

مقایسه روش رگرسیون بردار پشتیبان معمولی و گروهی برای پیش‌بینی قیمت سهام

سید محمد حسینی^{۱*}، استادیار، مجید ابتیاع^۲، کارشناس ارشد ریاضی مالی، محمدرضا آریافر^۳، دانشجوی دکتری اقتصاد

^۱ و ^۲ گروه پژوهشی هوش مصنوعی گهر - دانشکده علوم پایه - دانشگاه آیت الله بروجردی (ره) - بروجرد - ایران -

sm.hoseini@abru.ac.ir

^۲ گروه اقتصاد - دانشکده علوم اداری و اقتصادی - دانشگاه فردوسی مشهد - مشهد - ایران - m.ariafar@mail.um.ac.ir

چکیده: انباشت سرمایه و نگهداری صحیح از دارایی نقش به‌سزایی را در رشد اقتصادی کشورها ایفا می‌کند. نقش مؤثر این عامل و اهمیت آن را به‌روشنی در سیستم اکثر کشورها به‌ویژه نظام‌های سرمایه‌داری می‌توان مشاهده نمود. بازار سرمایه یکی از مناسبترین جایگاه‌ها به منظور جذب دارایی‌های اندک و به کارگیری آنها باهدف رشد و ارتقا یک شرکت و همچنین افزایش دارایی شخصی سرمایه‌گذاران است. از این رو، سرمایه‌گذاری درست، به افراد کمک می‌کند تا با مدیریت صحیح بخش مازاد درآمدهای خود، بتوانند به سرمایه لازم جهت رسیدن به اهداف‌شان دست پیدا کنند. اساساً قیمت سهام، غیرخطی و آشوبناک است، بنابراین سرمایه‌گذاری در بازار بورس، ریسک بالایی به‌همراه دارد. برای به‌حداقل رساندن این ریسک و کاهش مخاطرات، یک سیستم کارآمد نیاز است که قادر باشد حرکت قیمت سهام در آینده را با دقت بالایی پیش‌بینی نماید. در این راستا، مدل‌های یادگیری ماشین عملکرد مناسبی در زمینه پیش‌بینی با دقت بالا را دارا هستند. الگوریتم‌های یادگیری ماشین توانایی مدل نمودن سیستم‌های غیرخطی و پیچیده را دارند. در این پژوهش به کمک رویکرد گروهی جدید مبتنی بر رگرسیون بردار پشتیبان یک مدل جمعی یا ترکیبی ساخته می‌شود که از دقت و سرعت بالایی در پیش‌بینی قیمت سهام برخوردار است. سهم‌های مورد استفاده در این پژوهش ده سهام شاخص ساز فولاد، شستا، فارس، فملی، شپنا، شتران، شبندر، خودرو، خساپا و پارس است که در بازه زمانی ۱۴۰۰/۰۱/۰۷ تا ۱۴۰۳/۰۴/۳۱ در نظر گرفته شده است. بر اساس معیارهای مختلف، نتایج به دست آمده برتری روش رگرسیون بردار پشتیبان گروهی بر روش معمولی رگرسیون بردار پشتیبان و روش جنگل تصادفی را نشان می‌دهد.

واژه‌های کلیدی: یادگیری ماشین، رگرسیون بردار پشتیبان، روش گروهی، بازار سهام، پیش‌بینی قیمت،

* سید محمد حسینی، sm.hoseini@abru.ac.ir

Comparison of single and ensemble support vector regression model for stock price prediction

Sayyed Mohammad Hoseini ^{1*}, Assistant professor , Majid Ebita², Master of Financial Mathematics, Mohammad Reza Ariafar ³, PhD Student of Economics

^{1, 2} Gahar Artificial Intelligence Research Group, Faculty of Basic Sciences, Ayatollah Boroujerdi University , Boroujerd, Iran, sm.hoseini@abru.acir

³ Dept. of Economics, Faculty of Economics and Administrative Sciences, Ferdowsi University of Mashhad, Mashhad, Iran, m.ariafar@mail.um.ac.ir

Abstract: Capital accumulation and proper asset management play a crucial role in the economic growth of countries. The significant impact of this factor, as well as its importance, can be clearly observed in the system of most countries, especially capitalist systems. The capital market is one of the most appropriate places to attract small assets and use them in order to grow and improve a company and also increase the personal assets of investors. Therefore, correct investment helps individuals to manage the surplus of their incomes properly, enabling them to attain the necessary capital to achieve their goals. Essentially, stock prices are nonlinear and chaotic, hence investing in the stock market carries a high risk. To minimize this risk and reduce hazards, an efficient system is required that can predict future stock price movements with high accuracy. Given the importance of this issue, machine learning models have shown suitable performance in this area. Machine learning algorithms have the ability to model nonlinear and complex systems. In this research, a novel collective approach based on Support Vector Regression constructs a composite or hybrid model that enjoys high accuracy and speed in predicting stock prices. The stocks used in this research are the ten index-making shares of FOLD, TAMN, PKLJ, MSMI, PNES, PTEH, PNBA, IKCO, SIPA and PARS, considered in the period from 2021-03-27 to 2024-07-21. Based on various criteria, the results demonstrate the superiority of the collective Support Vector Regression method over the conventional Support Vector Regression method and the Random Forest approach.

Keywords: Regression support vector, stock market, price prediction, stock index maker, Random Forest

* Sayyed Mohammad Hoseini, sm.hoseini@abru.ac.ir

۱. مقدمه

دانش فنی و تخصص، راهکارهایی برای پیش‌بینی دقیق‌تر و مدیریت بهینه‌ی ریسک در بازارهای مالی ایجاد کنند. از این جهت، ترکیب مدل‌های مختلف و به‌کارگیری ابزارهایی همچون بگینگ^۱ و رگرسیون بردار پشتیبان^۲ (SVR) می‌تواند بهبودهای مهمی در عملکرد پیش‌بینی و مدیریت سرمایه حاصل کند.

به‌طور کلی، برای پیش‌بینی قیمت سهام دو نوع روش وجود دارد. یکی از رویکردها، استفاده از روش‌های آماری سنتی مثل میانگین ساده و رگرسیون است. این روش‌ها زمانی کارآمدند که بین متغیرهای مستقل و وابسته ارتباط و همبستگی مشخصی داشته باشند. بازار سهام به‌شدت نامنظم و آشوبناک است، بنابراین نیاز به یک سیستم کارآمد جهت پیش‌بینی دقیق قیمت سهام وجود دارد. مزیت عمده‌ی سیستم‌های هوشمند پیش‌بینی، دقت بالاتر آن‌هاست.

در دهه‌های اخیر، با پیشرفت روزافزون در حوزه فناوری و گسترش بازارهای مالی، پیش‌بینی قیمت سهام به یکی از چالش‌های اساسی در زمینه سرمایه‌گذاری تبدیل شده است. یکی از روش‌های مطرح در این زمینه، استفاده از مدل‌های رگرسیون بردار پشتیبان است در حال حاضر، باتوجه به پیشرفت چشمگیر

بازارهای مالی، به‌ویژه بازار سرمایه گسترده و کارآمد، به‌عنوان ابزاری برای تحقق رشد بلندمدت و پایداری اقتصادی، موردتوجه قرار می‌گیرند. تحقق این امر، نیازمند تخصیص بهینه منابع در سطح اقتصاد ملی است. بورس اوراق بهادار یکی از مسیرهای جذب سرمایه است. البته تحقق این امر، مستلزم جلب اعتماد سرمایه‌گذاران است. بازار سرمایه، امکان جذب اعتماد سرمایه‌گذاران را فراهم می‌کند و اوراق بهادار به‌عنوان ابزارهای متنوع جهت مدیریت ریسک و جمع‌آوری سرمایه‌های پراکنده استفاده می‌شوند. این رویکرد نقش مهمی در توسعه اقتصادی و پیشرفت کشورها ایفا می‌کند. گردش سرمایه در حال حاضر بیشترین حجم خود را از طریق بازارهای سرمایه دارد و عملکرد این بازار بر تحولات اقتصادی برخی کشورها با بازار سرمایه دارای شفافیت اطلاعات تأثیرگذار است [۱]. با این حال، عوامل ناشناخته و متعددی بر بازارهای اوراق بهادار تأثیر می‌گذارند که سبب عدم اطمینان در تصمیم‌گیری سرمایه‌گذاران شده است. به همین دلیل، تحقیقات و توسعه روش‌های پیش‌بینی دقیق و کاهش خطا در تحلیل بازارها از اهمیت ویژه‌ای برخوردارند [۲]. در این راستا، پژوهشگران و سرمایه‌گذاران در تلاشند تا با بهره‌گیری از

² Support Vector Regression

¹ Bootstrap Aggregating (Bagging)

سرمایه‌گذاران و تحلیلگران بازار ابزاری قوی ارائه می‌دهد تا در تصمیم‌گیری‌های سرمایه‌گذاری بهترین عملکرد را داشته باشند. همچنین، این رویکرد با مقاومت بیشتری در برابر نوسانات بازار و توانایی تعمیم به شرایط جدید، به‌عنوان یک ابزار قوی برای مدیریت ریسک در بازار سهام مطرح می‌شود. به‌طور کلی، استفاده از ترکیب مدل‌های مختلف با روش‌های گروهی یا جمعی برای پیش‌بینی قیمت سهام، به‌عنوان یک نوآوری مهم و کارآمد در زمینه تحلیل مالی و سرمایه‌گذاری معرفی می‌شود که می‌تواند به بهبود عملکرد مدل‌های پیش‌بینی و افزایش دقت تصمیم‌گیری‌های سرمایه‌گذاران کمک کند.

نتایج مطالعاتی که تا به اینجا بررسی شده‌اند نشان می‌دهند که روش گروهی باعث بهبود قابل توجهی در دقت پیش‌بینی‌ها نسبت به مدل‌های تکی SVR می‌شود. اما با وجود بهبودهای ایجاد شده توسط روش پیشنهادی، هنوز چالش‌هایی نیز وجود دارد که ممکن است تأثیر منفی بر عملکرد مدل داشته باشند. به‌عنوان مثال، یکی از چالش‌ها مدیریت داده‌های نوسانی بازار است، زیرا بازارهای مالی ممکن است تحت تأثیر عوامل ناپیوسته و پیچیده قرار گیرند که مدل را به چالش کشانده و دقت پیش‌بینی‌ها را کاهش دهند. همچنین، حساسیت مدل به تغییرات پارامترها نیز یک چالش مهم است، زیرا نیاز به تنظیم دقیق پارامترها برای حفظ عملکرد بهینه مدل وجود دارد.

در مراحل انجام این پژوهش، ابتدا با مرور ادبیات موجود در زمینه پیش‌بینی قیمت سهام و مدل‌های مورد استفاده، نقاط قوت و ضعف این روش‌ها شناسایی می‌شود. سپس، مدل پیشنهادی بر اساس ترکیب مدل‌های رگرسیون بردار پشتیبان (SVR) و تکنیک بگینگ ارائه شده است. سپس مراحل آموزش مدل و تولید پیش‌بینی‌ها شرح داده شدند. در ادامه، از معیارهای دقت پیش‌بینی، خطای میانگین مطلق (MAE) و خطای میانگین مربعات (MSE) برای ارزیابی عملکرد مدل استفاده شده است. در انتها ارزیابی جامع و دقیق عملکرد مدل انجام می‌شود.

تکنولوژی و سرعت بالای محاسبات، همچنین در دسترس بودن داده‌های گسترده در حوزه‌های مختلف، دیدگاه‌های علمی متمرکز بر داده، از جمله هوش مصنوعی و یادگیری ماشین، به‌عنوان رویکردهای اصلی در زمینه تحقیقات، توجه بسیاری از پژوهشگران را به خود جلب کرده است. این دیدگاه‌ها نه تنها در زمینه‌های علمی، بلکه در کاربردهای متعددی نیز وارد شده و نتایج قابل توجهی را به همراه داشته‌اند. ابزار SVR یکی از ابزارهای پر قدرت در پیش‌بینی قیمت سهام است که بر پایه ایده‌های ماشین بردار پشتیبان³ (SVM) برای مسائل پیش‌بینی مقادیر عددی پیوسته طراحی شده است. این مدل با توانایی مدیریت داده‌های پیچیده و ارائه پیش‌بینی دقیق، به‌عنوان یکی از ابزارهای مهم و مورد توجه در زمینه‌ی پیش‌بینی قیمت سهام شناخته می‌شود. بگینگ، یک تکنیک مؤثر در افزایش استحکام و دقت مدل‌های یادگیری ماشین است [۳]. این روش شامل نمونه‌گیری تصادفی از داده‌ها، آموزش مدل بر روی هر نمونه، و سپس ترکیب پیش‌بینی‌های این مدل‌ها است. این فرآیند منجر به بهبود دقت نهایی مدل می‌شود.

در این مقاله، یک روش جدید ترکیبی با استفاده از مدل‌های بگینگ و SVR ارائه شده است. در این روش پیشنهادی، ابتدا مدل SVR روی نمونه‌های مختلفی از داده‌ها آموزش داده می‌شود. سپس با استفاده از بگینگ، مدل‌های مختلفی ایجاد می‌شوند و هر کدام بر روی یک نمونه تصادفی از داده‌ها آموزش داده می‌شوند. این ترکیب از مدل‌ها منجر به بهبود دقت و مقاومت مدل در برابر نوسانات بازار می‌شود. روش پیشنهادی از قدرت پیش‌بینی SVR بهره می‌برد و با استفاده از بگینگ، دقت و پایداری پیش‌بینی‌ها را افزایش می‌دهد. ارزیابی عملکرد این روش با استفاده از معیارهای دقت پیش‌بینی، خطای میانگین مطلق (MAE) و خطای میانگین مربعات (MSE) صورت گرفته است [۴].

نوآوری اصلی این پژوهش در استفاده از ترکیب مدل‌های رگرسیون بردار پشتیبان (SVR) با تکنیک بگینگ برای پیش‌بینی قیمت سهام است. این رویکرد ترکیبی امکانات مدل SVR در مدیریت داده‌های پیچیده و پیش‌بینی دقیق را با افزایش دقت و پایداری مدل توسط بگینگ ترکیب می‌کند. این نوآوری به

³ Support Vector Machines

۲. پیشینه پژوهش

تاکنون، پیش‌بینی قیمت سهام موضوع بسیاری از پژوهش‌ها بوده است و رویکردها و تکنیک‌های مختلفی برای این منظور بررسی شده‌اند. در سال ۲۰۲۲ ایلیاس و همکاران [۵]، یک مدل ترکیبی را توسعه دادند که در آن رگرسیون بردار پشتیبان و فیلتر هودریک-پرسکات یکپارچه‌سازی می‌شوند و قابلیت پیش‌بینی قیمت سهام را افزایش می‌دهند. به طور مشابه، در سال ۲۰۱۴، هوآنگ و همکاران [۶] یک چارچوب را نشان دادند که رگرسیون بردار پشتیبان و تبدیل فوریه را برای پیش‌بینی قیمت سهام بر اساس داده‌های تاریخی یکپارچه می‌کند. این مطالعات نشان‌دهنده استفاده از رگرسیون بردار پشتیبان در مدل‌های پیش‌بینی قیمت سهام هستند.

در سال ۲۰۱۹، بشیر و همکاران [۷] بر اهمیت چالش‌هایی که در پیش‌بینی قیمت سهام وجود دارد تأکید کرد و اعلام کرد که در این زمینه موفقیت نسبی به دست آمده است. عدم قطعیت در سیاست‌های اقتصادی و تغییرات در قیمت کالاها می‌تواند تأثیرات غیرمنتظره‌ای بر بازارهای سهام داشته باشد که پیش‌بینی این تغییرات را دشوار می‌کند. روابط غیرخطی و پیچیده بین متغیرهای اقتصادی و قیمت سهام و واکنش‌های متفاوت بازارها به تغییرات مشابه، از جمله چالش‌های اصلی در مدل‌سازی و پیش‌بینی هستند که آنها در تحلیل خود بر آن تأکید کردند. همچنین، در سال ۲۰۲۰، پاسالیس و همکاران [۸] ویژگی‌های مختلفی را شناسایی کرده و از یک روش طبقه‌بندی مبتنی بر رگرسیون بردار پشتیبان (SVM) برای پیش‌بینی حرکات قیمت سهام استفاده کرده‌اند، که نمونه‌ای از اعمال تکنیک‌های یادگیری ماشین در این زمینه است. همچنین، ابتونگ و همکاران [۹] در همان سال به تحلیل و بهبود دقت پیش‌بینی قیمت سهام پرداختند.

روش‌های مبتنی بر یادگیری ماشین مانند جنگل تصادفی [۱۰] شبکه‌های پرسپترون چندلایه و ماشین‌های رگرسیون بردار پشتیبان [۱۱]، و همچنین تکنیک‌هایی مانند تبدیل فوریه [۶] و روش جنگل تصادفی^۴ و الگوریتم ماشین بردار پشتیبانی^۵ [۱۱] برای پیش‌بینی قیمت سهام استفاده شده است.

در این راستا، مطالعه لیو و همکاران [۱۳] به کارگیری یک مدل CAE-LSTM را نشان داده است که با استفاده از اطلاعات فنی، عملکرد بهتری در پیش‌بینی قیمت سهام به دست آورده است. این مطالعه به وضوح نشان می‌دهد که ترکیب شبکه‌های عصبی پیچشی و بازگشتی می‌تواند بهبود معناداری در دقت پیش‌بینی داشته باشد. همچنین، [۱۴] رستم و کیتاندانی در سال ۲۰۱۹ با تمرکز بر روی بازار سهام اندونزی، از رگرسیون بردار پشتیبان به همراه انتخاب ویژگی، نشان داده‌اند که این روش قابلیت پیش‌بینی بهتر در بازارهای خاصی را ارائه می‌دهد.

در پژوهشی توسط منادی و نجفی [۱۵]، با هدف ارتقای کیفیت انتخاب سبد سرمایه‌گذاری، استفاده از یک مدل پیشرفته رگرسیون بردار پشتیبان با خروجی‌های چندگانه مورد بررسی قرار گرفت. این رویکرد، که در جهت دستیابی به حداکثر بازده با کمترین ریسک است، بر تحلیل دقیق‌تر قیمت سهام تأکید دارد. استفاده از داده‌های مربوط به شرکت‌های برجسته S&P500 طی دوره‌ای معین نشان داد که این مدل، با در نظر گرفتن ارتباطات بین خروجی‌ها، عملکرد بهتری نسبت به روش‌های سنتی پیش‌بینی از خود نشان می‌دهد، خصوصاً در زمینه کارایی بر اساس معیار شارپ^۵.

دهقان و همکاران [۱۶] در پژوهشی، مدل‌های طبقه‌بندی مختلف برای دستیابی به یک مدل با کارایی و دقت بالا در پیش‌بینی قیمت‌گذاری کمتر از واقع سهام عرضه عمومی اولیه مورد بررسی

قرار دادند. نتایج نشان داده‌اند که مدل‌های ماشین بردار پشتیبان، شبکه بیزین دقت بالایی در پیش‌بینی قیمت‌گذاری کمتر از واقع دارند. همچنین، متغیرهای مهم و تأثیرگذار شامل رشد دارایی‌ها، دوره تصدی حسابرس، تخصص حسابرس در صنعت، نسبت فعالیت‌های تامین مالی، نسبت قیمت به سود هر سهم، بازده دارایی‌ها، نسبت فعالیت‌های عملیاتی، اندازه موسسه حسابرسی، فرصت‌های رشد و نوسانات قیمت سهام هستند.

خسروی‌نژاد و همکاران [۱۷] در مطالعه‌ای، ابتدا با استفاده از مدل سری زمانی و شبکه عصبی مصنوعی، متغیر هفتگی شاخص قیمت سهام در بورس اوراق بهادار تهران را برآورد کرده‌اند. این برآوردها برای دوره زمانی ۱۳۸۳ تا ۱۳۸۷ صورت گرفته و سپس قدرت پیش‌بینی دو مدل در دوره زمانی ۱۳۸۷ تا ۱۳۸۹ مورد آزمون و ارزیابی قرار گرفته است.

منجمی و همکاران [۱۸] یک مدل پیش‌بینی قیمت سهام با استفاده از شبکه‌های عصبی فازی و الگوریتم‌های ژنتیک طراحی کردند. نتایج حاصل این پژوهش مشخص می‌کند که ادغام شبکه‌های عصبی فازی و الگوریتم‌های ژنتیک به عنوان یک مدل ترکیبی، توانمندی بیشتری در پیش‌بینی قیمت سهام از نظر دقت و کارایی از خود نشان می‌دهد. این یافته می‌تواند به توسعه رویکردهای بهبود یافته در حوزه پیش‌بینی بازار سهام کمک کند و به سرمایه‌گذاران ابزاری قدرتمندتر برای تصمیم‌گیری در بازار سرمایه ارائه دهد.

رمضانی و عاملی [۱۹] در مطالعه‌ای از الگوریتم ژنتیک استفاده کردند. در بخش پیش‌بینی، از شبکه عصبی فازی به همراه دو ساختار مختلف (ممدانی^۶ و سوگنو^۷) برای پیش‌بینی قیمت سهام استفاده شده است. الگوریتم پیشنهادی با استفاده از داده‌های ۱۰

شرکت، هر کدام دارای ۷ ویژگی، ارزیابی شده است. نتایج نشان داده است که باتوجه به نوع شرکت‌ها، ویژگی‌ها و ساختار شبکه عصبی فازی ترکیبی، نتایج متفاوتی به دست می‌دهد. ارزیابی‌ها نشان می‌دهد که الگوریتم پیشنهادی با ساختار سوگنو با الگوریتم ژنتیک، در مقایسه با ساختار ممدانی، دارای عملکرد بهتری است، زیرا تعداد پارامترهای آموزش آن بیشتر است.

زمانی و همکاران [۲۰] در پژوهشی، از مدل‌های ریاضی بهینه‌سازی استفاده کردند تا با توجه به متغیرهایی از قبیل میانگین، واریانس، و چولگی، سبدي از سهام را پیش‌بینی کنند. نتایج تحقیق نشان می‌دهد که مدل‌های ارائه شده، در مقایسه با روش‌های سنتی و شاخص‌های بازار، بازدهی بالاتری را برای سرمایه‌گذاران فراهم می‌کنند. این مدل‌ها از ترکیب پتانسیل آتی سهام با اطلاعات گذشته برای تصمیم‌گیری بهینه‌تر استفاده می‌کنند، که به سرمایه‌گذاران ابزار قدرتمندتری برای تصمیم‌گیری در بازار بورس ارائه می‌دهد.

سادورسکی [۱۰] یک مطالعه مرتبط در خصوص استفاده از جنگل‌های تصادفی برای پیش‌بینی قیمت سهام ارائه کرده است. یافته‌های این مطالعه نشان می‌دهد که برای پیش‌بینی قیمت سهام، پیش‌بینی‌های جنگل‌های تصادفی دقیق‌تر از مدل‌های لاجیت باتوجه به بهره‌مندی از ظرفیت روش‌های گروهی است. [۲۱] انتی و همکاران در سال ۲۰۲۰ از یک روش گروهی، یعنی ماشین‌های رگرسیون بردار پشتیبان، برای پیش‌بینی کارایی بازار سهام استفاده کرده است. استفاده از الگوریتم ژنتیک برای بهینه‌سازی ابرپارامترهای SVM، رویکردی برای بهبود عملکرد پیش‌بینی را نشان می‌دهد [۲۲]. شائو و همکاران، ۲۰۱۹ از تحلیل طیفی تکین^۸ و ماشین‌های رگرسیون بردار پشتیبان برای

⁸ Singular spectrum analysis

⁶ Mamdani

⁷ Sogno

پیش‌بینی قیمت سهام، بر تلاش‌های مداوم برای توسعه و بهبود مدل‌های پیش‌بینی نیز تأکید دارند.

۳. یادگیری ماشین

امروزه، یادگیری ماشین به‌عنوان یک زمینه پژوهشی جدید و پرکاربرد در علوم کامپیوتر شناخته می‌شود. این حوزه با مسائل و مفاهیم گوناگونی همچون پیش‌بینی قیمت سهام، هوش مصنوعی، امور پزشکی و غیره، ارتباط نزدیکی دارند. الگوریتم‌های یادگیری ماشین به‌طور کلی به دو دسته باناظر و بدون ناظر تقسیم می‌شوند. هدف در یادگیری ماشین باناظر، یافتن نگاشت بین ورودی‌ها و خروجی‌ها است و برچسب‌ها برای هر زوج ورودی و خروجی مشخص هستند. این دسته از الگوریتم‌ها بر روی داده‌های برچسب‌خورده قابل‌استفاده هستند و در پروژه‌هایی که بر اساس تفکیک دقیق دسته‌ها انجام می‌شوند، مورد استفاده قرار می‌گیرند. در مقابل، یادگیری ماشین بدون ناظر، الگوهایی از داده‌ها را بدون نیاز به برچسب یا خروجی مشخص می‌کند. هدف در این دسته از الگوریتم‌ها، کشف الگوهای مخفی و مفهومی در داده‌هاست. این نوع یادگیری در مواردی که برچسب‌ها دسترسی محدود دارند یا دسترس‌ناپذیر هستند، بسیار مورد توجه قرار می‌گیرد. ترکیب این دو نوع یادگیری ماشین، به دلیل قابلیت کاربردی وسیع در زمینه‌های مختلف و توانایی پیش‌بینی دقیق، موضوع مهمی در پژوهش‌های امروزی است.

۴. جنگل تصادفی

تجزیه و تحلیل و پیش‌بینی قیمت سهام استفاده کرده‌اند. مطالعه‌ی مقیاس‌های زمانی مختلف و عوامل مؤثر بر قیمت‌های سهام، زمینه‌ای را برای چالش‌های پیش‌بینی فراهم می‌کند [۱۴]. رستم و کینتاندانی در سال ۲۰۱۹ به طور خاص بر روی رگرسیون بردار پشتیبان برای پیش‌بینی قیمت سهام اندونزی با انتخاب ویژگی تمرکز دارد. این تحقیق نشان‌دهنده‌ی کاربرد رگرسیون بردار پشتیبان در بازاری خاص است و همچنین کنترل ویژگی‌های مرتبط در آن بررسی شده است. کلهری و حسینی [۲۳] از یک روش مبتنی بر گراف، به کمک الگوریتم کوچک‌ترین درخت پوشا برای خوشه‌بندی سهام شرکت‌های بورس اوراق بهادار تهران استفاده کردند.

تحقیقات موجود از روش‌های گروهی مانند جنگل‌های تصادفی و ماشین‌های رگرسیون بردار پشتیبان و همچنین شبکه‌های عصبی مانند LSTM، برای پیش‌بینی قیمت سهام استفاده کرده‌اند. این مطالعات همچنین چالش‌های مرتبط با داده‌های بازار سهام در مقیاس‌های زمانی مختلف را برجسته می‌کنند و نیاز به انتخاب بهینه‌سازی مدل را برای بهبود دقت پیش‌بینی نشان می‌دهند. منابع انتخاب‌شده، پایه‌ای را برای رویکرد ترکیبی پیشنهادی رگرسیون بردار پشتیبان با استفاده از بگینگ فراهم می‌کنند.

در مجموع، تنوع در رویکردها و تکنیک‌های مورد استفاده در حوزه پیش‌بینی قیمت سهام یک چشم‌انداز پویا از پژوهش‌ها را ارائه می‌دهد. این مطالعات نه تنها به نقدهای اساسی مرتبط با چالش‌های پیش‌بینی قیمت سهام پرداخته‌اند؛ بلکه همچنین به راهکارهای نوآورانه و بهبود مدل‌های پیش‌بینی اشاره کرده‌اند. این توسعه‌ها و پیشرفت‌ها نشان‌دهنده تعهد محققان به ارتقاء دائمی در زمینه پیش‌بینی قیمت سهام هستند. بررسی پژوهش‌های پیشین در زمینه پیش‌بینی قیمت سهام نشان می‌دهد که این حوزه شامل طیف گسترده‌ای از روش‌ها و تکنیک‌های مختلف است. این مطالعات علاوه بر اشاره به چالش‌ها و پیچیدگی‌های مرتبط با

به‌خوبی با داده‌های پیچیده سازگاری دارد و عملکرد بهتری نسبت به روش‌های سنتی مانند رگرسیون یا درخت تصمیم به‌دست می‌آورد [۲۴].

۵. الگوریتم رگرسیون بردار پشتیبان

ماشین بردار پشتیبان یکی از الگوریتم‌های باناظر است که در سال ۱۹۹۸ توسط وپنیک کشف شد. به‌طور کلی مدل‌های SVM به دو گروه رگرسیون ماشین بردار پشتیبان و مدل دسته‌بند ماشین بردار پشتیبان تقسیم می‌شوند. از مدل دسته‌بندی جهت دسته‌بندی داده‌ها و از مدل رگرسیون جهت پیش‌بینی استفاده می‌شود. اگر مجموعه‌ای از نقاط داده به فرم $\{(x_1, y_1), \dots, (x_i, y_i)\}$ نشان داده شود که در آن $x_i \in R$ ورودی و $y_i \in R$ خروجی متناظر آن باشد، فرم کلی رگرسیون بردار پشتیبان بدین شرح است [۲۴]:

$$\min \frac{1}{2} \|w\|^2,$$

$$\text{s. t. } |y - \hat{y}| \leq \varepsilon,$$

که در آن w بردار وزن ویژگی‌ها است. در ادامه، هسته‌های مختلف قابل استفاده در رگرسیون ماشین بردار پشتیبان ارایه شده است:

۱. هسته گاوسی: هسته گاوسی به‌عنوان یکی از رایج‌ترین هسته‌ها در مدل‌های پیش‌بینی شناخته می‌شود. عملکرد برتر مدل‌های SVR با هسته گاوسی نسبت به دیگر هسته‌ها نشان‌دهنده توانایی بالاتر در دسته‌بندی داده‌های پیچیده و غیرخطی است. توزیع گاوسی این هسته اختلافات بین داده‌ها را به‌خوبی توزیع می‌کند و از آن برای مدل‌سازی روابط پیچیده استفاده می‌شود.

۲. هسته خطی: هسته خطی مختص داده‌هایی است که رابطه خطی بین ویژگی‌ها و پیش‌بینی وجود دارد. این هسته به مدل امکان

جنگل تصادفی^۹ یکی از الگوریتم‌های مؤثر در حوزه یادگیری ماشین است که برای مسائل دسته‌بندی و پیش‌بینی استفاده می‌شود. این الگوریتم به‌صورت مجموعه‌ای از درختان تصمیم عمل می‌کند و از رای‌گیری تصمیم‌گیری هر درخت برای تصمیم نهایی استفاده می‌کند. الگوریتم جنگل تصادفی به دلیل ویژگی‌های زیر جذابیت زیادی دارد:

۱. تصادفی بودن ویژگی‌ها: در هر مرحله از ساخت هر درخت تصمیم، تعدادی از ویژگی‌ها به صورت تصادفی انتخاب می‌شوند. این اقدام منجر به افزایش تنوع در مدل می‌گردد و از بروز بیش‌فرضی (اطلاعات غلط یا غیرقابل اعتماد) به داده‌ها جلوگیری می‌نماید. این رویکرد معمولاً در الگوریتم‌های یادگیری ماشین مانند روش جنگل تصادفی مورد استفاده قرار می‌گیرد.

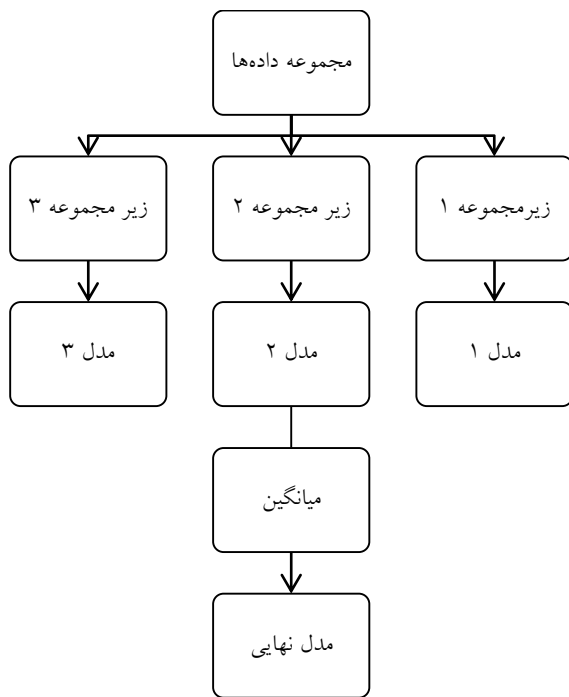
۲. استفاده از بوت‌استرپ نمونه‌ها: در هر مرحله از جنگل تصادفی، برای ساخت هر درخت، نمونه‌هایی از مجموعه‌ی داده‌ای به صورت بوت‌استرپ (انتخاب با جایگذاری) انتخاب می‌شوند. این باعث متنوع شدن داده‌ها و کاهش اثر تغییرات تصادفی در مجموعه‌ی داده‌ای می‌شود.

۳. رای‌گیری (میانگین‌گیری): نتایج حاصل از تمام درختان تصمیم به صورت میانگین محاسبه می‌شوند. این کار باعث مقاومت بیشتر مدل در مقابل داده‌های نویزی و پرت می‌شود.

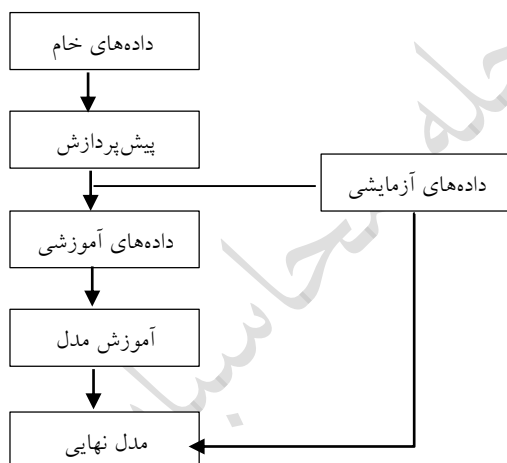
۴. پیچیدگی کم: درختان تصمیم فرآیندهای قابل فهم و قابل تفسیر هستند. جنگل تصادفی با ترکیب این درختان به‌عنوان یک مدل مجموعه‌ای به دست می‌آید که همچنان قابل فهم است.

الگوریتم جنگل تصادفی در انواع مختلفی از مسائل کاربرد دارد، از جمله پیش‌بینی قیمت سهام که در متن پیشین به آن اشاره شده است. باتوجه به ویژگی‌های یادگیری تصادفی، این الگوریتم

^۹ Random Forest



شکل ۱ مدل مفهومی رگرسیون بردار پشتیبان گروهی



شکل ۲ مدل مفهومی نحوه ورود داده‌ها

۶. الگوریتم رگرسیون بردار پشتیبان گروهی

ایده‌ی رگرسیون بردار پشتیبان گروهی از جنگل تصادفی رگرسیون نشأت می‌گیرد. در این روش هم ابتدا بر اساس داده‌ها و ایده‌ی بوت‌استرپ، چندین رگرسیون بردار پشتیبان تشکیل می‌شود و به هر رگرسیون بردار پشتیبان یک زیرمجموعه تصادفی

یادگیری روابط خطی میان ویژگی‌ها را می‌دهد. در شرایطی که داده‌ها دارای روابط خطی هستند، مدل‌های با هسته خطی به خوبی عمل می‌کنند.

۳. هسته چندجمله‌ای: هسته چندجمله‌ای برای مدل‌سازی روابط غیرخطی و پیچیده بین ویژگی‌ها مناسب است. این هسته به مدل امکان می‌دهد تا روابط چندجمله‌ای با ویژگی‌ها را تخمین بزند و برای داده‌های غیرخطی مؤثر باشد. اهمیت بالایی در مدل‌سازی روابط پیچیده و غیرخطی دارد.

۴. هسته سیگموئید: هسته‌های سیگموئید در زبان اصطلاحات ریاضی و آماری به مجموعه‌ای از الگوها یا توابع احتمالاتی اشاره دارد که به طور کلی به شکل منحنی‌های «سیگموئید» یا «S-shaped» شناخته می‌شوند. این الگوها معمولاً در مدل‌های آماری و ریاضی برای نمایش توزیع‌های احتمالاتی، توابع فعال‌سازی در شبکه‌های عصبی، و یا در توصیف رفتار و تغییرات در زمان استفاده می‌شوند.

یک هسته سیگموئید به طور معمول به عنوان یک تابع منحنی مانند S از یک متغیر مستقل به یک متغیر وابسته تعریف می‌شود. این توابع اغلب خاصیت‌هایی مانند صعود و نزول یکنواخت در ابتدا و پایان، معکوس‌پذیری و تغییر ناگهانی دارند. از جمله معروف‌ترین موارد استفاده از هسته‌های سیگموئید، مدل‌های رگرسیون لجستیک^{۱۰} در آمار و یادگیری ماشین می‌باشند. در این مدل، تابع لجستیک که یک هسته سیگموئید است، برای پیش‌بینی احتمال وقوع یک رویداد یا کلاس بر اساس متغیرهای مستقل استفاده می‌شود.

در بازه $[-1, 1]$ قرار می‌گیرند. در این مقاله، از روش زیر برای استانداردسازی استفاده شده است:

$$NRM = \frac{X - \mu}{\sigma},$$

که در فرمول بالا X داده مورد نظر، μ میانگین و σ انحراف معیار را تصریح می‌کند.

۹. طراحی مدل

ابتدا داده‌ها به دو مجموعه آموزشی و آزمایشی تقسیم می‌شود. اغلب پژوهشگران، ۸۰ درصد داده‌ها را به عنوان داده‌های آموزشی و ۲۰ درصد باقی‌مانده را به عنوان داده‌های آزمایشی در نظر می‌گیرند. تعداد نمونه‌های هر یک از سهم‌ها ۷۰۰ روز است که ۵۶۰ روز اول به منظور آموزش و ۱۴۰ روز دیگر برای آزمایش در نظر گرفته می‌شود. سپس به کمک روش جستجوی شبکه‌ای، پارامترهای مهم مدل‌ها تخمین زده می‌شوند و بعد از اعمال پارامترهای بهینه، مدل آموزش داده می‌شود. در انتها، عملکرد مدل‌ها بر اساس نتایج داده‌های آزمایشی و معیارهای ارزیابی که به شرح زیر است، بررسی می‌شود:

(۱) میانگین مربعات خطا با رابطه‌ی محاسبه می‌شود:

$$MSE = \frac{1}{N} \sum_{i=1}^N (y_i - \hat{y}_i)^2.$$

(۲) جذر میانگین مربعات خطا با رابطه‌ی تعیین می‌شود و در واقع جذر MSE است:

$$RMSE = \sqrt{\frac{1}{N} \sum_{i=1}^N (y_i - \hat{y}_i)^2}.$$

(۳) میانگین خطای مطلق که با رابطه‌ی زیر به دست می‌آید:

از داده‌ها با جای‌گذاری و اندازه‌ی یکسان ارایه می‌شود. سپس مدل رگرسیون بردار پشتیبان مختلفی آموزش داده می‌شود و در نهایت بین مدل‌های رگرسیون بردار پشتیبان آموزش داده‌شده، میانگین‌گیری می‌شود و میانگین خروجی‌ها به عنوان مدل نهایی انتخاب می‌شود [۲۵، ۲۶].

۷. روش‌شناسی و شناخت داده‌ها

پیاده‌سازی مدل پیشنهادی با استفاده از برنامه‌نویسی در محیط نرم‌افزار آناکوندا و زبان برنامه‌نویسی پایتون با بهره‌گیری از SciKit-learn انجام شده است. سهم‌های استفاده شده در این پژوهش عبارت‌اند از: شرکت فولاد مبارکه سپاهان (فولاد)، شرکت سرمایه‌گذاری تأمین اجتماعی (شستا)، شرکت صنایع پتروشیمی خلیج فارس (فارس)، شرکت ملی صنایع مس ایران (فملی)، شرکت پالایش نفت اصفهان (شپنا)، پالایش نفت تهران (شتران)، پالایش نفت بندرعباس (شبندر)، ایران خودرو (خودرو)، سایپا (خسپا) و پتروشیمی پارس (پارس) که در محدوده‌ی زمانی ۱۴۰۰/۰۱/۰۷ تا ۱۴۰۳/۰۴/۳۱ است که دلیل این انتخاب تأثیر گذاری ۱۰ سهم (شاخص ساز) در ۳ سال گذشته بر روی شاخص بورس اوراق بهادار بوده‌اند^{۱۱}. این داده‌ها شامل ویژگی‌های بالاترین قیمت، پایین‌ترین قیمت، قیمت باز شدن و قیمت بسته‌شدن سهام مذکور در بازه‌ی زمانی مورد نظر است.

۸. آماده‌سازی داده‌های ورودی

در این گام به نرمال‌سازی داده‌ها پرداخته می‌شود. نرمال کردن داده‌ها، فرایندی است که در آن داده‌ها به فرم مناسبی برای الگوریتم‌ها تبدیل می‌شوند. برای نرمال کردن داده‌ها روش‌های مختلفی وجود دارد که به طور کلی پس از استانداردسازی، داده‌ها

جنگل تصادفی قرار می‌گیرد. هسته‌های گاوسی، چندجمله‌ای و سیگموئید عملکرد متوسطی از خود نشان داده‌اند.

جدول ۱. ارزیابی و مقایسه روش پیشنهادی و سایر روش‌های بر روی داده‌های نماد فولاد.

روش	R^2	MAE	RMSE	MSE
معمولی SVR	گاوسی	۰٫۹۳۳ ۳	۰٫۰۵۱ ۲	۰٫۰۵۲۴ ۰٫۰۰۲۷
	خطی	۰٫۹۹۸ ۷	۰٫۰۰۴ ۶	۰٫۰۰۷۱ ۰٫۰۰۰۰۰۵
	چندجمله‌ای	۰٫۹۳۵ ۶	۰٫۰۵۰ ۳	۰٫۰۵۱۵ ۰٫۰۰۲۶
	سیگموئید	۰٫۸۷۸ ۱	۰٫۰۶۹ ۴	۰٫۰۷۰۸ ۰٫۰۰۰۵
گروهی SVR	گاوسی	۰٫۹۲۸ ۹	۰٫۰۵۲ ۹	۰٫۰۵۴۱ ۰٫۰۰۲۹
	خطی	۰٫۹۹۸ ۷	۰٫۰۰۴ ۶	۰٫۰۰۷۲ ۰٫۰۰۰۰۰۵
	چندجمله‌ای	۰٫۹۲۷ ۱	۰٫۰۵۳ ۵	۰٫۰۵۴۷ ۰٫۰۰۰۳
	سیگموئید	۰٫۸۸۱ ۱	۰٫۰۶۸ ۳	۰٫۰۶۹۹ ۰٫۰۰۴۸
	درخت تصمیم	۰٫۹۹۲ ۹	۰٫۰۱۲ ۸	۰٫۰۱۷ ۰٫۰۰۰۰۲
جنگل تصادفی	۰٫۹۹۷ ۰	۰٫۰۰۸ ۴	۰٫۰۱۱ ۴	۰٫۰۱۱ ۰٫۰۰۰۰۱

جدول ۲. ارزیابی و مقایسه روش پیشنهادی و سایر روش‌های بر روی داده‌های نماد شستا.

روش	R^2	MAE	RMSE	MSE
معمولی SVR	گاوسی	۲٫۰۶۴۴ -	۰٫۰۲۷ ۹	۰٫۰۳۰۱ ۰٫۰۰۰۰۹
	خطی	۰٫۹۹۵۴ -	۰٫۰۱۱ ۹	۰٫۰۰۱۱ ۰٫۰۰۰۰۰۰۱
	چندجمله‌ای	۲۷٫۸۶۶ -	۰٫۰۹۲ ۴	۰٫۰۹۲۴ ۰٫۰۰۸۵
	سیگموئید	۰٫۸۶۲۲ -	۰٫۰۱۹ ۳	۰٫۰۲۳۴ ۰٫۰۰۰۰۵
SVR	گاوسی	۲٫۰۷۲۳ -	۰٫۰۲۶ ۹	۰٫۰۳۰۱ ۰٫۰۰۰۰۹

$$MAE = \frac{\sum_{i=1}^N |y_i - \hat{y}_i|}{N}$$

۴) معیار ارزیابی R^2 که ارتباط بین دو متغیر را اندازه‌گیری می‌کند و به ضریب تعیین شناخته می‌شود همچنین هرچقدر این معیار به یک نزدیک‌تر باشد مدل عملکرد بهتری خواهد داشت.

که در آنها N ، y_i و $\hat{y}_i$ به ترتیب تعداد کل مشاهدات، مقدار واقعی و مقدار پیش‌بینی شده‌ی متناظر را نمایش می‌دهند، همچنین هر اندازه مقدار این معیارها به صفر نزدیک باشند، مدل از عملکرد بهتری برخوردار است.

۱۰. برآورد مدل و نتایج

در این بخش، مدل پیشنهادی بر روی داده‌های آموزشی هر یک از ده نماد مورد نظر اعمال و نتیجه‌ی ارزیابی پیش‌بینی مدل بر روی داده‌های آزمایشی در جداول ۱ تا ۱۰ گزارش شده است. همچنین به‌منظور مقایسه، مدل‌های حاصل از روش SVR معمولی، درخت تصمیم و جنگل تصادفی نیز بر روی داده‌های آموزشی به‌دست آمدند و نتایج ارزیابی پیش‌بینی آنها در جداول ۱ تا ۱۰ آورده شده است.

باتوجه به نتایج به‌دست آمده در مورد نماد فولاد که در جدول ۱ گزارش شده است، عملکرد SVR معمولی با هسته‌ی خطی و همچنین روش SVR گروهی با هسته‌ی خطی با تفاوت بسیار اندکی دارای بهترین عملکرد نسبت به سایر روش‌ها و هسته‌ها هستند و پس از آنها روش جنگل تصادفی که یک روش گروهی است در رتبه دوم عملکرد قرار می‌گیرد. همچنین در مورد تمام هسته‌ها، SVR گروهی نسبت به SVR معمولی در برخی موارد بهبود ایجاد کرده است.

در مورد نماد شستا که نتایج ارزیابی در جدول ۲ آورده شده است، روش SVR گروهی با هسته خطی بهترین عملکرد را دارد و با اختلاف خوبی، SVR معمولی با هسته‌ی خطی و بعد از آن روش

سیگماید	۰,۹۹۷	۰,۰۳۷	۰,۰۴۶۸	۰,۰۰۲
درخت تصمیم	۰,۹۹۸	۰,۰۳۱	۰,۰۳۸۸	۰,۰۰۱
جنگل تصادفی	۰,۹۹۸	۰,۰۲۵	۰,۰۳۱۱	۰,۰۰۰
	۳	۷	۵	۹

جدول ۴. ارزیابی و مقایسه روش پیشنهادی و سایر روش‌های بر روی داده‌های نماد فملی.

روش	R^2	MAE	RMSE	MSE
گوسی	۰,۹۴۹۱	۰,۰۳۵	۰,۰۳۸۴	۰,۰۰۱
خطی	۰,۹۹۶۲	۰,۰۰۷	۰,۰۱۰۳	۰,۰۰۰
چندجمله‌ای	۰,۹۴۰۷	۰,۰۳۹	۰,۰۴۱۴	۰,۰۰۱
سیگماید	۰,۸۶۸	۰,۰۰۶	۰,۰۶۱۸	۰,۰۰۳
گوسی	۰,۹۳۹	۰,۰۳۹	۰,۰۴۲	۰,۰۰۱
خطی	۰,۹۹۶۲	۰,۰۰۷	۰,۰۱۰۳	۰,۰۰۰
چندجمله‌ای	۰,۹۳۱	۰,۰۴۲	۰,۰۴۴۷	۰,۰۰۲
سیگماید	۰,۸۶۸۵	۰,۰۰۶	۰,۰۶۱۷	۰,۰۰۳
درخت تصمیم	۰,۹۸۶۷	۰,۰۱۴	۰,۰۱۸۹	۰,۰۰۰
جنگل تصادفی	۰,۰۹۸۹	۰,۰۱۳	۰,۰۱۷۷	۰,۰۰۰
	۷	۱	۳	۳

بهترین عملکرد هستند. همچنین هسته گوسی و پس از آن جنگل قرار می‌گیرند.

نماد فملی که نتایج ارزیابی روش‌ها بر روی داده‌های آن در جدول ۴ آمده است نیز همچنان SVR معمولی و گروهی با هسته خطی بهترین عملکرد را از خود نشان می‌دهند و بعد از آنها جنگل تصادفی و درخت تصمیم قرار می‌گیرند.

نتایج ارزیابی روش‌ها بر روی داده‌های نماد شپنا که در جدول ۵ آمده است نشان می‌دهد که SVR گروهی با هسته خطی بهترین

خطی	۰,۹۹۶۸	۰,۰۰۰۶	۰,۰۰۰۹	۰,۰۰۰۰۰۰۰۰
چندجمله‌ای	۲۷,۸۶۶	۰,۰۹۲	۰,۰۹۲۴	۰,۰۰۰۸۵
سیگماید	-۱,۰۰۶	۰,۰۰۲	۰,۰۲۴۳	۰,۰۰۰۰۵
درخت تصمیم	۰,۹۸۲۳	۰,۰۰۱	۰,۰۰۲۲	۰,۰۰۰۰۰۰۰۵
جنگل تصادفی	۰,۹۹۲۷	۰,۰۰۱	۰,۰۰۱۴	۰,۰۰۰۰۰۰۰۲
	۴	۷	۱	۹

در مورد سایر هسته‌ها، اختلاف اندکی بین SVR معمولی و SVR گروهی وجود دارد. جنگل تصادفی نیز به عنوان یک مدل گروهی، عملکرد خوبی در پیش‌بینی سهام شستا دارد. از این رو، استفاده از مدل‌های ماشین بردار گروهی با هسته خطی و جنگل تصادفی به نظر می‌رسد که دارای عملکرد بهتری برای پیش‌بینی سهام شستا هستند.

در جدول ۳ که نتایج ارزیابی روش‌های مورد مطالعه در این مقاله بر روی داده‌های نماد فارس را در خود دارد، SVR معمولی و SVR گروهی با هسته‌های خطی و چندجمله‌ای دارای

جدول ۳. ارزیابی و مقایسه روش پیشنهادی و سایر روش‌های بر روی داده‌های نماد فارس.

روش	R^2	MAE	RMSE	MSE
گوسی	۰,۹۹۹	۰,۰۱۹	۰,۰۲۴۷	۰,۰۰۰
خطی	۰,۹۹۹	۰,۰۱۴	۰,۰۱۹۶	۰,۰۰۰
چندجمله‌ای	۰,۹۹۹	۰,۰۱۷	۰,۰۲۲۴	۰,۰۰۰
سیگماید	۰,۹۹۷	۰,۰۳۷	۰,۰۴۶۶	۰,۰۰۲
گوسی	۰,۹۹۹	۰,۰۱۹	۰,۰۲۴۳	۰,۰۰۰
خطی	۰,۹۹۹	۰,۰۱۴	۰,۰۲۰	۰,۰۰۰
چندجمله‌ای	۰,۹۹۹	۰,۰۱۴	۰,۰۱۹۳	۰,۰۰۰
	۲	۱	۹	۵

عملکرد را دارد و پس از آن با اختلاف نسبتاً خوبی SVR معمولی با هسته خطی قرار می‌گیرد.

جدول ۵. ارزیابی و مقایسه روش پیشنهادی و سایر روش‌های بر روی

داده‌های نماد شپنا.

روش	R^2	MAE	RMSE	MSE
معمولی SVR	گاوسی	۰,۹۹۴	۰,۰۳۴	۰,۰۰۱
	خطی	۰,۹۹۹	۰,۰۱۰	۰,۰۰۰
	چندجمله‌ای	۰,۹۹۷	۰,۰۲۸	۰,۰۰۰
	سیگموئید	۰,۹۹۶	۰,۰۲۶	۰,۰۰۱
	گاوسی	۰,۹۹۳	۰,۰۳۹	۰,۰۰۲
گروهی SVR	خطی	۰,۹۹۹	۰,۰۰۹	۰,۰۰۰
	چندجمله‌ای	۰,۹۹۸	۰,۰۲۰	۰,۰۰۰
	سیگموئید	۰,۹۹۷	۰,۰۲۴	۰,۰۰۰
	درخت تصمیم	۰,۹۹۶	۰,۰۲۱	۰,۰۰۱
SVR	جنگل تصادفی	۰,۹۹۵	۰,۰۱۸	۰,۰۰۱

جدول ۶. ارزیابی و مقایسه روش پیشنهادی و سایر روش‌های بر روی

داده‌های نماد خودرو.

روش	R^2	MAE	RMSE	MSE
معمولی SVR	گاوسی	۰,۹۹۳	۰,۰۱۹	۰,۰۰۰
	خطی	۰,۹۹۲	۰,۰۲	۰,۰۰۰
	چندجمله‌ای	۰,۹۹۱	۰,۰۲۳	۰,۰۰۱
	سیگموئید	۰,۹۲۷	۰,۰۸۳	۰,۰۰۸
SVR	گاوسی	۰,۹۹۳	۰,۰۲۱	۰,۰۰۰

خطی	۰,۹۹۳	۰,۰۱۸	۰,۰۲۶۳	۰,۰۰۰
چندجمله‌ای	۰,۹۹۰	۰,۰۲۶	۰,۰۳۳۶	۰,۰۰۱
سیگموئید	۰,۹۶۰	۰,۰۶۰	۰,۰۶۶۸	۰,۰۰۴
درخت تصمیم	۰,۹۷۹	۰,۰۳۶	۰,۰۴۷۹	۰,۰۰۲
جنگل تصادفی	۰,۹۷۷	۰,۰۳۹	۰,۰۵۰۵	۰,۰۰۲

در مورد نماد خودرو، براساس نتایج مندرج در جدول ۶ به‌وضوح برتری SVR گروهی با هسته خطی بر سایر روش‌های ارایه شده در این مقاله مشاهده می‌شود و بعد از آن روش‌های SVR گروهی و معمولی با هسته گاوسی و چندجمله‌ای قرار می‌گیرند. در مورد این نماد، جنگل تصادفی و درخت تصمیم در رده‌های آخر قبل از هسته‌ی سیگموئید قرار می‌گیرند.

نمادهای پارس و شبندر در جدول ۷ و ۸ به‌کمک مدل SVR معمولی و گروهی با هسته خطی برترین عملکرد را رقم زده‌اند. همچنین جنگل تصادفی و درخت تصمیم نیز با اختلاف قابل

جدول ۷. ارزیابی و مقایسه روش پیشنهادی و سایر روش‌های بر روی

داده‌های نماد پارس.

روش	R^2	MAE	RMSE	MSE
معمولی SVR	گاوسی	۰,۰۷۹	۰,۰۷۹۵	۰,۰۰۶۳
	خطی	۰,۹۷۳۱	۰,۰۰۰۲	۰,۰۰۰۰۰۰۰۰
	چندجمله‌ای	۰,۹۷۳۱	۰,۰۰۰۲	۰,۰۰۰۰۰۰۰۰
	سیگموئید	۰,۹۷۳۱	۰,۰۰۰۲	۰,۰۰۰۰۰۰۰۰
گروهی SVR	گاوسی	۰,۰۸	۰,۰۸۰۲	۰,۰۰۶۴
	خطی	۰,۹۷۳۴	۰,۰۰۰۲	۰,۰۰۰۰۰۰۰۰
	چندجمله‌ای	۰,۹۷۳۴	۰,۰۰۰۲	۰,۰۰۰۰۰۰۰۰

با هسته گاوسی و چندجمله‌ای بهترین عملکرد را از خود نشان داده‌اند.

جدول ۹. ارزیابی و مقایسه روش پیشنهادی و سایر روش‌های بر روی داده‌های نماد شتران.

روش	R^2	MAE	RMSE	MSE
معمولی SVR	گاوسی	۰.۸۸۱ ۸	۰.۰۹۷ ۸	۰.۰۱۳
	خطی	۰.۹۹۹ ۳	۰.۰۰۵ ۲	۰.۰۰۰۰۰۰۶
	چندجمله‌ای	۰.۹۵۳ ۷	۰.۰۷۰ ۵	۰.۰۰۵۱
	سیگموئید	۰.۹۰۵ ۶	۰.۰۹۴ ۱	۰.۰۱۰۴
گروهی SVR	گاوسی	۰.۸۵۹ ۴	۰.۰۹۹ ۴	۰.۰۱۵۵
	خطی	۰.۹۹۹ ۳	۰.۰۰۵ ۲	۰.۰۰۰۰۰۰۶
	چندجمله‌ای	۰.۹۷۲ ۷	۰.۰۵۳ ۷	۰.۰۰۰۳
	سیگموئید	۰.۹۲۷ ۸	۰.۰۸۲ ۶	۰.۰۰۷۹
	درخت تصمیم	۰.۸۲۰ ۹	۰.۰۹۸ ۱	۰.۰۱۹۷
	جنگل تصادفی	۰.۷۹۲ ۱	۰.۱۰۴ ۴	۰.۰۲۲۹

جدول ۱۰. ارزیابی و مقایسه روش پیشنهادی و سایر روش‌های بر روی داده‌های نماد خسپا.

روش	R^2	MAE	RMSE	MSE
معمولی SVR	گاوسی	۰.۹۸۸ ۵	۰.۰۳۰ ۶	۰.۰۰۱ ۳
	خطی	۰.۹۹۲ ۱	۰.۰۲۱ ۱	۰.۰۰۰ ۹
	چندجمله‌ای	۰.۹۸۸ ۹	۰.۰۲۸ ۶	۰.۰۰۱ ۳
	سیگموئید	۰.۹۰۳ ۲	۰.۱۰۱ ۶	۰.۰۱۱ ۷

سیگموئید	۱۳۲۲,۴ -	۰.۰۵۳ ۳	۰.۰۵۶۲	۰.۰۰۳۱
درخت تصمیم	۰.۸۹۵۲	۰.۰۰۰ ۴	۰.۰۰۰۵	۰.۰۰۰۰۰۰۰۲
جنگل تصادفی	۰.۹۵۴۷	۰.۰۰۰ ۲	۰.۰۰۰۳	۰.۰۰۰۰۰۰۰۱

جدول ۸. ارزیابی و مقایسه روش پیشنهادی و سایر روش‌های بر روی داده‌های نماد شبندر.

روش	R^2	MAE	RMSE	MSE
معمولی SVR	گاوسی	۰.۹۷۵ ۹	۰.۰۲۶ ۳	۰.۰۰۰ ۹
	خطی	۰.۹۹۴ ۴	۰.۰۰۸ ۲	۰.۰۰۰ ۲
	چندجمله‌ای	۰.۹۳۸ ۸	۰.۰۴۶ ۱	۰.۰۰۲ ۳
	سیگموئید	۰.۵۹۱ ۸	۰.۱۱۰ ۲	۰.۰۱۵ ۷
گروهی SVR	گاوسی	۰.۹۶۸ ۹	۰.۰۳۰ ۲	۰.۰۰۱ ۱
	خطی	۰.۹۹۴ ۵	۰.۰۰۸ ۲	۰.۰۰۰ ۲
	چندجمله‌ای	۰.۹۶۷ ۷	۰.۰۳۲ ۱	۰.۰۰۱ ۲
	سیگموئید	۰.۶۴۲ ۷	۰.۱۰۲ ۴	۰.۰۱۳ ۸
	درخت تصمیم	۰.۹۸۴ ۷	۰.۰۱۷ ۲	۰.۰۰۰ ۵
	جنگل تصادفی	۰.۹۸۷ ۶	۰.۰۱۵ ۵	۰.۰۰۰ ۴

توجهی در رده‌های بعدی قرار می‌گیرند. در مورد نماد شبندر شاهد بهبود عملکرد SVR گروهی نسبت به SVR معمولی هستیم.

نمادهای شتران و خسپا نیز همچون اکثر نمادهای فوق بهترین عملکرد را با SVR معمولی و گروهی توسط هسته‌ی خطی کسب کرده‌اند. در مورد نماد شتران پس از هسته‌ی خطی، SVR گروهی با هسته چندجمله‌ای و نماد خسپا، روش SVR معمولی و گروهی

گروهی SVR	گاوسی	۰,۹۸۹	۰,۰۲۹	۰,۰۳۶۵	۰,۰۰۱	۳
	خطی	۰,۹۹۲	۰,۰۲۱	۰,۰۳۰۸	۰,۰۰۰	۹
	چندجمله‌ای	۰,۹۸۶	۰,۰۳۳	۰,۰۳۹۸	۰,۰۰۱	۵
	سیگموید	۰,۹۲۶	۰,۰۸۷	۰,۰۹۴۷	۰,۰۰۸	۹
درخت تصمیم		۰,۹۶۷	۰,۰۴۴	۰,۰۶۲۹	۰,۰۰۳	۹
جنگل تصادفی		۰,۹۸۲	۰,۰۳۳	۰,۰۴۶	۰,۰۰۲	۱

این مدل می‌تواند به‌عنوان یک ابزار قوی در تصمیم‌گیری‌های سرمایه‌گذاری مورد استفاده قرار گیرد و مسیرهای پژوهشی جدیدی را برای توسعه آینده این حوزه ارائه دهد. با اینکه این مقاله سعی در پاسخ به چالش‌های مهم در حوزه پیش‌بینی قیمت سهام داشته است، اما هنوز فرصت‌های بسیاری برای تحقیقات بیشتر و بهبود روش‌های پیشین وجود دارد. در این مقاله، یک روش جدید برای پیش‌بینی قیمت سهام با استفاده از ترکیب مدل‌های SVR با بگینگ معرفی شده است.

۱۲. تجزیه و تحلیل کلی نتایج

بر اساس نتایج به‌دست‌آمده از تجزیه و تحلیل داده‌های سهام مختلف شامل فولاد، شستا، فارس، فملی، شپنا، شتران، شبندر، خودرو، خساپا و پارس با استفاده از مدل‌های SVR معمولی و گروهی و جنگل تصادفی، می‌توان به نتایج و توصیه‌های زیر دست‌یافت:

مدل‌های SVR معمولی و گروهی با هسته خطی در بیشتر موارد، به‌عنوان مدل‌های با عملکرد برتر هستند.

پیشنهادها برای مطالعات آتی:

برای تحقیقات آینده، می‌توان پژوهش بیشتر در زمینه بهینه‌سازی پارامترها و افزودن ویژگی‌های جدید به مدل را مدنظر قرار داد. این اقدامات می‌توانند به بهبود عملکرد مدل و افزایش دقت در پیش‌بینی‌ها کمک کنند. همچنین، تطبیق مدل با محیط‌های پویا و نوسانی بازار یک چالش مهم است که ممکن است تأثیر مستقیم بر قابلیت پیش‌بینی مدل داشته باشد. ترکیب بگینگ و SVR در این مقاله به‌خوبی نشان داده است که چگونه این چالش می‌تواند با استفاده از چندین مدل و ترکیب پیش‌بینی‌ها کاهش یابد. این روش نه تنها دقت بالاتری در پیش‌بینی ارائه می‌دهد؛ بلکه مدل را مقاومت بیشتری در برابر تغییرات سریع بازار به ارمغان می‌آورد [۲۷].

نتایج اعمال روش پیشنهادی بر روی ده نماد مورد بررسی در این مقاله نشان می‌دهند که SVR معمولی و گروهی با هسته خطی در اکثر نمادها بهترین عملکرد در پیش‌بینی را به خود اختصاص داده‌اند و بنابراین استفاده از این هسته در روش پیشنهادی توصیه می‌شود. لازم به ذکر است که این توصیه بر مبنای مشاهدات این ده نماد در بازه زمانی مورد نظر در این مقاله است و این نتیجه‌گیری قطعی و حتمی نخواهد بود. از طرف دیگر، نتایج نشان می‌دهند که روش SVR معمولی با SVR گروهی از لحاظ معیارهای ارزیابی مطرح شده در این مقاله با اختلاف خیلی کمی از یکدیگر در اکثر موارد فاصله دارند که قابل چشم‌پوشی است. از آنجا که به‌کمک فناوری پردازش موازی، هزینه‌ی محاسباتی روش SVR گروهی می‌تواند به‌طور قابل ملاحظه‌ای از روش SVR معمولی کمتر باشد لذا استفاده از روش SVR گروهی با معیارهای ارزیابی تقریباً مشابه روش معمولی به‌صرفه‌تر و ارزشمندتر خواهد بود.

۱۱. نتیجه‌گیری و ارائه پیشنهاد

روش پیشنهادی با استفاده از ترکیب بگینگ و SVR نشان داده است که چگونه می‌تواند دقت پیش‌بینی قیمت سهام را افزایش دهد. این روش نه تنها در مقابل تغییرات بازار مقاومت بیشتری ایجاد می‌کند، بلکه قابلیت تعمیم به شرایط جدید را نیز داراست.

مقاله‌ی حاضر پیشنهادهای ارزشمندی مطرح کردند، کمال سپاسگزاری را دارند.

مراجع

[1] R. Oprea and O. Stoica, "Capital Markets Integration and Economic Growth," *Montenegrin Journal of Economics*, vol. 14, no. 3, pp. 23–35, 2018. doi: [10.14254/1800-5845/2018.14-3.2](https://doi.org/10.14254/1800-5845/2018.14-3.2).

[۲] م. نوشیار و ع. قاسمی مرزبالی، «پیش‌بینی روزانه قیمت برق با روش مبتنی بر ماشین یادگیری شدید، سیستم پیش‌پردازش‌کننده و الگوریتم بهبود یافته‌ی کلونی جستجوی ویروس»، *هوش محاسباتی در مهندسی برق*، جلد ۱۰، شماره ۲، صص ۷۳–۸۶، ۱۳۹۸. doi: [10.22108/isee.2019.116627.1215](https://doi.org/10.22108/isee.2019.116627.1215)

[۳] س. باجلان، س. فلاح‌پور و ن. دانا، «پیش‌بینی روند تغییرات قیمت سهم با استفاده از ماشین بردار پشتیبان وزن‌دهی شده و انتخاب ویژگی هیبرید به منظور ارائه استراتژی معاملاتی بهینه»، *راهبرد مدیریت مالی*، جلد ۴، شماره ۳، صفحات ۱۲۱–۱۴۸، ۱۳۹۵. doi: [10.22051/jfm.2016.2575](https://doi.org/10.22051/jfm.2016.2575)

[۴] ح. ر. کردلویی و ف. تیموری، «بررسی مقایسه‌ای توان مدل‌های ترکیب گوسی و ماشین بردار پشتیبان در تشخیص و پیش‌بینی حباب قیمتی»، *مهندسی مالی و مدیریت اوراق بهادار*، جلد ۲۳، شماره ۶، صفحات ۷۹–۱۰۴، ۲۰۱۵. doi: [20.1001.1.22519165.1394.6.23.5.3](https://doi.org/20.1001.1.22519165.1394.6.23.5.3)

[5] Q. Ilyas, K. Iqbal, S. Ijaz, A. Mehmood, and S. Bhatia, "A Hybrid Model to Predict Stock Closing Price Using Novel Features and a Fully Modified Hodrick–Prescott Filter," *Electronics*, vol. 11, p. 3588, 2022. doi: [10.3390/electronics11213588](https://doi.org/10.3390/electronics11213588).

[6] H. Huang, W. Zhang, G. Deng, and J. Chen, "Predicting Stock Trend Using Fourier

پژوهش‌های آتی می‌تواند به سمت بهینه‌سازی پارامترهای مدل، ارتقای تکنیک‌های بگینگ، و استفاده از داده‌های نوآورانه برای آموزش مدل ادامه یابند. همچنین، بررسی تأثیر ترکیب مدل‌های دیگر به جای SVR نیز می‌تواند یک جنبه جالب برای پژوهش‌های آینده باشد و به سرمایه‌گذاران امکان می‌دهد تا با در نظر گرفتن تغییرات وضعیت بازار، تصمیمات بهتری اتخاذ کنند.

در تحقیقات آتی، می‌توان مدل‌های دیگری را نیز در مقایسه با مدل‌های حاصل از این مطالعه در نظر گرفت. از جمله مدل‌های عمیق یادگیری و شبکه‌های عصبی می‌تواند موضوعات مناسبی برای بررسی باشند. تحلیل تفصیلی‌تر نتایج و تأثیرپذیری از عوامل مختلف مانند حجم معاملات، شاخص‌های اقتصادی، و اخبار بازار می‌تواند به بهترین استفاده از این مدل‌ها کمک کند.

انجام تحقیقات بر روی دوره‌های آتی بازار سهام به منظور پیش‌بینی عملکرد مدل‌ها در آینده و تعیین احتمال موفقیت آنها می‌تواند مفید باشد. تلاش برای بهبود عملکرد مدل‌ها از طریق بهینه‌سازی پارامترها، افزایش حجم داده‌ها، و استفاده از تکنیک‌های پیشرفته یادگیری ماشین می‌تواند تحقیقات آینده را بهبود بخشد.

مطالعات آینده می‌توانند معیارهای جدیدی را در تجزیه و تحلیل ارزیابی عملکرد مدل‌ها معرفی کنند تا نتایج بهتر و مطمئن‌تری حاصل شود. تحقیقات بیشتر در این حوزه می‌تواند به بهبود روش‌ها و افزایش دقت پیش‌بینی در بازار سهام کمک کند و تصمیم‌گیری‌های مبتنی بر داده را تسهیل نماید.

سپاسگزاری

نویسندگان این مقاله از همفکری اعضای هیات تحریریه مجله علمی محاسبات نرم و داوران محترمی که در تکمیل و اصلاح

[12] P. Samad, S. Mutalib, and S. Rahman, "Analytics of stock market prices based on machine learning algorithms," *Indonesian Journal of Electrical Engineering and Computer Science*, vol. 16, no. 2, pp. 1050–1058, 2019. doi: [10.11591/ijeecs.v16.i2.pp1050-1058](https://doi.org/10.11591/ijeecs.v16.i2.pp1050-1058).

[13] H. Liu, L. Qi, and M. Sun, "Short-Term Stock Price Prediction Based on CAE-LSTM Method," *Wireless Communications and Mobile Computing*, vol. 2022, pp. 1–7, 2022. doi: [10.1155/2022/4809632](https://doi.org/10.1155/2022/4809632).

[14] Z. Rustam and P. Kintandani, "Application of Support Vector Regression in Indonesian Stock Price Prediction with Feature Selection Using Particle Swarm Optimisation," *Modelling and Simulation in Engineering*, vol. 2019, pp. 1–5, 2019. doi: [10.1155/2019/8962717](https://doi.org/10.1155/2019/8962717).

[۱۵] م. منادی و ا. نجفی، «مدل بهینه سازی سبد سرمایه گذاری مبتنی بر پیش بینی با استفاده از رگرسیون بردار پشتیبان»، *مهندسی مالی و مدیریت اوراق بهادار*، جلد ۵۵، شماره ۱۴، صفحات ۲۵۲–۲۷۰، ۱۴۰۲.

[۱۶] ب. دهقان خانقاهی، ج. بحری ثالث، س. جبارزاده کنگرلویی، و ع. آشتاب، «بررسی مقایسه ای دقت پیش بینی مدل های ماشین بردار پشتیبان، شبکه بیزین و سی فایو در پیش بینی قیمت گذاری کمتر از واقع شرکت های پذیرفته شده در بورس اوراق بهادار و فرابورس»، *مهندسی مالی و مدیریت اوراق بهادار*، جلد ۱۱، شماره ۴۴، صفحات ۹۵–۱۱۳، ۱۳۹۹.

[۱۷] ع. خسروی نژاد و م. شعبانی صدر پیشه، «ارزیابی مدل های خطی و غیرخطی در پیش بینی شاخص قیمت سهام در بورس اوراق بهادار تهران»، *اقتصاد مالی*، جلد ۲۷، شماره ۸، صفحات ۵۱–۶۴، ۱۳۹۳.

[۱۸] س. ا. ح. منجمی، م. ابزری، و ع. رعیتی شوازی، «پیش بینی قیمت سهام در بازار بورس اوراق بهادار با استفاده از شبکه ی

Transform and Support Vector Regression," in *Proceedings of the 2014 IEEE 17th International Conference on Computational Science and Engineering (CSE)*, pp. 213–216, 2014. doi: [10.1109/CSE.2014.70](https://doi.org/10.1109/CSE.2014.70).

[7] S. A. Basher, A. A. Haug, and P. Sadorsky, "The impact of economic policy uncertainty and commodity prices on CARB country stock market volatility," *MPRA Paper 96577*, University Library of Munich, Germany, 2019. [Online]. Available: <https://ideas.repec.org/p/pra/mprapa/96577.html>

[8] N. Passalis, A. Tefas, J. Kannianen, M. Gabbouj, and A. Iosifidis, "Temporal Bag-of-Features Learning for Predicting Mid Price Movements Using High Frequency Limit Order Book Data," *IEEE Transactions on Emerging Topics in Computational Intelligence*, pp. 1–12, 2018. doi: [10.1109/TETCI.2018.2872598](https://doi.org/10.1109/TETCI.2018.2872598).

[9] M. Obthong, N. Tantisantiwong, W. Jeamwathanachai, and G. Wills, "A Survey on Machine Learning for Stock Price Prediction: Algorithms and Techniques," in *Proceedings of the 12th International Joint Conference on Computational Intelligence (IJCCI 2020)*, pp. 63–71, 2020. doi: [10.5220/0009340700630071](https://doi.org/10.5220/0009340700630071).

[10] P. Sadorsky, "A Random Forests Approach to Predicting Clean Energy Stock Prices," *Journal of Risk and Financial Management*, vol. 14, no. 2, p. 48, 2021. doi: [10.3390/jrfm14020048](https://doi.org/10.3390/jrfm14020048).

[11] B. Egüz, F. E. Çorbacı, and T. Kaya, "Stock Price Prediction of Turkish Banks Using Machine Learning Methods," in *Intelligent and Fuzzy Techniques for Emerging Conditions and Digital Transformation - Proceedings of the INFUS 2021 Conference*, C. Kahraman et al., Eds., vol. 308, pp. 222–229, Springer Science and Business Media Deutschland GmbH, 2022. doi: [10.1007/978-3-030-85577-2_26](https://doi.org/10.1007/978-3-030-85577-2_26).

[۲۵] س. م. حسینی، م. ابتیاع، و ر. خوجانی، «یک روش گروهی مبتنی بر بگینگ SVM برای مساله‌ی رتبه‌بندی اعتباری»، *محاسبات نرم*، در حال انتشار. doi: [10.22052/scj.2024.246330.1063](https://doi.org/10.22052/scj.2024.246330.1063).

[26] A. Smola and B. Schölkopf, "A Tutorial on Support Vector Regression," *Statistics and Computing*, vol. 14, no. 3, pp. 199–222, 2004. doi: [10.1023/B:STCO.0000035301.49549.88](https://doi.org/10.1023/B:STCO.0000035301.49549.88).

[27] M. Sedighi, H. Jahangirnia, M. Gharakhani, and S. Farahani Fard, "A Novel Hybrid Model for Stock Price Forecasting Based on Metaheuristics and Support Vector Machine," *Data*, vol. 4, no. 2, p. 75, 2019. doi: [10.3390/data4020075](https://doi.org/10.3390/data4020075).

عصبی فازی و الگوریتم‌های ژنتیک و مقایسه‌ی آن با شبکه‌ی عصبی مصنوعی»، *فصلنامه علمی پژوهشی اقتصاد مقداری*، جلد ۶، شماره ۳، صفحات ۱–۲۶، ۱۳۸۸. doi: [10.22055/jqe.2009.10697](https://doi.org/10.22055/jqe.2009.10697).

[۱۹] م. رمضانی و ا. عاملی، «پیش‌بینی قیمت سهام با استفاده از شبکه عصبی فازی مبتنی بر الگوریتم ژنتیک و مقایسه با شبکه عصبی فازی»، *فصلنامه تحقیقات مدل سازی اقتصادی*، جلد ۶، شماره ۲۲، صفحات ۶۱–۹۱، ۱۳۹۴. doi: [10.18869/acadpub.jemr.6.22.61](https://doi.org/10.18869/acadpub.jemr.6.22.61).

[۲۰] م. زمانی، س. و. افسر تقفی، و ا. بیات، «سیستم خبره پیش‌بینی قیمت سهام و بهینه‌سازی سبد سهام با استفاده از شبکه‌های عصبی فازی، مدل‌سازی فازی و الگوریتم ژنتیک»، *مهندسی مالی و مدیریت اوراق بهادار*، جلد ۵، شماره ۲۱، صفحات ۱۰۷–۱۳۰، ۱۳۹۳.

[21] I. Nti, A. Adekoya, and B. Weyori, "Efficient Stock-Market Prediction Using Ensemble Support Vector Machine," *Open Computer Science*, vol. 10, pp. 154–163, 2020. doi: [10.1515/comp-2020-0199](https://doi.org/10.1515/comp-2020-0199).

[22] A. Li, Q. Wei, Y. Shi, and Z. Liu, "Research on stock price prediction from a data fusion perspective," *Data Science in Finance and Economics*, vol. 3, pp. 230–250, 2023. doi: [10.3934/DSFE.2023014](https://doi.org/10.3934/DSFE.2023014).

[۲۳] ف. کلهری و س. م. حسینی، «خوشه‌بندی قیمت سهام با استفاده از الگوریتم کوچک‌ترین درخت پوشا»، *محاسبات نرم*، در حال انتشار. doi: [10.22052/scj.2024.253526.1183](https://doi.org/10.22052/scj.2024.253526.1183).

[۲۴] م. ابتیاع، س. م. حسینی، و ر. خوجانی، «رتبه‌بندی اعتباری مشتریان بانک به کمک روش جدید گروهی بر پایه ماشین بردار پشتیبان: مطالعه موردی بانک پاسارگاد»، *محاسبات نرم*، جلد ۱۰، شماره ۲، صفحات ۲–۱۵، ۱۴۰۰. doi: [10.22052/scj.2022.243227.1016](https://doi.org/10.22052/scj.2022.243227.1016).