌های ذکر شده، مدل‌های پیشرفته‌تری نیز وجود دارند که از جمله آن‌ها می‌توان به مدل‌های مبتنی بر معماری ترنسفورمر اشاره کرد، که ابزارهایی مانند ChatGPT بر اساس آن‌ها ساخته شده‌اند. این مدل‌ها به دلیل نیاز به منابع پردازشی زیاد و حجم بالای داده، از ابتدا در این پژوهش آموزش داده نشده‌اند. با این حال، از آخرین نسخه مدل (GPT در زمان نگارش این مقاله) با دسترسی API و با رویکرد تنظیم دقیق استفاده شده که نتایج آن در ادامه آمده است.

برای رسیدن به مدل نهایی، جدیدترین مدل‌های متناسب با مسئله پیاده‌سازی، آموزش و تنظیم دقیق شده‌اند. طبق نتایج به‌دست‌آمده، مدل تنظیم دقیق شده LLM و مدل Bi-GRU بهترین عملکرد را در سناریوهای مختلف ارائه داده‌اند. با توجه به اینکه مدل LLM فقط از طریق API قابل دسترسی است و به صورت جعبه سیاه عمل می‌کند، مدل Bi-GRU به عنوان بهترین مدل مورد تشریح قرار گرفته است.

۱-۵-۴- مدل Bi-GRU

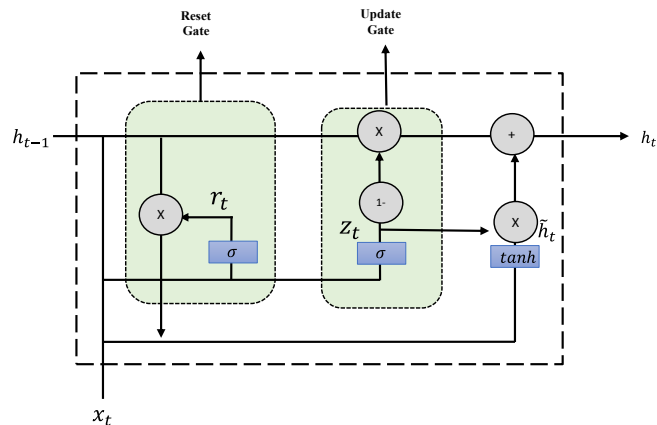
مدل انتخاب شده GRU دوطرفه یا Bi-GRU است که در زمره‌ی شبکه‌های عصبی تکرارپذیر^۱ قرار می‌گیرد. به طور کلی مدل‌های عصبی تکرارپذیر برای مدل‌سازی مقادیر یا مسائل وابسته به زمان بسیار کارا هستند. مزیت این مدل نسبت به مدل Bi-LSTM که در اکثر پژوهش‌های دیگر استفاده شده، سرعت آموزش بالاتر آن است (Seabe و همکاران، ۲۰۲۳). همچنین مدل پیشنهادی به دلیل ترکیب و استفاده از انواع داده در فرآیند آموزش و پیش‌بینی، ذیل تکنیک‌های تلفیق داده^۲ قابل طبقه‌بندی است (Liu و همکاران، ۲۰۲۰).

ساختار کلی مدل استفاده شده در شکل ۱۱ نمایش داده شده است. بردارهای عددی بعد از پردازش در مرحله‌ی نرمال‌سازی عددی و تولید خروجی جاسازی به مدل GRU وارد می‌شوند. این مدل دارای تعدادی مرحله

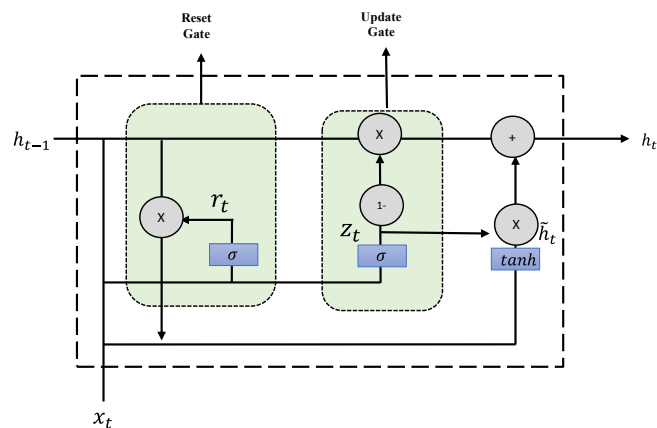
1 Recurrent Models

2 Data Fusion

یا واحد^۱ است. نمایی از این واحدها که در آن دروازه‌های بازنشانی و بروزرسانی ترسیم شده در



شکل ۱۰ نمایش داده شده است.

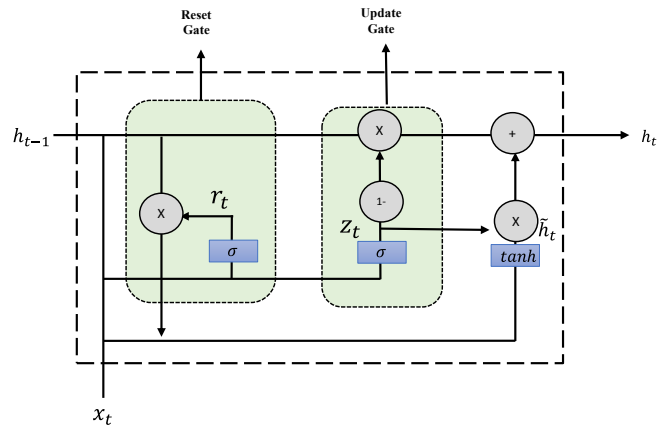


شکل ۱۰ - یک واحد پردازشی در مدل GRU

منبع: یافته‌های پژوهشگر.

در اینجا این مراحل به اندازه پنجره‌ی وابستگی زمانی یا همان f می‌باشد. مدل پیشنهادی در هر مرحله دادگان مربوط به f روز گذشته را دریافت کرده و براساس آن وزن‌ها یا پارامترهای مدل را بروز می‌کند. لذا برای هر

پیش‌بینی در زمان t ، مقادیر ورودی در بازه‌ی $[t - f + 1, \dots, t - 1]$ ، را به مدل می‌دهیم. مقادیر مشخص شده



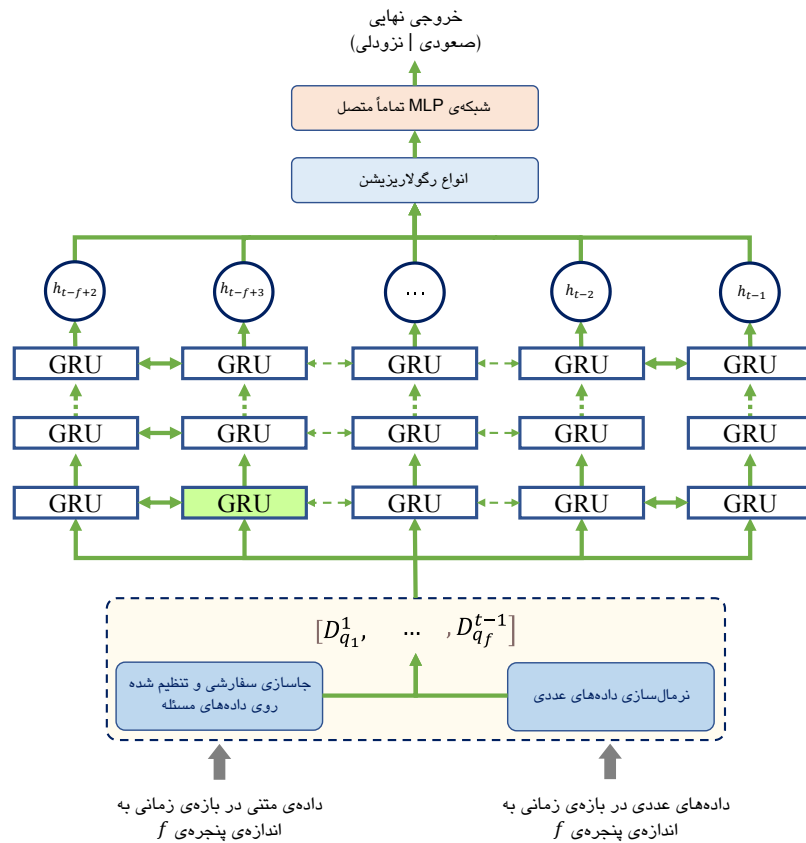
در

شکل ۱۰، به ازای هر واحد در معماری GRU با استفاده از روابط زیر محاسبه می‌شوند:

$$\begin{aligned} z_t &= \sigma(W_z \cdot [h_{t-1}, x_t]) \\ r_t &= \sigma(W_r \cdot [h_{t-1}, x_t]) \\ \tilde{h}_t &= \tanh(W \cdot [r_t * h_{t-1}, x_t]) \\ h_t &= (1 - z_t) * h_{t-1} + z_t * \tilde{h}_t \end{aligned} \quad (5)$$

در هر ایپاک وزن‌ها براساس میزان خطای مدل و در جهت کاهش آن با استفاده از الگوریتم بهینه‌سازی بروز می‌شوند. بروزرسانی وزن‌های مدل تا جایی ادامه می‌یابد که مقدار خطای مدل به ازای چند ایپاک متوالی کاهش نیافته و یا به آستانه‌ی مورد قبول تعریف شده برسد. به طور کلی تابع خطا از تفاضل مقدار پیش‌بینی شده توسط مدل و مقدار برچسب واقعی محاسبه می‌شود. تابع خطا در اینجا در رابطه‌ی ۵ نمایش داده شده است. این تابع به تابع کراس‌آنتروپی^۱ معروف بوده و به دلیل ویژگی‌های محاسباتی مناسب در فرآیند آموزش، در مسائل کلاسیک بسیار مرسوم است. در این رابطه $\hat{y}^t$ پیش‌بینی مدل و y^t برچسب حقیقی دیتا در زمان t است. این تابع توسط بهینه‌ساز ADAM در فرآیند آموزش بهینه می‌شود.

$$Loss(t) = \frac{1}{T} \sum_{t=1}^T (y^t \log(\hat{y}^t) + (1 - y^t) \log(1 - \hat{y}^t)) \quad (5)$$



شکل ۱۱ - مدل پیشنهادی مبتنی بر تلفیق داده و یادگیری عمیق
منبع: یافته‌های پژوهشگر.

برای جلوگیری از بیش‌برازش^۱ مدل، از تکنیک‌های رگولاریزیشن^۲ استفاده شده است. این تکنیک‌ها کمک می‌کنند تا مدل در فرآیند بهینه‌سازی و پیدا کردن وزن‌های بهینه، از افتادن در بهینه‌های محلی^۳ یا ایجاد بایاس روی بخشی از داده‌ها جلوگیری کند. دو تکنیک رایج در این زمینه، دوراندازی^۴ و نرمال‌سازی دسته‌ای^۵ هستند که در این مدل نیز به کار رفته‌اند. در لایه آخر مدل GRU، برای تجمیع خروجی‌ها از یک لایه MLP کاملاً متصل استفاده

- 1 Overfitting
- 2 Regularization
- 3 Local Optimum
- 4 Dropout
- 5 Batch Normalization

شده است تا خروجی نهایی در حالت طبقه‌بندی به شکل یک بردار به اندازه تعداد کلاس‌ها و در حالت رگرسیون به صورت یک عدد پیوسته باشد. تابع فعال‌سازی در حالت طبقه‌بندی، تابع سافت‌مکس^۱ است که خروجی دودویی^۲ را تعیین می‌کند، و در حالت رگرسیون از یک تابع خطی ساده استفاده شده است.

بردار ورودی به مدل به ازای هر روز، شامل خروجی جاسازی به ازای هر سند (خبر) و مقادیر متغیرهای بدست آمده از مرحله‌ی مهندسی ویژگی است. با قراردادن بردار حاصل از جاسازی متنی، به ازای هر سند یا خبر در روز t برداری به صورت $[e_1^t, \dots, e_K^t | v_1^t, \dots, v_K^t]$ خواهیم داشت که در آن e_i^t مقدار هر درایه از بردار جاسازی به طول K ، و v_i بیانگر مقدار نرمال شده‌ی یک متغیر عددی از مجموعه‌ی R متغیر عددی است. این بردار را به صورت D_p^t نمایش داده و آن را به صورت بردار داده برای خبر p در روز t تفسیر می‌کنیم. با این کدینگ از مسئله، برای هر روز، به تعداد اخبار موجود از منابع مختلف، نمونه خواهیم داشت. برای ساخت داده‌های مربوط به روز $t - 1$ ، متغیرهای عددی را عیناً استفاده کرده و برای بخش بردار جاسازی، میانگین جاسازی‌های تمامی اخبار منتشره در $t - 1$ را قرار می‌دهیم.

۲-۵-۴- پیکربندی مدل Bi-GRU

منظور از پیکربندی^۳، جزئیات مربوط به معماری و هایپرپارامترهای مدل است. مدل مورد استفاده از خانواده‌ی Bi-GRU یا اصطلاحاً GRU دوطرفه است. مزیت این مدل نسبت به مدل‌های یک‌طرفه، مدل‌سازی بهتر روابط داده‌های متوالی است. یکی از مهم‌ترین مشخصات در پیکربندی هر مدل، اندازه ورودی است که متناسب با اندازه‌ی بردارهای داده در مجموعه‌های داده است. در مدل اول، به دلیل آنکه ورودی متنی به مدل داده نمی‌شود، اندازه‌ی بردار ورودی صرفاً شامل تمامی داده‌های عددی موجود از متغیرهای انتخاب شده است. تعداد این متغیرها ۲۱ عدد می‌باشد که ۴ مورد از آنها مربوط به داده‌های روزانه‌ی نرخ ارز، ۷ مورد مربوط به داده‌های منتج شده و نگاشت شده از شاخص‌های اقتصادی، و ده مورد بیانگر مرجع خبر است.

در مدل دوم که داده‌های متنی نیز در کنار داده‌های عددی مورد استفاده قرار می‌گیرند، به ازای هر داده‌ی متنی یک بازنمایی برداری با ابعاد ۷۶۸ خواهیم داشت. این بردار در کنار ۲۱ متغیر دیگر، بردار ورودی به طول ۷۸۹ مؤلفه را ایجاد می‌کند.

در معماری مدل ۱ تا ۳ لایه از واحدهای GRU به صورت چیده شده روی هم در نظر گرفته شده است. همچنین در هر لایه، ۳۲ تا ۱۲۸ واحد پردازشی GRU در نظر گرفته شده است.

به منظور جلوگیری از بیش‌براش مدل، از تکنیک دوراندازی استفاده شده و پارامتر دوراندازی به اندازه ۰.۲ تا ۰.۵ وزن‌ها در نظر گرفته شده است.

1 SoftMax
2 Binary
3 Configuration

به دنبال پشته‌ای از لایه‌های GRU، ۱ تا ۲ لایه‌ی تماماً متصل قرار داده شده که در هر لایه ۶۴ الی ۲۵۶ نورون در نظر گرفته شده است. همچنین تابع فعال‌سازی برای لایه‌های پنهان تابع ReLU، برای لایه‌ی خروجی در حالت رگرسیون تابع خطی، و برای لایه‌ی خروجی در حالت کلاس‌بند تابع سیگموید لحاظ شده است. به منظور آموزش مدل، اندازه Batch متناسب با GPU و محدودیت‌های سرعت و حافظه، ۳۲ یا ۶۴ انتخاب شده و نرخ یادگیری در بازه‌ی $[1e-3, 1e-4]$ تنظیم شده است. برای آموزش کامل مدل حداقل ۵۰ و حداکثر حدود ۱۰۰ اپیاک لازم است. عدد دقیق براساس وضعیت همگرایی مدل مشخص می‌شود. برای بهینه‌سازی تابع هزینه در مدل، از بهینه‌ساز Adam با پارامترهای $\beta_1 = 0.9$ و $\beta_2 = 0.999$ استفاده شده است. تابع خطا برای حالت رگرسیون تابع میانگین مربعات خطا و برای حالت کلاس‌بندی، تابع کراس‌آنتروپی محاسبه شده است. به منظور جلوگیری از بیش‌برازش مدل، علاوه بر تکنیک دوراندازی ذکر شده، از رگولاریزیشن L2 (کاهش وزن‌ها) در لایه‌های تماماً متصل خروجی استفاده شده است.

۳-۵-۴- کدینگ مسئله‌ی کلاس‌بندی

منظور از کدینگ در اینجا، تبدیل مقادیر پیوسته نرخ ارز در خروجی به مقادیر گسسته است. این رویکرد اهمیت ویژه‌ای دارد زیرا در بسیاری از موارد و کاربردهای واقعی، میزان تغییرات نرخ ارز مهم‌تر از مقدار دقیق آن پس از تغییر است. بر همین اساس، هفت کلاس برای تبدیل مقادیر پیوسته نرخ ارز به مقادیر گسسته در خروجی تعریف شده‌اند.

- مقدار تغییرات کمتر از ۰.۵٪ در جهت مثبت یا منفی (یک کلاس - معادل با بدون تغییر)
- تغییرات بین ۰.۵٪ تا ۱٪ هم در جهت مثبت و هم در جهت منفی (دو کلاس)
- تغییرات بین ۱٪ تا ۳٪ هم در جهت مثبت و هم در جهت منفی (دو کلاس)
- تغییرات بیش از ۳٪ هم در جهت مثبت و هم در جهت منفی (دو کلاس)

لذا به ازای هر نمونه‌ی ورودی، یکی از ۷ کلاس فوق به نمونه نسبت داده می‌شود. همانطور که در شکل ۵ نشان داده شده است، برای نمایش برچسب هر نمونه از کدینگ وان-هات استفاده شده است. لذا هفت ستون به ازای هر نمونه داریم که مقادیر ستون متناظر با کلاسی که نمونه به آن تعلق دارد برابر با ۱ بوده، و مابقی مقادیر صفر هستند.

۴-۵-۴- مدل زبانی بزرگ

در این حالت، از مدل زبانی بزرگ GPT-4o شرکت OpenAI که با داده‌های آموزشی تنظیم دقیق شده، استفاده شده است. داده‌ها از طریق API به سرورهای OpenAI ارسال شده و تنظیم دقیق انجام شده است. دسترسی به مدل تنها از طریق واسطه‌های تعیین‌شده ممکن است. پس از تنظیم دقیق، با استفاده از فرمان‌های مشخص، مدل خروجی مورد نظر را تولید می‌کند. نمونه‌ای از این فرمان در جدول ۱ نشان داده شده است.

جدول ۱ - قالب فرمان استفاده شده برای تنظیم دقیق مدل مبتنی بر LLM

###Instruction: Given the following economic indicators and related news articles, predict the next day's percentage change in the USD/IRR exchange rate. ###Economic Indicators: -GDP Growth: {gdp_growth} % -Inflation Rate: {inflation_rate} % -Unemployment Rate: {unemployment_rate} % -Trade Balance: {trade_balance} billion USD -Public Debt: {public_debt}% of GDP -Oil Price: {oil_price} USD per barrel -Foreign Exchange Reserves: {fx_reserves} billion USD -Money Supply Growth (M2): {money_supply_growth}% ###News Articles: "{news_article_1}" ... ###Target: The percentage change in the USD/IRR exchange rate for the next day is {target}%.
--

منبع: یافته های پژوهشگر.

۴-۶- عملکرد مدل ها در پیش بینی نرخ ارز

به منظور ارزیابی درست و استاندارد، از مجموعه داده‌ی آزمون^۱ استفاده شده است. این داده‌ها پیش‌تر توسط هیچ یک از مدل‌ها عیناً دیده نشده‌اند. البته در فرآیند آموزش، داده‌های مشابه با آنها در آموزش مدل‌ها استفاده شده است. همانطور که پیش‌تر نیز ذکر گردید، حدود ۱۰٪ از داده‌های موجود در مجموعه‌ی داده به عنوان داده‌ی آزمون در نظر گرفته شده‌اند. این داده‌ها به صورت تصادفی و به نحوی انتخاب شده‌اند که کل بازه‌ی زمانی موجود در مجموعه‌ی داده را پوشش دهند. این مسئله از آن جهت حائز اهمیت است که خروجی گزارش شده از عملکرد مدل، سوگیری^۲ روی بخش خاصی از داده‌ها نداشته باشد و بتوان نتایج حاصله از مدل را به درستی تعمیم داد.

1 Test Set

2 Bias

متأسفانه مقایسه‌ی مدل‌های پیشنهادی با سایر روش‌های پیشنهاد شده در مقالات میسر نیست. دلایل اصلی این موضوع عبارتند از:

- (۱) مجموعه داده‌های مورد استفاده، مکانیزم‌های نرمال‌سازی و پالایش داده‌ها یکسان نیست.
 - (۲) مجموعه‌ی داده‌ها و الگوریتم‌های پیاده‌سازی شده‌ی سایر مقالات در دسترس نیست.
 - (۳) پارامترها و هایپرپارامترهای مربوط به هر مدل، عموماً در مقالات ذکر نشده‌اند.
- با توجه به موارد فوق، در این پژوهش تمامی مدل‌های مدرن و مرتبط پیاده‌سازی شده و نتایج آنها گزارش شده است.
- به منظور جلوگیری از سوگیری مدل روی بخش خاصی از داده، از روش اعتبارسنجی متقابل^۱ استفاده کرده‌ایم. در این روش، مجموعه‌های آموزش و آزمون به تعداد مشخص (مثلاً ۵ بار) به صورت کاملاً تصادفی با توزیع یکنواخت انتخاب می‌شوند و کل فرآیند آموزش و آزمون به این تعداد از ابتدا انجام می‌شود. نتیجه‌ی گزارش‌شده‌ی نهایی، میانگین تمام نتایج بدست آمده در همه‌ی حالت‌هاست.

۱-۶-۴- معیارهای ارزیابی

در صورتی که پیش‌بینی نرخ ارز به عنوان یک مسئله‌ی رگرسیون مطرح باشد، از خطای میانگین مربعات (MSE^۲) و همچنین خطای میانگین مطلق (MAE^۳) استفاده شده است. این روابط به صورت زیر هستند:

$$MAE = \frac{\sum_{i=1}^n |y_i - \hat{y}_i|}{n} \quad (۶)$$

$$MSE = \frac{1}{n} \sum_{i=1}^n (y_i - \hat{y}_i)^2 \quad (۷)$$

در صورتی که مسئله را به عنوان یک مسئله‌ی کلاس‌بندی تفسیر کنیم، عموماً از معیار ارزیابی صحت برای ارزیابی کارایی مدل‌های پیش‌بینی نرخ ارز استفاده شده است. تعریف این معیار در رابطه‌ی (۸) نمایش داده شده است.

$$Accuracy = \frac{TP + TN}{TP + FP + TN + FN} \quad (۸)$$

تفسیر پارامترهای فوق براساس ماتریس درهمی^۴ است که در نمایش داده شده است.

1 Cross-Validation
2 Mean Square Error
3 Mean Absolute Error
4 Confusion Matrix

		برچسب حقیقی	
		مثبت (P)	منفی (N)
پیش‌بینی مدل	مثبت (P)	به درستی مثبت TP	به اشتباه مثبت FP
	منفی (N)	به اشتباه منفی FN	به درستی منفی TN

شکل ۱۲ - ماتریس درهمی

منبع: یافته های پژوهشگر.

لازم به ذکر است که با وجود آنکه در **Error! Reference source not found.** ماتریس درهمی برای یک مسئله‌ی دو کلاسه تشریح شده است، اما این مفهوم قابل تعمیم به مسائل چند کلاسه است. علاوه بر این معیار، معیار استاندارد دیگر مورد استفاده، امتیاز $F1^1$ است. این معیار در واقع دو معیار دیگر مطرح در مسائل کلاسیک یعنی دقت^۲ و فراخوان یا یادآوری^۳ را با یکدیگر ترکیب می‌کند. معیارهای دقت و یادآوری براساس جدول درهمی به صورت زیر تعریف می‌شوند:

$$Precision = \frac{TP}{TP + FP} \quad Recall = \frac{TP}{TP + FN} \quad (9)$$

معیار دقت نشان می‌دهد که از میان تمام نمونه‌هایی که مدل به عنوان مثبت پیش‌بینی کرده، چه تعداد واقعاً مثبت بوده‌اند. معیار یادآوری نیز مشخص می‌کند که از میان تمامی نمونه‌های مثبت، مدل چه تعداد را به درستی شناسایی کرده است. این دو معیار معمولاً رابطه معکوس دارند؛ به طوری که بهبود یکی اغلب منجر به کاهش دیگری می‌شود. برای تراز کردن این دو معیار، از متریک امتیاز $F1$ استفاده می‌شود که میانگینی هندسی از دقت و یادآوری است و به صورت زیر محاسبه می‌شود.

$$F1 = 2 \times \frac{Precision \times Recall}{Precision + Recall}$$

- 1 F1 Score
2 Precision
3 Recall

برای گزارش عملکرد یک مدل، استفاده از هر دو متریک صحت و امتیاز F1 ضروری است. این موضوع به توزیع نامتوازن کلاس‌ها در مجموعه‌های داده بازمی‌گردد. در داده‌هایی با ناترازی کلاس‌ها، تکیه صرف بر متریک صحت می‌تواند گمراه‌کننده باشد. برای مثال، در تشخیص بیماری، بیشتر نمونه‌ها ممکن است سالم باشند و اگر مدل همه را سالم پیش‌بینی کند، صحت بالایی (مثلاً ۹۷٪) به دست می‌آید، اما عملکرد واقعی مدل ضعیف خواهد بود. بنابراین، امتیاز F1 که عدم توازن کلاس‌ها را نیز در نظر می‌گیرد، باید گزارش شود.

۲-۶-۴- ارزیابی مدل‌ها در پیش‌بینی نرخ ارز

به منظور بررسی فرضیه‌ی تأثیرگذاری اخبار در پیش‌بینی نرخ ارز، مدل‌ها در دو حالت بدون استفاده از اخبار، و با استفاده از اخبار آموزش داده‌ایم. نتایج بدست آمده برای حالت رگرسیون در جدول ۲، و برای حالت کلاس‌بندی در جدول ۳ نمایش داده شده است.

جدول ۲ - نتایج عملکرد مدل عنوان رگرسیون در پیش‌بینی نرخ ارز

	بدون استفاده از داده‌ی متنی		با استفاده از داده‌ی متنی	
	MSE	MAE	MSE	MAE
Bi-GRU	0.032	0.112	0.014	0.078
ARIMA	0.050	0.135	---	---
Prophet	0.045	0.135	---	---
Fine-tuned LLM	0.027	0.104	0.012	0.070

منبع: یافته‌های پژوهشگر.

جدول ۳ - نتایج عملکرد مدل‌ها به عنوان کلاس‌بند در پیش‌بینی نرخ ارز

	بدون استفاده از داده‌ی متنی		با استفاده از داده‌ی متنی	
	Accuracy	F1-Score	Accuracy	F1-Score
Bi-GRU	0.78	0.76	0.86	0.84
ARIMA	0.63	0.61	---	---
Prophet	0.66	0.63	---	---
Fine-tuned LLM	0.82	0.80	0.88	0.87

منبع: یافته‌های پژوهشگر.

همانطور که مشاهده می‌شود در نظر گرفتن اخبار تأثیر قابل ملاحظه‌ای در بهبود عملکرد پیش‌بینی نرخ ارز دارد. در صدر جدول مدل زبانی بزرگ تنظیم دقیق شده قرار دارد. علی‌رغم داشتن بهترین نتیجه، این مدل مشکلاتی نیز دارد؛ در ساده‌ترین حالت عملکرد این مدل مشابه یک جعبه‌ی سیاه است که ما از داخل آن و نحوه‌ی عملکرد آن، اطلاع دقیقی نداریم. این مسئله تفسیرپذیری خروجی مدل را کم می‌کند. به عبارت دیگر، در صورتی که مدل پیش‌بینی‌های خوبی داشته باشد، همه‌چیز به خوبی پیش می‌رود. اما وقتی پیش‌بینی‌ها با کاهش دقت مواجه

شوند، نمی‌توان به طور دقیق مدل را عیب‌یابی کرده و دلایل افت عملکرد را متوجه شد. این مسئله در مدل‌های دیگر که پیاده‌سازی و کدمنبع آنها در دسترس است، وجود ندارد. مسئله‌ی دیگر، حفاظت و صیانت از داده‌هاست. در بسیاری از کاربردها، داده‌های مورد استفاده دارای ارزش سازمانی و ارگانی بوده و ارسال آنها به یک مرجع خارجی، خلاف سیاست‌های آن سازمان است. لذا با وجود عملکرد بهتر مدل مبتنی بر LLM، اما در استفاده از آن باید ملاحظات سازمانی و اهداف بلندمدت را لحاظ کرد.

در مقاوم دوم، مدل Bi-GRU قرار دارد. همچنین استفاده از داده‌ی متنی در تمامی سناریوها منجر به بهبود کارایی مدل و کاهش خطا شده است (در مدل‌های آماری چون ARIMA و Prophet امکان استفاده از داده‌ی متنی به صورت مستقیم وجود نداشت). نکته‌ی حائز اهمیت در اینجا در رابطه با روش‌های آماری ARIMA و Prophet است. این روش‌ها ذاتاً برای کلاس‌بندی طراحی نشده‌اند. برای داشتن مقایسه‌ای از کارایی این روش‌ها، نتایج بدست آمده در آن‌ها را مطابق با کدینگ برچسب که پیش‌تر توضیح داده شد، تبدیل به خروجی کلاس‌بند شده‌اند.

۵- آزمون‌ها و نتایج

ایده‌ی این پژوهش ایجاد یک روند سیستماتیک به منظور تحلیل پدیده‌های کیفی (چون اخبار منتشره) بر پدیده‌های کمی (مانند نرخ ارز) است. از آنجا که مسئله‌ی نرخ ارز و بسیاری از مسائل دیگر، مسائلی چندعامله بوده و در بعضی از مراجع در زمره‌ی سیستم‌های پویا^۱ طبقه‌بندی می‌شوند، نیاز به رویکردی داریم که بتوان در آن تمامی عوامل دخیل در یک پدیده را تا حد امکان مشارکت داد. برای این منظور مدل پیش‌بینی کننده‌ی نرخ ارز که از انواع ورودی داده پشتیبانی می‌کند در بخش قبل معرفی شد. با توجه به نتایج گزارش شده به وضوح مشخص شد که اخبار منتشره در کنار سایر عوامل اقتصادی بر نرخ ارز تأثیرگذار است.

جدول ۴ - نمادهای منابع خبری

نماد	منبع خبری	نماد	منبع خبری
A	صنعت-معدن-تجارت	F	دنای اقتصاد
B	ایران اقتصادی	G	بانک مرکزی
C	ایسنا	H	بی‌بی‌سی فارسی
D	تابناک	I	صدای آمریکا فارسی
E	آسیا	J	خبرآنلاین

منبع: یافته‌های پژوهشگر.

سؤال دیگری که در اینجا مطرح است آن است که منابع خبری به چه میزان در نرخ ارز مؤثرند؟ به طور دقیق‌تر، می‌خواهیم بررسی کنیم که آیا وجود خبری از یک منبع مشخص، یا ترکیبات مختلفی از منابع، به طور معناداری در

خروجی مدل مؤثر است، یا خیر. این مسئله از آن جهت حائز اهمیت است که ممکن است وزن یک منبع خبر به نحوی بالا باشد که بر رفتار بازار غلبه داشته باشد. به منظور انجام تحلیل حساسیت برای منابع خبری، سه رویکرد متداول است که در ادامه ذکر شده است. نگاشت منابع خبری به نمادها در جداول به صورت جدول ۴ است. رویکرد اول، رویکرد مطالعه فرسایشی^۱ است. در این روش مدل را به ازای تمامی منابع خبری آموزش می‌دهیم و آن را به عنوان مدل پایه در نظر می‌گیریم. سپس یکی از منابع را از مجموعه‌ی آموزش حذف می‌کنیم (مابقی ویژگی‌ها و منابع خبری را نگه می‌داریم) و مدل جدیدی آموزش می‌دهیم. مدل جدید را براساس شاخص‌های موجود (صحت، امتیاز F1، خطای MAE یا MSE) را با مدل پایه مقایسه می‌کنیم. این مقایسه نشان می‌دهد که حذف هریک از منابع خبری به چه میزان می‌تواند در خروجی مؤثر باشد. در جدول ۵، نتایج این رویکرد نمایش داده شده است.

جدول ۵ - تغییرات خروجی مدل در صورت حذف منابع خبری

منبع خبری محدوف	MSE	Δ_{MSE}	MEA	Δ_{MAE}	Accuracy	$\Delta_{Accuracy}$	F1-Score	Δ_{F1}
---	0.014	---	0.078	---	0.86	---	0.84	---
A	0.016	+0.002	0.084	+0.006	0.84	-0.02	0.82	-0.02
B	0.015	+0.001	0.080	+0.002	0.85	-0.01	0.83	-0.01
C	0.014	≈ 0	0.078	≈ 0	0.86	≈ 0	0.84	0
D	0.014	≈ 0	0.078	≈ 0	0.86	≈ 0	0.84	0
E	0.014	≈ 0	0.078	≈ 0	0.86	≈ 0	0.84	0
F	0.018	+0.004	0.088	+0.010	0.83	-0.03	0.80	-0.04
G	0.016	+0.002	0.082	+0.004	0.84	-0.02	0.82	-0.02
H	0.015	+0.001	0.080	+0.002	0.85	-0.01	0.83	-0.01
I	0.014	≈ 0	0.078	≈ 0	0.86	≈ 0	0.84	≈ 0
J	0.015	+0.001	0.079	+0.001	0.85	-0.01	0.83	-0.01

منبع: یافته‌های پژوهشگر.

مطابق با جدول ۵، حذف منبع خبر F بیشترین تأثیر منفی را در خروجی مدل داشته است؛ به طوری که منجر به افزایش MSE به اندازه ۰.۰۰۴ و MAE به اندازه‌ی ۰.۰۱۰ شده، در حالیکه نرخ صحت کلاس‌بندی به اندازه‌ی ۳٪ و امتیاز F1 به اندازه‌ی ۰.۰۴ کاهش داشته‌اند. حذف منابع خبری C، D و E تأثیر چندانی روی عملکرد مدل نداشته و حذف منابع دیگر تأثیرات اندکی روی خروجی مدل داشته است. تأثیرگذاری بیشتر برخی منابع خبری را می‌توان به طرق مختلفی تفسیر کرد:

1 Ablation Study

- ممکن است در طول زمان و به دلیل سابقه‌ی یک منبع خبری، میزان اعتماد به صحت تحلیل‌ها یا اخبار منتشره از آن منبع کاهش یا افزایش یافته باشد.
- ممکن است برخی از منابع ذاتاً خبرهای بامحتواتری منتشر کنند. به عنوان مثال منابعی که صرفاً خبری را بدون تحلیل یا بازنویسی آن، بازنشر می‌کنند، عموماً تأثیر کمی در خروجی مدل داشته‌اند.

رویکرد دوم استفاده از رویکردهای وزن‌دهی به ویژگی‌ها به منظور تفسیر اثرگذاری آنهاست. برای این منظور روش‌های مختلفی چون ¹SHAP، DeepLIFT، GradientSHAP، KernelSHAP، گرادیان‌های مجتمع² و پس‌انتشار هدایت شده³ وجود دارد {Lundberg, 2017 #179}. برخی از این روش‌ها مانند SHAP سربار محاسباتی بالایی داشته و برای مدل‌های ساده‌تر مناسب هستند. روش‌های دیگر چون گرادیان مجتمع از نظر محاسباتی کاراتر بوده و برای شبکه‌های عصبی بهینه شده‌اند. در اینجا از این روش استفاده شده است. نحوه‌ی محاسبه میزان مشارکت هر ویژگی در خروجی در رابطه‌ی زیر نمایش داده شده است:

$$IG_i(x) = (x_i - x'_i) \times \int_{\alpha=0}^1 \frac{\partial F(x' + \alpha(x - x'))}{\partial x_i} d\alpha \quad (9)$$

در این رابطه، x ، مقادیر ورودی‌های حقیقی، x' ، مقادیر ورودی‌های پایه (مثلاً همه‌ی مقادیر صفر)، $F(x)$ ، تابع پیش‌بینی مدل (تولید خروجی مدل براساس ورودی)، $\frac{\partial F}{\partial x_i}$ گرادیان (درجه‌ی تغییر) خروجی مدل نسبت به ویژگی ورودی x_i ، و α یک فاکتور برای مقیاس‌دهی⁴ بین ۰ و ۱ است.

با اعمال این روش روی مدل Bi-GRU ای که داریم، می‌توانیم تخمینی از مشارکت⁵ هر ویژگی (شامل منابع خبری) در پیش‌بینی‌های مدل داشته باشیم. برای آنکه تأثیر هر منبع خبری را مشاهده کرد، کدینگ منابع مختلف خبری به صورت وان-هات انجام شده است.

نتایج این بررسی در جدول ۶ نمایش داده شده است. براساس مقادیر حاصل از روش IG می‌توان متوجه شد که منبع خبری F بالاترین مشارکت در خروجی مدل را داشته است. این مشاهده با تأثیر حذف این منبع و بیشترین تغییرات در متریک‌های ارزیابی مطابقت دارد. لذا می‌توان گفت منابع C، D، E و I کمترین مشارکت در خروجی مدل را داشته و منابع خبری F و A مهم‌ترین مشارکت‌کنندگان در خروجی هستند.

1 Shapley Additive Explanations

2 Integrated Gradients

3 Guided Backpropagation

4 Scaling

5 Contribution

جدول ۶ - میزان مشارکت منابع خبری در پیش‌بینی نرخ ارز براساس روش گرادیان مجتمع (IG)

مشارکت براساس IG	عنوان منبع خبر	کد منبع
0.15	صنعت-معدن-تجارت	A
0.10	ابرار اقتصادی	B
0.03	ایسنا	C
0.02	تابناک	D
0.01	آسیا	E
0.18	دنیای اقتصاد	F
0.09	بانک مرکزی	G
0.07	بی‌بی‌سی فارسی	H
0.04	صدای آمریکا فارسی	I
0.06	خبرآنلاین	J

منبع: یافته‌های پژوهشگر.

رویکرد سوم آن است که به صورت تصادفی فیلد مربوط به منبع اخبار را تغییر دهیم و مجدداً مدل را آموزش دهیم. به این روش اختلاط^۱ داده‌ها نیز گفته می‌شود. به عبارت دیگر، محتوای خبر همان است، اما منبع خبر به صورت تصادفی تغییر داده شده است.

برای بررسی این ایده، سه حالت جابجایی عناصر را در نظر گرفته‌ایم. در حالت اول و براساس مشاهدات قبلی، منابع خبری A، B و F را با یکدیگر جابجا کرده‌ایم. دلیل این انتخاب آن است که در مشاهدات قبلی، این منابع خبری تأثیر بیشتری داشته‌اند. در حالت دوم، منابع C، D، E و I را مخلوط کرده‌ایم زیرا تأثیرات آنها نیز بر خروجی مدل مشابه بوده است. و در نهایت، تمامی منابع را با یکدیگر به صورت تصادفی جابجا کرده‌ایم. نتایج این رویکرد در جدول ۷ نمایش داده شده است.

با توجه به نتایج بدست آمده می‌توان گفت به هم زدن منابع خبری A، B و F منجر به افت عملکرد متوسطی در حدود افزایش MSE تا ۰.۰۰۳+، کاهش دقت به میزان ۳٪ و کاهش F1 به اندازه‌ی ۰.۰۳- می‌شود. به هم ریختن منابع خبری C، D، E و I هیچ تغییر قابل توجهی ایجاد نمی‌کند؛ این مسئله نشان دهنده‌ی سهم کم آنها در عملکرد مدل است. مخلوط کردن همه منابع منجر به بزرگترین افت عملکرد می‌شود به طوری که MSE به اندازه‌ی ۰.۰۰۵+ افزایش، MAE به اندازه‌ی ۰.۰۱۲+ افزایش، و نرخ صحت به میزان ۵٪ کاهش می‌یابد! نتایج تأیید می‌کنند که خروجی و عملکرد مدل کاملاً به منابع خبری وابستگی دارد به طوری که حتی اگر خبری را عیناً با همان محتوا، خبرگزاری یا پایگاه دیگری منتشر کند، تأثیر متفاوتی در خروجی دارد.

جدول ۷ - نتایج سناریوهای جابجایی منابع خبری

سناریو	MSE	Δ_{MSE}	MEA	Δ_{MAE}	Accuracy	$\Delta_{Accuracy}$	F1-Score	Δ_{F1}
بدون جابه‌جایی (مدل پایه)	0.014	--	0.078	--	0.86	--	0.84	--
جابه‌جایی منابع A، B و F	0.017	+0.003	0.085	+0.007	0.83	-0.03	0.81	-0.03
جابه‌جایی منابع C، D، E، I	0.014	≈ 0	0.078	≈ 0	0.85	-0.01	0.83	-0.01
جابه‌جایی تمامی منابع	0.019	+0.005	0.090	+0.012	0.81	-0.05	0.79	-0.05

* منبع: یافته‌های پژوهشگر.

۶- جمع‌بندی

در این پژوهش با استفاده از مدل‌های یادگیری ماشین، روشی عملیاتی برای دخیل کردن اخبار منتشره در مدل‌سازی نرخ ارز در ایران ارائه گردید. رویکرد پیشنهادی براساس معماری و خط لوله‌ی استاندارد یادگیری ماشین از ورود داده‌ی خام تا ارزیابی خروجی نهایی مدل طراحی شده است. این روش شامل پنج مرحله‌ی پیش‌پردازش داده‌های متنی، عددی و دسته‌ای، بازنمایی برداری داده‌ها، مهندسی ویژگی‌ها، آموزش مدل‌ها و ارزیابی و تحلیل نهایی است. در مرحله‌ی بازنمایی داده‌های متنی، از یک جاسازی سفارشی‌شده برای داده‌های اقتصادی استفاده شده است. این ایده با رفع محدودیت‌های روش‌های قبلی، چالش کار با داده‌های متنی در مدل‌های عددی و اقتصادی را به خوبی برطرف می‌نماید.

با تکنیک تلفیق انواع داده‌ها، مسئله برای دو حالت کلاس‌بندی و رگرسیون حل شده است. به منظور رسیدن به بهترین مدل، انواع مختلفی از مدل‌های بروز یادگیری ماشین و مدل‌های کلاسیک چون ARIMA و Prophet پیاده‌سازی و آزموده شده‌اند. برای هر مدل دو نسخه آموزش داده شده که در یکی از آنها از داده‌ی خبری در کنار داده‌های عددی استفاده شده و در دیگری، صرفاً داده‌ی عددی مبنای مدل‌سازی نرخ ارز قرار گرفته است. نتایج بدست آمده کاملاً تأیید می‌کند که استفاده از داده‌های متنی در هر دو حالت کلاس‌بندی و رگرسیون تأثیر کاملاً مثبتی در بهبود کارایی مدل در پیش‌بینی و مدل‌سازی نرخ ارز داشته است.

در ادامه با سه روش اثر منبع منتشرکننده‌ی خبر بررسی شده است. براساس نتایج بدست آمده منبع منتشرکننده‌ی خبر نیز کاملاً در نرخ ارز مؤثر است؛ به نحوی که وقتی صرفاً کد منابع خبری (و نه محتوای اخبار) را جابجا کردیم، کارایی مدل با افت مواجه شد. این بدان معناست که حتی اگر خبری کاملاً مشابه توسط منابع مختلف منتشر شود، تأثیر متفاوتی بر نرخ ارز خواهد داشت. در ادامه به منظور کمی‌سازی تأثیر هر منبع خبری بر نرخ ارز، از رویکرد گرادیان یکپارچه (IG) برای تخمین میزان مشارکت هر منبع در خروجی استفاده شد. ضرایب بدست آمده در این بخش را می‌توان درجه‌ای از تأثیرگذاری یا اهمیت یک منبع خبری در نرخ ارز تفسیر نمود.

۷- کارهای آینده

معماری جامع طراحی شده امکان پیاده‌سازی و ارزیابی استاندارد انواع روش‌ها و ایده‌ها را فراهم ساخته است. به طور کلی اقدامات زیر برای بهبود و یا گسترش این پژوهش پیشنهاد می‌گردد:

- پیاده‌سازی و آزمون مدل‌های نسبتاً سنگین مبتنی بر معماری ترنسفورمر در صورت وجود منابع پردازشی کافی
- افزودن داده‌های منتشره در کانال‌ها، پلتفرم‌ها و شبکه‌های اجتماعی به مجموعه‌ی داده و غنی‌تر کردن مدل
- بررسی استفاده از مدل‌های زبانی ساین متوسط و حتی بزرگ که به صورت متن باز در دسترس هستند.
- استفاده از جاسازی‌های پیشرفته‌تر در خط لوله‌ی پردازش و تنظیم دقیق آنها برای زبان فارسی
- توسعه‌ی مدل پیشنهادی به منظور استفاده در سایر بازارها، مانند بازارهای سهام
- تهیه‌ی و انتشار مجموعه داده‌های استاندارد به منظور تسهیل فرآیند تحقیق و مدل‌سازی توسط سایر محققین
- ایجاد تابلوهای گزارش کارایی در پلتفرم‌های استاندارد به منظور گزارش نتایج عملکرد مدل‌های مختلف به صورت مجتمع با امکان مقایسه بین آنها

فهرست منابع

- ابرار-اقتصادی. (۲۰۲۳). روزنامه ابرار اقتصادی. Retrieved 20 July from <https://www.abrarnews.com>.
- ابونوری، ا.، خانعلی‌پور، ا. & عباسی، ج. (۲۰۰۹). اثر اخبار بر نوسانات نرخ ارز در ایران: کاربردی از خانواده ARCH. پژوهشنامه بازرگانی، ۵۰، ۱۲۰-۱۰۱.
- آسیا. (۲۰۲۳). روزنامه‌ی آسیا. Retrieved 20 July from <https://www.asianews.ir>.
- ایسنا. (۲۰۲۳). بخش سیاسی-اقتصادی پایگاه خبری ایسنا. Retrieved 20 July from <https://www.isna.ir/service/Economy>.
- آقامیری، س.، دامن کشیده، م. & هادی نژاد، م. (۲۰۲۱). بررسی عوامل موثر بر نرخ ارز در ایران از دیدگاه پست کینزین‌ها، فصلنامه اقتصاد مالی، ۵۴، ۲۰۹-۲۳۸.
- بیات، ن. (۲۰۱۸). پیش بینی نرخ ارز با استفاده از نقشه‌های خودسازمان ده بازگشتی. اقتصاد و تجارت نوین، ۱۳، ۸۴-۵۵.
- بی‌بی‌سی-فارسی. (۲۰۲۳). بخش اقتصادی-سیاسی بی‌بی‌سی فارسی. Retrieved 20 July from <https://www.bbc.com/persian/topics/cl819mvlllqt>.
- تابناک. (۲۰۲۳). بخش سیاسی-اقتصادی پایگاه خبری تابناک. Retrieved 20 July from <https://www.tabnak.ir/fa/economic>.

- تقوی، م. & خدام، م. (۲۰۱۱). بررسی تطبیقی کارآمدی نظریه های ارزی در پیش بینی تغییرات نرخ ارز در بازار تبادلات بین المللی ارز. دانش مالی تحلیل اوراق بهادار (مطالعات مالی)، ۹، ۱۹۲-۱۴۷.
- حسین زاده یوسف آباد، س. حقیقت، ع. (۲۰۱۳). اثر سیاست پولی بر نرخ ارز در ایران با استفاده از الگوی خودهمبسته با وقفه توزیع شده (ARDL)، فصلنامه اقتصاد مالی، ۲۵، ۱۲۳-۱۴۶.
- حسن، خ. & احمد، م. (۲۰۱۲). مدل سازی و پیش بینی نرخ ارز بر اساس معادلات دیفرانسیل تصادفی. تحقیقات اقتصادی، ۱۰۰ (۴۷)، ۱۲۹-۱۴۴.
- خاشعی، م. بیجاری، م. مخاطب رفیعی، ف. & فریمه، (۲۰۱۲). پیش بینی نرخ ارز با بکارگیری مدل های ترکیبی پرسپترون های چندلایه (MLPs) و طبقه بندی کننده های عصبی احتمالی (PNNs). نشریه علمی پژوهشی روش های عددی در مهندسی، ۳۲ (۱)، ۱-۱۴.
- خبرآنلاین. (۲۰۲۳). بخش سیاسی-اقتصادی پایگاه خبری ایسنا. Retrieved 20 July from <https://www.khabaronline.ir/service/Economy>
- دمیری، م. سعیدی، پ. دیده خانی، ح. & عباسی، ا. (۲۰۲۰). مدل سازی نوسانات نرخ ارز با رویکرد پویایی های سیستم. مهندسی مالی و مدیریت اوراق بهادار، ۱۱ (۴۳)، ۲۲۰-۲۴۴.
- دنیای-اقتصاد. (۲۰۲۳). روزنامه ی دنیای اقتصاد. Retrieved 20 July from <https://donya-e-eqtasad.com>
- زاهدی، ی. رضایی، ن. & نجاری، و. (۲۰۲۳). حباب قیمتی و تاثیر متغیرهای اقتصادی بر نرخ ارز در بازارهای مالی ایران با استفاده از روش های ARIMA و TAR، فصلنامه اقتصاد مالی، ۶۴، ۲۲۳-۲۴۵.
- رحیمی بروجردی، ع. (۲۰۰۰). نظام مطلوب ارزی و تنظیم و پیش بینی نرخ ارز برای اقتصاد ایران. پژوهش های اقتصادی ایران، ۵، ۴۰-۴۵.
- زرانژاد، م. فقه مجیدی، ع. & رضایی، ر. ا. (۲۰۰۸). پیش بینی نرخ ارز با استفاده از شبکه های عصبی مصنوعی و مدل ARIMA. اقتصاد مقداری (بررسی های اقتصادی)، ۱۹، ۱۳۰-۱۰۷.
- شریف مقدم، ش. & هاشمی، س. ا. (۲۰۱۸). پیش بینی نرخ ارز یورو به دلار با تکنیک شبکه عصبی مصنوعی. مهندسی مالی و مدیریت اوراق بهادار، ۹ (۳۷)، ۳۹۹-۴۱۳.
- شیرازی، همایون، & نصراللهی. (۲۰۱۴). مدل های پولی و پیش بینی نرخ ارز در ایران: از تئوری تا شواهد تجربی. فصلنامه سیاست های مالی و اقتصادی، ۱ (۴)، ۵-۲۴.
- صدای-امریکا-فارسی. (۲۰۲۳). بخش اقتصادی-سیاسی صدای امریکای فارسی. Retrieved 20 July from <https://ir.voanews.com/z/1066>
- صنعت-معدن-تجارت. (۲۰۲۳). روزنامه ی صنعت-معدن-تجارت. Retrieved 20 July from <https://www.smtnews.ir>
- طیبی، س. موحدنیا، ن. & کاظمینی، م. (۲۰۰۸). به کارگیری شبکه های عصبی مصنوعی در پیش بینی متغیرهای اقتصادی و مقایسه آن با روش های اقتصادسنجی: پیش بینی روند نرخ ارز در ایران. مهندسی صنایع و مدیریت (شریف ویژه علوم مهندسی)، ۴۳، ۱۰۴-۱۹۹.

- فتاحی، ش.، احمدی، آ.، & اکرم میرزائی، ع. (۲۰۱۳). مقایسه دقت روش الگوریتم ژنتیک با روش های دیگر پیش بینی های نرخ ارز. مطالعات و سیاست های اقتصادی ۹۶(۹۶)، ۱۱۳-۱۳۶.
- فراهانی، گ.، عادلی، امیدعلی، & قربانی. (۲۰۲۰). تاثیر نااطمینانی سیاست های اقتصادی بر نوسانات نرخ ارز با استفاده از رویکرد مدل خودهمبسته با وقفه های توزیعی غیرخطی (NARDL). مدل سازی اقتصادسنجی، ۱۴۷-۱۷۱، (۴)۵.
- مرزبان، د. ح.، جواهری، ب.، بهنام، & اکبریان. (۲۰۰۵). یک مقایسه بین مدل های اقتصادسنجی ساختاری، سری زمانی و شبکه عصبی برای پیش بینی نرخ ارز. مجله تحقیقات اقتصادی، ۴۰(۲).
- منصورى گرگرى، ح.، & خداوىسى، ح. (۲۰۱۹). پیش بینی نرخ ارز: مقایسه الگوهای رشد لجستیک با الگوهای رقیب. اقتصاد و الگوسازی، ۱۰، ۱۴۱-۱۷۹.
- هاشمی دیزج، ع.، حاضری نیری، ه.، & پوروحدانی، ر. (۲۰۲۰). مقایسه ی عملکرد مدل های شبکه های عصبی مصنوعی برای پیش بینی نرخ ارز در ایران. دوفصلنامه علمی مطالعات و سیاست های اقتصادی، ۷(۲)، ۵۳-۸۰.
- یارمحمدی، م.، & محمودوند، ر. (۲۰۱۵). پیش بینی نرخ ارز با استفاده از روش تحلیل مجموعه ی مقادیر تکین. ۱۴۶-۱۳۷، ۱۸.
- Bojanowski, P., Grave, E., Joulin, A., & Mikolov, T. (2017). Enriching word vectors with subword information. Transactions of the association for computational linguistics, 5, 135-146.
- Bonbast. (۲۰۲۳). Live exchange rates in Iran's free market. Retrieved 20 July from <https://www.bonbast.com>
- Brown, T., Mann, B., Ryder, N., Subbiah, M., Kaplan, J. D., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., & Askell, A. (2020). Language models are few-shot learners. Advances in neural information processing systems, 33, 1877-1901.
- Chan, S. W., & Chong, M. W. (2017). Sentiment analysis in financial texts. Decision Support Systems, 94, 53-64.
- Danielsson, S., & Gramer, A. (2023). Predicting Forex Rates using Sentiment Analysis on Financial Articles.
- Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2018). Bert: Pre-training of deep bidirectional transformers for language understanding. arXiv preprint arXiv:1810.04805.
- Dridan, R., & Oepen, S. (2012). Tokenization: Returning to a long solved problem—a survey, contrastive experiment, recommendations, and toolkit—. Proceedings of the 50th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers),
- Engle, R. F., & Ng, V. K. (1993). Measuring and testing the impact of news on volatility. The journal of finance, 48(5), 1749-1778.
- Farahani, M., Gharachorloo, M., Farahani, M., & Manthouri, M. (2021). Parsbert: Transformer-based model for persian language understanding. Neural Processing Letters, 53, 3831-3847.
- Feuerriegel, S., & Prendinger, H. (2016). News-based trading strategies. Decision Support Systems, 90, 65-74.
- Galperti, S., & Cerigioni, F. (2023). Listing Specs: The Effect of Framing Attributes on Choice. Journal of the European Economic Association, jvad032.

- Hagenau, M., Liebmann, M., & Neumann, D. (2013). Automated news reading: Stock price prediction based on financial news using context-capturing features. *Decision Support Systems*, 55(3), 685-697 .
- Heaton, J. (2016). An empirical analysis of feature engineering for predictive modeling. *SoutheastCon 2016* ,
- Jee, W., & Hyun, M. (2023). “10,000 Available” or “10% Remaining”: The Impact of Scarcity Framing on Ticket Availability Perceptions in the Secondary Ticket Market. *Behavioral Sciences*, 13(4), 338 .
- Jiménez-Jiménez, F., & Rodero-Cosano, J. (2023). Conditioning competitive behaviour in experimental Bertrand markets through contextual frames. *Journal of Behavioral and Experimental Economics*, 103, 101987 .
- Josse, J., & Husson, F. (2012). Handling missing values in exploratory multivariate data analysis methods. *Journal de la société française de statistique*, 153(2), 79-99 .
- Kapoor, A., Gulli, A., Pal, S., & Chollet, F. (2022). Deep Learning with TensorFlow and Keras: Build and deploy supervised, unsupervised, deep, and reinforcement learning models. Packt Publishing Ltd .
- Khan, W., Daud, A., Khan, K., Muhammad, S., & Haq, R. (2023). Exploring the frontiers of deep learning and natural language processing: A comprehensive overview of key challenges and emerging trends. *Natural Language Processing Journal*, 100026 .
- Li, Y., & Yang, T. (2018). Word embedding for understanding natural language: a survey. *Guide to big data applications*, 83-104 .
- Liebmann, M., Orlov, A. G., & Neumann, D. (2016). The tone of financial news and the perceptions of stock and CDS traders. *International Review of Financial Analysis*, 46, 159-175 .
- Liu, J., Li, T., Xie, P., Du, S., Teng, F., & Yang, X. (2020). Urban big data fusion based on deep learning: An overview. *Information Fusion*, 53, 123-133 .
- Lopez-Lira, A., & Tang, Y. (2023). Can chatgpt forecast stock price movements? return predictability and large language models. *arXiv preprint arXiv:2304.07619* .
- Lunsford, K. G. (2020). Policy language and information effects in the early days of Federal Reserve forward guidance. *American Economic Review*, 110(9), 2899-2934 .
- McDonald, B. (2012). Loughran and mcdonald financial sentiment dictionary. In. *MEX*. صرافی ملی. (۲۰۲۳). Retrieved 20 July from <https://www.mex.co.ir/>
- Mikolov, T., Chen, K., Corrado, G., & Dean, J. (2013). Efficient estimation of word representations in vector space. *arXiv preprint arXiv:1301.3781* .
- Mohamadi, S., Mujtaba, G., Le, N., Doretto, G., & Adjero, D. A. (2023). ChatGPT in the Age of Generative AI and Large Language Models: A Concise Survey. *arXiv preprint arXiv:2307.04251* .
- Mohammed, A. H., & Ali, A. H. (2021). Survey of bert (bidirectional encoder representation transformer) types. *Journal of Physics: Conference Series* ,
- Morales, E. F., & Escalante, H. J. (2022). A brief introduction to supervised, unsupervised, and reinforcement learning. In *Biosignal processing and classification using computational learning and intelligence* (pp. 111-129). Elsevier .
- Mundell, R. (1963). M.(1963): Capital Mobility and Stabilization Policy Under Fixed and Flexible Exchange Rates. *Canadian Journal of Economics and Political Science*, 29(4), 472-485 .
- Nargesian, F., Samulowitz, H., Khurana, U., Khalil, E. B., & Turaga, D. S. (2017). Learning Feature Engineering for Classification. *Ijcai* ,
- Nassirtoussi, A. K., Aghabozorgi, S., Wah, T. Y., & Ngo, D. C. L. (2015). Text mining of news-headlines for FOREX market prediction: A Multi-layer Dimension Reduction Algorithm with semantics and sentiment. *Expert Systems with Applications*, 42(1), 306-324 .

- Nayak, A. S., Kanive, A. P., Chandavekar, N., & Balasubramani, R. (2016). Survey on pre-processing techniques for text mining. *International Journal of Engineering and Computer Science*, 5(6), 16875-16879 .
- Pennington, J., Socher, R., & Manning, C. D. (2014). Glove: Global vectors for word representation. *Proceedings of the 2014 conference on empirical methods in natural language processing (EMNLP)* ,
- Potdar, K., Pardawala, T. S., & Pai, C. D. (2017). A comparative study of categorical variable encoding techniques for neural network classifiers. *International journal of computer applications*, 175(4), 7-9 .
- Sarzynska-Wawer, J., Wawer, A., Pawlak, A., Szymanowska, J., Stefaniak, I., Jarkiewicz, M., & Okruszek, L. (2021). Detecting formal thought disorder by deep contextualized word representations. *Psychiatry Research*, 304, 114135 .
- Seabe, P. L., Moutsinga, C. R. B., & Pindza, E. (2023). Forecasting cryptocurrency prices using LSTM, GRU, and bi-directional LSTM: a deep learning approach. *Fractal and Fractional*, 7(2), 203 .
- Semiromi, H. N., Lessmann, S., & Peters, W. (2020). News will tell: Forecasting foreign exchange rates based on news story events in the economy calendar. *The North American Journal of Economics and Finance*, 52, 101181 .
- Shilpa, B., & Shambhavi, B. (2021). Combined deep learning classifiers for stock market prediction: integrating stock price and news sentiments. *Kybernetes*, 52(3), 748-773 .
- Shynkevich, Y., McGinnity, T. M., Coleman, S. A., & Belatreche, A. (2016). Forecasting movements of health-care stock prices based on different categories of news articles using multiple kernel learning. *Decision Support Systems*, 85, 74-83 .
- Singh, D., & Singh, B. (2020). Investigating the impact of data normalization on classification performance. *Applied Soft Computing*, 97, 105524 .
- Singh, J., & Gupta, V. (2016). Text stemming: Approaches, applications, and challenges. *ACM Computing Surveys (CSUR)*, 49(3), 1-46 .
- Sousa, M. G., Sakiyama, K., de Souza Rodrigues, L., Moraes, P. H., Fernandes, E. R., & Matsubara, E. T. (2019). BERT for stock market sentiment analysis. *2019 IEEE 31st international conference on tools with artificial intelligence (ICTAI)* ,
- Sun, A., Lachanski, M., & Fabozzi, F. J. (2016). Trade the tweet: Social media text mining and sparse matrix factorization for stock market prediction. *International Review of Financial Analysis*, 48, 272-281 .
- TGJU. نرخ ارز. (۲۰۲۳). Retrieved 20 July from <https://www.tgju.org>
- Van de Kauter, M., Breesch, D., & Hoste, V. (2015). Fine-grained analysis of explicit and implicit sentiment in financial news articles. *Expert Systems with Applications*, 42(11), 4999-5010 .
- Villamil, L., Bausback, R., Salman, S., Liu, T. L., Horn, C., & Liu, X. (2023). Improved stock price movement classification using news articles based on embeddings and label smoothing. *arXiv preprint arXiv:2301.10458* .

nal linguistics, 5, 135-146 .

Bonbast. (۲۰۲۳). Live exchange rates in Iran's free market. Retrieved 20 July from <https://www.bonbast.com>

Brown, T., Mann, B., Ryder, N., Subbiah, M., Kaplan, J. D., Dhariwal, P., Neelakantan, A., Shyam, P., Sastry, G., & Askell, A. (2020). Language models are few-shot learners. *Advances in neural information processing systems*, 33, 1877-1901 .

- Chan, S. W., & Chong, M. W. (2017). Sentiment analysis in financial texts. *Decision Support Systems*, 94, 53-64 .
- Chong, D., & Druckman, J. N. (2007). Framing theory. *Annu. Rev. Polit. Sci.*, 10, 103-126 .
- Danielsson, S., & Gramer, A. (2023). Predicting Forex Rates using Sentiment Analysis on Financial Articles .
- Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2018). Bert: Pre-training of deep bidirectional transformers for language understanding. *arXiv preprint arXiv:1810.04805* .
- Dharshing, S., Hille, S. L., & Wüstenhagen, R. (2017). The influence of political orientation on the strength and temporal persistence of policy framing effects. *Ecological Economics*, 142, 295-305 .
- Dridan, R., & Oepen, S. (2012). Tokenization: Returning to a long solved problem—a survey, contrastive experiment, recommendations, and toolkit—. *Proceedings of the 50th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)* ,
- Engle, R. F., & Ng, V. K. (1993). Measuring and testing the impact of news on volatility. *The journal of finance*, 48(5), 1749-1778 .
- Farahani, M., Gharachorloo, M., Farahani, M., & Manthouri, M. (2021). Parsbert: Transformer-based model for persian language understanding. *Neural Processing Letters*, 53, 3831-3847 .
- Feuerriegel, S., & Prendinger, H. (2016). News-based trading strategies. *Decision Support Systems*, 90, 65-74 .
- Galperti, S., & Cerigioni, F. (2023). Listing Specs: The Effect of Framing Attributes on Choice. *Journal of the European Economic Association*, jvad032 .
- Hagenau, M., Liebmann, M., & Neumann, D. (2013). Automated news reading: Stock price prediction based on financial news using context-capturing features. *Decision Support Systems*, 55(3), 685-697 .
- Heaton, J. (2016). An empirical analysis of feature engineering for predictive modeling. *SoutheastCon 2016* ,
- Jee, W., & Hyun, M. (2023). “10,000 Available” or “10% Remaining”: The Impact of Scarcity Framing on Ticket Availability Perceptions in the Secondary Ticket Market. *Behavioral Sciences*, 13(4), 338 .
- Jiménez-Jiménez, F., & Rodero-Cosano, J. (2023). Conditioning competitive behaviour in experimental Bertrand markets through contextual frames. *Journal of Behavioral and Experimental Economics*, 103, 101987 .
- Josse, J., & Husson, F. (2012). Handling missing values in exploratory multivariate data analysis methods. *Journal de la société française de statistique*, 153(2), 79-99 .
- Kahneman, D., & Tversky, A. (1979). Prospect theory. *Econometrica*, 47 .(۲)
- Kapoor, A., Gulli, A., Pal, S., & Chollet, F. (2022). *Deep Learning with TensorFlow and Keras: Build and deploy supervised, unsupervised, deep, and reinforcement learning models*. Packt Publishing Ltd .
- Khan, W., Daud, A., Khan, K., Muhammad, S., & Haq, R. (2023). Exploring the frontiers of deep learning and natural language processing: A comprehensive overview of key challenges and emerging trends. *Natural Language Processing Journal*, 100026 .
- Kühberger, A. (1998). The influence of framing on risky decisions: A meta-analysis. *Organizational behavior and human decision processes*, 75(1), 23-55 .
- Levin, I. P., Schneider, S. L., & Gaeth, G. J. (1998). All frames are not created equal: A typology and critical analysis of framing effects. *Organizational behavior and human decision processes*, 76(2), 149-188 .
- Li, Y., & Yang, T. (2018). Word embedding for understanding natural language: a survey. *Guide to big data applications*, 83-104 .

- Liebmann, M., Orlov, A. G., & Neumann, D. (2016). The tone of financial news and the perceptions of stock and CDS traders. *International Review of Financial Analysis*, 46, 159-175 .
- Liu, J., Li, T., Xie, P., Du, S., Teng, F., & Yang, X. (2020). Urban big data fusion based on deep learning: An overview. *Information Fusion*, 53, 123-133 .
- Lopez-Lira, A., & Tang, Y. (2023). Can chatgpt forecast stock price movements? return predictability and large language models. *arXiv preprint arXiv:2304.07619* .
- Lunsford, K. G. (2020). Policy language and information effects in the early days of Federal Reserve forward guidance. *American Economic Review*, 110(9), 2899-2934 .
- McDonald, B. (2012). Loughran and mcdonald financial sentiment dictionary. In. MEX صرافی ملی. (۲۰۲۳). Retrieved 20 July from <https://www.mex.co.ir/>
- Mikolov, T., Chen, K., Corrado, G., & Dean, J. (2013). Efficient estimation of word representations in vector space. *arXiv preprint arXiv:1301.3781* .
- Mohamadi, S., Mujtaba, G., Le, N., Doretto, G., & Adjero, D. A. (2023). ChatGPT in the Age of Generative AI and Large Language Models: A Concise Survey. *arXiv preprint arXiv:2307.04251* .
- Mohammed, A. H., & Ali, A. H. (2021). Survey of bert (bidirectional encoder representation transformer) types. *Journal of Physics: Conference Series* ,
- Morales, E. F., & Escalante, H. J. (2022). A brief introduction to supervised, unsupervised, and reinforcement learning. In *Biosignal processing and classification using computational learning and intelligence* (pp. 111-129). Elsevier .
- Mundell, R. (1963). Capital Mobility and Stabilization Policy Under Fixed and Flexible Exchange Rates. *Canadian Journal of Economics and Political Science*, 29(4), 472-485 .
- Nargesian, F., Samulowitz, H., Khurana, U., Khalil, E. B., & Turaga, D. S. (2017). Learning Feature Engineering for Classification. *Ijcai* ,
- Nassirtoussi, A. K., Aghabozorgi, S., Wah, T. Y., & Ngo, D. C. L. (2015). Text mining of news-headlines for FOREX market prediction: A Multi-layer Dimension Reduction Algorithm with semantics and sentiment. *Expert Systems with Applications*, 42(1), 306-324 .
- Nayak, A. S., Kanive, A. P., Chandavekar, N., & Balasubramani, R. (2016). Survey on pre-processing techniques for text mining. *International Journal of Engineering and Computer Science*, 5(6), 16875-16879 .
- Pennington, J., Socher, R., & Manning, C. D. (2014). Glove: Global vectors for word representation. *Proceedings of the 2014 conference on empirical methods in natural language processing (EMNLP)* ,
- Potdar, K., Pardawala, T. S., & Pai, C. D. (2017). A comparative study of categorical variable encoding techniques for neural network classifiers. *International journal of computer applications*, 175(4), 7-9 .
- Sarzynska-Wawer, J., Wawer, A., Pawlak, A., Szymanowska, J., Stefaniak, I., Jarkiewicz, M., & Okruszek, L. (2021). Detecting formal thought disorder by deep contextualized word representations. *Psychiatry Research*, 304, 114135 .
- Seabe, P. L., Moutsinga, C. R. B., & Pindza, E. (2023). Forecasting cryptocurrency prices using LSTM, GRU, and bi-directional LSTM: a deep learning approach. *Fractal and Fractional*, 7(2), 203 .
- Semiromi, H. N., Lessmann, S., & Peters, W. (2020). News will tell: Forecasting foreign exchange rates based on news story events in the economy calendar. *The North American Journal of Economics and Finance*, 52, 101181 .
- Shilpa, B., & Shambhavi, B. (2021). Combined deep learning classifiers for stock market prediction: integrating stock price and news sentiments. *Kybernetes*, 52(3), 748-773 .
- Shynkevich, Y., McGinnity, T. M., Coleman, S. A., & Belatreche, A. (2016). Forecasting movements of health-care stock prices based on different categories of news articles using multiple kernel learning. *Decision Support Systems*, 85, 74-83 .

- Singh, D., & Singh, B. (2020). Investigating the impact of data normalization on classification performance. *Applied Soft Computing*, 97, 105524 .
- Singh, J., & Gupta, V. (2016). Text stemming: Approaches, applications, and challenges. *ACM Computing Surveys (CSUR)*, 49(3), 1-46 .
- Sousa, M. G., Sakiyama, K., de Souza Rodrigues, L., Moraes, P. H., Fernandes, E. R., & Matsubara, E. T. (2019). BERT for stock market sentiment analysis. 2019 IEEE 31st international conference on tools with artificial intelligence (ICTAI) ,
- Sun, A., Lachanski, M., & Fabozzi, F. J. (2016). Trade the tweet: Social media text mining and sparse matrix factorization for stock market prediction. *International Review of Financial Analysis*, 48, 272-281 .
- TGJU. نرخ ارز. (۲۰۲۳). Retrieved 20 July from <https://www.tgju.org>
- Tversky, A., & Kahneman, D. (1981). The framing of decisions and the psychology of choice. *science*, 211(4481), 453-458 .
- Van de Kauter, M., Breesch, D., & Hoste, V. (2015). Fine-grained analysis of explicit and implicit sentiment in financial news articles. *Expert Systems with Applications*, 42(11), 4999-5010 .
- Villamil, L., Bausback, R., Salman, S., Liu, T. L., Horn, C., & Liu, X. (2023). Improved stock price movement classification using news articles based on embeddings and label smoothing. *arXiv preprint arXiv:2301.10458* .
- Wang, X.-T., Rao, L., & Zheng, H. (2016). Framing effects: Behavioral dynamics and neural basis. *Neuroeconomics*, 145-165 .

**Evaluating the "framing effect" on the exchange rate
in Iran using machine learning methods**

Elmira Asle Roosta¹

Alireza Erfani²

Abdolmohammad Kashian³

Received: 2025-Dec-29

Accepted: 2026-Feb-28

Abstract

This study examines the sharp fluctuations in Iran's exchange rate, which are often triggered by significant political-economic news. It aims to determine whether the impact of such news can be systematically measured. Using a deep learning approach, the research presents a tool for predicting exchange rates and evaluates the influence of news and news sources. Data from 10 major domestic and international news sources between 2014 and 2023, alongside exchange rate and economic indicators, were fed into the model. Results show that including news enhances prediction accuracy in both regression and classification modes. Moreover, the model reveals that the impact of news varies based on both its content and the source reporting it, emphasizing the role of news sources in shaping market reactions.

Keywords: Exchange Rate Forecasting, News Impact, News Sources, Machine Learning, Artificial Intelligence.

JEL Classification: C45, C51, C53, C55, G17

⁴ Department of Economics, Faculty of Economics, Management and Administrative Sciences, Semnan University, Semnan. Iran. e.asleroosta@semnan.ac.ir

² Department of Economics, Faculty of Economics, Management and Administrative Sciences, Semnan University, Semnan. Iran. (Corresponding author) aerfani@semnan.ac.ir

³ Department of Economics, Faculty of Economics, Management and Administrative Sciences, Semnan University, Semnan. Iran. a.m.kashian@profs.semnan.ac.ir