



Financial Management Perspective

Journal homepage: <https://jfmp.sbu.ac.ir/>



Original Article

Predicting Iran's Yield Curve: Combining Factor Model with Machine Learning Approach

Saeed Mohammadiaghdam*
Moslem Peymany Foroushany**
Meysam Amiry***
Mohammad Bahrani****

Abstract

Introduction: The yield curve is a key analytical tool in economics, offering vital insights into market expectations regarding monetary policy, economic conditions, and inflation across various time horizons. It also plays a critical role in fiscal policymaking, financial institution modeling, and investment decisions such as asset valuation and risk management. Despite its importance, the analysis and forecasting of the yield curve have received limited attention in Iran. This becomes especially significant in the context of chronic inflation, currency volatility, international sanctions, and dependence on oil revenues. The present study aims to forecast the risk-free government bond yield curve in Iran. To this end, a two-dimensional forecasting approach is employed across both time and maturity dimensions, allowing for the simultaneous analysis of the term structure and its dynamic behavior over time.

Method: Among the various approaches to yield curve forecasting, the Dynamic Nelson-Siegel (DNS) factor model is adopted as the foundational framework due to its interpretability, dimensionality reduction capabilities, and its ability to summarize the curve through three latent factors: level, slope, and curvature. These factors have well-established economic and financial interpretations, providing a meaningful

Received: Feb. 28, 2025

Accepted: June. 2, 2025

* Ph, D Student of Financial Engineering, Department of Finance and Banking, Faculty of Management and Accounting, Allameh Tabataba'i University, Tehran, Iran. (**Corresponding Author**), Email: Saeed.aghdam@yahoo.com

** Associate Prof., Department of Finance and Banking, Faculty of Management and Accounting, Allameh Tabataba'i University, Tehran, Iran.

*** Assistant Prof., Department of Finance and Banking, Faculty of Management and Accounting, Allameh Tabataba'i University, Tehran, Iran

**** Assistant Prof., Department of Computer Science, Faculty of Computer, Statistics, Mathematics, Allameh Tabataba'i University, Tehran, Iran

basis for strategic and policy-level decision-making. Using data from Iranian Islamic Treasury Bills (ITBs), this study forecasts the factors mentioned above using a range of models, including the Vector Autoregressive-GARCH (VAR-GARCH) model as a classical baseline, gradient boosting algorithms as shallow machine learning models, and deep learning architectures such as Convolutional-Recurrent Long Short-Term Memory (Conv-LSTM) networks and Gated Recurrent Units (GRU). These models differ in terms of complexity, interpretability, data requirements, computational demands, and their capacity to capture linear or nonlinear relationships.

Results and Discussion: The empirical results reveal that the VAR-GARCH model outperforms others in forecasting the level factor, largely due to its autoregressive structure, which is better suited for modeling stable long-term trends. Conversely, deep learning models underperform in predicting the level factor due to limited data availability and difficulty in capturing persistent trends. However, for the slope and curvature factors—more influenced by short- and medium-term fluctuations—deep learning models demonstrate superior performance, owing to their ability to capture complex nonlinear temporal patterns. In contrast, traditional statistical models exhibit limitations in handling such fluctuations due to rigid assumptions. Subsequently, the predicted factors were integrated into the DNS model, and the accuracy of the reconstructed yield curve was evaluated using the Root Mean Square Error (RMSE). The results indicate that no single model dominates in predicting all three factors simultaneously. Therefore, a hybrid model strategy, in which each factor is forecasted by the most accurate model, leads to enhanced reconstruction performance. This approach is also theoretically consistent with the DNS model's assumption of factor independence. The optimal configuration was achieved when the level factor was predicted using either VAR-GARCH or Conv-LSTM, the slope factor using GRU, and the curvature factor using either VAR-GARCH or a gradient boosting algorithm, resulting in a reconstruction error of approximately 0.5%.

Conclusion: This study introduces an accurate and data-driven framework for yield curve forecasting in the Iranian financial market by leveraging the Dynamic Nelson-Siegel model. Unlike previous studies that primarily relied on classical approaches such as VAR, this research integrates both shallow and deep machine learning models. In the first stage, these models were evaluated based on their ability to predict the DNS factors. The VAR-GARCH model was found to be most effective for forecasting the level factor, while deep learning models were more accurate in forecasting slope and curvature. In the second stage, the reconstructed yield curve, based on the predicted factors, was assessed using RMSE. The findings suggest that a tailored combination of models for each factor—specifically, VAR-GARCH or Conv-LSTM for level, GRU for slope, and VAR-GARCH or gradient boosting for curvature—results in the highest forecasting accuracy, with a reconstruction error below 0.5%.



Keywords: Yield Curve; Factor Model; Machine Learning; Deep Learning; Fixed Income Securities.

How to Cite: Mohammadiaghdam, S. , Peymany Foroushany, M. , Amiry, M. and bahrani, M. (2025). Predicting Iran's Yield Curve: Combining Factor Model with Machine Learning Approach. *Financial Management Perspective*, 15(1), 9-39. doi:10.48308/jfmp.2025.238972.1476 (In Persian).

چشم انداز مدیریت مالی، ۱۴۰۴، دوره ۱۵، شماره ۱، صص ۹-۳۹، doi: [10.48308/jfmp.2025.238972.1476](https://doi.org/10.48308/jfmp.2025.238972.1476)



نوع مقاله: پژوهشی

پیش‌بینی منحنی بازده ایران: ترکیب مدل عاملی با رویکرد یادگیری ماشین

سعید محمدی اقدم*
مسلم پیمانی فروشانی**
میثم امیری***
محمد بحرانی****

چکیده

هدف: منحنی بازده یکی از ابزارهای کلیدی در تحلیل‌های اقتصادی به‌شمار می‌رود که نقش مهمی در تفسیر انتظارات بازار نسبت به سیاست‌های پولی، وضعیت اقتصادی و تورم در بازه‌های زمانی مختلف ایفا می‌کند. این منحنی همچنین در حوزه‌هایی چون سیاست‌گذاری مالی، مدل‌سازی کسب‌وکار نهادهای مالی و تصمیم‌گیری‌های سرمایه‌گذاری مانند ارزش‌گذاری دارایی‌ها و مدیریت ریسک کاربرد فراوانی دارد. با وجود اهمیت بالای موضوع، پیش‌بینی و تحلیل منحنی بازده در ایران کمتر مورد توجه قرار گرفته است درحالی‌که اقتصاد ایران با چالش‌هایی مانند تورم مزمن، نوسانات ارزی، تحریم‌ها و وابستگی به درآمدهای نفتی مواجه است. هدف این پژوهش، پیش‌بینی منحنی بازده اوراق دولتی بدون ریسک در ایران است. در این راستا، پیش‌بینی

تاریخ پذیرش: ۱۴۰۴/۰۳/۱۲

تاریخ دریافت: ۱۴۰۳/۱۲/۱۰

* دانشجوی دکتری مهندسی مالی، گروه مالی و بانکداری، دانشکده مدیریت و حسابداری، دانشگاه علامه طباطبائی (ره)، تهران، ایران. (نویسنده مسئول)
Email: Saeed.aghdam@yahoo.com

** دانشیار گروه مالی و بانکداری، دانشکده مدیریت و حسابداری، دانشگاه علامه طباطبائی (ره)، تهران، ایران.

*** استادیار گروه مالی و بانکداری، دانشکده مدیریت و حسابداری، دانشگاه علامه طباطبائی (ره)، تهران، ایران.

**** استادیار، گروه رایانه، دانشکده آمار، ریاضی و رایانه، دانشگاه علامه طباطبائی (ره)، تهران، ایران.

با توجه به دو بعد زمان و سررسید انجام شد به طوریکه ضمن بررسی رفتار بازده اوراق با سررسید مختلف در هر زمان، روند تغییرات هر سررسید در طول زمان نیز تحلیل شد.

روش: با وجود توسعه روش‌های مختلف برای پیش‌بینی منحنی بازده، مدل عاملی نلسون-سیگل پویا به دلیل تفسیرپذیری بالا، کاهش ابعاد و توانایی خلاصه‌سازی منحنی در سه عامل کلیدی سطح، شیب و انحنا، به عنوان چارچوب پایه برآورد انتخاب شد. این عوامل به دلیل دلالت‌های اقتصادی و مالی مشخص، نقشی مهم در تصمیم‌گیری‌های سیاستی و راهبردی ایفا می‌کنند. در این پژوهش، با استفاده از داده‌های اسناد خزانه اسلامی در بازار سرمایه ایران، تلاش شد تا عوامل مذکور با مجموعه مدل‌ها از جمله مدل خود رگرسیون برداری-گارچ (به عنوان مدل مبنا) و سایر مدل‌ها ذیل یادگیری ماشین مانند الگوریتم مبتنی بر تقویت گرادینان به عنوان مدل سطحی و مدل‌های شبکه عصبی پیچشی - حافظه طولانی کوتاه-مدت و واحد بازگشتی دارای دروازه به عنوان مدل یادگیری عمیق پیش‌بینی شوند. در نهایت با جایگذاری مقادیر پیش‌بینی شده سه عامل در معادله نلسون-سیگل پویا منحنی بازده آینده بازسازی گردد. شایان ذکر است که هر یک از مدل‌های برآورد از نظر پیچیدگی، تفسیرپذیری، نیازهای داده‌ای، الزامات محاسباتی و نوع روابط (خطی یا غیرخطی)، ویژگی‌هایی متفاوت دارند.

یافته‌ها: یافته‌های پژوهش نشان می‌دهد که مدل خودرگرسیون برداری-گارچ در پیش‌بینی عامل سطح عملکرد بهتری نسبت به سایر مدل‌ها دارد. این برتری به دلیل ساختار خودرگرسیو این مدل است که برای تحلیل روندهای پایدار و طولانی مدت مناسب‌تر عمل می‌کند. در مقابل، مدل‌های یادگیری عمیق به دلیل محدودیت داده و ضعف در شناسایی روندهای طولانی مدت، دقت کمتری در پیش‌بینی این عامل داشته‌اند. اما در مورد عامل‌های شیب و انحنا که بیشتر تحت نوسانات کوتاه‌مدت و میان‌مدت قرار دارند، مدل‌های یادگیری عمیق عملکرد بهتری نسبت به مدل‌های سنتی از خود نشان داده‌اند. این برتری به توانایی آن‌ها در درک الگوهای پیچیده و غیرخطی در طول زمان بازمی‌گردد، درحالی‌که مدل‌های آماری کلاسیک به دلیل مفروضات سخت‌گیرانه در مواجهه با چنین نوساناتی دچار خطا می‌شوند. در مرحله بعد، پیش‌بینی سه عامل در معادله نلسون-سیگل پویا جای‌گذاری شده و دقت بازسازی منحنی بازده با معیار ریشه میانگین مربعات خطا سنجیده شد. نتایج نشان داد که هیچ‌یک از مدل‌ها به تنهایی برتری مطلق در پیش‌بینی هر سه عامل را ندارند. بنابراین، استفاده از ترکیب بهینه از مدل‌ها - به گونه‌ای که هر عامل توسط مدلی با کمترین خطا پیش‌بینی شود - می‌تواند دقت بازسازی منحنی بازده را افزایش دهد و این رویکرد با ساختار مدل نلسون-سیگل، مبتنی بر فرض استقلال عامل‌ها از یکدیگر، نیز سازگار است. نتایج نشان داد در صورتیکه عامل سطح با مدل خود رگرسیون برداری - گارچ یا شبکه عصبی پیچشی - حافظه طولانی کوتاه‌مدت، شیب با واحد بازگشتی دارای دروازه و انحنا با مدل خود رگرسیون برداری-گارچ یا الگوریتم مبتنی بر تقویت گرادینان برآورد شوند به بهترین نتایج یعنی کمترین انحراف از واقعیت معادل حدود نیم درصد منجر خواهد شد.

نتیجه‌گیری: این پژوهش با هدف ارائه مدلی برای پیش‌بینی منحنی بازده در بازار مالی ایران انجام شد. از این رو، مدل نلسون-سیگل پویا انتخاب شد که منحنی بازده را در قالب سه عامل سطح، شیب و انحنا مدل‌سازی می‌کند. این تحقیق از مجموعه مدل‌های سنجی و یادگیری ماشین جهت برآورد استفاده کرد. در مرحله نخست، عملکرد مدل‌ها در پیش‌بینی عامل‌های نلسون-سیگل ارزیابی شد. نتایج نشان داد که مدل خودرگرسیون برداری-گارچ برای پیش‌بینی عامل سطح عملکرد برتری دارد، در حالی که مدل‌های یادگیری عمیق در پیش‌بینی عامل‌های شیب و انحنا، که نوسانات کوتاه‌مدت و میان‌مدت دارند، دقیق‌تر عمل کردند. در مرحله دوم، دقت بازسازی منحنی بازده بر اساس عامل‌های پیش‌بینی شده سنجیده شد. یافته‌ها نشان داد که بهترین ترکیب برای پیش‌بینی منحنی زمانی حاصل می‌شود که عامل سطح با مدل خودرگرسیون برداری - گارچ یا شبکه عصبی پیچشی - حافظه طولانی کوتاه مدت، عامل شیب با واحد بازگشتی دارای دروازه و عامل انحنا با خودرگرسیون برداری-گارچ یا الگوریتم مبتنی بر تقویت گرادینان پیش‌بینی شود که منجر به خطای بازسازی کمتر از نیم درصد خواهد شد.

کلیدواژه‌ها: منحنی بازده؛ مدل عاملی؛ یادگیری ماشین؛ یادگیری عمیق؛ اوراق بهادار با درآمد ثابت.

استناد دهی: محمدی اقدم، سعید، پیمانی فروشانی، مسلم، امیری، میثم و بحرانی، محمد. (۱۴۰۴). پیش‌بینی منحنی بازده ایران: ترکیب مدل عاملی با رویکرد یادگیری ماشین. چشم‌انداز مدیریت مالی، ۱۵(۱)، ۹-۳۹.

۱. مقدمه

منحنی بازده به عنوان نگاشتی از بازده اوراق بهادار با درآمد ثابت دارای سررسیدهای مختلف است که ریسک، نقد شوندگی و مشخصه مالیاتی مشابه دارند [۴۶]. این منحنی می تواند در تفسیر انتظارات بازار از سیاست پولی، فعالیت اقتصادی و انتظارات تورم در افق زمانی کوتاه، میان و طولانی مدت استفاده شود و نقش کلیدی در اقتصاد ایفا کند؛ به طوری که منحنی بازده به عنوان یک مؤلفه مهم در پیش بینی نوسانات اقتصادی است که اطلاعات قابل استخراج از آن می تواند جهت اجرای سیاست پولی صحیح توسط بانک مرکزی به کار گرفته شود. به علاوه این متغیر امکان پیش بینی رشد اقتصادی طولانی مدت را به واسطه ترسیم انتظارات سیاست گذاران از رشد اقتصادی آینده در خصوص مصرف یا سرمایه گذاری منابع، فراهم می کند [۹]. علاوه بر این، بر اساس مطالعات تجربی و نظری، منحنی بازده این قابلیت را دارد که رکود یا رونق آینده اقتصاد را توصیف و پیش بینی کند. نقش این منحنی صرفاً به حوزه اقتصاد کلان محدود نمی شود، بلکه بازارهای مالی نیز از آن بهره مند می شوند. به عنوان نمونه، مطابق اسناد بانک مرکزی اتحادیه اروپا، منحنی بازده که بر پایه انتظارات بازار از نرخ بهره و تورم آینده شکل می گیرد، اطلاعات ارزشمندی در زمینه قیمت گذاری دارایی ها ارائه می دهد که می تواند بر تصمیمات مرتبط با مصارف اثرگذار باشد. همچنین، منحنی بازده ابزاری ارزشمند برای دستیابی به اهداف بانک مرکزی از جمله ثبات مالی، انسجام بازارهای مالی و تحلیل عملیات بازار باز محسوب می شود [۵۳]. علاوه بر این، تحلیل منحنی بازده در زمینه کسب و کارهای مالی، به ویژه صنعت بانکداری، از اهمیت بالایی برخوردار است. دلیل این اهمیت، وابستگی شدید مدل درآمد بانک ها به نرخ بهره برای کسب سود از حاشیه خالص بهره است؛ به طوری که تغییرات نرخ بهره می تواند به طور مستقیم بر درآمد و پایداری بانک ها تأثیر بگذارد. نبود پیش بینی صحیح از روند آتی منحنی بازده، علاوه بر ایجاد اختلال در مدیریت بهینه تطابق سررسید دارایی ها و بدهی ها، زمینه ساز بروز انواع ریسک ها، از جمله ریسک اعتباری خواهد شد. شواهد تجربی نیز نشان می دهد که در اقتصادهای توسعه یافته و حتی در حال توسعه نیز، بانک های بزرگ، سیاست گذاران و نهادهای مالی از تحلیل منحنی بازده به عنوان ابزاری کلیدی برای ارزیابی وضعیت آینده اقتصاد، پیش بینی تورم و تدوین سیاست های بهینه پولی بهره می گیرند [۲]. دلایل یاد شده نشان دهنده ضرورت ارائه پیش بینی از منحنی بازده به طور خاص در اقتصادهای نوظهور، از جمله اقتصاد ایران، است.

تا به امروز، روش ها و مدل های متعددی برای مدل سازی منحنی بازده ارائه شده اند. با این حال، به دلیل ویژگی هایی همچون قابلیت تفسیرپذیری بهتر، خلاصه سازی اطلاعات و کاهش ابعاد داده های ورودی از طریق تعریف عامل ها به عنوان نماینده های منحنی بازده، مدل های عاملی^۱ نسبت به سایر روش ها از برتری برخوردار هستند. در میان این مدل ها، مدل نلسون-سیگل به عنوان یکی از شناخته شده ترین و پرکاربردترین روش ها است که منحنی بازده را بر اساس سه مؤلفه اصلی شامل سطح، شیب و انحنا برآورد می کند [۴۹]. این مدل، به دلیل سادگی در تفسیر و انعطاف پذیری، به طور گسترده توسط نهادهای سیاست گذاری اصلی مانند بانک های مرکزی در سراسر جهان به کار گرفته می شود. به علاوه، بسیاری از مطالعات، تأثیر منحنی بازده در اقتصاد کلان، بازارهای مالی و بازار کالا را بر مبنای همین سه عامل کلیدی (سطح، شیب و انحنا) تحلیل می کنند. عامل سطح منعکس کننده وضعیت طولانی مدت نرخ های بهره و شرایط کلی اقتصاد، از جمله انتظارات تورمی طولانی مدت و نرخ بهره سیاستی است افزایش این عامل نشان دهنده افزایش نرخ بهره در تمام سررسیدها (کوتاه مدت، میان مدت و طولانی مدت) است. این عامل در سیاست گذاری پولی، برآورد هزینه تأمین مالی آتی، تصمیم گیری های مربوط به سرمایه گذاری های طولانی مدت و ارزیابی ثبات اقتصادی

اهمیت ویژه‌ای دارد [۸]. عامل شیب که به‌عنوان تفاوت میان نرخ بهره‌های کوتاه‌مدت و طولانی مدت تعریف می‌شود، در پیش‌بینی چرخه‌های اقتصادی (رکود - رونق) نقش اساسی ایفا می‌کند. به‌عنوان نمونه، شیب مثبت (نرخ بهره طولانی‌مدت بالاتر از نرخ بهره کوتاه‌مدت) معمولاً نشانه‌ای از رشد یا بهبود بالقوه فعالیت‌های اقتصادی، در قالب افزایش مصرف یا سرمایه‌گذاری، تلقی می‌شود [۲۷]؛ عامل انحنا تغییرات کوتاه‌مدت نرخ بهره‌های میان‌مدت را بازتاب می‌دهد و نشان‌دهنده نوسانات و نا اطمینانی در بازار است و برای پیش‌بینی نحوه تعدیل کوتاه‌مدت فعالیت‌های اقتصادی در واکنش به تغییرات سیاست پولی فعلی کاربرد دارد [۱]؛ برای مثال، انحنا یا بالا اغلب نشان‌دهنده نا اطمینانی و بی‌ثباتی بیشتر در اقتصاد است. در چنین شرایطی، فعالان بازار انتظار تغییر سیاست یا بروز یک شوک موقت اقتصادی را دارند، که این می‌تواند به شکل کاهش سرمایه‌گذاری میان‌مدت، تعویق در اعطای تسهیلات مالی میان‌مدت یا تأخیر در اجرای برنامه‌ها و تأمین مالی پروژه‌های توسعه‌ای شرکت‌ها بروز یابد. تغییرات منحنی بازده، علاوه بر تأثیرگذاری در سطح کلان اقتصاد، می‌تواند بر سایر بازارها، از جمله بازار کالاها و به‌ویژه قیمت آن‌ها نیز اثرگذار باشد. این اثرگذاری عمدتاً از طریق دو کانال انتظارات تورمی و هزینه تأمین مالی صورت می‌گیرد. عامل‌ها در این تحلیل از محتوای اطلاعاتی برخوردارند. به‌عنوان مثال، بروز یک شوک مثبت در مؤلفه سطح نشان‌دهنده افزایش انتظارات تورمی در طولانی مدت، رشد هزینه‌های تأمین مالی، کاهش سرمایه‌گذاری و درنهایت افزایش قیمت کالاهاست یا اگر مؤلفه شیب تحت شوک فزاینده قرار گیرد بیانگر کاهش سرمایه‌گذاری و تضعیف چرخه تجاری یا رکود خواهد بود [۴]؛ می‌تواند نشان‌دهنده نا اطمینانی‌های اقتصادی باشد و باعث کاهش سرمایه‌گذاری، به‌ویژه در کالاهایی که نوسانات قیمتی بالایی دارند، شود. روند تغییر عامل‌ها در تحلیل و اتخاذ تصمیمات سرمایه‌گذاری بازارهای مالی نیز حائز اهمیت است. هر یک از این عوامل، یعنی سطح، شیب، و انحنا، نقش متمایزی در پیش‌بینی و مدیریت سرمایه‌گذاری ایفا می‌کنند: عامل سطح به‌طور عمده در تخمین قیمت دارایی‌ها و برآورد هزینه تأمین مالی ناشران کاربرد دارد؛ شیب منحنی بازده به دلیل ارتباط با پیش‌بینی چرخه‌های اقتصادی، در تصمیم‌گیری‌های مربوط به تخصیص دارایی‌ها میان طبقات مختلف دارایی‌ها (مانند سهام، اوراق قرضه، یا صنایع خاص) مؤثر است؛ و انحنا نیز به دلیل ارتباط با نوسانات، می‌تواند به‌عنوان معیاری برای ارزیابی ریسک در نظر گرفته شود. تأثیر این عامل‌ها از طریق چندین کانال مختلف قابل‌بررسی است. یکی از این کانال‌ها، تأثیر بر نرخ تنزیل و ارزش‌گذاری دارایی‌ها (اعم از سهام و اوراق) بر مبنای رویکرد جریان‌های نقدی تنزیل شده است. این تأثیر به دو شکل قابل مشاهده است: نخست، تأثیر مستقیم بر نرخ تنزیل و ارزش‌گذاری دارایی‌ها؛ و دوم، تأثیر غیرمستقیم بر جریان‌های نقدی انتظاری شرکت‌ها، که ناشی از تغییرات هزینه تأمین مالی به‌واسطه نوسانات منحنی بازده است. کانال دوم، به تأثیر تغییرات نرخ بهره بر بازده سهام از طریق باز تنظیم پرتفوی سرمایه‌گذاران اشاره دارد. برای مثال، در شرایط کاهش نرخ بهره، مدیران پرتفوی معمولاً با فروش اوراق و خرید سهام، ترکیب سبد سرمایه‌گذاری خود را تغییر می‌دهند؛ این تغییرات می‌تواند منجر به افزایش قیمت سهام شود. افزون بر این، هر یک از عامل‌های منحنی بازده می‌توانند بازده یا نوسانات را به دارایی‌ها منتقل کنند. بنابراین، پیش‌بینی دقیق این عامل‌ها در مدیریت پرتفوی، تصمیم‌گیری در زمینه تخصیص دارایی‌ها و مدیریت ریسک، اهمیت بالایی دارد [۶۴]. ترکیب شواهد موجود نشان می‌دهد که ارائه پیش‌بینی‌های قابل‌اعتماد از منحنی بازده مبتنی بر عامل‌های اصلی، می‌تواند سیاست‌گذاران را در اتخاذ تصمیمات مناسب پیرامون نرخ بهره، مدیریت تورم و حفظ ثبات اقتصادی یاری کند و بالاترین عایدی را برای فعالان بازار داشته باشد. این پیش‌بینی‌ها باید دارای دقت و کارایی بالا بوده و به‌طور مستمر به‌روزرسانی شوند، زیرا صحت آن‌ها به‌عنوان یک عامل کلیدی، نقش تعیین‌کننده‌ای در تصمیم‌گیری‌های سیاست‌گذاران و سرمایه‌گذاران در بازارهای مالی ایفا می‌کند.

اهمیت پیش‌بینی منحنی بازده و تدوین تصمیمات و راهبردهای متناسب در اقتصادهای نوظهوری مانند ایران، که با چالش‌هایی نظیر تورم، بازارهای کمتر توسعه‌یافته و کم‌عمق، عدم بلوغ کافی، و مجموعه‌ای از شوک‌ها از جمله

تحریم‌ها، نوسانات ارزی، مشکلات ساختاری، وابستگی به درآمدهای نفتی و سایر محدودیت‌ها مواجه‌اند، به‌مراتب بالاتر است. در چنین شرایطی، پیش‌بینی دقیق و قابل‌اتکا از منحنی بازده می‌تواند در سطوح مختلف تصمیم‌گیری - از سیاست‌گذاری کلان تا اجرا و نظارت - نقشی کلیدی ایفا کند. در سطح کلان، این پیش‌بینی‌ها به سیاست‌گذاران کمک می‌کند تا سیاست‌های پولی و مالی بهینه، پیش‌بینی چرخه‌های اقتصادی، مدیریت ترازنامه بانک مرکزی و سایر بانک‌ها، زمان‌بندی و قیمت‌گذاری مناسب برای انتشار اوراق دولتی، اجرای مؤثر عملیات بازار باز، تعیین سقف بهینه قیمت‌ها در بورس کالا (در شرایطی که قیمت‌گذاری ضروری است)، و هدایت نرخ‌های تأمین مالی در سایر بازارها (مانند بازار اوراق شرکتی یا تأمین مالی جمعی) را طراحی و اجرا کنند. علاوه بر این، پیش‌بینی منحنی بازده می‌تواند در تصمیم‌سازی برای نهادهای قانون‌گذار، بانک مرکزی و رگولاتورها مؤثر باشد و اتخاذ تصمیمات دقیق‌تر و کم‌ریسک‌تر را در سطوح مختلف (از تصمیم‌گیری‌های کلان تا تصمیمات خرد توسط فعالان بازار) تسهیل کند. در غیاب چنین پیش‌بینی‌هایی، در شرایط اقتصادی نامطمئن، احتمال اتخاذ تصمیمات غیر بهینه و در نتیجه بروز آثار منفی و زیان‌بار بر اقتصاد افزایش می‌یابد.

پیش‌بینی منحنی بازده در سطح خرد نیز آثار مثبتی به همراه دارد، به‌گونه‌ای که نبود آن در اقتصادهای نوظهوری، مانند ایران، می‌تواند منجر به تخصیص ناکارا و غیر بهینه سرمایه، قیمت‌گذاری نادرست ابزارهای مالی، مداخلات غیراصولی دولت و درنهایت، بی‌ثباتی اقتصادی و اختلال در روند رشد اقتصادی شود.

در این مقاله، با توجه به اهمیت مدل عاملی، این مدل به‌عنوان مبنای اصلی انتخاب‌شده و تلاش می‌شود تا با بهره‌گیری از مجموعه روش‌ها، به‌ویژه تکنیک‌های یادگیری ماشین، بهترین تخمین و پیش‌بینی ممکن از عامل‌های مؤثر در شکل‌گیری منحنی بازده ارائه گردد؛ به‌گونه‌ای که نتایج به‌دست‌آمده کمترین میزان انحراف از واقعیت را داشته باشند. استفاده از روش‌های مبتنی بر یادگیری ماشین به دلیل مزایای متعدد آن‌هاست، از جمله: افزایش تنوع و حجم داده‌های ورودی برای تحلیل، عدم نیاز به اعمال مفروضات متعدد برخلاف رویکردهای سنتی، محدودیت رویکرد سنتی در مدل‌سازی روابط غیرخطی و آموزش آن‌ها، ویژگی‌های خاص داده‌های مالی مانند چولگی، دنباله‌های پهن و ساختار وابستگی با پیچیدگی بالا و همچنین تعداد زیاد عوامل مؤثر در ارزش‌گذاری ابزارهای مالی و پیچیدگی آن‌ها. این چالش‌ها و نیازها باعث شدند که روش‌های یادگیری ماشین و رویکردهای زیرمجموعه آن به‌عنوان ابزارهایی مؤثر برای تخمین نتایج و حل مسائل در این حوزه مطرح شوند [۵۲] و مقاله کنونی نیز از آن‌ها استفاده کند.

ساختار مقاله به شرح زیر است: ابتدا ادبیات نظری پژوهش تشریح می‌شود که شامل معرفی مدل عاملی نلسون-سیگل پویا، مدل خود رگرسیون برداری-گارچ، مدل شبکه عصبی پیچشی-حافظه طولانی کوتاه‌مدت، الگوریتم تقویت گرادیان، و مدل واحد بازگشتی دارای دروازه است. سپس به بررسی پیشینه تجربی پژوهش و مرور مطالعات مرتبط پرداخته می‌شود. در ادامه، بخش روش تحقیق به تشریح داده‌های اولیه، مفروضات و جزئیات هر یک از مدل‌ها اختصاص یافته است. در بخش یافته‌ها، مقادیر برآوردشده عامل‌ها، مقایسه مقادیر پیش‌بینی‌شده و واقعی هر یک از عامل‌ها در هر مدل، و ارزیابی توان پیش‌بینی مدل‌ها بر اساس شاخص‌های ارزیابی مشخص شده، ارائه می‌شود. درنهایت، بخش پایانی مقاله به جمع‌بندی نتایج، ارائه پیشنهادها، و فهرست منابع اختصاص دارد.

۲. مبانی نظری و پیشینه پژوهش

مدل‌سازی منحنی بازده به‌طور کلی به سه دسته اصلی تقسیم می‌شود: مدل‌های مبتنی بر رگرسیون، مدل‌های تجربی^۱ و مدل‌های تعادلی^۲. در مدل‌های رگرسیونی، منحنی بازده با برازش یک تابع ریاضی بر داده‌های بازده و

1 Empirical Models

2 Equilibrium Models

سررسیدهای مختلف به دست می‌آید؛ به عبارتی، تلاش می‌شود رابطه بین بازده اوراق و سررسیدهای آن‌ها به صورت یک تابع ریاضی مشخص شود. مدل تجربی معمولاً شکل خاصی از تابع تنزیل^۱ منحنی بازده را فرض می‌کنند که شامل ۲ نوع روش تکراری باز نمونه‌گیری^۲ و خانواده نلسون سیگل هستند که مطابق رویکرد اول آن منحنی بازده به صورت گام‌به‌گام و بر مبنای قیمت اوراق استخراج می‌شود درحالی‌که در مورد دوم با استفاده از سه عامل اصلی، یعنی سطح، شیب و انحنا، منحنی بازده تخمین زده می‌شود [۴۹]. مدل‌های تعادلی به‌عنوان سومین گروه، در چارچوب مدل‌های با راه‌حل با فرم بسته (راه‌حل تحلیل صریح) قرار می‌گیرند در این مدل‌ها، مشابه رویکردهای تجربی، تلاش می‌شود قیمت اوراق برازش شود، اما به دلیل انعطاف‌پذیری محدود این مدل‌ها، دقت لازم در برازش کامل و دقیق منحنی بازده فراهم نمی‌شود. در نهایت، بر اساس یافته‌های مطالعات، مدل‌های تجربی نلسون سیگل پویا، نلسون سیگل بدون آربیتراژ^۴ و سایر مدل‌های مرتبط، به‌عنوان مناسب‌ترین روش‌ها برای مدل‌سازی منحنی بازده شناخته می‌شوند [۱۳].

مدل‌های عاملی از جهت اهمیت و کاربرد در مقایسه با سایر مدل‌ها توسعه بیشتری داشتند. مشهورترین آن‌ها خانواده نلسون سیگل است که توسط نلسون (۱۹۸۷) ارائه و توسط بلیس (۱۹۹۶) با اضافه کردن یک پارامتر کاهنده مازاد توسعه داده شد نلسون بر این باور بود که مدل پیشنهادی می‌تواند ۹۶ درصد از تغییرات بازده اوراق در سررسیدهای مختلف را توضیح دهد [۴۹، ۶]. اسونسون (۱۹۹۵) مدل نلسون سیگل را توسعه داده به طوری‌که یک عامل اصلی و یک عامل کاهنده اضافه کردند مقاله آن‌ها بر مبنای کمینه‌سازی خطاهای قیمت و بازده عمل می‌کند و داده‌های اوراق خزانه و اوراق قرضه دارای کوپن را به‌عنوان ورودی می‌پذیرد [۶۲]. دیپولد و لی (۲۰۰۶) عامل‌های ذاتی را از حالت ایستا به حالت متحرک در طول زمان تغییر دادند. مدل ارائه‌شده، تطابق بالایی با ویژگی‌های شناخته‌شده منحنی بازده دارند و توانسته‌اند پیش‌بینی‌های دقیقی برای ساختار زمانی بازده‌ها در افق‌های کوتاه و طولانی مدت ارائه دهند. به‌طور خاص، نتایج پیش‌بینی در افق‌های طولانی مدت به‌طور قابل‌توجهی بهتر از پیش‌بینی‌های استاندارد و مدل‌های مرجع عمل کرده است [۲۳]. رزنده و فریرا (۲۰۰۸) یک مدل توسعه‌یافته از مدل نلسون سیگل شامل ۵ عامل - دو عامل شیب - را معرفی کردند، پیش‌بینی‌ها در افق‌های زمانی متفاوت و با روش‌های متنوع انجام شد نتایج نشان داد مدل ۵ عاملی لزوماً عامل بهبود قدرت پیش‌بینی نیست [۵۷]. کاپمن و دیگران (۲۰۰۷) به دنبال بهبود مدل پویای نلسون-سیگل برای تحلیل و پیش‌بینی ساختار زمانی نرخ بهره بودند در این راستا آن‌ها تغییرات نوسانات نرخ بهره در طول زمان را با یک فرآیند گارچ مدل‌سازی کردند که منجر به بهبود قابل‌توجه برازش شد [۳۶]. کریستنسن و دیگران (۲۰۰۷) و (۲۰۰۹) به بررسی مدل تعمیم‌یافته اسونسون از می‌پردازد و تلاش دارد نسخه‌ای از این مدل را ارائه دهد که با شرط نبود فرصت آربیتراژ در طول زمان سازگار باشد نتایج نشان می‌دهد که مدل پیشنهادی علاوه بر قابلیت اعمال شرط عدم آربیتراژ، از نظر برازش درون‌نمونه‌ای نیز عملکرد مناسبی دارد [۲۱، ۲۰]. در ادامه اولاه (۲۰۱۷) با تمرکز بر افزودن مؤلفه‌ی نوسان‌پذیری زمانی به معادله مشاهده منحنی بازده با استفاده از مدل ای گارچ، تلاش دارد تا دقت پیش‌بینی مدل در بازه‌های زمانی مختلف افزایش یابد نتایج نشان داد مدل‌های استاندارد نلسون-سیگل در افق‌های زمانی کوتاه بهتر عمل می‌کنند، اما مدل‌های توسعه‌یافته با مؤلفه‌ای گارچ برای افق‌های زمانی میان‌مدت و طولانی مدت کارایی بیشتری دارند. هر یک از این رویکردها در بازارهای مختلف مورد استفاده قرار گرفته‌اند [۶۳]؛ از جمله توسط لیتون و همکاران (۲۰۰۱) در بازار اوراق بهادار ایالات متحده [۳۸]، چو و همکاران (۲۰۰۹) در بازار اوراق قرضه دولتی تایوان [۱۹]، کنگ (۲۰۱۲) در بازار اوراق

1 Discount function

2 Bootstrap model

3 Closed-Form Solution

4 Arbitrage Free Nelson Siegel

دولتی کره جنوبی [۳۲]، گوگاس و همکاران (۲۰۱۵) با بهره‌گیری از مدل‌های پارامتریک و یادگیری ماشین در ایالات متحده [۳۲] و همچنین لورنس (۲۰۱۶) و نگی (۲۰۲۰) در بازار اوراق بهادار اتریش [۳۹، ۴۷]. در سال‌های اخیر، کاربرد یادگیری ماشین در توسعه این مدل‌ها با استقبال فزاینده‌ای همراه بوده است.

همان‌طور که پیش‌تر اشاره شد، ساختار عاملی توصیف تجربی دقیق‌تری از منحنی بازده ارائه می‌دهد؛ به گونه‌ای که کلیه اطلاعات مربوط به اوراق بهادار در قالب چند متغیر یا عامل خلاصه می‌شود و از مدلی بهره می‌برد که شامل تعداد محدودی عامل به همراه ضرایب مرتبط با سرسیدهای مختلف است. دومین عامل جذابیت این مدل، مزایای آماری آن است؛ به این معنا که با کاهش ابعاد مسئله از حالت پیچیده و غیرقابل حل به مدلی با ابعاد پایین‌تر، امکان استخراج اطلاعات ارزشمند فراهم می‌شود. سومین دلیل برتری این مدل، به ملاحظات اقتصادی و مالی بازمی‌گردد؛ چراکه صرف ریسک و بازده انتظاری ابزارهای مالی را می‌توان بر پایه تعداد محدودی از عامل‌های اصلی ریسک محاسبه کرد [۱۳]. در عمل نیز، همانند جنبه‌های نظری، این مدل توسط بانک‌های مرکزی کشورهای مختلف مورد استفاده قرار گرفته و مقادیر عامل‌ها به صورت روزانه منتشر می‌شود. دلیل انتخاب مدل نلسون-سیگل پویا از میان سایر مدل‌های توسعه‌یافته، عملکرد بهتر آن نسبت به مدل‌هایی مانند اسونسون است؛ چراکه با کاهش تعداد پارامترهای برآوردی، محاسبات ساده‌تری دارد و در عین حال، در پیش‌بینی بازده در افق‌های زمانی کوتاه‌مدت، دقت بالاتری ارائه می‌دهد. افزون بر این، در مقایسه با مدل نلسون-سیگل بدون آربیتراژ، از تفسیرپذیری بهتری در تحلیل پویایی زمانی منحنی بازده برخوردار است [۳]. مزیتی که مدل‌های بدون آربیتراژ علی‌رغم دقت بالاتر در برازش، فاقد آن هستند [۵۵]. بنابراین، مدل نلسون-سیگل پویا به دلیل پایداری برآوردها، قابلیت اتکا و تفسیرپذیری بهتر، به عنوان مدل پایه در این مطالعه انتخاب شده است [۶۵].

پس از انتخاب مدل مبنا، گام بعدی شامل تعیین مدل و روش مناسب برای برآورد و پیش‌بینی عامل‌هاست. نخستین رویکرد مورد استفاده، مدل‌های خود رگرسیون برداری هستند که در آن‌ها عامل‌های مربوطه با بهره‌گیری از روش‌های آماری از داده‌های نرخ بهره استخراج می‌شوند [۲۳]. با این حال، در سال‌های اخیر، سایر رویکردها از جمله روش‌های مبتنی بر یادگیری ماشین مانند ماشین بردار پشتیبان^۱ یا جنگل تصادفی^۲ به واسطه توانایی آن‌ها در محاسبه روابط غیرخطی و الگوهای پیچیده داده‌های مالی در پیش‌بینی منحنی بازده نیز مورد توجه قرار گرفته‌اند. اما چرا مدل‌های یادگیری ماشین می‌توانند نسبت به مدل‌های پایه مزایای بیشتری ارائه دهند؟ نخست آنکه، مدل‌های آماری سنتی عمدتاً مبتنی بر رگرسیون خطی متعارف هستند، درحالی‌که مدل‌های یادگیری ماشین قادرند انواع روابط خطی و غیرخطی را پوشش دهند. دوم آنکه، یادگیری ماشین از انعطاف‌پذیری بیشتری در فرایند پیش‌بینی برخوردار است؛ این مدل‌ها امکان بهره‌گیری از یک مدل منفرد یا مدل‌های ترکیبی را فراهم می‌کنند و نسبت به مدل‌های رگرسیونی، وابستگی کمتری به مفروضات دارند. در نهایت، مدل‌های یادگیری ماشین توانایی پردازش داده‌های حجیم و انجام محاسبات سنگین را دارا هستند؛ قابلیت‌هایی که در مدل‌های آماری یا وجود ندارد یا مستلزم صرف زمان زیادی است، درحالی‌که در بازارهای مالی، به دلیل اهمیت تصمیم‌گیری سریع، سرعت پردازش اهمیت بالایی دارد [۱۳]. البته باید توجه داشت که هر روش مزایا و محدودیت‌های خاص خود را دارد؛ مدل‌های سنتی تفسیرپذیری و سادگی بیشتری دارند، درحالی‌که رویکردهای یادگیری ماشین از دقت پیش‌بینی بالاتری در مواجهه با داده‌های پیچیده و انبوه برخوردارند البته باید دقت داشت منجر به بیش‌برازش نشوند. بنابراین، حذف هر یک از این دو رویکرد ممکن است به از دست رفتن بخشی از اطلاعات مفید منجر شود.

1 Support Vector Machine

2 Random Forest

ذکر این نکته ضروری است که منحنی بازده شامل پیش‌بینی منحنی است نه صرفاً یک متغیر؛ بدین معنا که در فرآیند پیش‌بینی این منحنی، باید به‌طور هم‌زمان به هم‌بستگی میان بازده‌ها در سررسیدهای مختلف و نحوه تغییر بازده‌ها در طول زمان توجه شود. بنابراین، ساختار مسئله به‌صورت دویعدی شامل زمان و سررسید است [۲۳]؛ به این معنا که داده‌ها نه تنها شامل بازده اوراق با سررسیدهای مختلف هستند، بلکه زمان نیز به‌عنوان یک بُعد دیگر وارد مدل می‌شود؛ یعنی بازده اوراق برای سررسیدهای گوناگون در دوره‌های زمانی متفاوت مورد بررسی قرار می‌گیرد. این مسئله به‌مراتب پیچیده‌تر از حالتی است که تنها یک سری زمانی مربوط به نرخ بازده یک ورقه خاص پیش‌بینی می‌شود و مستلزم در نظر گرفتن سطوح مختلفی از اطلاعات و پیچیدگی روابط میان مؤلفه‌ها، ناشی از ابعاد بالای مسئله در مدل‌سازی است. در نتیجه، نیاز به پیش‌بینی پویا در طول زمان وجود دارد که به پویایی عوامل مؤثر بر منحنی بازده اشاره دارد [۲۲].

مدل عاملی نلسون - سیگل پویا: مدل نلسون سیگل به‌عنوان یکی از اولین مدل‌های عاملی، منحنی بازده را مبتنی بر سه مؤلفه اصلی سطح^۱، شیب^۲ و انحنا^۳ برآورد می‌کند. سطح، نشان‌دهنده نرخ بهره یا تعادل طولانی مدت، نرخ‌های مورد انتظار و میانگین بازدهی سررسیدهای مختلف در منحنی بازده است. شیب، بیانگر تفاوت میان نرخ‌های بهره طولانی مدت و کوتاه‌مدت بوده و جهت تغییرات نرخ بهره را منعکس می‌سازد؛ عاملی که به‌طور معمول چرخه‌های اقتصادی را نشان می‌دهد [۴۹]. انحنا نیز مؤلفه‌ی غیرخطی منحنی بازده را توصیف می‌کند که ناحیه‌ی میانی منحنی را در برمی‌گیرد و در تعیین میزان کشیدگی یا خمیدگی منحنی بازده نقش دارد؛ ویژگی‌ای که در پیش‌بینی نوسانات و نا اطمینانی‌های اقتصادی کاربرد دارد [۲۳].

دیبولد و لی پس از معرفی مدل نلسون-سیگل و بررسی نتایج حاصل از آن، به این نتیجه رسیدند که برای دستیابی به برآوردهای دقیق‌تر و نتایج قابل‌اعتماد، پارامترهای این مدل باید در طول زمان تغییرپذیر باشند. بر همین اساس، آن‌ها نسخه پویای مدل نلسون-سیگل را به‌صورت زیر ارائه کردند.

$$y_t(\tau) = l_{1t} + S_{2t} \left(\frac{1 - e^{-\lambda\tau}}{\lambda\tau} \right) + C_{3t} \left(\frac{1 - e^{-\lambda\tau}}{\lambda\tau} - e^{-\lambda\tau} \right) + v_t(\tau) \quad (۱)$$

تابع فوق سری زمانی خطی از برازش متغیر y_t بر عامل‌های سطح (l_{1t})، شیب (S_{2t}) و انحناست (C_{3t}). $\lambda^{۵۴}$ یک متغیر ثابت و τ سررسید است. برای برآورد عامل‌ها، دو رویکرد تک‌گامه و دوگامه معرفی شده‌اند. در رویکرد دوگامه، ابتدا عامل‌های سطح، شیب و انحنا با استفاده از رگرسیون حداقل مربعات ساده محاسبه می‌شوند و سپس پویایی سری‌زمانی این عامل‌ها با بهره‌گیری از مدل‌های مختلف، از جمله مدل خود رگرسیون برداری، برآورد می‌گردد. در رویکرد دوم، یعنی رویکرد تک‌گامه، مدل تو سعه‌یافته نلسون-سیگل در قالب نمایش فضا-حالت^۶ ارائه می‌شود. در این چارچوب، دیبولد و همکارانش برازش هم‌زمان نرخ‌های بهره و تخمین عامل‌ها را بر پایه روش حداکثر درست‌نمایی و با استفاده از نرخ‌های بهره خروجی فیلتر کالمن انجام می‌دهند [۳۵]. بولدر و لی بر این باورند که رویکرد برآورد دوگامه نسبت به رویکرد تک‌گامه از عملکرد بهتری برخوردار است [۷]؛ در نتیجه در این مقاله از رویکرد دوگامه استفاده شد.

مدل خود رگرسیون برداری-گارج^۷: اولین مدلی که برای برآورد عامل‌های مدل نلسون-سیگل پویا مورد استفاده قرار گرفت، مدل خود رگرسیون برداری است. در این مقاله نیز این مدل به‌عنوان مبنایی برای مقایسه توان پیش‌بینی

1 Level

2 Slope

3 Curvature

۴ در مقاله اولیه نلسون سیگل ۱۹۸۷ m نماینده سررسید و $1/t$ نماینده مقدار ثابت λ است.

۵ این متغیر شکل ضرایب عامل‌ها را کنترل کرده و تعیین می‌کند که عامل‌ها در سررسید مختلف چگونه بازده را تحت تأثیر قرار می‌دهند به علاوه نشان می‌دهد در کدام سررسید عامل انحنا حداکثر خواهد شد.

6 State-space representation

7 Vector Autoregression

مدل‌های یادگیری ماشین انتخاب شده است. مدل خود رگرسیون برداری یک مدل آماری است که با هدف برآورد روابط خطی میان متغیرهای سری زمانی چندمتغیره به کار می‌رود، به گونه‌ای که هر متغیر نه تنها به مقادیر وقفه دار خود، بلکه به مقادیر وقفه دار سایر متغیرها نیز وابسته است [۲۹].

در ادامه، با استفاده از نتایج و خروجی آزمون‌های مدل خود رگرسیون برداری، مدل گارچ با هدف مدل‌سازی نوسانات متغیر در طول زمان معرفی شد. این مدل امکان تغییر واریانس شرطی جزء اخلاص در طول زمان را فراهم می‌سازد و فرض می‌کند نوسانات جاری به مقادیر گذشته جزء اخلاص و همچنین نوسانات قبلی وابسته هستند [۲۶]. مدل ترکیبی خود رگرسیون برداری - گارچ^۱ با ادغام دو مدل یادشده و با هدف مدل‌سازی هم‌زمان ارتباطات پویا میان متغیرهای سری زمانی چندمتغیره و نوسانات متغیر آن‌ها در طول زمان ارائه شده است. بخش خود رگرسیون برداری این مدل، به بررسی وابستگی درونی و روابط بین متغیرها می‌پردازد، درحالی‌که مؤلفه گارچ، نوسانات زمانی جزء اخلاص حاصل از تابع خود رگرسیون برداری را مدل‌سازی می‌کند [۴۲].

الگوریتم مبتنی بر تقویت گرادیان^۲: الگوریتم مبتنی بر تقویت گرادیان، به عنوان روشی کارآمد شناخته می‌شود که فرایند تقویت گرادیان را در قالبی از ترکیب چندین درخت تصمیم پیاده‌سازی می‌کند [۱۴]. در این الگوریتم، سری‌های زمانی به صورت ترکیبی از درخت‌های تصمیم مدل‌سازی می‌شوند، به گونه‌ای که مقدار پیش‌بینی شده برای یک متغیر وابسته برابر با مجموع خروجی‌های K درخت تصمیم است، به طوری که در هر مرحله، خطای مدل به تدریج کاهش می‌یابد. بر اساس سازوکار این الگوریتم، مدل به صورت مرحله به مرحله ساخته می‌شود، به گونه‌ای که در هر مرحله یک درخت تصمیم جدید ایجاد می‌شود تا خطای پیش‌بینی‌های مرحله‌ی قبل را کاهش دهد. این فرایند با استفاده از گرادیان تابع هزینه انجام می‌پذیرد [۵۰]. این الگوریتم قادر است روابط میان ویژگی‌ها را مدیریت کرده و در انجام پیش‌بینی‌ها، از جمله پیش‌بینی قیمت سهام، عملکرد قابل قبولی از خود نشان دهد [۱۶].

مدل ترکیبی شبکه عصبی پیچشی - حافظه طولانی کوتاه-مدت^۳: یکی از مدل‌های اصلی معرفی شده در این مقاله ترکیبی دو معماری شبکه عصبی پیچشی و حافظه طولانی کوتاه‌مدت است. این ترکیب با هدف بهره‌گیری هم‌زمان از وابستگی‌های زمانی^۴ و مکانی^۵ طراحی شده و در مسائل مختلفی همچون طبقه‌بندی تصاویر، ویدئوها و داده‌های سری زمانی کاربرد دارد. در این مدل، ابتدا لایه‌های شبکه پیچشی به کار گرفته می‌شوند تا ویژگی‌های کلیدی داده‌های مکانی استخراج شوند؛ سپس این ویژگی‌ها به مدل حافظه طولانی کوتاه‌مدت وارد می‌شوند تا وابستگی‌های زمانی یا ترتیبی داده‌ها تحلیل و مدل‌سازی شوند [۶۶]. در مجموع، نتایج مطالعات پیشین نشان می‌دهد که استفاده از مدل‌های یادگیری عمیق نظیر شبکه عصبی پیچشی و حافظه طولانی کوتاه‌مدت در پیش‌بینی سری‌های زمانی مالی موفقیت‌آمیز بوده است. با این حال، مدل شبکه عصبی پیچشی به دلیل ماهیت اصلی خود در تمرکز بر استخراج ویژگی، در مقایسه با مدل حافظه طولانی کوتاه‌مدت دقت پیش‌بینی کمتری در داده‌های سری زمانی عددی دارد. از سوی دیگر، در استخراج مؤثرترین ویژگی‌ها نسبت به شبکه عصبی پیچشی ضعف‌هایی دارد. بنابراین، ساخت یک مدل ترکیبی که بتواند از مزایای هر یک از این مدل‌ها برای جبران ضعف‌هایشان بهره گیرد، رویکردی منطقی برای افزایش دقت پیش‌بینی سری‌های زمانی خواهد بود. در این معماری، داده‌های سری زمانی به گونه‌ای شکل می‌گیرند که با ساختار ورودی شبکه عصبی پیچشی و سپس حافظه طولانی کوتاه‌مدت سازگار باشند. این مدل از دو لایه اصلی

1 VAR-GARCH

2 XGBoost

3 Convolutional Neural Networks - Long Short-Term Memory (CNN-LSTM)

4 Spatial

5 Temporal

تشکیل شده است: لایه شبکه عصبی پیچشی که مسئول استخراج ویژگی‌های کلیدی از داده‌های سری زمانی پردازش شده است، و لایه حافظه طولانی کوتاه‌مدت که وظیفه پیش‌بینی نهایی را بر عهده دارد.

اجزای اصلی این معماری شامل: لایه ورودی، لایه پیچشی تک‌بعدی، لایه تجمیعی^۱، لایه پنهان^۲ مدل حافظه طولانی کوتاه‌مدت و در نهایت لایه اتصال کامل^۳ برای تولید خروجی پیش‌بینی است.

فرآیند آموزش و پیش‌بینی معماری فوق به شرح زیر است:

✓ مرحله آموزش اولیه: فرآیند آموزش با ورود داده‌های آموزشی آغاز می‌شود. در این مرحله، داده‌های سری زمانی برای آموزش مدل وارد می‌شوند. سپس، پارامترهای شبکه شامل وزن‌ها و بایاس‌ها در لایه‌های مختلف مدل مقداردهی اولیه می‌شوند.

✓ استخراج ویژگی‌ها در لایه شبکه عصبی پیچشی: داده‌های ورودی ابتدا از طریق لایه‌های پیچشی و تجمیعی در بخش شبکه عصبی پیچشی عبور می‌کنند. در این مرحله، ویژگی‌های مهم از داده‌های سری زمانی استخراج شده و به عنوان خروجی به لایه حافظه طولانی کوتاه‌مدت منتقل می‌شوند. فرآیند استخراج ویژگی عمدتاً در این مرحله انجام می‌شود.

✓ پیش‌بینی در شبکه عصبی طولانی کوتاه‌مدت: خروجی حاصل از شبکه عصبی پیچشی وارد شبکه طولانی کوتاه‌مدت می‌شود. در این بخش، عملیات پیش‌بینی بر روی مقادیر سری زمانی انجام شده و خروجی شبکه طولانی کوتاه‌مدت به لایه اتصال کامل^۴ ارسال می‌شود تا مقدار نهایی پیش‌بینی تولید شود. پس از آن، خطای پیش‌بینی با مقادیر واقعی مقایسه شده و میزان خطا محاسبه می‌گردد.

✓ بررسی شرط توقف آموزش: در این مرحله، نتایج ارزیابی مدل بررسی می‌شود تا مشخص شود که آیا شرط توقف آموزش فراهم شده است یا خیر. این شرط معمولاً شامل رسیدن به تعداد دوره‌های مشخص و کاهش خطای پیش‌بینی به کمتر از یک آستانه از پیش تعیین شده است.

✓ به‌روزرسانی وزن‌ها و تکرار آموزش: اگر شرط توقف محقق نشده باشد، خطای محاسبه شده به لایه‌های قبلی برگشت داده می‌شود که تحت عنوان فرآیند پس انتشار خطا^۵ شناخت می‌شوند و وزن‌ها و بایاس‌ها اصلاح می‌گردند. سپس فرآیند آموزش مجدداً از مرحله اول تکرار می‌شود. در غیر این صورت، آموزش به پایان رسیده و پیکربندی نهایی مدل ذخیره می‌شود.

✓ آزمایش مدل آموزش دیده: در این مرحله، مدل آموزش دیده با استفاده از داده‌های آزمون مورد ارزیابی قرار می‌گیرد. داده‌های آزمون وارد مدل ذخیره شده شده و خروجی پیش‌بینی نهایی به دست می‌آید.

✓ ارزیابی دقت پیش‌بینی: برای اندازه‌گیری دقت مدل انواع سنج‌ها از جمله ریشه میانگین مربعات خطا مورد استفاده قرار می‌گیرد تا میزان انحراف بین مقادیر واقعی و مقادیر پیش‌بینی شده سنجیده شود [۶۶].

مدل واحد بازگشتی دارای دروازه^۵: مدل واحد بازگشتی دارای دروازه یکی از انواع شبکه‌های عصبی بازگشتی است که به عنوان جایگزینی مدل حافظه طولانی کوتاه‌مدت معرفی شده است. هر دو مدل با هدف مدل‌سازی وابستگی‌های زمانی در داده‌ها توسعه یافته‌اند، با این تفاوت که مدل واحد بازگشتی دارای دروازه به دلیل ساختار ساده‌تر و نیاز به محاسبات کمتر برای به‌روزرسانی وضعیت‌های داخلی، از کارایی محاسباتی بالاتری برخوردار است، همچنان توانایی مقابله با مشکل ناپدید شدن گرادیان را، که از مزایای مدل حافظه طولانی کوتاه‌مدت است،

1 Pooling Layer

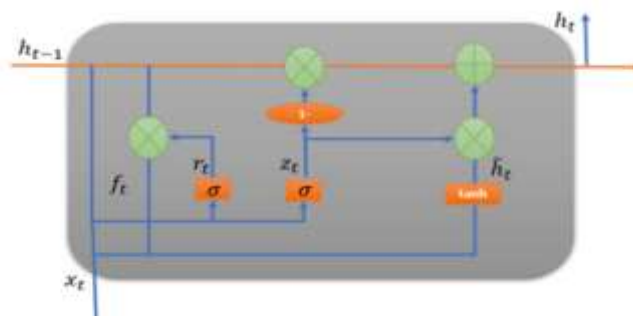
2 Hidden Layer

3 Full Connected Layer

4 Back-Propagation

5 Gated Recurrent Unit

نیز حفظ می‌کند. این مدل شامل ۲ دروازه اصلی به‌روزرسانی^۱ و باز تنظیم^۲ است. دروازه به‌روزرسانی مشخص می‌کند چه میزان از اطلاعات گذشته در گام بعدی حفظ شود و دروازه باز تنظیم تعیین می‌کند چه میزان از اطلاعات قبلی باید کنار گذاشته شود. این مدل با ترکیب دروازه ورودی^۳ و فراموشی^۴ در قالب یک دروازه به‌روزرسانی ساختار واحد بازگشتی دارای دروازه را ساده‌تر کرده و به کاهش تعداد پارامترها و پیچیدگی محاسباتی منجر شده است. با وجود این سادگی، در بسیاری از کاربردها، به‌ویژه در پیش‌بینی سری‌های زمانی، عملکرد بهتری نسبت مدل حافظه طولانی کوتاه‌مدت به از خود نشان داده است [۱۸].



شکل ۱ نحوه عملکرد واحد بازگشتی دارای دروازه - منبع [۴۵]

تصویر فوق نشان‌دهنده ساختار داخلی یک سلول مربوط به واحد بازگشتی دارای دروازه است. این سلول‌ها به‌عنوان توابع محاسباتی هستند که باهدف کنترل مکانیسم آموزش در سلول‌ها استفاده می‌شوند.

$$Z_t = \sigma(x_t w^z + h_{t-1} U^z + b_z) \quad (۲)$$

$$r_t = \sigma(x_t w^r + h_{t-1} U^r + b_r) \quad (۳)$$

$$h_t^- = \tanh(r_t * h_{t-1} U + x_t W + b) \quad (۴)$$

$$h_t = (1 - z_t) * h_t^- + Z_t * h_{t-1} \quad (۵)$$

در فرمول‌های بالا، W و w^z w^r سنجه‌های وزن برای بردار ورودی است. U U^r U^z سنجه وزن مربوط به گام زمانی پیشین و b b_r b_z بایاس مربوطه است. σ تابع سیگموئید لاجیت و r_t نشان‌دهنده باز تنظیم، Z_t نشان‌دهنده دروازه به‌روزرسانی و h_t^- نشان‌دهنده لایه پنهان منتخب است. [۴۵]

پیشینه پژوهش: بیشتر مطالعات مبتنی بر یادگیری ماشین بر بازارهایی مانند سهام، رمز ارز و شاخص‌ها متمرکز بوده‌اند، درحالی‌که استفاده از این رویکرد در بازار اوراق با درآمد ثابت و پیش‌بینی منحنی بازده، هنوز بسیار محدود و کمتر موردتوجه قرار گرفته است. تحقیقات اولیه از مدل‌های سری زمانی کلاسیک از جمله خود رگرسیون برداری استفاده کردند از جمله همیلتون (۱۹۹۴) که البته این منطق توسط برخی محققان از جمله سوئل (۲۰۱۱) به‌واسطه ناکارآمدی مدل‌های کلاسیک در توصیف سری‌های زمانی مالی رد شد [۲۹،۵۹]. لیترمن و شوئمن (۱۹۹۱) با استفاده از تحلیل عناصر اصلی به این نتیجه رسیدند که ۳ عامل در ساختار زمانی نرخ بهره دارای بیشترین تأثیر هستند [۴۵] که این یافته با تحقیقات نش (۲۰۱۲) چاو و چو (۲۰۲۲)، اپرا (۲۰۲۲) برای سایر کشورها تأیید شد [۴۸،۵۴]. مالیک و میشرا (۲۰۱۹) این مقاله با تمرکز بر بازار هند، در پی توسعه روشی برای پیش‌بینی نرخ‌های بهره با سرسیدهای مختلف و همچنین آزمون فشار آن‌ها است. مطابق یافته این پژوهش مدل خود رگرسیون-میانگین متحرک در دوره‌های پرتنش و معمولی عملکرد مناسبی داشته است. آن‌ها با استفاده از رویکرد تحلیل مؤلفه اساسی به این نتیجه رسیدند که

1 Update Gate

2 Reset Gate

3 Input Gate

4 Forget Gate

عناصر اصلی ۸۸ درصد از نوسانات منحنی بازده را نشان می‌دهند [۴۴]. دومین دسته از مدل‌ها از رویکرد آماری تحلیل مبتنی بر داده تابعی^۱ استفاده کردند که به عنوان رویکردی نا پارامتریک توسط کالدريا و تورنت (۲۰۱۷) باهدف پیش‌بینی منحنی بازده آمریکا استفاده شد که البته یافته‌ها کاملاً قابل اتکا نبود و نتوانستند توان عملکردی این طبقه مدل‌ها را تأیید کنند [۱۰]. سومین دسته مدل‌ها توسط کانسکی و دیگران (۲۰۰۸) کانسکی و تیمونین (۲۰۱۰) ارائه شد آن‌ها در پژوهش خود از تحلیل آماره فضایی^۲ استفاده کردند تا منحنی بازده را در فضای دوبعدی نگاشت کنند هدف اصلی آن‌ها، تحلیل ساختارهای زمانی و خوشه‌بندی این منحنی‌ها برای درک بهتر از مدل‌ها و قابلیت پیش‌بینی آن‌هاست. آن‌ها به این نتیجه رسیدند تحلیل پارامترهای مدل نلسون-سیگل و ساختار زمانی و خوشه‌بندی آن‌ها، به سنجش کارایی این مدل برای پیش‌بینی کمک می‌کند [۳۱،۳۰]. سمبسون و داس (۲۰۱۷) برای مدل‌سازی از فرآیند گاوسی پویا استفاده کرد و به این نتیجه رسیدند رویکرد سری زمانی چند متغیره در سررسید کوتاه‌مدت نسبت به فرآیند گاوسی پویا عملکرد بهتری دارد درحالی‌که رویکرد دوم در سررسید میان و طولانی مدت از کارایی بیشتری برخوردار است [۵۸]. این روش توسط پیمنتال و دیگران (۲۰۲۲) با چارچوب رگرسیون کوانتایل ترکیب شد. آن‌ها به این نتیجه رسیدند که این مدل قابلیت توصیف بالای اکثر سررسیدها را داراست؛ در سالیان اخیر استفاده از یادگیری ماشین در پیش‌بینی منحنی بازده تشدید شده است [۵۶]. کرشنزو و دیگران (۲۰۱۸) از رویکرد خود رمزگذار حذف‌کننده نویز برای استخراج مشخصه‌های منحنی بازده اوراق شرکتی با نقد شوندگی پایین استفاده کردند و به نتایج قابل دفاع دست یافتند [۳۴]. نانز و دیگران (۲۰۱۹) از پرسپترون چندلایه با مشخصه‌های مختلف برای پیش‌بینی منحنی بازده اروپا استفاده کردند. این مدل با استفاده از مهم‌ترین ویژگی‌ها، در مجموع بهترین عملکرد را در پیش‌بینی منحنی بازده در افق‌های زمانی مختلف داشته است [۵۱]. سویمون و دیگران (۲۰۲۰) در ۲ مقاله خود به دنبال توسعه یک مدل پیش‌بینی نرخ بهره طولانی مدت ژاپن با استفاده از روش‌های مختلف یادگیری ماشین است؛ به‌ویژه با در نظر گرفتن تأثیرات بازارهای نرخ بهره و ارز خارجی (مانند آمریکا و اروپا) بر بازار ژاپن. آن‌ها در مطالعه دوم خود از مدل‌های یادگیری عمیق مختلف مانند مدل حافظه طولانی کوتاه‌مدت برای پیش‌بینی منحنی بازده استفاده کردند و بهبود نتیجه را گزارش کردند [۶۱،۶۰]. بیانچی و دیگران (۲۰۲۱) در مقاله خود به دنبال بررسی قابلیت پیش‌بینی بازده اوراق قرضه با استفاده از روش‌های یادگیری ماشین، به‌ویژه شبکه‌های عصبی و مدل‌های درختی و مقایسه عملکرد آن‌ها با مدل‌های سنتی مبتنی بر نرخ بهره. بودند. آن‌ها نشان دادند شبکه‌های عصبی شواهد آماری قوی‌تری برای پیش‌بینی بازده اوراق نسبت به مدل‌های سنتی ارائه می‌دهند [۵]. کافمن و دیگران (۲۰۲۲) مدل فضایی - حالت گوسی خطی را با مدل‌های شبکه عصبی ترکیب کردند تا ضرایب عامل‌های منحنی بازده را برآورد کنند. مطابق نتایج، ارزیابی عملکرد مدل با استفاده از داده‌های تاریخی ۱۴ ساله از منحنی بازده برزیل نشان می‌دهد که مدل پیشنهادی پیش‌بینی خارج از نمونه دقیق‌تری نسبت به روش‌های سنتی مانند مدل پویای نلسون-سیگل و دیگر توسعه‌های آن ارائه می‌دهد [۳۳]. کوستیرا (۲۰۲۴) در مقاله خود ابتدا با استفاده از تحلیل مؤلفه‌های اصلی، منحنی بازده را به مؤلفه‌های اصلی تجزیه کردند. سپس مؤلفه‌ها را با استفاده از مدل حافظه طولانی کوتاه‌مدت پیش‌بینی کردند و از طریق بازترکیب، منحنی بازده نهایی را بازسازی کردند. آن‌ها به این نتیجه رسیدند استفاده مستقیم از مدل حافظه طولانی کوتاه‌مدت دقیق‌ترین پیش‌بینی‌ها را دارد و مدل خود رگرسیون معمولاً عملکرد ضعیف‌تری دارد. از طرف دیگر مدل ترکیبی می‌تواند روی مؤلفه‌های خاصی از منحنی بازده مانند سطح یا شیب تمرکز کند [۳۷].

^۱ Functional Data Analysis: به عنوان یک شاخه از آمار است که داده‌ها را با هدف استخراج اطلاعات در خصوص منحنی، سطوح و سایر متغیرها

تحلیل می‌کند.

^۲ Spatial Statistics: به عنوان یک شاخه ای از آمار است که شامل مجموعه ابزارهایی است که برای شناسایی و تحلیل توزیع فضایی، بررسی الگوها و فرایندها و روابط مکانی بین عوارض در نقشه‌ها به کار گرفته می‌شود.

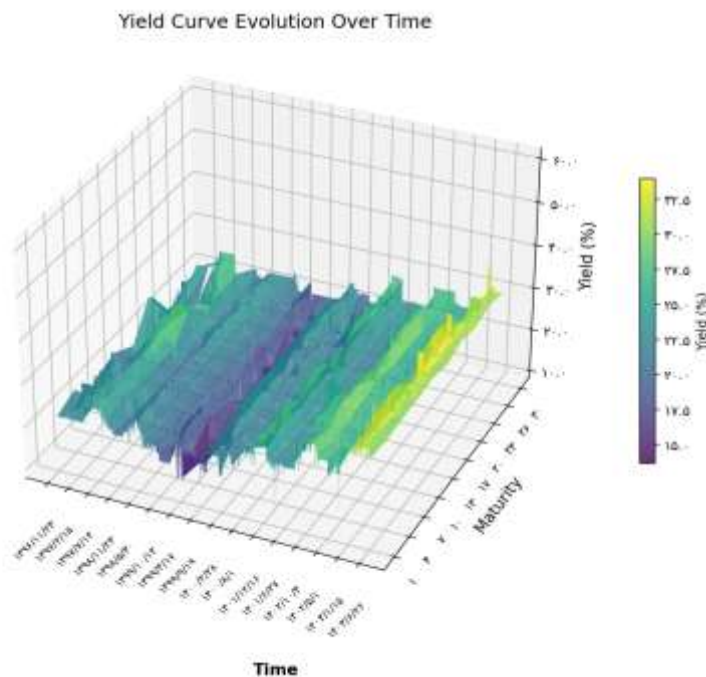
با این حال، برخی از مطالعات، تحلیل منحنی بازده را صرفاً به نرخ بهره یا بازده محدود کرده‌اند و با نادیده گرفتن مؤلفه زمان تا سررسید، آن را به صورت تک بعدی و بدون در نظر گرفتن ساختار دوبعدی آن بررسی کرده‌اند. در این رویکرد، تنها سررسید به عنوان متغیر مستقل لحاظ شده و از مدل‌های مختلف صرفاً برای پیش‌بینی بازده یک دارایی خاص استفاده شده است به عنوان نمونه، دنیس و موریسون (۲۰۰۷) در پژوهش خود با بهره‌گیری از چارچوب مدل سازی فضای حالت، فیلتر کالمن و رگرسیون مبتنی بر شبکه عصبی، بازده اوراق خزانه ۱۰ ساله در سه کشور منتخب را تخمین زده‌اند و به این نتیجه رسیدند که شبکه عصبی می‌تواند جایگزینی قابل اعتمادتر نسبت به مدل‌های سنتی باشد [۲۴]. همان‌طور که پیش‌تر نیز اشاره شد، پیش‌بینی دقیق و معنادار منحنی بازده مستلزم توجه به ساختار دوبعدی آن و لحاظ هم‌زمان هر دو مؤلفه (زمان و بازده) است؛ در غیر این صورت، مدل از نظر تحلیلی و تفسیرپذیری با محدودیت مواجه خواهد شد.

با وجود اهمیتی که برای منحنی بازده وجود دارد و پیش‌تر نیز به آن اشاره شد، تاکنون مطالعه‌ای قابل اتکا در ایران، مبتنی بر مدل عاملی و در چارچوب مسئله تحقیق حاضر، انجام نشده است. این در حالی است که اقتصاد ایران به عنوان یک اقتصاد نوظهور، به دلیل شرایط خاصی همچون تحریم‌ها، نوسانات ارزی و نفتی، نبود بازار بدهی کارآمد و سایر عوامل، نیازمند برآورد و پیش‌بینی منحنی بازده است تا بتواند در اتخاذ سیاست‌های بهینه عملکرد مناسبی داشته باشد. در این مقاله تلاش شده است متناسب با داده‌های اقتصاد ایران، از طیف متنوعی از مدل‌های پیش‌بینی - از مدل‌های سری زمانی چندمتغیره گرفته تا روش‌های یادگیری ماشین سطحی و عمیق - در قالب مدل عاملی نلسون-سیگل پویا بهره گرفته شود تا میزان قدرت پیش‌بینی این رویکردها مشخص گردد. ترکیب مدل عاملی با روش‌های متنوع یادگیری ماشین از نوآوری‌های این پژوهش به شمار می‌رود و در مجموع، گامی در جهت توسعه مدل‌های پیش‌بینی عاملی پویا محسوب می‌شود. افزون بر این، موضوع پژوهش تاکنون در فضای تحقیقاتی ایران مورد بررسی قرار نگرفته است.

۳. روش‌شناسی پژوهش

پژوهش حاضر با رویکردی کمی انجام و بر پایه منطق تحلیل سری‌های زمانی استوار است؛ در این راستا، از ترکیب مدل عاملی با مجموعه‌ای از مدل‌های رگرسیونی و روش‌های یادگیری ماشین بهره گرفته شده است.

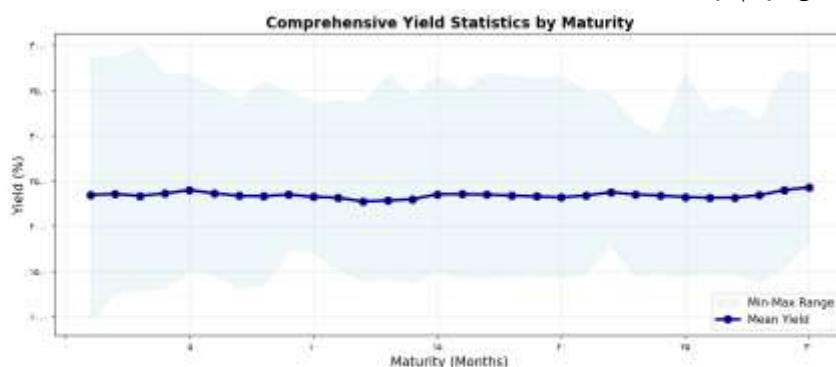
داده ورودی (جامعه آماری): داده‌های اولیه این پژوهش شامل سری زمانی قیمت پایانی کلیه اسناد خزانه اسلامی منتشرشده در بازار سرمایه ایران طی بازه زمانی شهریور ۱۳۹۴ تا شهریور ۱۴۰۳ است که در مجموع ۱۸۸ عنوان اوراق خزانه اسلامی را در برمی‌گیرد. به منظور افزایش دقت و اتکای محاسبات، در گام نخست اوراق سخاب با سررسید کمتر از یک سال که صرفاً در سال اول منتشرشده و به دلیل آشنایی محدود بازار با این ابزار و عدم تداوم انتشار، فاقد کشف قیمت مناسب بودند، از مجموعه داده‌ها حذف شدند. در ادامه و پس از بررسی‌های انجام شده، مشخص شد اوراق با سررسید بیش از ۳۰ ماه از نظر داده‌ای کفایت لازم را ندارند؛ بنابراین، دامنه سررسید انتخابی بین یک تا ۳۰ ماه در نظر گرفته شد. در نتیجه، مجموعه داده نهایی شامل ۱۵۸۳ ردیف (روز معاملاتی) در بازه زمانی ۱۳۹۶/۱۱/۲۱ تا ۱۴۰۳/۶/۲۸ و برای سررسیدهای یک‌ماهه تا ۳۰ ماهه با فواصل یک‌ماهه تنظیم شده است. مقادیر هر دوره سررسید نیز به صورت روزانه ثبت شده‌اند. در تصویر زیر، تاریخچه منحنی بازده اسناد خزانه اسلامی قابل مشاهده است:



نمودار ۱ سابقه منحنی بازده

ابزار و روش گردآوری داده‌ها: داده‌های مورد استفاده از وب سایت فرابورس ایران و شرکت مدیریت فناوری بورس تهران گردآوری شده و پردازش‌های اولیه آن‌ها با استفاده از نرم‌افزار اکسل صورت گرفته است. همچنین، تحلیل‌های اقتصادسنجی با بهره‌گیری از نرم‌افزار ایویوز و مدل‌های یادگیری ماشین با استفاده از زبان برنامه‌نویسی پایتون انجام شده‌اند.

مراحل پردازش: پس از دریافت قیمت اوراق منتشر شده، بازده تا سررسید هر یک از اوراق به صورت سری زمانی محاسبه شد. به منظور ترسیم منحنی بازده، بعد دوم داده‌ها یعنی "زمان تا سررسید" نیز ایجاد گردید. برای این منظور، بازه‌های زمانی با فاصله‌های یک‌ماهه از ۱ تا ۳۰ ماه تا سررسید در نظر گرفته شد. خروجی این مرحله، یک ماتریس دوبعدی از بازده تا سررسید اوراق در قالب زمان و سررسید است که نمایانگر سری زمانی منحنی بازده روزانه اوراق خزانه اسلامی می‌باشد. در ادامه، میانگین بازده برای سررسیدهای مختلف به همراه محدوده تغییرات آن‌ها (بازه بین حداکثر و حداقل بازده) ارائه شده است.



نمودار ۲ متوسط و محدوده حداقل - حداکثر نرخ بازده در هر سررسید

در ادامه، روند تاریخی عامل‌های مدل نلسون-سیگل پویا با استفاده از روش حداقل مربعات معمولی و بر پایه ماتریس بازده اوراق به‌دست‌آمده از مرحله پیشین برآورد می‌شوند. خروجی این مرحله، شامل سری‌های زمانی تاریخی

سه عامل سطح، شیب و انحنا است که در ادامه توسط مجموعه مدل‌ها از جمله خود رگرسیون برداری-گارچ، الگوریتم مبتنی بر تقویت گرادیان ذیل یادگیری سطحی و مدل‌های شبکه عصبی پیچشی - حافظه طولانی کوتاه‌مدت و واحد بازگشتی دارای دروازه ذیل یادگیری عمیق پیش‌بینی خواهند شد.

در پایان نیز، با جایگذاری مقادیر پیش‌بینی شده سه عامل سطح، شیب و انحنا در معادله نلسون-سیگل پویا منحنی بازده آینده بازسازی می‌گردد و دقت این برآورد نیز با استفاده از معیار ریشه میانگین مربعات خطا مورد ارزیابی قرار خواهد گرفت.

اولین مدل خود رگرسیون برداری-گارچ با وقفه (۱،۱) است که در سایر مطالعات نیز استفاده شده است. مدل پایه خود رگرسیون برداری با وقفه ۲ (بر اساس مشخصه سنج تعداد وقفه ۱) است که پس از بررسی آزمون‌های مختلف مشخص شد جزء اخلاص در تمام حالات دارای ناهمسانی واریانس است و به هیچ‌وجه قابل حل نیست. در نتیجه مدل مذکور با مدل گارچ (۱،۱) به عنوان مدل مکمل جهت مدل‌سازی واریانس ترکیب شد. توضیحات مربوطه در بخش یافته پژوهش به طور کامل آمده است.

مدل دوم بر پایه الگوریتم تقویت گرادیان طراحی شده است. این مدل از ترکیب مجموعه‌ای از درخت‌های تصمیم با رویکرد گرادیان تقویت شده بهره می‌برد و قادر است روابط غیرخطی میان متغیرها را به خوبی در نظر بگیرد. به دلیل ساختار خاص خود، این مدل از بروز بیش برآزش جلوگیری می‌کند. برخلاف مدل‌های یادگیری عمیق، این روش فاقد قابلیت مهندسی خودکار ویژگی‌هاست و همچنین ارتباط زمانی بین داده‌ها ۲ را مستقیماً لحاظ نمی‌کند. بنابراین، در کاربرد آن برای پیش‌بینی سری‌های زمانی، لازم است از راهکارهایی نظیر مدیریت آرایه‌ها و تعریف پنجره‌های متحرک^۳ استفاده شود. داده‌های ورودی این مدل شامل سری زمانی عامل‌هایی هستند که با استفاده از مدل نلسون-سیگل پویا برآورد شده‌اند. در ابتدا، داده‌ها با نسبت‌های ۷۰، ۱۵ و ۱۵ درصد به مجموعه‌های آموزش، اعتبارسنجی و آزمون تقسیم شدند. در ادامه داده آموزش پیش از ورود به مدل با استفاده از آزمون Z نرمال‌سازی می‌شوند تا مقیاس ویژگی‌ها قابل مقایسه گردد، فرآیند یادگیری ماشین با سرعت و دقت بیشتری همگرا شود و از تأثیر نامتناسب ویژگی‌هایی با مقادیر بزرگ‌تر جلوگیری شود. در پایان برای ارزیابی عملکرد مدل، از معیارهای ریشه میانگین مربعات خطا و میانگین قدر مطلق خطا در خصوص داده بخش آزمون استفاده شده است. مقادیر هاپرپارامترها از طریق رویکرد جست‌وجوی شبکه‌ای به دست آمده که نتایج به شرح زیر است:

جدول ۱ مقادیر و توضیحات هاپرپارامترها در مدل الگوریتم تقویت گرادیان

مقدار	هاپرپارامتر	توضیح
۱۰۰	تعداد درخت‌ها	تعداد کل درخت‌های تصمیم که مدل باید بسازد. مقدار بیشتر ممکن است دقت را بالا ببرد، اما خطر بیش برآزش را هم افزایش می‌دهد.
۲	حداکثر عمق درخت	حداکثر عمقی که هر درخت می‌تواند رشد کند. درخت‌های عمیق‌تر می‌توانند الگوهای پیچیده‌تری یاد بگیرند اما ممکن است مدل دچار بیش برآزش شود.
۰.۸	نسبت نمونه برداری از داده‌ها	درصدی از داده‌های آموزشی که برای آموزش هر درخت استفاده می‌شود (بین ۰ و ۱). مقدار کمتر باعث کاهش هم‌وابستگی درخت‌ها و کمک به جلوگیری از بیش برآزش می‌شود.
۰.۸	نسبت ویژگی‌ها در هر درخت	درصد ویژگی‌ها که برای ساخت هر درخت تصادفی انتخاب می‌شوند. کاهش این مقدار باعث افزایش تنوع در درخت‌ها می‌شود که به کاهش بیش برآزش کمک می‌کند.

1 Lag Length Criteria

2 Temporal Dependency

3 Rolling windows

دو مدل نخست در چارچوب یادگیری عمیق تعریف می‌شوند. دلیل انتخاب این دو مدل، منطق و ساختار معماری آن‌هاست که مبتنی بر زنجیره‌ای از اطلاعات و مدیریت حافظه است؛ ویژگی‌ای که با ماهیت مسئله‌ی مورد مطالعه یعنی مدل‌سازی منحنی بازده و وابستگی اطلاعات در طول زمان، کاملاً هم‌راستا و سازگار است. علاوه بر این، برخلاف مدل‌هایی که بر پایه‌ی الگوریتم‌های بهینه‌سازی گرادیانی تقویتی طراحی شده‌اند، این مدل‌ها از قابلیت مهندسی خودکار ویژگی‌ها بهره‌مند هستند.

مدل سوم، ترکیبی از شبکه عصبی پیچشی و حافظه طولانی کوتاه‌مدت است. بخش شبکه عصبی پیچشی به تحلیل الگوهای مکانی یا روابط کوتاه‌مدت می‌پردازد، در حالیکه حافظه طولانی کوتاه‌مدت بر جنبه‌ی زمانی و روابط طولانی مدت تمرکز دارد. داده‌های ورودی، مشابه سایر مدل‌ها، شامل سری‌های زمانی عامل‌های استخراج‌شده از مدل نلسون-سیگل پویای برآوردی است. این داده‌ها در ابتدا به سه مجموعه‌ی آموزش، اعتبارسنجی و آزمون تقسیم شده‌اند داده آموزش با استفاده از آزمون Z نرمال‌سازی شد. تابع فعال‌سازی^۱ مورد استفاده تابع نمایی خطی مقیاس بندی شده^۲ و تابع بهینه‌سازی نیز AdamW است. مسئله بیش برآزش یک موضوع مهم است که جهت جلوگیری از بروز آن چند راهکار از جمله نرخ حذف تصادفی نورون‌ها^۳، نرمال‌سازی دسته‌ای^۴، توقف زود هنگام در صورتیکه تابع زیان اعتبارسنجی آخرین گام از میانگین خود بیشتر شود البته با لحاظ دوره ۱۰ گامه صبر، لحاظ شد. برای دستیابی به بهترین نرخ یادگیری^۵، یک برنامه تعریف شده و باهدف جلوگیری از انفجار گرادیان از برش (قطع) گرادیان^۶ باهدف کاهش گرادیان و جلوگیری از انفجار استفاده شد. مشابه سایر مدل‌ها ارزیابی مدل با استفاده از سنج‌های ریشه میانگین مربعات خطا و میانگین قدر مطلق خطا عملکرد مدل انجام شد. مشخصات مدل از جمله اندازه دسته^۷ معادل ۶۴ و تعداد دوره آموزش^۸ ۳۰۰ است. مقادیر هایپر پارامترها نیز از طریق رویکرد جست‌وجوی شبکه‌ای به دست آمده که نتایج به شرح زیر است:

جدول ۲ مقادیر و توضیحات هایپر پارامترها در مدل شبکه عصبی پیچشی - حافظه طولانی کوتاه‌مدت

مقدار	هایپر پارامتر	توضیح
۳۲	اندازه نورون لایه پنهان ۱۰	تعداد نورون‌ها در لایه پنهان حافظه طولانی کوتاه‌مدت، مقدار بیشتر آن به مدل توانایی بیشتری برای یادگیری الگوهای پیچیده می‌دهد اما ممکن است باعث بیش برآزش شود.
۱	تعداد لایه پنهان در هر دو مدل	تعداد لایه‌های حافظه طولانی کوتاه‌مدت و شبکه عصبی پیچشی. لایه‌های بیشتر می‌توانند روابط پیچیده‌تری را یاد بگیرند، اما یادگیری را دشوارتر و زمان‌برتر می‌کنند.
۳	اندازه کرنل (هسته) در شبکه عصبی پیچشی ۱۱	اندازه فیلترهای پیچشی در شبکه عصبی پیچشی، اندازه کوچک‌تر (مثل ۳×۳) برای استخراج ویژگی‌های محلی مناسب است، اندازه بزرگ‌تر (مثل ۵×۵) اطلاعات گسترده‌تری را پوشش می‌دهد.
۳۲	تعداد فیلتر در شبکه عصبی پیچشی	تعداد فیلترهای لایه شبکه عصبی پیچشی، هر فیلتر ویژگی خاصی از داده را استخراج می‌کند. افزایش تعداد فیلترها به معنای استخراج ویژگی‌های متنوع‌تر است.
۰,۲	نرخ حذف تصادفی نورون‌ها ۱۲	درصد نورون‌هایی که در طول آموزش به صورت تصادفی غیرفعال می‌شوند تا از بیش برآزش جلوگیری شود. معمولاً بین ۰,۱ تا ۰,۵ تنظیم می‌شود.

1 Activation Function

2 Scaled Exponential Linear Unit

3 Dropout Rate

4 Batch Normalization

5 Learning Rate

6 Gradient Clipping

7 Batch_Size

8 Epoch

^۹ تعداد اجرای کامل فرآیند آموزش است، یک دوره یعنی مدل یک بار تمام نمونه‌های آموزش را دیده است؛ به عنوان مثال این مدل ۳۰۰ بار یک کتاب را خوانده است تا به طور کامل مفهوم آنرا متوجه شود. تعداد بیشتر دوره امکان برآورد مناسب پارامترها را فراهم میکند و امکان فراگیری الگوهای ظریف را خواهد داشت و بالقوه به صحت پیش بینی بالاتری دست خواهد یافت اما می‌تواند به بیش برآزش منجر شود زمان بیهوده صرف آموزش شود.

10 Hidden Size

11 Kernel-size

12 Dropout_Rate

مدل نهایی، واحد بازگشتی دارای دروازه است و داده‌های ورودی آن، مشابه سایر مدل‌ها، شامل سری زمانی عامل‌های برآورد شده از مدل نلسون-سیگل پویا هستند. داده‌ها در ابتدا به سه بخش آموزش، اعتبارسنجی و آزمون تقسیم شده‌اند و در ادامه داده آموزش با استفاده از آزمون Z نرمال‌سازی شد. برای افزایش قابلیت تعمیم مدل و جلوگیری از بیش‌برازش، همانند سایر مدل‌ها، از چندین تکنیک استفاده شده است؛ از جمله نرخ حذف تصادفی نورون‌ها برابر با ۰/۱، نرمال‌سازی دسته‌ای و استفاده از مکانیسم توقف زود هنگام در صورتی که مقدار تابع زیان اعتبارسنجی در گام آخر از میانگین آن در دوره‌ای ۱۰ مرحله‌ای بیشتر شود. تابع فعال‌سازی تابع نمایی خطی مقیاس بندی شده و تابع بهینه‌سازی نیز AdamW است. در خصوص دستیابی به بهترین نرخ یادگیری، یک برنامه تعریف شد و مشابه سایر مدل‌ها از سنج‌های ریشه میانگین مربعات خطا و میانگین قدر مطلق خطا جهت سنجش عملکرد مدل استفاده شد. مقادیر هایپرپارامترها نیز از طریق رویکرد جست‌وجوی شبکه‌ای به دست آمده که نتایج به شرح زیر است:

جدول ۳ مقادیر و توضیحات هایپرپارامترها در واحد بازگشتی دارای دروازه

مقدار	هایپر پارامتر	توضیح
۱۲۸	اندازه نورون لایه پنهان ۱	تعداد نورون‌ها در لایه پنهان مدل، مقدار بیشتر آن به مدل توانایی بیشتری برای یادگیری الگوهای پیچیده می‌دهد اما ممکن است باعث بیش‌برازش شود.
۱	تعداد لایه پنهان مدل	لایه‌های بیشتر می‌توانند روابط پیچیده‌تری را یاد بگیرند، اما یادگیری را دشوارتر و زمان‌برتر می‌کنند.
۰,۱	نرخ حذف تصادفی نورون‌ها ۲	درصد نورون‌هایی که در طول آموزش به صورت تصادفی غیرفعال می‌شوند تا از بیش‌برازش جلوگیری شود. معمولاً بین ۰,۱ تا ۰,۵ تنظیم می‌شود.

سنجه ارزیابی مدل‌ها: جهت ارزیابی مدل‌ها از دو سنجه خطای ریشه میانگین مربعات خطا^۳ و میانگین قدر مطلق خطا^۴ استفاده شده است که فرمول آن‌ها به شرح زیر است:

$$RMSE = \sqrt{\frac{\sum (y_i - y_i^A)^2}{N}} \quad (۶)$$

$$MAE = \frac{1}{2} \sum_{i=1}^n |x_i - xI| \quad (۷)$$

پیش‌پردازش داده‌ها و محاسبه بازده تا سررسید در نرم‌افزار اکسل، محاسبات خود رگرسیون برداری-گارچ در نرم‌افزار ایویوز ۱۳ و برآورد در مدل‌های الگوریتم مبتنی بر تقویت گرادیان، شبکه عصبی پیچشی - حافظه طولانی کوتاه مدت و مدل واحد بازگشتی دارای دروازه توسط زبان برنامه‌نویسی پایتون و با استفاده از مجموعه کتابخانه یا فرمول‌ها صورت پذیرفته است.

۴. تحلیل داده‌ها و یافته‌ها

پس از محاسبه بازده تا سررسید، از این مقادیر برای برآورد عامل‌های مدل نلسون-سیگل پویا استفاده شد. به منظور جلوگیری از برآورد مقادیر غیرواقع بینانه، همانند رویه رایج در مطالعات پیشین، در ابتدا محدوده‌ای برای تخمین عامل‌ها تعریف گردید. در این راستا، مدل با استفاده از مقادیر اولیه برآورد شد و سپس با در نظر گرفتن حداقل و حداکثر مقادیر به همراه انحراف معیار مشخص، دامنه‌ای برای عامل‌ها تعیین شد؛ با این حال، برای مؤلفه‌ی سطح نیازی به اعمال محدودیت وجود نداشت. در گام بعد، با به کارگیری ماسک برای مدیریت داده‌های مفقود، عامل‌ها با استفاده از روش حداقل مربعات معمولی تخمین زده شدند. جهت دستیابی به مقادیر بهینه چندین متد بهینه‌سازی با

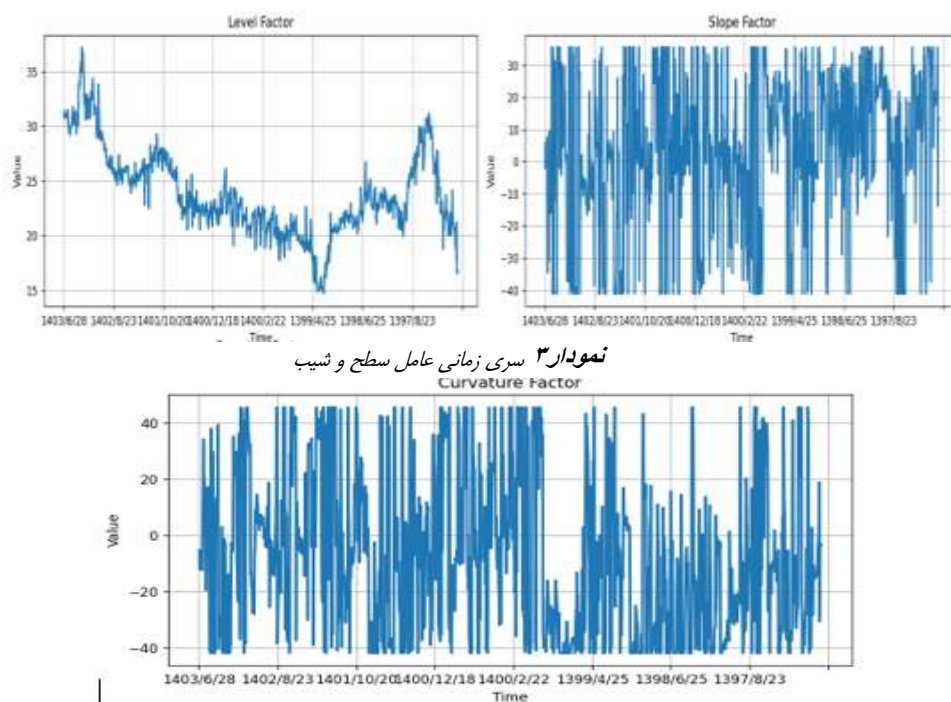
1 Hidden Size

2 Dropout_Rate

3 RMSE

4 MAE

یکدیگر مقایسه شدند که در نهایت بر اساس منطق ریشه میانگین مربعات خطا متد الگوریتم برنامه‌ریزی مربعات کمینه ترتیبی^{۲۱} به عنوان متد بهینه انتخاب شد. این فرآیند به صورت روزانه تکرار شد تا در نهایت، سری زمانی عامل‌ها به دست آید. اگرچه امکان برآورد این عامل‌ها در نرم‌افزار اکسل نیز وجود داشت، اما به منظور بهره‌گیری از انعطاف‌پذیری بیشتر در تنظیمات مدل، به‌ویژه در انتخاب روش‌های بهینه سازی برای دستیابی به بهترین پارامترها، تخمین مدل‌ها در محیط زبان برنامه‌نویسی پایتون انجام گرفت. نتایج حاصل، سری زمانی عامل‌ها را مطابق با نمودارهای ۳ و ۴ نشان می‌دهد:



در نهایت آمار توصیفی از عامل‌های برآورد شده مطابق مدل به شرح جدول ۴ ارائه شده است:

جدول ۴ آمار توصیفی عامل‌های محاسبه شده

عنوان	سطح	شیب	انحنا
میانگین	۲۳/۸۷	۵/۸۰	-۸/۶۳
انحراف معیار	۳/۹۵	۱۸/۹۱	۲۴/۱۳
حداقل	۱۳/۳۱	-۴۱/۱۱	-۴۱/۷۳
چارک اول	۲۱/۲۲	-۲/۲۹	-۲۸/۵۹
میانه	۲۳/۰۹	۳/۷۰	-۷/۱۹
چارک سوم	۲۶/۲۲	۱۹/۶۰	۳/۱۵
حداکثر	۳۷/۲۸	۳۵/۵۵	۴۵/۵۰

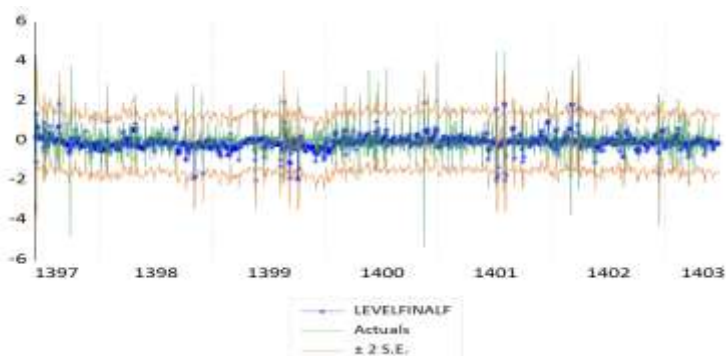
بر اساس جدول ۴، انحراف معیار مؤلفه‌های شیب و انحنا نسبتاً بالا بوده و این عامل‌ها دارای مقادیر مثبت و منفی هستند؛ بنابراین مدل‌سازی و پیش‌بینی آن‌ها نسبت به مؤلفه سطح، که تنها شامل مقادیر مثبت با پراکندگی کمتر است، به مراتب دشوارتر خواهد بود.

در ادامه، نتایج پیش‌بینی هر یک از عامل‌ها توسط مدل‌های مختلف ارائه شده است.

¹ The Sequential Least Squares Programming algorithm

^۲ این متد به دنبال حل مسائل بهینه‌سازی غیرخطی است و از اطلاعات گزاردیان استفاده می‌کند تا بهترین راه حل را ارائه دهد.

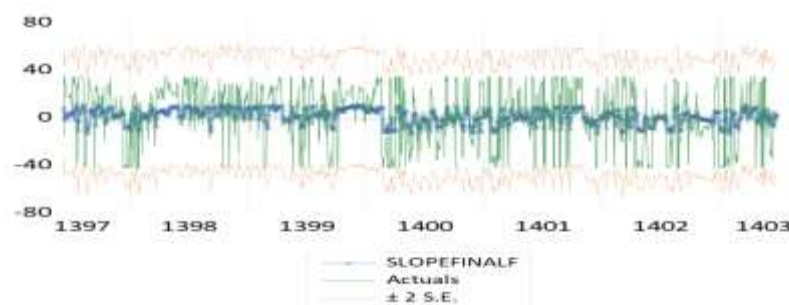
مدل خود رگرسیون برداری: برای اجرای مدل خود رگرسیون برداری، ابتدا لازم است وضعیت مانایی متغیرها بررسی شود. با ترسیم نمودار عامل‌ها و به‌ویژه با بهره‌گیری از آزمون دیکی-فولر، مشخص شد که عامل اول (سطح) مانا نیست، درحالی‌که دو عامل دیگر یعنی شیب و انحنا، فاقد مشکل مانایی بوده و از این نظر قابل استفاده در مدل هستند. مطابق نتیجه، ابتدا الگوی تصحیح خطای برداری^۱ انتخاب شد. در این راستا ضمن برآورد بهترین وقفه، از آزمون جوهانسون جهت انتخاب مرتبه صحیح هم انباشستگی استفاده شد و در ادامه ضمن بررسی نکات مربوط به ضرایب مؤلفه رابطه طولانی مدت و مؤلفه تصحیح خطای^۲ روابط کوتاه‌مدت مشخص شد دو مورد از آن‌ها به‌طور معنادار مثبت هستند که عامل تشدید فاصله از میانگین طولانی مدت و در نتیجه به وضعیت بدتر مؤلفه‌ها به‌مرور زمان منجر خواهد شد. علاوه بر این، با انجام آزمون‌های مختلف بر روی جزء اخلاص مشخص شد که مدل منتخب از نظر خودهمبستگی سریالی مشکلی ندارد، اما با مسئله ناهمسانی واریانس مواجه است. همچنین، آزمون علیت گرنجر نشان داد که هیچ‌یک از مؤلفه‌ها عامل علیت سایر مؤلفه‌ها نیستند. بر اساس نتایج به‌دست آمده، برای رفع مشکل نا مانایی سطح، از راهکار دوم یعنی انجام تفاضل استفاده شد و پس از تکرار آزمون دیکی-فولر برای عامل سطح، پایایی آن تأیید گردید. با بررسی سنجح طول وقفه، وقفه ۲ به‌عنوان وقفه مناسب برای مدل انتخاب شد. در ادامه، آزمون‌های مربوط به جزء اخلاص نشان داد که مشابه الگوی تصحیح خطای برداری، مسئله ناهمسانی واریانس همچنان باقی است و باوجود اتخاذ چندین راهکار از جمله تغییر طول وقفه، به دلیل شدت زیاد مشکل، برطرف نشد. همچنین، آزمون علیت گرنجر بیانگر وجود علیت معنادار و قابل تأیید متغیرها نسبت به سطح و شیب است. در نهایت جهت انتخاب مدل بهینه بین الگوی تصحیح خطای برداری و خود رگرسیون برداری از عامل‌های AIC, SC, Log likelihood استفاده شد که با تفاوت فراوان مدل خود رگرسیون برداری نتایج بهتری داشت. در نهایت مدل ترکیبی خود رگرسیون برداری-گارچ با انتخاب مدل خود رگرسیون برداری به‌عنوان تابع میانگین و گارچ به‌عنوان تابع واریانس انتخاب شد. شایان ذکر است جهت انتخاب تابع واریانس از انواع مدل‌ها مانند EGARCH, TGARCH با وقفه‌های مختلف استفاده شد که بر اساس معناداری ضرایب و انجام آزمون جزء اخلاص و پایداری نتایج، در نهایت مدل گارچ با وقفه (۱،۱) انتخاب گردید. در ادامه نمودارهای مقایسه مقادیر پیش‌بینی شده و واقعی برای هر یک از ۳ عامل آمده است:



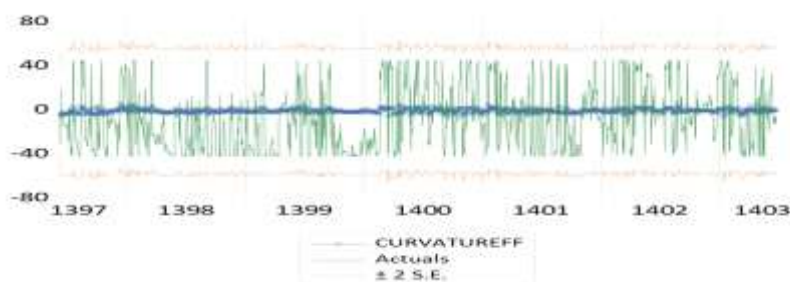
نمودار ۵ مقایسه مقادیر واقعی و پیش‌بینی سطح

1 VECM

2 Error Correction Term

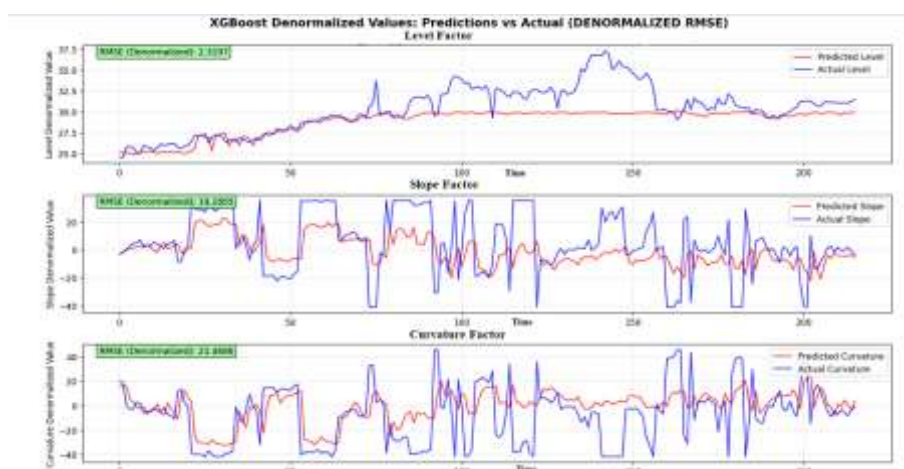


نمودار ۶ مقایسه مقادیر واقعی و پیش‌بینی شیب



نمودار ۷ مقایسه مقادیر واقعی و پیش‌بینی انحنا

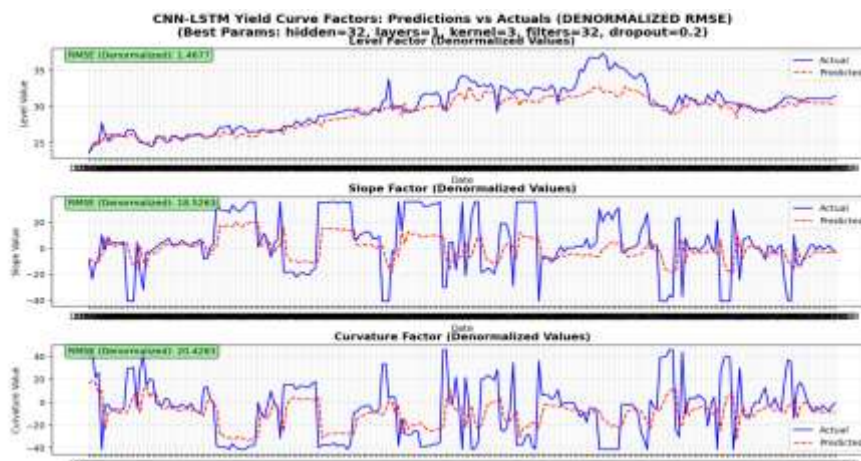
مدل الگوریتم مبتنی بر تقویت گرادیان: ورودی مدل مانند سایر مدل‌ها، سری زمانی تاریخی عامل‌های برآورد شده مدل نلسون-سیگل پویاست. ابتدا داده‌ها مدیریت و تبدیل شده و پس از نرمال‌سازی، برای آموزش و استفاده در مدل منتخب یادگیری ماشین آماده می‌شوند. در نهایت عامل‌ها برآورد شده‌اند که نتایج مقایسه مقادیر واقعی و پیش‌بینی هر یک از عامل‌ها در نمودار ۸ نشان داده شده است.



نمودار ۸ مقایسه مقادیر واقعی و پیش‌بینی عامل‌ها ذیل مدل مبتنی بر تقویت گرادیان

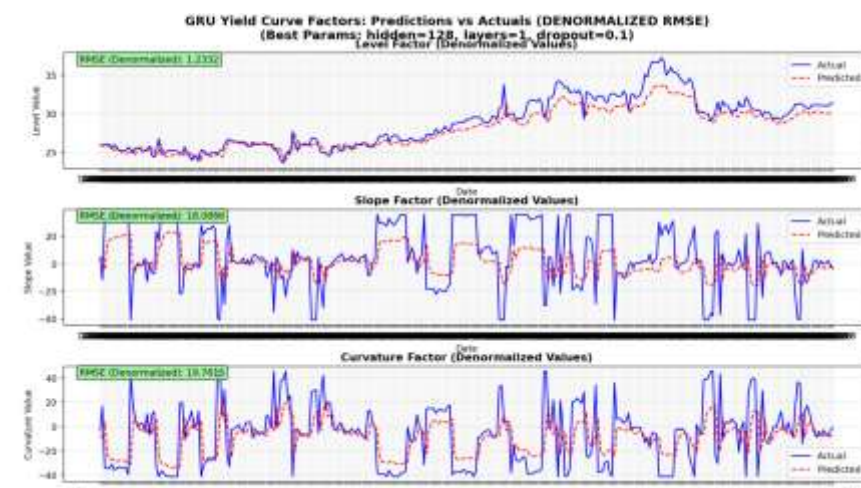
مدل شبکه عصبی پیچشی - حافظه طولانی کوتاه مدت: ورودی مدل، سری زمانی تاریخی عامل‌های برآورد شده مدل نلسون-سیگل پویاست. ابتدا داده‌های اولیه به منظور سازگاری با لایه تک‌بعدی شبکه عصبی پیچشی، تغییر قالب داده شدند و سپس برای پردازش و استخراج ویژگی‌های داخلی، وارد این لایه شدند. در مرحله بعد، به منظور نرمال‌سازی خروجی لایه پیچشی، تثبیت فرآیند آموزش و بهبود عملکرد مدل، از روش نرمال‌سازی دسته‌ای استفاده شد و تابع فعال‌سازی تابع نمایی خطی مقیاس بندی شده است که برای ایجاد مؤلفه غیرخطی در مدل به کار گرفته شد. سپس داده‌ها در لایه حافظه طولانی کوتاه مدت قرار گیرند. این لایه، آرایه‌های ورودی را

پردازش و خروجی تولید می کند و به طور همزمان، مؤلفه نرخ حذف تصادفی نورون ها جهت مقابله با بیش برآزش، به صورت تصادفی بخشی از واحدهای خروجی را در طول آموزش غیرفعال می کند. در نهایت، داده ها از یک لایه کاملاً متصل عبور کرده و خروجی نهایی تولید می شود. هدف کلی از طراحی این مدل ترکیبی، آموزش الگوهای داخلی داده ها با استفاده از شبکه عصبی پیچشی برای استخراج ویژگی ها و یادگیری وابستگی های طولانی مدت توسط لایه حافظه طولانی کوتاه مدت بود. در ادامه، تصویری از نتایج مدل بر اساس مقایسه مقادیر پیش بینی شده با مقادیر واقعی هر عامل در نمودار شماره ۹ ارائه شده است.



نمودار ۹ مقایسه مقادیر واقعی و پیش بینی عامل ها ذیل مدل شبکه عصبی پیچشی - حافظه طولانی کوتاه _ مدت

مدل واحد بازگشتی دارای دروازه: این مدل علاوه بر حفظ مزایای حافظه طولانی کوتاه مدت، از ساختار ساده تری برخوردار است؛ به طوری که دو دروازه برای مدیریت اطلاعات تعریف شده اند تا میزان فراموشی اطلاعات گذشته و نگهداری اطلاعات جاری کنترل شود. ورودی مدل مشابه سایر مدل ها است و داده های اولیه پس از نرمال سازی، پیش پردازش می شوند. این مدل شامل یک لایه بازگشتی واحد دارای دروازه است که به دنبال آن لایه نرمال سازی دسته ای و دو لایه خطی قرار دارند؛ لایه اول خطی به همراه تابع فعال سازی Tanh برای افزودن مؤلفه غیرخطی و لایه دوم به منظور پیش بینی نهایی استفاده می شود. در نمودار ۱۰، مقادیر پیش بینی شده در مقایسه با مقادیر واقعی هر عامل نشان داده شده است.



نمودار ۱۰ مقایسه مقادیر واقعی و پیش بینی عامل ها ذیل مدل واحد بازگشتی دارای دروازه

مقایسه مدل‌ها : در پایان، با بهره‌گیری از دو معیار ریشه میانگین مربعات خطا و میانگین قدر مطلق خطا، عملکرد پیش‌بینی عامل‌ها توسط مدل‌های مختلف مورد مقایسه قرار گرفت که نتایج آن در جدول شماره ۵ ارائه شده است:

جدول ۵ مقایسه توان پیش‌بینی انواع مدل‌ها

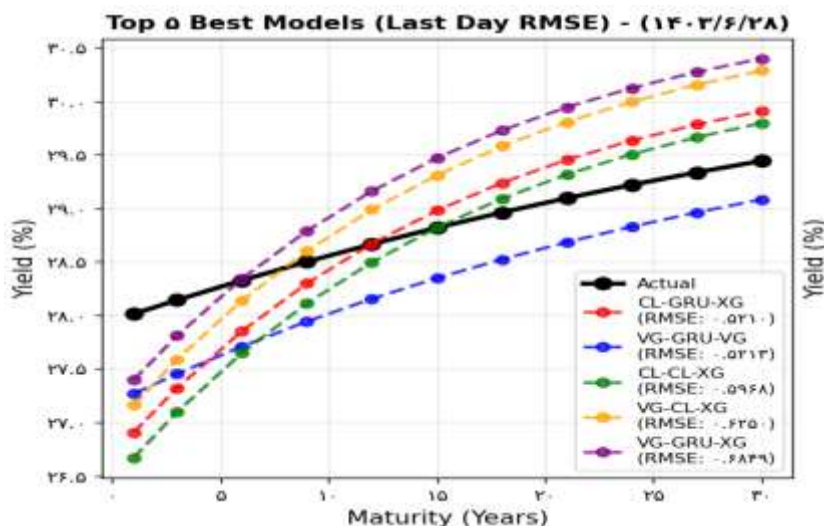
مدل	ریشه میانگین مربعات خطا	میانگین قدر مطلق خطا
سطح		
خود رگرسیون برداری-گارچ (۱،۱)	۰/۶۹۰	۰/۴۴۹
واحد بازگشتی دارای دروازه	۱/۲۳۳	۰/۹۳۹
شبکه عصبی پیچشی - حافظه طولانی کوتاه‌مدت	۱/۴۶۷	۱/۰۲۰
الگوریتم مبتنی بر تقویت گرادیان	۲/۳۱۹	۱/۵۷۹
شیب		
خود رگرسیون برداری-گارچ (۱،۱)	۲۱/۶۰۰	۱۷/۱۱۰
واحد بازگشتی دارای دروازه	۱۸/۰۸۸	۱۳/۱۱۴
شبکه عصبی پیچشی - حافظه طولانی کوتاه‌مدت	۱۸/۵۲۶	۱۳/۴۳۰
الگوریتم مبتنی بر تقویت گرادیان	۱۹/۲۸۵	۱۵/۰۰۷
انحنا		
خود رگرسیون برداری-گارچ (۱،۱)	۲۷/۹۰۰	۲۳/۹۶۰
واحد بازگشتی دارای دروازه	۱۹/۷۶۱	۱۳/۱۲۲
شبکه عصبی پیچشی - حافظه طولانی کوتاه‌مدت	۲۰/۴۲۶	۱۳/۹۸۹
الگوریتم مبتنی بر تقویت گرادیان	۲۱/۴۶۸	۱۶/۱۴۷

مطابق جدول ۵، یافته‌ها نشان می‌دهد که در مورد عامل سطح، که دارای روندی طولانی مدت است، مدل سنتی خود رگرسیون برداری -گارچ عملکرد بهتری نسبت به سایر مدل‌ها از خود نشان داد. این برتری را می‌توان به ماهیت خودرگرسیو و توانایی این مدل در مدل‌سازی واریانس شرطی نسبت داد، که سبب می‌شود برای داده‌های دارای روند پایدار و نوسانات ساختاری مناسب‌تر باشند. در مقابل، مدل‌های یادگیری عمیق مانند واحد بازگشتی دارای دروازه و شبکه عصبی پیچشی - حافظه طولانی کوتاه‌مدت در مدل‌سازی این عامل عملکرد ضعیف‌تری داشتند؛ که احتمالاً به دلیل دشواری معماری‌های شبکه عصبی در شناسایی روندهای طولانی مدت در حالت محدودیت داده است.

در خصوص دو عامل شیب و انحنا، که معمولاً در افق‌های زمانی کوتاه‌مدت و میان‌مدت دارای نوسانات موقتی هستند، مدل‌های یادگیری عمیق عملکرد بهتری نسبت به مدل سنتی خود رگرسیون برداری - گارچ داشته‌اند. علت این برتری آن است که مدل‌های آماری کلاسیک با ساختار خطی و مفروضات محدودکننده، در برابر رفتارهای غیر ایستا یا نوسانات پیچیده کوتاه‌مدت، دچار بایاس می‌شوند. حتی در مقایسه با مدل‌های یادگیری سطحی، مدل‌های یادگیری عمیق نتایج دقیق‌تری ارائه داده‌اند. به‌طور خاص، مدل واحد بازگشتی دارای دروازه در هر دو عامل شیب و انحنا نسبت به سایر مدل‌ها عملکرد برتری داشته است. این نشان می‌دهد که معماری‌های مبتنی بر حافظه در یادگیری عمیق، توانایی بالایی در شناسایی و پیش‌بینی روابط غیرخطی و الگوهای پیچیده و متغیر در طول زمان را دارند.

در مرحله دوم، مقادیر پیش‌بینی شده برای سه عامل سطح، شیب و انحنا در معادله نلسون-سیگل پویا جای‌گذاری شدند تا منحنی بازده بازسازی گردد، و دقت این بازسازی با استفاده از معیار ریشه میانگین مربعات خطا در مقایسه با مقادیر واقعی مورد ارزیابی قرار گرفت. بر اساس نتایج جدول ۵ نتایج نشان داد که هیچ‌یک از مدل‌ها به‌تنهایی قادر

به ارائه بهترین عملکرد برای هر سه عامل به صورت همزمان نیست. بنابراین، استفاده از ترکیبی از مدل‌ها به گونه‌ای که هر عامل به صورت مستقل و توسط مدلی با کمترین میزان خطا برآورد شود، می‌تواند دقت بازسازی منحنی بازده را بهبود بخشد؛ چرا که این رویکرد با ساختار مدل نلسون-سیگل، مبتنی بر فرض استقلال عامل‌ها از یکدیگر، نیز سازگار است. با این حال، باید توجه داشت که چنین ترکیبی لزوماً منجر به مدل بهینه در بازسازی نهایی منحنی بازده نخواهد شد و باید تمام حالات ممکن بررسی شود. با بررسی نتایج حاصل از ۶۴ حالت، ۵ ترکیب برتر به شرح تصویر زیر برای آخرین روز ارائه شده است:



نمودار ۱۱ پیش‌بینی منحنی بازده آخرین روز بر اساس ۵ مدل برتر

بر اساس نمودار ۱۱، استفاده از ترکیبی از مدل‌ها به گونه‌ای که هر عامل به صورت مستقل و توسط مدلی با کمترین میزان خطا تخمین زده شود، دقت بازسازی منحنی بازده را در مقایسه با حالتی که تمامی عامل‌ها با یک مدل واحد برآورد شوند، افزایش داده است؛ با این حال، این ترکیب لزوماً به بهترین نتیجه منجر نشده و ترکیب بهینه مدل‌ها برای بازسازی منحنی، متفاوت بوده است. علت این امر به ضرایب متفاوت هر عامل در فرمول نلسون-سیگل بازمی‌گردد، چراکه این ضرایب به صورت تابعی از سررسید در معادله ظاهر می‌شوند؛ به طوری که ضریب عامل سطح در کلیه سررسیدها مقدار ثابتی دارد، در حالی که ضریب عامل شیب در سررسیدهای کوتاه مدت بیشترین تأثیر را داشته و عامل انحنا در سررسیدهای میان مدت نقش برجسته‌تری ایفا می‌کند. افزون بر این، در برخی موارد ممکن است مدلی که در پیش‌بینی هر سه عامل دارای خطاهای متوسط است، به دلیل جهت‌گیری متضاد این خطاها، اثرات آن‌ها را خنثی کرده و در نهایت خطای کلی منحنی بازده را کاهش دهد. در مقابل، مدلی که عامل‌ها را با دقت بالاتر پیش‌بینی می‌کند اما خطاهایی با جهت‌گیری مشابه تولید می‌نماید، می‌تواند منجر به انباشت خطا و افزایش انحراف کل منحنی شود. در ادامه خطای مربوط به تمام حالات ممکن به شرح جدول ۶ آمده است. مطابق جدول مذکور بهترین نتیجه مربوط به ترکیبی است که سطح با مدل خود رگرسیون برداری - گارچ یا شبکه عصبی پیچشی - حافظه طولانی کوتاه مدت، شیب با واحد بازگشتی دارای دروازه و انحنا با مدل خود رگرسیون برداری - گارچ یا الگوریتم مبتنی بر تقویت گرادیان برآورد شوند که انحراف از واقعیت معادل حدود نیم درصد است.

جدول ۶ میزان ریشه میانگین مربعات خطا منحنی بازده آخرین روز در تمام حالات

سطح	شیب	انحنا			
		شبکه عصبی پیچشی - حافظه طولانی کوتاه مدت	واحد بازگشتی دارای دروازه	خود رگرسیون برداری-گارچ (۱،۱)	الگوریتم مبتنی بر تقویت گرادیان
شبکه عصبی پیچشی - حافظه طولانی کوتاه مدت	شبکه عصبی پیچشی - حافظه طولانی کوتاه مدت	۲/۸۸۴	۲/۵۸۵	۱/۱۷۴	۰/۵۹۶
	واحد بازگشتی دارای دروازه	۱/۷۲۹	۲/۴۲۸	۱/۰۰۴	۰/۵۲۱
	خود رگرسیون برداری-گارچ (۱،۱)	۱/۹۸۰	۲/۰۹۵	۳/۰۶۷	۳/۰۶۷
	الگوریتم مبتنی بر تقویت گرادیان	۳/۸۱۸	۳/۵۲۸	۲/۱۶۹	۱/۴۲۰
واحد بازگشتی دارای دروازه	شبکه عصبی پیچشی - حافظه طولانی کوتاه مدت	۳/۳۵۴	۳/۰۵۸	۱/۶۴۸	۰/۸۸۵
	واحد بازگشتی دارای دروازه	۳/۱۹۷	۲/۸۹۹	۱/۴۷۹	۰/۷۳۵
	خود رگرسیون برداری-گارچ (۱،۱)	۱/۷۱۱	۱/۷۶۸	۲/۶۱۱	۳/۴۵۸
	الگوریتم مبتنی بر تقویت گرادیان	۴/۲۹۳	۴/۰۰۴	۲/۶۴۰	۱/۸۳۵
خود رگرسیون برداری-گارچ (۱،۱)	شبکه عصبی پیچشی - حافظه طولانی کوتاه مدت	۲/۴۰۴	۲/۱۰۳	۰/۶۹۲	۰/۶۲۵
	واحد بازگشتی دارای دروازه	۲/۲۵۱	۱/۹۴۹	۰/۵۲۱	۰/۶۸۴
	خود رگرسیون برداری-گارچ (۱،۱)	۲/۳۲۷	۲/۴۸۲	۳/۵۴۱	۴/۴۱۹
	الگوریتم مبتنی بر تقویت گرادیان	۳/۳۲۹	۳/۰۳۸	۱/۶۹۰	۱/۰۴۸
الگوریتم مبتنی بر تقویت گرادیان	شبکه عصبی پیچشی - حافظه طولانی کوتاه مدت	۳/۵۲۴	۳/۲۲۷	۱/۸۱۸	۱/۰۲۴
	واحد بازگشتی دارای دروازه	۳/۳۶۶	۳/۰۶۸	۱/۶۴۹	۰/۸۶۵
	خود رگرسیون برداری-گارچ (۱،۱)	۱/۶۳۸	۱/۶۶۸	۲/۴۴۹	۳/۲۸۹
	الگوریتم مبتنی بر تقویت گرادیان	۴/۴۶۴	۴/۱۷۵	۲/۸۰۹	۱/۹۹۱

۵. بحث و نتیجه‌گیری

منحنی بازده نمایانگر رابطه بین بازده اوراق بهادار با درآمد ثابت و سررسیدهای مختلف است؛ اوراقی که از نظر ریسک، نقد شوندگی و ویژگی‌های مالیاتی مشابه‌اند. این منحنی ابزار مهمی برای تحلیل انتظارات بازار نسبت به سیاست‌های پولی، روند فعالیت‌های اقتصادی و پیش‌بینی تورم در افق‌های زمانی کوتاه‌مدت، میان‌مدت و طولانی مدت محسوب می‌شود. همچنین در تدوین سیاست‌های مالی دولت، مدیریت کسب‌وکار نهادهای مالی مانند بانک‌ها و اتخاذ تصمیمات سرمایه‌گذاری از قبیل ارزش‌گذاری دارایی‌ها، مدیریت پرتفوی و کنترل ریسک کاربرد دارد. در این میان، ارائه مدلی که بتواند پیش‌بینی دقیق، به‌روز و کارآمدی از منحنی بازده ارائه دهد، به‌ویژه در اقتصادهای نوظهوری مانند ایران، اهمیت مضاعفی دارد. با این حال، تاکنون این موضوع در ایران موردتوجه شایسته قرار نگرفته و مقاله حاضر درصدد است تا این خلأ را پوشش دهد.

بر اساس مطالعات صورت گرفته، مجموعه‌ای از مدل‌ها و رویکردها برای پیش‌بینی منحنی بازده معرفی شده‌اند. در این میان، مدل عاملی نلسون-سیگل پویا به دلیل ویژگی‌هایی نظیر قابلیت تفسیرپذیری، انعطاف‌پذیری بالا، توانایی تلخیص اطلاعات و کاهش ابعاد داده‌های ورودی، و همچنین کاربرد گسترده آن توسط اکثر بانک‌های مرکزی جهان،

به عنوان گزینه‌ی مناسب انتخاب شده است. این مدل، منحنی بازده را بر مبنای سه مؤلفه‌ی اصلی سطح، شیب و انحنا برآورد می‌کند. در اغلب مطالعات پیشین، از مدل خود رگرسیون برداری برای برآورد این عامل‌ها و پیش‌بینی منحنی بازده به عنوان مدلی مبنا استفاده شد؛ اما در این پژوهش، علاوه بر مدل مذکور، مجموعه‌ای از روش‌های یادگیری ماشین -سطحی یا عمیق- نیز به کار گرفته شده است. جهت انجام محاسبات، ابتدا بازده تا سررسید اوراق خزانه اسلامی استخراج و سپس عامل‌های مدل نلسون-سیگل پویا برآورد شده‌اند. در ادامه پیش‌بینی از عامل‌ها با استفاده از چندین مدل از جمله خود رگرسیون برداری-گارچ (۱،۱)، واحد بازگشتی دارای دروازه، الگوریتم تقویت گرادیان و مدل ترکیبی شبکه عصبی پیچشی-حافظه طولانی کوتاه مدت جهت پیش‌بینی انجام شد. به منظور ارزیابی عملکرد مدل‌ها در پیش‌بینی هر یک از سه عامل، از دو شاخص میانگین قدر مطلق خطا و ریشه میانگین مربعات خطا استفاده شد. یافته‌ها حاکی از آن است که در پیش‌بینی عامل سطح که دارای روندی طولانی مدت است، مدل سنتی خود رگرسیون برداری-گارچ عملکرد بهتری نسبت به سایر مدل‌ها ارائه داده است. این برتری را می‌توان به ماهیت خودرگرسیو این مدل و توانایی آن در مدل‌سازی واریانس شرطی نسبت داد، که آن را برای داده‌های با روند پایدار و نوسانات ساختاری مناسب‌تر می‌سازد. در مقابل، مدل‌های یادگیری عمیق نظیر واحد بازگشتی دارای دروازه و شبکه عصبی پیچشی-حافظه طولانی کوتاه مدت در مدل‌سازی این عامل عملکرد ضعیف‌تری نشان داده‌اند که احتمالاً به دلیل چالش این شبکه‌ها در شناسایی روندهای طولانی مدت در شرایط محدودیت داده است.

در رابطه با دو عامل شیب و انحنا نیز که معمولاً در افق‌های زمانی کوتاه مدت و میان مدت با نوسانات موقت همراه‌اند، مدل‌های یادگیری عمیق نسبت به مدل سنتی خود رگرسیون برداری-گارچ عملکرد بهتری نشان داده‌اند. این برتری ناشی از آن است که مدل‌های آماری کلاسیک به دلیل ساختار خطی و مفروضات محدودکننده خود، در مواجهه با رفتارهای غیر ایستا و نوسانات پیچیده کوتاه مدت دچار سوگیری می‌شوند. حتی در مقایسه با مدل‌های یادگیری سطحی نیز، مدل‌های یادگیری عمیق نتایج دقیق‌تری ارائه داده‌اند. به ویژه، مدل بازگشتی دارای دروازه در هر دو عامل شیب و انحنا نسبت به سایر مدل‌ها عملکرد بهتری ثبت کرد.

در مرحله دوم، مقادیر پیش‌بینی شده برای سه عامل سطح، شیب و انحنا در معادله نلسون-سیگل پویا جای‌گذاری شدند تا منحنی بازده بازسازی گردد، و دقت این بازسازی با استفاده از معیار ریشه میانگین مربعات خطا در مقایسه با مقادیر واقعی مورد ارزیابی قرار گرفت. نتایج نشان داد که هیچ‌یک از مدل‌ها به تنهایی قادر به ارائه بهترین عملکرد برای هر سه عامل به صورت هم‌زمان نیست. بنابراین، استفاده از ترکیبی از مدل‌ها به گونه‌ای که هر عامل به صورت مستقل و توسط مدلی با کمترین میزان خطا برآورد شود، می‌تواند دقت بازسازی منحنی بازده را بهبود بخشد؛ چرا که این رویکرد با ساختار مدل نلسون-سیگل، مبتنی بر فرض استقلال عامل‌ها از یکدیگر، نیز سازگار است. با این حال، باید توجه داشت که چنین ترکیبی لزوماً منجر به مدل بهینه در بازسازی نهایی منحنی بازده نخواهد شد. در نتیجه، پس از برآورد تمام حالات ممکن، بهترین دقت در بازسازی منحنی بازده زمانی حاصل می‌شود که سطح با مدل خود رگرسیون برداری - گارچ یا شبکه عصبی پیچشی - حافظه طولانی کوتاه مدت، شیب با واحد بازگشتی دارای دروازه و انحنا با مدل خود رگرسیون برداری-گارچ یا الگوریتم مبتنی بر تقویت گرادیان برآورد شوند که انحراف از واقعیت معادل حدود نیم درصد است.

پیشنهادهای محدودیت‌ها : در زمینه مطالعات آتی، چند پیشنهاد قابل طرح است: نخست، به منظور افزایش تفسیرپذیری نتایج و بهره‌برداری سیاستی، از متغیرهای اقتصادی به عنوان ورودی مدل استفاده شود. دوم، از روش‌های یادگیری ماشین به صورت مستقیم برای پیش‌بینی مدل بهره گرفته شود. در نهایت، پیشنهاد می‌شود سایر الگوریتم‌های یادگیری ماشین به صورت ترکیبی با مدل‌های عاملی به کار گرفته شوند.

در خصوص محدودیت‌های پژوهش، باید توجه داشت که در سطح جهانی، منحنی بازده بر پایه اوراق خزانه‌داری دولت که با سازوکار مشخص و در بازه‌های زمانی معین منتشر می‌شوند، و همچنین با اتکای بر اجرای عملیات بازار باز شکل می‌گیرد؛ اما در ایران، اوراق اسناد خزانه اسلامی درواقع مطالبات پیمانکاران هستند که انتشار و عرضه آن‌ها طبق برنامه زمان‌بندی مشخصی انجام نمی‌شود. علاوه بر این، به دلیل نبود عملیات بازار باز، بازار با نوسانات قیمتی بالا و حجم معاملات پایین مواجه است. همچنین، وجود شکست ساختاری در مدل از دیگر محدودیت‌های موجود به شمار می‌رود. با این حال، با توجه به اهمیت موضوع، پیش‌بینی منحنی بازده همچنان ضرورتی جدی و قابل‌پیگیری دارد.

تعارض منافع. برای ارائه مطالب و نگارش این مقاله هیچ‌گونه کمک مالی از هیچ فرد، نهاد و سازمانی دریافت نشده است و نتایج و دستاوردهای این مقاله به نفع یا ضرر سازمان یا فردی خاص نخواهد بود. حضور نویسندگان در این پژوهش به عنوان شاهدی بی‌طرف ولی متخصص بوده است و نویسندگان هیچ‌گونه تعارض منافی ندارند.

Reference

- Argyropoulos, E., & Tzavalis, E. (2016). Forecasting economic activity from yield curve factors. *The North American Journal of Economics and Finance*, 36, 293-311.
- Ang, A., Piazzesi, M., & Wei, M. (2006). What does the yield curve tell us about GDP growth? *Journal of Econometrics*, 131(1-2), 359-371. <https://doi.org/10.1016/j.jeconom.2005.04.004>
- Bank for international settlements (2005) Zero-coupon yield curves: technical documentation. Monetary and economic department. No. 25.
- Bayaa, Y., & Qadan, M. (2024). The shape of the Treasury yield curve and commodity prices. *International Review of Financial Analysis*, 94, 103311
- Bianchi, M Büchner, A Tamoni (2021) Bond Risk Premiums with Machine Learning, *The Review of Financial Studies*, volume 34, p. 1046 – 1089.
- Bliss, R. R. (1996). Testing term structure estimation methods (No. 96-12a). Working Paper.
- Bolder, D.J. and Liu, S. (2007). Examining simple joint macroeconomic and term-structure models: a practitioner's perspective. *Bank of Canada working paper no. 2007-49*.
- Borio, C., Gambacorta, L., & Hofmann, B. (2017). The influence of monetary policy on bank profitability. *International finance*, 20(1), 48-63.
- Boukhatem, J., & Sekouhi, H. (2017). What does the bond yield curve tell us about Tunisian economic activity? *Research in International Business and Finance*, 42, 295-303.
- Caldeira, J., Torrent, H., (2017). Forecasting the US term structure of interest rates using nonparametric functional data analysis. *Journal of Forecasting* 36 (1), 56–7.
- Castellani, M., Santos, E. A. d., (2006). Forecasting long-term government bond yields: an application of statistical and AI models. ISEG, Departamento de Economia
- Castro-Iragorri, J Ramirez .(2021) Forecasting Dynamic Term Structure Models with Autoencoders. Documentos De Trabajo 019431.
- Cao, Y. (2023). Forecast Yield Curve of China's Government Bond with Machine Learning. .
- CERNA, S., GUYEUX, C., ARCOLEZI, H. H., COUTURIER, R., and ROYER, G. (2020). A comparison of lstm and xgboost for predicting _remen interventions. *In World Conference on Information Systems and Technologies*, pages 424{434. Springer.
- Chae S.C., Choi S.Y. (2022), Analysis of the term structure of major currencies using principal component analysis and autoencoders, *Axioms*, 11(3), 135.
- CHEN, T. and GUESTRIN, C. (2016). Xgboost: A scalable tree boosting system. *In Proceedings of the 22nd acm sigkdd international conference on knowledge discovery and data mining*, pages 785-794.
- Cho, K.; van Merriënboer, B.; Bahdanau, D.; Bengio, Y. (2014) On the Properties of Neural Machine Translation: Encoder-Decoder Approaches. arXiv:1409.1259

18. Cho, K., van Merriënboer, B., Gulcehre, C., Bahdanau, D., Bougares, F., Schwenk, H., & Bengio, Y. (2014). Learning phrase representations using RNN encoder-decoder for statistical machine translation. *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 1724-1734. <https://aclanthology.org/D14-1179>
19. Chou, J., Su, Y., Tang, H., & Chen, C. (2009). Fitting the term structure of interest rates in illiquid market: Taiwan experience. *Investment Management and Financial Innovations*, 6(1), 101–116.
20. Christensen, J., Diebold, F., & Rudebusch, G. (2009). An arbitrage-free generalized nelson-siegel term structure model. *The Econometrics Journal*, 12(3), 33–64. <https://doi.org/10.1111/j.1368-423X.2008.00267.x>
21. Christensen, J., Diebold, F., & Rudebusch, G. (2007). The afne arbitrage-free class of nelson-siegel term structure models. Working Paper 13611, *National Bureau of Economic Research*. <https://doi.org/10.3386/w13611>
22. Cox, J. C., Ingersoll, J. E., & Ross, S. A. (1985). A theory of the term structure of interest rates. *Econometrica*, 53(2), 385-407. <https://doi.org/10.2307/1911242>
23. Diebold, F. X., & Li, C. (2006). Forecasting the term structure of government bond yields. *Journal of Econometrics*, 130(2), 337-364. <https://doi.org/10.1016/j.jeconom.2005.04.014>
24. Dunis, C. L., Morrison, V., (2007). The economic value of advanced time series methods for modelling and trading 10-year government bonds. *European Journal of Finance* 13 (4), 333–352.
25. Enders, W.,(2014). *Applied Econometric Time Series*, 4th Edition. Wiley Series in Probability and Statistics. Wiley.
26. Engle, R. F. (1982). Autoregressive conditional heteroskedasticity with estimates of the variance of United Kingdom inflation. *Econometrica*, 50(4), 987-1007. <https://doi.org/10.2307/1912773>
27. Estrella, A., & Hardouvelis, G. A. (1991). The term structure as a predictor of real economic activity. *The journal of Finance*, 46(2), 555-576.
28. Gogas, P., Papadimitriou, T., Matthaïou, M., & Chrysanthidou, E. (2015). Yield curve and recession forecasting in a machine learning framework. *Computational Economics*, 45, 635–645. <https://doi.org/10.1007/s10614-014-9432-0>
29. Hamilton, J. D. (1994). *Time series analysis*. Princeton University Press.
30. Kanevski, M., Maignan, M., Pozdnoukhov, A., Timonin, V., (2008). Interest rates mapping. *Physica A: Statistical Mechanics and its Applications* 387 (15), 15 3897–3903.
31. Kanevski, M., Timonin, V., (2010). Machine learning analysis and modeling of interest rate curves. *European Symposium on Artificial Neural Networks, Computational Intelligence and Machine Learning*, ESANN.
32. Kang, K. (2012). Forecasting the term structure of korean government bond yields using the dynamic nelson-siegel class models. *Asia-Pacific Journal of Financial Studies*, 41(6), 765–787. <https://doi.org/10.1111/ajfs.12000>
33. Kauffmann P. C., Takada, H. H., Terada, A. T. & Stern, J. M. (2022). Learning ForecastEfficient Yield Curve Factor Decompositions with Neural Networks. *Econometrics*, 10.
34. Kirczenow, G., Fathi, A. & Davison, M. (2018). Machine Learning for Yield Curve Feature Extraction: Application to Illiquid Corporate Bonds. *Arxiv Preprint Arxiv:1806.01731*.
35. Knapp, B. (2020). *Macroeconomic factors in interest rate modelling* (Doctoral dissertation, Wien).
36. Koopman, S., Mallee, M., & Van der Wel, M. (2007). Analyzing the term structure of interest rates using the dynamic Nelson-Siegel model with time-varying parameters. *Tinbergen Institute Discussion Papers* 07-095/4, Tinbergen Institute
37. Kostyra, T. P. (2024). Forecasting the yield curve for Poland with the PCA and machine learning. *Bank i Kredyt*, 56(4), 459-478.
38. Linton, O., Mammen, E., Nielsen, J., & Tanggaard, C. (2001). Yield curve estimation by kernel smoothing methods. *Journal of Econometrics*, 105(1), 185–223. [https://doi.org/10.1016/S0304-4076\(01\)00075-6](https://doi.org/10.1016/S0304-4076(01)00075-6)
39. Litterman R., Scheinkman J. (1991), Common factors affecting bond returns, *Journal of Fixed Income*, 1(1), 54–61.

40. Lorenčič, E. (2016). Testing the performance of cubic splines and Nelson-Siegel model for estimating the zero-coupon yield curve. *Naše gospodarstvo/Our Economy*, 62(2), 42–50. <https://doi.org/10.1515/ngoe-2016-0011>
41. Lu, W., Li, J., Li, Y., Sun, A., & Wang, J. (2020). A CNN-LSTM-based model to forecast stock prices. *Complexity*, 2020(1), 6622927.
42. Lütkepohl, H., & Krätzig, M. (2004). Applied time series econometrics. Cambridge University Press.
43. Lynn, H.M.; Pan, S.B.; Kim, P. A (2019) Deep Bidirectional GRU Network Model for Biometric Electrocardiogram Classification Based on Recurrent Neural Networks. *IEEE Access*, 7, 145395–145405
44. Mallick A.K., Mishra A.K. (2019), Interest rates forecasting and stress testing in India: a PCA-ARIMA approach, *Palgrave Communications*, 5(1), 1–17.
45. Mateus, B. C., Mendes, M., Farinha, J. T., Assis, R., & Cardoso, A. M. (2021). Comparing LSTM and GRU models to predict the condition of a pulp paper press. *Energies*, 14(21), 6958.
46. Mishkin, F. S. (2019). *The economics of money, banking, and financial markets, Global edition. 12th Edition*, Pearson Education.
47. Nagy, K. (2020). Term structure estimation with missing data: Application for emerging markets. *The Quarterly Review of Economics and Finance*, 75, 347–360. <https://doi.org/10.1016/j.qref.2019.04.002>
48. Nath G.C. (2012), Estimating term structure changes using principal component analysis in the Indian sovereign bond market, SSRN 2075635
49. Nelson, C. R., & Siegel, A. F. (1987). Parsimonious modeling of yield curves. *Journal of Business*, 60(4), 473–489. <https://doi.org/10.1086/296368>
50. Nichani, R., Gasmi, L., Laiche, N., & Kabou, S. (2024). Optimizing financial time series predictions with hybrid ARIMA, LSTM, and XGBoost Models. *Studies in Engineering and Exact Sciences*, 5(2), e11188-e11188.
51. Nunes, M., Gerding, E., McGroarty, F., & Niranjana, M. (2019). A comparison of multitask and single-task learning with artificial neural networks for yield curve forecasting. *Expert Systems with Applications*, 119, 362-375.
52. Nunes, M. (2022). Machine learning in fixed income markets: forecasting and portfolio management (Doctoral dissertation, University of Southampton).
53. Nyman-Andersen, P. (2018). Yield curve modelling and a conceptual framework for estimating yield curves: evidence from the European Central Bank's yield curves (No. 27). *ECB Statistics Paper*.
54. Oprea A. (2022), The use of Principal Component Analysis (PCA) in building yield curve scenarios and identifying relative-value trading opportunities on the Romanian government bond market, *Journal of Risk and Financial Management*, 15(6), 247.
55. Piene, F. B., & Vedvik, J. O. (2020). Forecasting the US Treasury Yield Curve using Targeted Diffusion Indices (Master's thesis, Handelshøyskolen BI).
56. Pimentel R., Rissstad M., Westgaard S. (2022), Predicting interest rate distributions using PCA & quantile regression, *Digital Finance*, 4(4), 291–311.
57. Rezende, R., & Ferreira, M. (2008). Modeling and forecasting the Brazilian term structure of interest rates by an extended Nelson-Siegel class of models: A quantile autoregression approach. resreport, *Escola Brasileira de Economia e Finanças*.
58. Sambasivan, R., Das, S. (2017). A statistical machine learning approach to yield curve forecasting. In: *International Conference on Computational Intelligence in Data Science*, ICCIDS. IEEE, pp. 1–6.
59. Sewell, M. (2011). Characterization of financial time series. UCL Department of Computer Science,
60. Suimon, Y., Sakaji, H., Izumi, K. & Matsushima, H. (2020a). Autoencoder-based ThreeFactor Model for the Yield Curve of Japanese Government Bonds and A Trading Strategy. *Journal of Risk and Financial Management*, 13, 82-103.
61. Suimon, Y., Sakaji, H., Izumi, K., Shimada, T. & Matsushima, H. (2020b). Japanese Interest Rate Forecast Considering the Linkage of Global Markets using Machine Learning Methods. *International Journal of Smart Computing and Artificial Intelligence*, 4, 1-17.
62. Svensson, L. E. (1995). Estimating forward interest rates with the extended Nelson & Siegel method. *Sveriges Riksbank Quarterly Review*, 3(1), 13-26.

63. Ullah, W. (2017). Term structure forecasting in afne framework with time-varying volatility. *Stat Methods Appl*, 26, 453–483. <https://doi.org/10.1007/s10260-017-0378-y>
64. Umar, Z., Yousaf, I., & Aharon, D. Y. (2021). The relationship between yield curve components and equity sectorial indices: evidence from China. *Pacific-Basin Finance Journal*, 68, 101591.
65. Walstrøm R, Paraschiv F, Schürle M. (2022). A comparative analysis of parsimonious yield curve models with a focus on the Nelson-Siegel, Svensson, and Bliss versions. *Computational Economics* 59:967–1004
66. Widiputra, H., Mailangkay, A., & Gautama, E. (2021). Multivariate CNN-LSTM Model for Multiple Parallel Financial Time-Series Prediction. *Complexity*, 2021(1), 9903518.
67. Zhang, X., Song, Y., & Liu, Y. (2018). Deep learning for time series forecasting: A survey. *In Proceedings of the 2018 International Conference on Machine Learning and Data Mining* (pp. 45-59). Springer. https://doi.org/10.1007/978-3-319-99068-4_6.