

مدل‌سازی فرآیند در کوره گندله‌سازی با استفاده از روش‌های هوشمند: مطالعه موردی در شرکت معدنی و صنعتی گل‌گهر

محدثه رضایی^{۱*}، حسین نظام‌آبادی‌پور^۲، رضا خاکسارپور^۳

^۱ دانشجوی دکتری بخش مهندسی برق، دانشکده فنی و مهندسی، دانشگاه شهید باهنر کرمان، کرمان، ایران

^۲ استاد بخش مهندسی برق، دانشکده فنی و مهندسی، دانشگاه شهید باهنر کرمان، کرمان، ایران

^۳ کارشناسی ارشد، مجتمع معدنی و صنعتی گل‌گهر سیرجان، کرمان، ایران

چکیده

فرآیند گندله‌سازی یکی از مراحل کلیدی زنجیره تولید فولاد است که در آن کنسانتره سنگ‌آهن به گندله‌های قابل مصرف در کوره بلند تبدیل می‌شود. مرحله پخت در کوره، به دلیل تأثیر مستقیم توزیع دما و فشار بر ویژگی‌های فیزیکی و مکانیکی گندله، یکی از عوامل اصلی در کیفیت محصول است. هدف این پژوهش، توسعه مدلی داده‌محور برای پیش‌بینی رفتار کوره گندله‌سازی شرکت معدنی و صنعتی گل‌گهر بر پایه داده‌های عملیاتی واقعی است. بدین منظور، متغیرهای ورودی و خروجی فرآیند جمع‌آوری و پس از پیش‌پردازش، انتخاب ویژگی با سه روش امتیاز-اف، اطلاعات متقابل و ضریب همبستگی پیرسون انجام شد. سپس مدل‌های رگرسیون خطی چندگانه، جنگل تصادفی، k-نزدیک‌ترین همسایه و شبکه عصبی پرسپترون چندلایه در چارچوب چندورودی-چندخروجی برای پیش‌بینی همزمان ۳۶ متغیر دما و ۳۱ متغیر فشار آموزش داده شدند. نتایج نشان داد روش انتخاب ویژگی مبتنی بر اطلاعات متقابل موجب بهبود عملکرد مدل‌های غیرخطی شده و شبکه عصبی چندلایه بالاترین دقت را با میانگین ضریب تعیین ۹۱/۵۰ درصد برای دما و ۸۹/۶۲ درصد برای فشار به دست آورد. یافته‌ها بیانگر آن است که استفاده از رویکردهای یادگیری داده‌محور می‌تواند رفتار پیچیده کوره را با دقت بالا مدل کرده و بستر مناسبی برای طراحی سامانه‌های کنترلی هوشمند و بهینه‌سازی شرایط پخت فراهم سازد.

مقاله پژوهشی

تاریخ دریافت:

۱۴۰۴/۰۵/۳۱

تاریخ پذیرش:

۱۴۰۴/۰۸/۰۴

کلیدواژه‌ها:

مدل‌سازی هوشمند، یادگیری ماشین، انتخاب ویژگی، شبکه عصبی پرسپترون چندلایه، کوره گندله‌سازی

نویسنده مسئول:

mo.rezaei@eng.uk.ac.ir

doi : 10.22034/ABMIR.2025.23568.1156

E-ISSN: [2821-2037](#)

© 2025. Published by Yazd University This is an open access article

under the CC BY 4.0 License (<https://creativecommons.org/licenses/by/4.0/>).



۱- مقدمه

محیطی به حدود ۱۳۰۰ درجه سانتی‌گراد در مرحله پخت افزایش یافته و سپس در نواحی خنک‌کن تا حدود ۲۰۰ درجه سانتی‌گراد کاهش می‌یابد. کیفیت گندله‌های تولید شده، به شدت تحت تأثیر نسبت ترکیب مواد خام، ارتفاع بستر و شرایط حرارتی نظیر دما و فشار در بخش‌های مختلف کوره قرار دارد. تغییرات ناخواسته در این پارامترها می‌تواند منجر به کاهش استحکام مکانیکی، افزایش مصرف انرژی و بروز مشکلات زیست‌محیطی شود [۱۱].

در عمل، کنترل دما و فشار اغلب بر تجربه اپراتورها متکی است که دستیابی به تعادل حرارتی پایدار را دشوار می‌سازد. از این‌رو، استفاده از مدل‌های پیش‌بینی مبتنی بر داده، برای ارزیابی اثر تغییرات ورودی پیش از اجرای واقعی، ضروری به نظر می‌رسد. در سال‌های اخیر، مطالعات متعددی از هوش مصنوعی و یادگیری ماشین برای مدل‌سازی و کنترل فرآیندهای پیچیده صنعتی از جمله گندله‌سازی استفاده کرده‌اند. برای مثال، مدل‌های یادگیری عمیق در پیش‌بینی اندازه گندله، کنترل فرآیند پخت و بهینه‌سازی مصرف انرژی به‌کار رفته‌اند [۱۲]، و شبکه‌های عصبی پیچشی^۷ (CNN) نیز برای پیش‌برخظ اندازه گندله توسعه یافته‌اند [۱۳]. نتایج این پژوهش‌ها نشان داده که استفاده از فناوری‌های یادگیری ماشین می‌تواند دقت پیش‌بینی و کارایی فرآیندهای معدنی همچون تف‌جوشی و گندله‌سازی را افزایش داده و هم‌زمان با بهبود بهره‌وری انرژی، به پایداری زیست‌محیطی کمک کند [۱۴-۱۶].

در حوزه مدل‌سازی فرآیندهای حرارتی، دو رویکرد اصلی در ادبیات علمی وجود دارد: مدل‌های مکانیکی/فیزیکی^۸ و مدل‌های داده‌محور^۹. مدل‌های فیزیکی بر مبنای قوانین انتقال حرارت، جرم و واکنش‌های سینتیکی توسعه یافته‌اند و می‌توانند رفتار فرآیند را با دقت بالا توصیف کنند. برای نمونه، در [۸]، مدلی ریاضی برای فرآیند استحکام گندله‌ها بر پایه شبکه مستقیم ارائه شده که توزیع دما، نرخ خشک شدن و واکنش‌های اکسیداسیون را شبیه‌سازی

سنگ آهن به عنوان اصلی‌ترین ماده اولیه در تولید فولاد، نقش حیاتی در صنایع معدنی و متالورژی ایفا می‌کند. این ماده معدنی حدود ۵ درصد از پوسته زمین را تشکیل می‌دهد و عمدتاً به صورت اکسیدهای هماتیت و مغنتیت یافت می‌شود [۱]. فرآیند استخراج و فرآوری سنگ آهن شامل مراحل خردایش، تغلیظ و جداسازی ناخالصی‌ها است که در نهایت به تولید کنسانتره سنگ آهن منجر می‌شود. کنسانتره، که حاوی بیش از ۶۵ درصد آهن است، برای استفاده مستقیم در کوره‌های احیای فولاد مناسب نیست، زیرا ذرات ریز آن باعث ایجاد گرد و غبار، اختلال در جریان گاز و مشکلات حمل‌ونقل می‌گردد [۲، ۳]. بنابراین، گندله‌سازی^۱ به عنوان یک فرآیند کلیدی برای تبدیل کنسانتره به گلوله‌های کروی ضروری است و مواد مناسب برای احیای مستقیم یا کوره بلند را فراهم می‌کند [۴، ۵]. این فرآیند که از دهه ۱۹۵۰ میلادی آغاز شده، امروزه به یکی از گلوگاه‌های بهبود کیفیت و بهره‌وری در زنجیره تولید فولاد تبدیل شده است [۶].

در فرایند گندله‌سازی، کنسانتره با مواد افزودنی مانند بنتونیت، سنگ آهک و آب مخلوط شده و در دیسک‌ها یا استوانه‌های چرخان به گلوله تبدیل می‌شود، سپس در کوره پخت حرارت می‌بیند تا استحکام مکانیکی لازم را کسب کند [۷]. روش‌های اصلی گندله‌سازی شامل شبکه-کوره-خنک‌کن^۲ و شبکه مستقیم^۳ است که هر کدام مزایا و چالش‌های خاصی دارند [۸]. برای مثال، روش شبکه مستقیم که در شرکت معدنی و صنعتی گل‌گهر مورد استفاده قرار می‌گیرد، با کاهش خرد شدن گندله‌ها و تنظیم دقیق عملیات حرارتی، کیفیت بالاتری ارائه می‌دهد [۹].

کوره گندله‌سازی مورد مطالعه در این پژوهش، از نوع شبکه مستقیم و با طول ۱۵۶ متر، شامل نواحی خشک‌کن دمشی^۴ (UDD) و مکشی^۵ (DDD)، پیش‌گرمایش^۶ (PH)، پخت^۷ (F و AF) و خنک‌کن^۸ (C1 و C2) است [۱۰]. در این فرآیند، دما از شرایط

⁷ Firing

⁸ Cooling

⁹ Convolutional Neural Network (CNN)

¹⁰ Mechanistic/Physical Models

¹¹ Data-Driven Models

¹ Pelletizing

² Grate-Kiln Cooler

³ Straight-Grate

⁴ Updraft Drying

⁵ Downdraft Drying

⁶ Pre-Heating

- طرح و ارزیابی مدل‌های چندورودی-چندخروجی^۴ (MIMO) برای پیش‌بینی هم‌زمان دما و فشار در نواحی مختلف کوره.
 - ارائه تحلیل عملیاتی (پیشنهاد برای کاربرد در سیستم‌های کنترلی صنعتی) و راهنمایی برای انتخاب مدل بر اساس محدودیت‌های محاسباتی و مقیاس خروجی.
- ادامه مقاله بدین شرح سازمان‌دهی شده است: در بخش دوم مرور ادبیات موضوعی و در بخش سوم مفاهیم پایه‌ای که برای هدف این پژوهش مورد توجه هستند، بیان می‌شوند. در بخش چهارم روش پیشنهادی و در بخش پنجم آزمایش‌ها و نتایج بدست آمده ارائه می‌شود. در بخش ششم نتیجه‌گیری و پیشنهادات برای تحقیقات آینده بیان می‌شود.

۲- مفاهیم پایه

برای درک بهتر روش پیشنهادی و تحلیل نتایج، آشنایی با مفاهیم اصلی صنعتی و علمی مرتبط با موضوع پژوهش ضروری است. در این بخش، ابتدا مروری بر فرآیند گندله‌سازی و عملکرد کوره پخت ارائه می‌شود و سپس مفاهیم یادگیری ماشین و شبکه عصبی پرسپترون چندلایه که در مدل‌سازی به کار رفته‌اند، توضیح داده خواهد شد.

۲-۱ فرآیند گندله‌سازی

گندله‌سازی فرآیندی است که در آن کنسانتره سنگ آهن با افزودنی‌هایی مانند بنتونیت، سنگ آهک و آب مخلوط شده و به شکل گلوله‌های کروی با قطر ۸ تا ۱۶ میلی‌متر در می‌آید. هدف از این عملیات، بهبود قابلیت حمل، کاهش گرد و غبار، و فراهم‌سازی شرایط بهینه برای فرآیندهای بعدی مانند احیای مستقیم یا استفاده در کوره بلند است.

پس از تشکیل گندله خام، این ذرات وارد کوره پخت می‌شوند تا با اعمال پروفیل حرارتی مناسب، استحکام مکانیکی و ویژگی‌های فیزیکی مورد نظر را به‌دست آورند. کوره‌های گندله‌سازی به دو دسته اصلی تقسیم می‌شوند: الف) روش شبکه-کوره-خنک‌کن، ب) روش شبکه مستقیم. در کوره‌های شبکه مستقیم، از جمله کوره مورد مطالعه در این پژوهش، فرآیند پخت در چند ناحیه متوالی

می‌کند. همچنین در [۱۷، ۱۸]، مدل‌های فیزیکی جامعی برای تحلیل واکنش‌های شیمیایی، اکسیداسیون مگنتیت و احتراق سوخت در نواحی مختلف کوره ارائه شده است. با وجود دقت بالا، این مدل‌ها نیازمند داده‌های آزمایشگاهی دقیق و توان محاسباتی زیاد هستند، از این رو در کاربردهای کنترل برخط چندان عملیاتی نیستند.

در مقابل، مدل‌های داده‌محور با استفاده از داده‌های واقعی کارخانه و الگوریتم‌های یادگیری ماشین قادرند رفتار سیستم را بدون نیاز به مدل‌سازی دقیق فیزیکی پیش‌بینی کنند. مطالعات متعددی [۱۹، ۲۰] نشان داده‌اند که رویکردهای یادگیری ماشین، از جمله شبکه‌های عصبی، رگرسیون جنگل تصادفی و روش‌های ترکیبی، می‌توانند توزیع دما در انواع کوره‌ها را با دقت بالاتری نسبت به مدل‌های ساده فیزیکی پیش‌بینی کنند و در کاربردهای بلادرنگ نیز قابل پیاده‌سازی‌اند.

بر این اساس، پژوهش حاضر بر استخراج مدل ورودی-خروجی برای کوره گندله‌سازی با استفاده از روش‌های هوش مصنوعی تمرکز دارد. اهداف شامل جمع‌آوری و پیش‌پردازش داده‌ها، تحلیل آماری روابط متغیرها، اعمال شیفت زمانی برای همگام‌سازی ورودی-خروجی، انتخاب ویژگی و مدل‌سازی با الگوریتم‌های یادگیری ماشین مانند رگرسیون خطی چندگانه^۱ (MLR)، رگرسیون جنگل تصادفی^۲ (RFR)، رگرسیون k -همسایه نزدیک‌تر^۳ (KNN) و شبکه‌های عصبی عمیق است. این رویکرد می‌تواند ابزار مفیدی برای کنترل پیشرفته فرآیند و دستیابی به گندله با کیفیت مطلوب فراهم کند.

نوآوری‌های اصلی این پژوهش را می‌توان در موارد زیر خلاصه نمود:

- استفاده از یک مجموعه داده یک ساله با گام زمانی یک دقیقه از کوره گندله‌سازی صنعتی به‌عنوان داده واقعی (۵۲۵،۵۴۰ نمونه) برای آموزش و اعتبارسنجی مدل‌ها.
- مقایسه نظام‌مند سه روش انتخاب ویژگی تک‌متغیره انتخاب-اف، اطلاعات متقابل و همبستگی پیرسون در ترکیب با سه الگوریتم رگرسیون (MLR, KNN, RFR) بر روی ۶۷ خروجی عملیاتی.

^۳ k-Nearest Neighbor Regression

^۴ Multiple Input Multiple Output

^۱ Multiple Linear Regression

^۲ Random Forest Regression

- جنگل تصادفی با قابلیت بالای مدل‌سازی روابط غیرخطی و تعیین اهمیت ویژگی‌ها [۲۵].
 - رگرسیون k -همسایه نزدیک‌تر به عنوان روشی ساده ولی کارآمد در پیش‌بینی [۲۶].
 - شبکه‌های عصبی مصنوعی برای یادگیری الگوهای پیچیده و غیرخطی داده‌ها [۲۷، ۲۸].
- این ترکیب از روش‌ها امکان تحلیل داده‌های چندمتغیره صنعتی را از جنبه‌های مختلف فراهم ساخته و موجب افزایش دقت و اطمینان نتایج می‌شود.

۳- روش پیشنهادی

با توجه به ماهیت پیچیده، غیرخطی و چندمتغیره‌ی فرآیند گندله‌سازی، استفاده از مدل‌های تحلیلی مبتنی بر قوانین فیزیکی به‌سختی امکان‌پذیر است. از این‌رو، در این پژوهش یک رویکرد داده‌محور نظارت‌شده مبتنی بر الگوریتم‌های یادگیری ماشین و یادگیری عمیق به‌منظور مدل‌سازی رفتار فرآیند و پیش‌بینی هم‌زمان فشار و دما در نواحی مختلف کوره پیشنهاد می‌شود. این روش به دلیل توانایی در یادگیری روابط غیرخطی و بهره‌گیری از داده‌های واقعی صنعتی برای مسائل پیچیده‌ی فرآیندی، مناسب‌ترین گزینه برای هدف این تحقیق است.

در مرحله نخست، کلان‌داده (۵۲۵،۵۴۰ نمونه) حاصل از سامانه پایش کارخانه گندله‌سازی شماره ۱ شرکت معدنی و صنعتی گل‌گهر گردآوری و پردازش شد. داده‌ها به دو گروه ورودی و خروجی تقسیم شدند، مجموعه‌ی ورودی شامل ۶۶ متغیر فرآیندی (۳۴ متغیر دمای مشعل‌های نواحی PH و F، ۲۳ متغیر درصد دمپر هواکش‌ها، ۴ متغیر خوراک ورودی، یک متغیر سرعت ماشین، سه متغیر ارتفاع بستر گندله و یک متغیر میانگین ارتفاع بستر) و مجموعه‌ی خروجی شامل ۶۷ متغیر (۳۱ فشار و ۳۶ دما در نواحی مختلف کوره) می‌باشد.

در فرآیند نگاشت پارامتری، ویژگی‌هایی مانند دمای مشعل‌ها، درصد بازبودن دمپرها، سرعت نوار نقاله، ارتفاع بستر و نرخ خوراک ورودی، به‌عنوان بردار ورودی مدل ($x = [x_1, x_2, \dots, x_n]$) در نظر گرفته شدند که n بیانگر تعداد متغیرهای ورودی است. در مقابل، بردار خروجی شامل فشار و دما در نواحی

انجام می‌شود. شکل (۱) طرح کلی کوره گندله‌سازی را نشان می‌دهد.

- خشک‌کن دمشی (UDD) و مکشی (UDD): تبخیر رطوبت گندله‌ها با جریان هوای گرم از پایین یا بالا. دمای خروجی این دو ناحیه حدود ۳۵۰ درجه سانتی‌گراد است.
 - پیش‌گرمایش (PH): افزایش تدریجی دما از ۳۵۰ درجه سانتی‌گراد تا حدود ۱۰۰۰ درجه سانتی‌گراد.
 - پخت (F و AF): رسیدن دما به حدود ۱۳۰۰ درجه سانتی‌گراد برای ایجاد استحکام نهایی. در AF با حفظ گرما و بدون مشعل، دما ثابت می‌ماند.
 - خنک‌کن (C1 و C2): کاهش دما تا حدود ۲۰۰ درجه سانتی‌گراد همراه با بازیافت حرارت و بازگرداندن هوای گرم به نواحی پخت و خشک‌کن، با هدف بهینه‌سازی مصرف انرژی.
- در انتهای این فرآیند، گندله‌هایی با ویژگی‌های مکانیکی و شیمیایی مطلوب برای مراحل بعدی تولید فولاد حاصل می‌شوند. شکل (۲) نمایی از کارخانه گندله‌سازی شماره ۱ شرکت معدنی و صنعتی گل‌گهر را نشان می‌دهد که داده‌های مورد استفاده در این پژوهش از کوره گندله‌سازی همین واحد استخراج شده‌اند.

۲-۲ الگوریتم‌های یادگیری ماشین

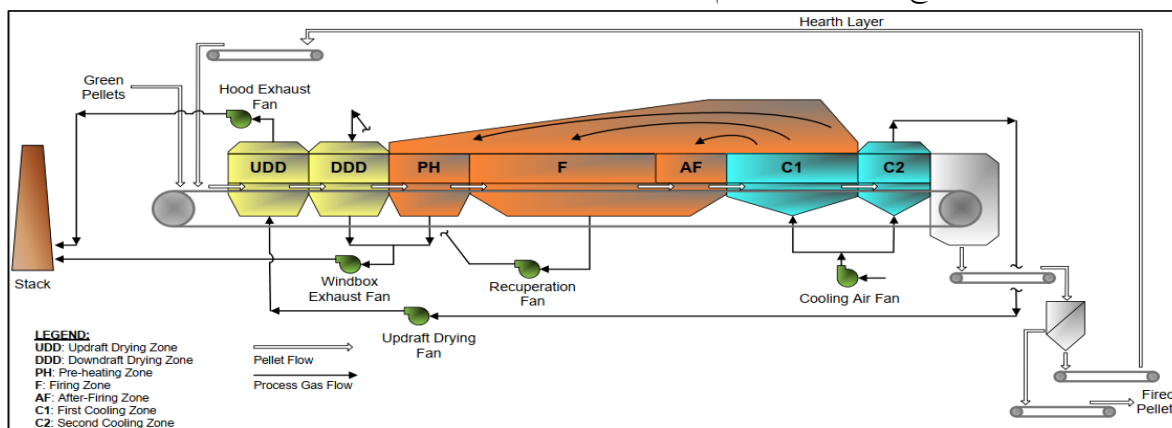
در این بخش، به معرفی مفاهیم نظری الگوریتم‌های مورد استفاده برای انتخاب ویژگی و مدل‌سازی داده‌های فرآیند صنعتی در این پژوهش پرداخته می‌شود. این مدل‌ها شامل روش‌های کلاسیک یادگیری ماشین و شبکه‌های عصبی مصنوعی هستند که هر یک دارای کاربردها و مزایای خاص خود در تحلیل داده‌های چندمتغیره صنعتی می‌باشند.

در این پژوهش برای انتخاب ویژگی از سه روش فیلتر آماری و برای مدل‌سازی داده‌های فرآیند صنعتی، از روش‌های مختلف یادگیری ماشین استفاده شده است. در مرحله انتخاب ویژگی، شاخص‌هایی مانند امتیاز-اف [۲۲]، اطلاعات متقابل [۲۳] و ضریب همبستگی پیرسون به‌کار رفتند تا متغیرهای مؤثر شناسایی شوند. برای مدل‌سازی، چند رویکرد مورد استفاده قرار گرفت:

- رگرسیون خطی چندگانه به عنوان روشی ساده برای تحلیل روابط خطی [۲۴].

روش پیشنهادی شامل دو مسیر مکمل است: مسیر (الف) برای انتخاب ویژگی‌های مؤثر با استفاده از سه الگوریتم فیلترینگ آماری، و مسیر (ب) مدل‌سازی رابطه‌ی بین ورودی‌ها و خروجی‌ها با بهره‌گیری از رویکردهای یادگیری ماشین و یادگیری عمیق. این ساختار به‌گونه‌ای طراحی شده که علاوه بر افزایش دقت پیش‌بینی، تفسیرپذیری و کارایی محاسباتی مدل نیز حفظ شود.

مختلف کوره $(Y = [y_1, y_2, \dots, y_m])$ است که m نشان‌دهنده‌ی تعداد متغیرهای خروجی است. هر ردیف از داده‌ها معرف یک نمونه از اندازه‌گیری‌های فرآیند در زمان معین بوده و در مجموع شامل N نمونه آموزشی است. بدین ترتیب، مدل‌های یادگیری ماشین و شبکه‌ی عصبی پرسپترون چندلایه به منظور یادگیری نگاشت غیرخطی بین ورودی‌ها و خروجی‌ها $(Y = f(X))$ آموزش داده شدند. این چارچوب، امکان مدل‌سازی هم‌زمان اثرات متقابل متغیرهای فرآیندی و پیش‌بینی پاسخ‌های چندگانه را فراهم می‌کند.



شکل (۱): نمای کلی کوره گندله‌سازی [۱]



شکل (۲): کارخانه گندله‌سازی شماره ۱ شرکت معدنی و صنعتی گل‌گهر.

فشار نواحی UDD و DDD که به ترتیب با نمادهای UDD_P و DDD_P نشان داده می‌شوند، نشان می‌دهد. بر اساس داده‌های این جدول، برای متغیر خروجی UDD_P، مشاهده می‌شود که از میان ۱۰ ویژگی برتر که توسط هر الگوریتم انتخاب شده‌اند، فقط یک ویژگی در هر سه فهرست مشترک است.

۳-۱ انتخاب ویژگی فیلتری

همان‌طور که گفته شد، در مسیر الف، انتخاب ویژگی با روش‌های فیلتری معرفی شده در بخش ۲-۲ انجام می‌شود. به ازای هر متغیر خروجی، ده ویژگی برتر هر معیار برای هر متغیر خروجی انتخاب شدند و جدول ۱ نتایج این انتخاب را برای متغیرهای خروجی

آن‌ها است. به‌ویژه، معیار اطلاعات متقابل متغیرهایی را آشکار ساخته که علی‌رغم نداشتن همبستگی خطی قوی، در رفتار کلی فرآیند اثرگذار هستند. استفاده‌ی هم‌زمان از این سه روش، تصویری جامع‌تر از ساختار داده‌ها فراهم کرده و مبنایی مناسب برای مدل‌سازی دقیق‌تر فراهم ساخته است.

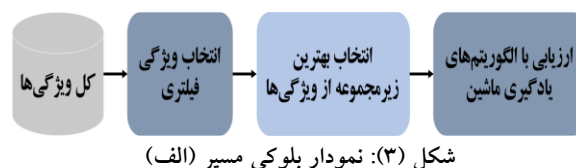
یعنی دو الگوریتم دیگر روی انتخاب‌های متفاوتی تمرکز کرده‌اند و تنها یک متغیر به‌طور هم‌زمان توسط هر سه روش مهم تشخیص داده شده است. اما در مورد متغیر خروجی DDD_P ، سه الگوریتم انتخاب ویژگی، ۸ ویژگی مشترک و تنها ۲ ویژگی متفاوت را انتخاب کرده‌اند. همان‌گونه که مشاهده می‌شود، تفاوت میان ویژگی‌های منتخب در سه معیار نشان‌دهنده‌ی دیدگاه‌های مکمل

جدول (۱): ۱۰ ویژگی منتخب توسط هر یک از سه روش انتخاب ویژگی برای متغیرهای خروجی DDD_P و UDD_P

ویژگی	DDD_P			UDD_P		
	همبستگی پرسون	اطلاعات متقابل	انتخاب-اف	همبستگی پرسون	اطلاعات متقابل	انتخاب-اف
۱	DDDtoHEX_D	DDDtoHEX_D	DDDtoHEX_D	DDDtoHEX_D	DDDtoHEX_D	DDDtoHEX_D
۲	F36_D	F22_D	F36_D	UDDtoHEX_D	Oct_D	UDDtoHEX_D
۳	F37_D	F37_D	F37_D	B14_T	F22_D	B3_T
۴	UDDtoHEX_D	F36_D	UDDtoHEX_D	B12_T	F37_D	F22_D
۵	GPTolM_FR	F42_D	GPTolM_FR	B27_T	F36_D	F12_D
۶	GPTolRS_FR	IM_S	GPTolRS_FR	Oct27B_D	F42_D	B6_T
۷	Num_Disks	F002_D	Num_Disks	Oct27A_D	F22toF12_D	B5_T
۸	F42_D	B33_T	F42_D	F22toF12_D	B3_T	B14_T
۹	IM_S	Oct_D	IM_S	GPTolM_FR	B6_T	B4_T
۱۰	B27_T	B31_T	B27_T	B11_T	B19_T	B2_T

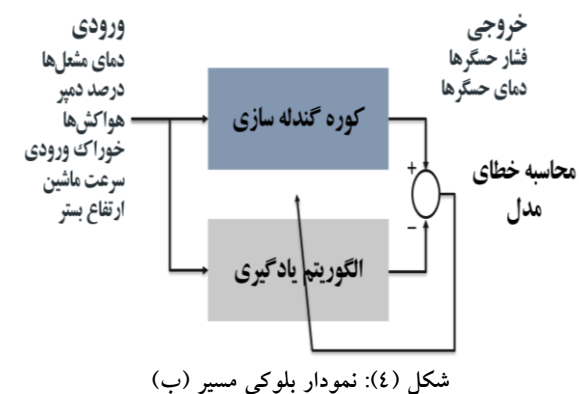
مدل‌سازی مجدداً انجام می‌شود. اما این بار همه متغیرهای ورودی بدون اعمال فیلتر انتخاب ویژگی در نظر گرفته شده‌اند. این کار با مدل‌های تک خروجی MLR ، KNR و RFR و نیز یک مدل مبتنی بر MLP به نام $MIMO$ انجام می‌شود. نمودار بلوکی مسیر (ب)، یعنی فرآیند مدل‌سازی داده‌محور در شکل ۴ نمایش داده شده است. در این چارچوب، متغیرهای ورودی وارد مدل یادگیری می‌شوند تا رابطه‌ی غیرخطی میان ورودی‌ها و خروجی‌های فرآیند یاد گرفته شود.

این مرحله برای ۶۷ متغیر خروجی موجود در مجموعه داده انجام شده و ۱۰ ویژگی برتر برای هر یک از آن‌ها شناسایی شده‌اند. از آنجا که برای ارزیابی مستقیم کیفیت الگوریتم‌های انتخاب ویژگی فیلتری معیار مشخص و مستقلی وجود ندارد، عملکرد آن‌ها به‌طور غیرمستقیم و از طریق تأثیرشان بر نتایج مدل‌های رگرسیونی سنجیده می‌شود. بدین منظور، از سه مدل MLR ، KNR و RFR استفاده شده است. نمودار بلوکی این مرحله در شکل (۳) نشان داده شده است.



۲-۳ مدل‌سازی

در بخش ۳-۱، مدل‌سازی با روش‌های سنتی یادگیری ماشین بر مبنای ویژگی‌های منتخب توسط سه الگوریتم انتخاب ویژگی انجام شد تا اثر نوع روش انتخاب ویژگی بر عملکرد مدل‌های رگرسیون بررسی شود. در این بخش، با هدف مقایسه و ارزیابی جامع‌تر،



MIMO_All انجام می‌شود. در ادامه، نتایج بدست آمده تحلیل و مقایسه می‌گردند تا نقاط قوت و ضعف هر رویکرد مشخص شود.

۴-۱ معیارهای ارزیابی عملکرد

برای سنجش دقت مدل‌های رگرسیونی، از سه معیار متداول استفاده شده است:

- میانگین مربعات خطا (MSE): نشان‌دهنده میانگین توان دوم اختلاف بین مقادیر واقعی و پیش‌بینی شده بوده و مقادیر کمتر بیانگر عملکرد بهتر مدل است [۲۹].
- میانگین قدر مطلق خطا (MAE): مشابه MSE، اما بر پایه قدر مطلق خطا محاسبه می‌شود و معیار ساده‌تری برای ارزیابی دقت پیش‌بینی است [۳۰].
- ضریب تعیین: که با نماد R^2 -Score نشان داده می‌شود، میزان تغییرات توضیح داده شده توسط مدل را نشان می‌دهد و بین صفر تا یک متغیر است. مقدار بالاتر نشان‌دهنده دقت و توانایی بالاتر مدل در تبیین داده‌هاست [۱۰].

در مسأله حاضر، از دو معیار MSE و MAE برای ارزیابی عملکرد مدل‌ها استفاده شده است. زیرا هر یک جنبه متفاوتی از خطا را نشان می‌دهند. معیار MSE به دلیل توان دوم کردن خطاها، نسبت به خطاهای بزرگ حساس‌تر است و برای تحلیل مدل‌هایی که هدف آن‌ها کاهش خطاهای شدید و پایدارسازی پیش‌بینی‌هاست (مانند مدل‌های مبتنی بر شبکه عصبی و جنگل تصادفی)، شاخص مناسب‌تری محسوب می‌شود. در مقابل، معیار MAE میانگین قدر مطلق خطاها را بدون تأکید بر مقدار آن‌ها ارائه می‌دهد و در سنجش دقت کلی پیش‌بینی‌ها (به‌ویژه برای متغیرهایی با مقیاس‌های مختلف) قابل تفسیرتر است.

۴-۲ پیش‌پردازش داده‌ها

داده‌های موجود، شامل داده‌های یک ساله کوره گندله‌سازی هستند که از تاریخ اول ژانویه ۲۰۲۱ الی ۳۱ دسامبر ۲۰۲۱ با گام زمانی یک دقیقه جمع‌آوری شده‌اند. در این پژوهش، ۶۶ متغیر به عنوان ویژگی‌های ورودی در نظر گرفته شده‌اند که شامل دمای مشعل‌ها، درصد دمپر هواکش‌ها، سرعت ماشین، ارتفاع بستر گندله و میزان خوراک ورودی هستند. در مقابل، ۶۷ متغیر خروجی شامل مقادیر اندازه‌گیری‌شده‌ی دما و فشار در نقاط مختلف کوره می‌باشند که

همان‌طور که در شکل مشاهده می‌شود، مدل یادگیری به‌صورت یک حلقه بازخوردی عمل می‌کند. به‌گونه‌ای که خروجی پیش‌بینی‌شده‌ی مدل با مقادیر واقعی اندازه‌گیری‌شده مقایسه و خطای مدل محاسبه می‌شود. سپس الگوریتم یادگیری با استفاده از این خطا، پارامترهای مدل را به‌روزرسانی می‌کند تا همگرایی به نداشت بهینه $Y = f(X)$ حاصل شود. این ساختار امکان بهبود تدریجی دقت پیش‌بینی و مدل‌سازی هم‌زمان اثرات متقابل میان متغیرهای فرآیندی را فراهم می‌کند.

در مدل‌های سنتی مانند MLR، KNR و RFR، به دلیل ساختار تک‌خروجی آن‌ها، لازم است برای پیش‌بینی هر یک از متغیرهای خروجی، یک مدل جداگانه آموزش داده شود. این امر موجب افزایش تعداد مدل‌های مورد نیاز، زمان آموزش و پیچیدگی محاسباتی می‌گردد. در مقابل، مدل MIMO قادر است به‌طور هم‌زمان تمام خروجی‌های فرآیند را پیش‌بینی کند. این ویژگی باعث کاهش تعداد مدل‌ها، بهره‌برداری بهتر از روابط متقابل بین خروجی‌ها و در بسیاری از موارد بهبود دقت پیش‌بینی می‌شود. مدل MIMO به صورت یک شبکه عصبی پرسپترون چندلایه طراحی شده است. این مدل‌سازی به سه صورت انجام شده است: MIMO-Temperature: آموزش یک مدل بر اساس تمام متغیرهای ورودی و پیش‌بینی هم‌زمان همه خروجی‌های دما، MIMO-Pressure: آموزش یک مدل بر اساس تمام متغیرهای ورودی و پیش‌بینی هم‌زمان همه خروجی‌های فشار، MIMO-All: آموزش یک مدل بر اساس تمام متغیرهای ورودی و پیش‌بینی هم‌زمان همه خروجی‌های فشار و دما.

۴-۳ آزمایش‌ها و نتایج

در این بخش، فرآیند ارزیابی مدل‌های پیشنهادی و تحلیل نتایج ارائه می‌شود. ابتدا تنظیمات آزمایش‌ها، شامل پیکربندی و پیش‌پردازش داده‌ها، جداسازی داده‌های آموزش و آزمون، معیارهای ارزیابی عملکرد و مقادیر پارامترهای هر الگوریتم برای هر مدل توضیح داده می‌شود. سپس عملکرد مدل‌ها مورد بررسی قرار می‌گیرد.

مقایسه بین مدل‌های سنتی (MLR، KNR، RFR) و مدل MIMO در ساختارهای MIMO-Temperature، MIMO-Pressure و

هدف از این مرحله، همسان‌سازی دامنه مقادیر متغیرهای مختلف است تا الگوریتم بتواند آن‌ها را به شکلی متعادل پردازش کند. از این رو، پیش از اجرای الگوریتم‌های انتخاب ویژگی نیز لازم است داده‌ها مقیاس‌بندی شوند. در روش نرمال‌سازی کمینه-بیشینه، به منظور تغییر مقیاس خطی ویژگی j -ام در یک محدوده دلخواه $[a, b]$ می‌توان از رابطه (۱) استفاده کرد [۳۱]:

$$\hat{x}_j^i = a + \frac{(x_j^i - x_j^{\min})(b - a)}{x_j^{\max} - x_j^{\min}} \quad (1)$$

در رابطه بالا، $x_j^{\min}$ و $x_j^{\max}$ به ترتیب برابر با کمینه و بیشینه مقادیر ویژگی j -ام و $\hat{x}_j^i$ مقدار مقیاس شده را نشان می‌دهد. در این پژوهش از نرمال‌سازی کمینه-بیشینه استفاده شده است و کلیه مقادیر متغیرهای ورودی و خروجی به بازه $[-1, 1]$ نگاشت شده‌اند.

لازم به ذکر است که نرمال‌سازی فقط برای آموزش مدل انجام شد. برای گزارش عملکرد، مقادیر پیش‌بینی شده توسط مدل، به مقیاس اصلی بازگردانده شده و سپس معیارهای ارزیابی عملکرد مانند MAE و MSE محاسبه شده‌اند.

۴-۲-۳ جداسازی داده‌های آموزش و آزمون

در این پژوهش، هدف مدل‌سازی، پیش‌بینی هم‌زمان خروجی‌های فرآیند بر اساس متغیرهای ورودی در همان لحظه است، نه پیش‌بینی روند زمانی آینده. در چنین مسائلی، مدل به دنبال کشف رابطه میان مقادیر ورودی و خروجی در هر نمونه از داده است، بدون آن‌که به توالی زمانی داده‌ها وابسته باشد. از آنجا که ورودی‌ها و خروجی‌ها مربوط به اندازه‌گیری‌های لحظه‌ای از وضعیت فرآیند هستند و در بازه‌های زمانی نسبتاً کوتاه ثبت شده‌اند، ترتیب وقوع داده‌ها در آموزش مدل اهمیتی ندارد و بر روابط فیزیکی میان متغیرها تأثیر نمی‌گذارد.

به همین دلیل، تقسیم داده‌ها به صورت تصادفی میان مجموعه‌های آموزش (۷۰ درصد) و آزمون (۳۰ درصد) انجام شد تا هر دو مجموعه شامل گستره کامل شرایط عملیاتی باشند. الگوریتم‌های مورد مقایسه با استفاده از داده‌های آموزشی آموزش دیده و سپس عملکرد آن‌ها بر روی داده‌های آزمون ارزیابی می‌شود.

مدل باید آن‌ها را بر اساس ورودی‌ها پیش‌بینی کند.

در مجموع، مجموعه داده شامل ۱۳۳ ویژگی (۶۶ ورودی و ۶۷ خروجی) و ۵۲۵۰۵۴۰ نمونه است. این داده‌ها پس از پاک‌سازی و پیش‌پردازش برای آموزش مدل‌های یادگیری ماشین استفاده شده‌اند. مراحل پیش‌پردازش مجموعه داده شامل مواجهه با مقادیر ناموجود، مقیاس‌بندی ویژگی با روش نرمال‌سازی کمینه-بیشینه در بازه $[-1, 1]$ و تقسیم به داده‌های آموزش و آزمون است.

۴-۲-۱ مواجهه با مقادیر ناموجود

برخی از داده‌ها دارای مقادیر ناموجود (NaN) هستند که لازم است پیش از تحلیل تکمیل شوند. این مقادیر عمدتاً در شرایطی مشاهده شده‌اند که تعداد دیسک‌های فعال کمتر از سه بوده است. جدول (۲) میانگین تعداد مقادیر ناموجود را برای داده‌های ورودی و خروجی ارائه می‌کند.

جدول (۲): میانگین تعداد مقادیر ناموجود به ازای هر متغیر مجموعه داده

نام مجموعه داده	میانگین تعداد مقادیر ناموجود در هر متغیر
داده‌های دمای مشعل	۵۸۳۳۶
داده‌های درصد دمپر هواکش‌ها	۵۸۳۲۸
داده‌های میزان خوراک ورودی	۵۸۹۹۹
داده‌های سرعت ماشین و ارتفاع بستر	۵۹۰۹۳
داده‌های فشار حسگرها	۵۸۹۸۱
داده‌های دمای حسگرها	۵۸۴۲۴

به عنوان نمونه، سطر اول این جدول به داده‌های دمای مشعل‌ها مربوط است. این داده‌ها شامل ۳۴ متغیر دماست که به طور میانگین، هر متغیر دارای ۵۸۳۳۶ مقدار ناموجود است. پیش از انجام تحلیل‌های آماری، این مقادیر باید تکمیل گردند. با توجه به اینکه تغییرات این متغیرها در طول زمان چندان زیاد نیست، در این پژوهش از روش جایگزینی مقدار میانگین هر متغیر برای پر کردن مقادیر ناموجود آن استفاده شده است.

۴-۲-۲ مقیاس‌بندی ویژگی

پیش از به‌کارگیری هر الگوریتم یادگیری ماشین، معمولاً مرحله‌ای تحت عنوان مقیاس‌بندی ویژگی‌ها بر روی داده‌ها انجام می‌شود.

۳-۴ تنظیم پارامترها

برای مرحله انتخاب ویژگی از ماژول‌های موجود در کتابخانه Scikit-learn پایتون استفاده شده است. به کمک ماژول SelectKBest این کتابخانه و انتخاب تابع رتبه‌دهی با امتیاز-اف، اطلاعات متقابل و همبستگی پیرسون، رتبه-های مربوط به هر ویژگی محاسبه شده و بر این اساس، به ازای هر متغیر خروجی کوره گندله‌سازی، ۱۰ ویژگی ورودی با بالاترین رتبه-ها انتخاب می‌شوند.

در ادامه لازم است مدل‌های رگرسیون MLR، KNR و RFR روی ویژگی‌های انتخاب شده برازش داده شده و مقادیر متغیرهای خروجی را پیش‌بینی کنند. این مدل‌ها با استفاده از کتابخانه cuML (بخشی از چارچوب RAPIDS) پیاده‌سازی شده‌اند که امکان اجرای الگوریتم‌ها بر بستر GPU را فراهم می‌کند. استفاده از این کتابخانه با هدف افزایش سرعت پردازش و مدیریت کارآمد حجم بالای داده‌ها انتخاب شده است. تنظیمات پارامترها برای الگوریتم‌های MLR، KNR و RFR در جدول (۳) ارائه شده است.

جدول (۳): تنظیمات پارامترهای مدل‌های MLR، KNR و RFR

الگوریتم	پارامترها
MLR	<code>algorithm = 'eig', fit_intercept = True, normalize = False</code>
KNR	<code>n_neighbors = 5, metric = 'euclidean', weights = 'uniform',</code>
RFR	<code>n_estimator = 100, split_criterion = 'mse', max_depth = 16</code>

مدل MIMO، با استفاده از لایه‌های پیش‌خور طراحی شده است. از آنجا که داده‌های ورودی شامل متغیرهای عددی مستقل و بدون ساختار مکانی هستند، استفاده از مدل‌های کانولوشنی در این مسئله توجیه فیزیکی و محاسباتی ندارد زیرا بین ویژگی‌ها رابطه مکانی یا همسایگی تعریف‌پذیر وجود ندارد. بنابراین از ساختار تماماً متصل برای همه لایه‌ها استفاده شده است. این شبکه دارای یک لایه ورودی با ۶۶ گره برای دریافت متغیرهای ورودی است. پس از لایه ورودی، ما از پنج لایه پنهان که به ترتیب شامل ۱۳۲، ۲۶۴، ۲۶۴، ۲۶۴، ۱۳۲ گره هستند استفاده کرده‌ایم. لایه خروجی در مدل‌های MIMO_All، MIMO_Pressure و

MIMO_Temperature به ترتیب شامل ۶۷ (تعداد کل متغیرهای خروجی)، ۳۱ (تعداد متغیرهای فشار) و ۳۶ (تعداد متغیرهای دما) گره است.

انتخاب ساختار پنج‌لایه با هدف افزایش توان مدل در یادگیری روابط غیرخطی میان ورودی‌ها و خروجی‌ها صورت گرفت. پیش از انتخاب این معماری، نسخه‌های ساده‌تر شبکه با دو و سه لایه پنهان نیز آزمایش شدند. اما با وجود زمان آموزش کوتاه‌تر، این مدل‌ها دقت پایین‌تر و نوسان بیشتری در خطای اعتبارسنجی داشتند. بنابراین ساختار پنج‌لایه به عنوان پیکربندی نهایی برگزیده شد.

تابع فعال‌ساز در لایه‌های پنهان ReLU و در لایه خروجی، خطی است. تابع زیان مورد استفاده، MSE بوده و بهینه‌ساز ADAM با نرخ یادگیری ۰/۰۰۰۱ برای بروزرسانی وزن‌های شبکه بکار گرفته شده است. در مرحله آموزش مدل پیشنهادی، اندازه دسته برابر با ۳۲ و اندازه مجموعه اعتبارسنجی برابر با بیست درصد داده‌های آموزشی انتخاب شده است. تعداد مراحل آموزش مدل نیز برابر با ۲۰۰ در نظر گرفته شده است. به منظور جلوگیری از بیش‌برازش، از سازوکار توقف زودهنگام استفاده گردید. بر این اساس، در صورتی که در طول ۱۰ درصد پایانی دوره‌های آموزشی، بهبودی معناداری در مقدار تابع زیان روی داده‌های اعتبارسنجی مشاهده نشود، فرآیند آموزش به‌طور خودکار متوقف شده و وزن‌های مدل متناظر با کمترین خطای اعتبارسنجی ذخیره می‌گردد.

۴-۴ نتایج

پس از آماده‌سازی داده‌ها و تنظیم پارامترهای مدل‌ها بر اساس توضیحات ارائه‌شده در بخش قبل، مدل‌های رگرسیون سنتی (MLR، KNR و RFR) و مدل‌های MIMO بر روی داده‌های آموزشی برازش داده شدند. سپس، عملکرد هر مدل بر روی داده‌های آزمون ارزیابی گردید. در این بخش، نتایج حاصل از اجرای مدل‌ها ارائه و مقایسه می‌شوند. تحلیل‌ها شامل بررسی دقت پیش‌بینی، ارزیابی شاخص‌های آماری عملکرد و مقایسه کیفی بین مدل‌های مختلف است تا بتوان کارآمدترین روش در پیش‌بینی متغیرهای خروجی کوره گندله‌سازی را شناسایی کرد.

۴-۱-۴ نتایج انتخاب ویژگی

در این بخش عملکرد سه الگوریتم انتخاب ویژگی تک متغیره امتیاز اف، اطلاعات متقابل و همبستگی پیرسون روی مجموعه داده کوره گندله‌سازی بررسی خواهد شد. به این منظور، از سه مدل MLR، KNR و RFR استفاده شده است. مدل‌های رگرسیون به ازای هر متغیر خروجی، با ویژگی‌های انتخاب شده توسط هر یک از سه الگوریتم انتخاب ویژگی فیلتری روی نمونه‌های مجموعه داده آموزش، اعمال شده و سپس عملکرد آن‌ها روی مجموعه داده آزمون بررسی خواهد شد. جدول (۴)، میانگین نتایج بدست آمده را برای معیارهای R^2 -Score، MAE و MSE را روی ۶۷ متغیر خروجی نشان می‌دهد.

لازم به ذکر است که در این پژوهش، معیارهای ارزیابی عملکرد بر مبنای مقادیر واقعی خروجی‌ها محاسبه شده‌اند. اگرچه دامنه مقادیر متغیرهای دما و فشار متفاوت است، اما هدف از میانگین‌گیری معیارها در اینجا مقایسه نسبی عملکرد مدل‌ها در پیش‌بینی کلیه خروجی‌ها بوده است، نه مقایسه دقیق بین متغیرهای هم‌مقیاس. در واقع، میانگین معیارها تنها برای ایجاد شاخصی تجمیعی جهت مقایسه کیفی بین الگوریتم‌ها به کار رفته و تمام مدل‌ها تحت شرایط یکسان آموزش و ارزیابی شده‌اند. بنابراین اختلاف دامنه‌ها بر مقایسه نهایی تأثیرگذار نیست.

بر اساس نتایج گزارش شده در جدول (۴)، از نظر میانگین R^2 -Score روی همه متغیرهای خروجی، بهترین دقت برای KNR با الگوریتم انتخاب ویژگی اطلاعات متقابل بدست آمده که برابر با ۸۸/۵۹ درصد است. رتبه‌های بعدی که مقادیر ۸۶/۵۲ درصد و ۸۲/۷۷ درصد هستند، به ترتیب توسط RFR با الگوریتم انتخاب ویژگی اطلاعات متقابل و KNR با امتیاز-اف کسب شده‌اند. همچنین الگوریتم RFR با امتیاز-اف رتبه چهارم R^2 -Score را بدست آورده که برابر با ۸۱/۹۶ درصد است.

این نتایج نشان می‌دهند که استفاده از الگوریتم انتخاب ویژگی اطلاعات متقابل در مجموع عملکرد بهتری نسبت به دو روش دیگر دارد، به‌ویژه در مدل‌های غیرخطی مانند KNR و RFR. در مقابل، الگوریتم همبستگی پیرسون در اکثر موارد کمترین مقدار R^2 -

Score را به دست آورده است که می‌تواند به محدودیت این روش در شناسایی روابط غیرخطی بین متغیرها نسبت داده شود.

جدول (۴): نتایج میانگین معیارهای ارزیابی مدل‌های رگرسیون روی ۶۷ خروجی در سه روش انتخاب ویژگی

		امتیاز-اف	اطلاعات متقابل	همبستگی پیرسون
MLR	R^2 -Score	۵۱/۲۵	۴۹/۹	۳۶/۸
	MAE	۸/۷۵	۸/۵۶	۹/۴۵
	MSE	۳۵۵/۸۴	۳۷۵/۵۳	۴۱۱/۳۴
KNR	R^2 -Score	۸۲/۷۷	۸۸/۵۹	۷۸/۳۴
	MAE	۳/۹۵	۳/۰۹	۴/۱۸
	MSE	۱۱۳/۵۸	۸۳/۵۵	۱۳۱/۵۸
RFR	R^2 -Score	۸۱/۹۶	۸۶/۵۲	۷۷/۰۲
	MAE	۴/۸۱	۴/۱	۵/۲
	MSE	۱۳۴/۹۳	۹۸/۹۷	۱۵۶/۴۴

از نظر معیار MAE، کمترین خطای میانگین مطلق (۳/۰۹) مربوط به مدل KNR با الگوریتم اطلاعات متقابل است، که بیانگر دقت بالاتر پیش‌بینی در این ترکیب است. همچنین در معیار MSE، همین ترکیب با مقدار ۸۳/۵۵ کمترین میانگین خطای مربعات را داشته و در نتیجه بهترین عملکرد کلی را از نظر کاهش خطای پیش‌بینی ارائه داده است. نمودارهای شکل (۵)، مقایسه عملکرد مدل‌های MLR، KNR و RFR را در سه روش انتخاب ویژگی (امتیاز-اف، اطلاعات متقابل، همبستگی پیرسون) از نظر معیارهای R^2 -Score، MAE و MSE به صورت بصری نشان می‌دهند.

به طور کلی، می‌توان نتیجه گرفت که در مسئله حاضر، ترکیب انتخاب ویژگی اطلاعات متقابل و مدل KNR بهترین تعادل را بین دقت پیش‌بینی و کاهش خطا ایجاد کرده است. این امر بیانگر آن است که بهره‌گیری از روش‌های انتخاب ویژگی که قادر به شناسایی روابط غیرخطی میان متغیرهای ورودی و خروجی هستند، می‌تواند نقش مؤثری در بهبود عملکرد مدل‌های یادگیری ماشین برای فرآیندهای صنعتی پیچیده ایفا کند.

۴-۴-۲ نتایج مدل‌سازی

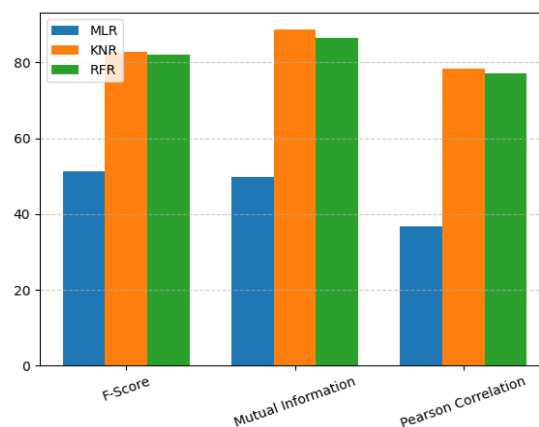
در این بخش، نتایج حاصل از مدل‌سازی با الگوریتم‌های MLR، KNR و RFR بدون مرحله انتخاب ویژگی و با در نظر گرفتن همه متغیرهای ورودی و همچنین مدل‌های MIMO ارائه خواهد شد. جدول (۵) نتایج حاصل از مدل‌سازی را نشان می‌دهد. لازم به ذکر است که در این جدول میانگین معیارهای ارزیابی روی همه خروجی‌ها نشان داده شده است.

جدول (۵): نتایج میانگین معیارهای ارزیابی عملکرد در مدل‌سازی

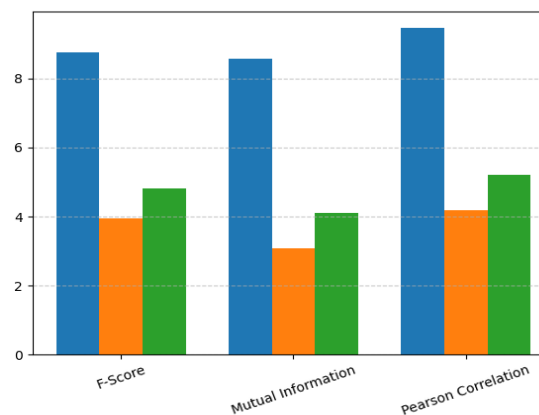
	R ² -Score	MAE	MSE
MLR	۶۴/۲۲	۶/۸۱	۲۴۳/۸۴
RFR	۸۹/۴۱	۳/۵۲	۵۲/۱۲
KNR	۹۱/۶۷	۲/۵	۷۲/۷۸
MIMO-Temperature	۹۱/۵۰	۵/۹۴	۱۱۴/۹۶
MIMO-Pressure	۸۹/۶۲	۰/۵۴	۰/۸۳
MIMO-All	۸۹/۳۴	۳/۷۶	۷۱/۵۶

بر اساس نتایج ارائه‌شده در جدول (۵)، مشاهده می‌شود که عملکرد مدل‌های غیرخطی نسبت به مدل خطی چندمتغیره (MLR) به‌طور محسوسی بهتر است. مدل MLR با مقدار R²-Score برابر با ۶۴/۲۲ و مقادیر بالای MAE و MSE نتوانسته است ساختارهای پیچیده موجود در داده‌ها را به‌خوبی مدل‌سازی کند و در نتیجه دقت پایینی دارد. در مقابل، دو مدل غیرخطی RFR و KNR توانسته‌اند با افزایش چشم‌گیر ضریب تعیین (به ترتیب ۸۹/۴۱ و ۹۱/۶۷) و کاهش خطاهای مطلق و مربعی، نتایج بسیار بهتری ارائه دهند. در میان این دو مدل، KNR از نظر خطای مطلق عملکرد بهتری دارد، در حالی که RFR در کاهش خطاهای بزرگ و پایدارسازی مدل از طریق MSE، برتری نسبی نشان می‌دهد.

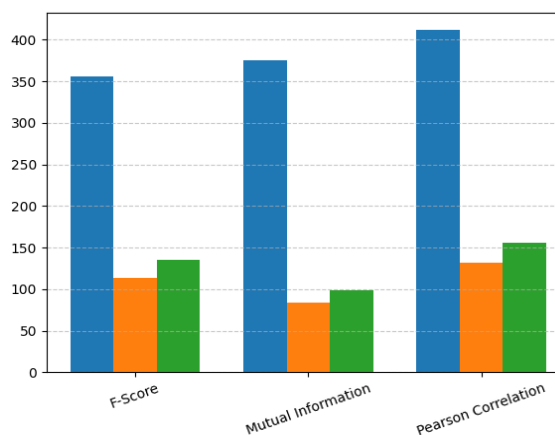
در مقایسه با این مدل‌ها، مدل‌های MIMO نتایج متفاوتی ارائه داده‌اند. مدل MIMO-Temperature از نظر R²-Score عملکردی مشابه KNR دارد، اما خطاهای آن (MAE و MSE) بالاتر است که می‌تواند ناشی از پیچیدگی پیش‌بینی همزمان چند خروجی باشد. در مقابل، مدل MIMO-Pressure با مقادیر بسیار پایین MAE و MSE (به ترتیب ۰/۵۴ و ۰/۸۳) و R²-Score بالا، به‌ظاهر بهترین عملکرد را نشان می‌دهد. با این حال، باید توجه داشت که دامنه تغییرات متغیرهای فشار به‌طور ذاتی بسیار



(الف)

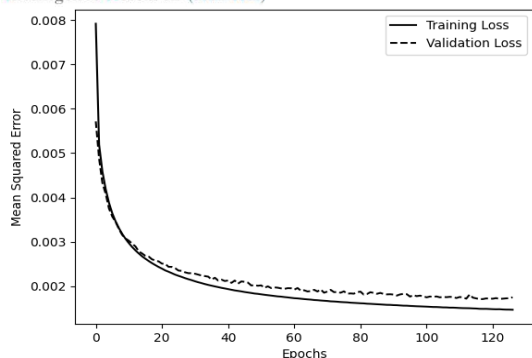


(ب)

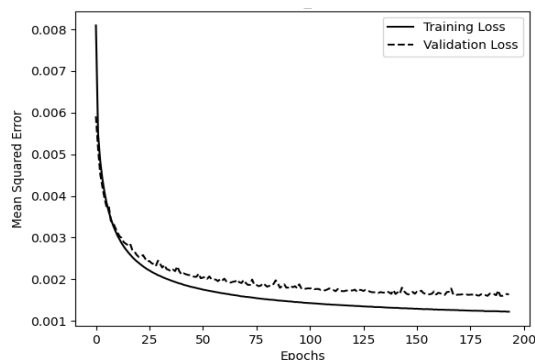


(ج)

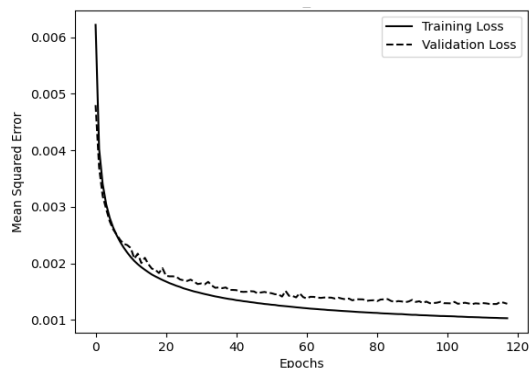
شکل (۵): نمودارهای مقایسه عملکرد مدل‌های MLR، KNR و RFR در سه روش انتخاب ویژگی با معیارهای الف) R²-Score، ب) MAE و ج) MSE



(الف)



(ب)



(ج)

شکل (۶): نمودارهای تابع زیان در سه مدل (الف)

MIMO-Pressure (ب)، MIMO-Temperature (ج) و MIMO-All

Pressure

در مقابل، RFR با پیچیدگی تقریبی $O(T \times m_{try} \times N \log N)$ است که در آن T تعداد درخت‌ها و m_{try} تعداد ویژگی‌هایی است که در هر گره به صورت تصادفی انتخاب می‌شوند [۳۴]. RFR به‌خاطر قابلیت موازی‌سازی ساخت درخت‌ها در کاربردهای صنعتی یک گزینه متعادل و محبوب است. در مورد مدل MIMO،

محدودتر از متغیرهای دما است و همین موضوع باعث کوچک شدن مقادیر خطای میانگین و مربعی در مدل MIMO-Pressure شده است. بنابراین، برتری ظاهری این مدل در معیارهای خطای الزاماً به معنای دقت بالاتر آن نسبت به سایر مدل‌ها نیست، بلکه بخشی از این نتایج به تفاوت مقیاس متغیرها باز می‌گردد. به طور کلی، انتخاب مدل باید متناسب با شرایط و محدودیت‌ها باشد:

■ اگر منابع محاسباتی محدود باشند و نیاز به پیاده‌سازی سریع وجود داشته باشد، مدل‌های ساده‌ای مانند MLR به دلیل سرعت بالا در آموزش و پیش‌بینی گزینه مناسبی هستند، هرچند دقت پایین‌تری دارند.

■ مدل KNR به دلیل نیاز به محاسبات فاصله با تمام داده‌های آموزشی، برای داده‌های بزرگ یا محیط‌های بلادرنگ مناسب نیست، اما در داده‌های کوچک و جایی که دقت بالاتر از سرعت اهمیت دارد می‌تواند کارآمد باشد.

■ مدل RFR تعادلی بین دقت و سرعت برقرار می‌کند و نسبت به KNR در داده‌های بزرگ مقیاس‌پذیرتر است، به همین دلیل در کاربردهای بلادرنگ برتری دارد.

■ مدل‌های MIMO مبتنی بر شبکه عصبی به منابع سخت‌افزاری قدرتمند (GPU و RAM بالا) نیاز دارند، اما وقتی چنین منابعی در دسترس باشد و تعداد خروجی‌ها زیاد باشد (مانند دما و فشار کوره)، به دلیل امکان یادگیری روابط پیچیده و پیش‌بینی همزمان چند خروجی گزینه‌ای بهینه محسوب می‌شوند. شکل (۶) نمودارهای تابع زیان مدل‌های MIMO را نشان می‌دهد.

از نظر هزینه محاسباتی، تفاوت قابل‌توجهی میان مدل‌های استفاده‌شده وجود دارد. مدل MLR از نظر پیچیدگی زمانی در حدود $O(N \times n^2 + n^3)$ برای n ویژگی و N نمونه آموزشی است [۳۲]. در عمل رگرسیون خطی معمولاً از نظر زمان اجرا و پیاده‌سازی سریع و سبک است، ولی سرعت آن نسبی است و به مقادیر N و n بستگی دارد. مدل KNR به دلیل نیاز به محاسبه فاصله بین هر نمونه آزمون و تمام داده‌های آموزشی، دارای پیچیدگی زمانی $O(n \times N)$ است [۳۳]. این الگوریتم برای استنتاج بلادرنگ با داده زیاد معمولاً مناسب نیست.

شرایط عملیاتی است. مدل‌های ساده مانند MLR برای شرایط با محدودیت منابع محاسباتی مناسب‌اند، RFR برای داده‌های بزرگ و کاربردهای بلندپایه کارایی بیشتری دارد، و مدل‌های MIMO در صورت دسترسی به منابع سخت‌افزاری قدرتمند بهترین گزینه برای پیش‌بینی همزمان چند خروجی محسوب می‌شوند.

نتایج این تحقیق می‌تواند مبنایی برای توسعه ابزارهای هوشمند کنترل فرآیند در صنعت فولاد باشد. در ادامه، توسعه روش‌های پیشرفته‌تر انتخاب ویژگی برای کاهش عدم قطعیت داده‌ها، بهره‌گیری از شبکه‌های عصبی مدرن نظیر LSTM، GRU و معماری‌های مبتنی بر توجه برای استخراج وابستگی‌های زمانی و مکانی، ترکیب مدل‌های فیزیکی و داده‌محور در قالب رویکردهای ترکیبی، اعتبارسنجی مدل‌ها در محیط واقعی و در نهایت طراحی مدل‌های هوشمندی که علاوه بر پیش‌بینی خروجی‌ها به بهینه‌سازی انرژی و کاهش آلاینده‌ها کمک کنند، از مهم‌ترین مسیرهای آینده این حوزه به شمار می‌روند.

۶- سپاس‌گزاری

این پژوهش برگرفته از پروژه‌ای پژوهشی است که با حمایت شرکت معدنی و صنعتی گل‌گهر انجام شده است. بدین‌وسیله از این شرکت به‌دلیل همکاری، ارائه داده‌های عملیاتی و پشتیبانی فنی در انجام این تحقیق صمیمانه قدردانی می‌شود.

References

- [1] F. Habashi, Handbook of extractive metallurgy. Wiley-Vch, 1997.
- [2] Saeedi and N. Sotoudeh, Iron Production. Isfahan University of Technology, Jihad Daneshgahi Press, 2006. [In Persian].
- [3] M. Alizadeh, M. Alizadeh, and H. Jeylan. "Creating optimal conditions for producing suitable green pellets from iron ore concentrates with low specific surface area using organic and mineral additives," Journal of Metallurgical Engineering, vol. 21, no. 3, pp. 216-224, 2018. [In Persian].
- [4] K. Meier, Y. Nabialahi, and A. Mansouri Aliabadi. Iron Ore Pelletizing. Shamloo Publishing, 2012. [In Persian].
- [5] B. Abazarpour, R. Hejazi, M. Saghaeian, and V. Sheikhzadeh, "Effect of iron ore particle size and shape on green pellet quality," in International Mineral Processing Congress, Moscow, 2018.
- [6] L. Lu, Iron ore: mineralogy, processing and environmental sustainability, Elsevier, 2015.
- [7] K. Meyer, Pelletizing of iron ores, Springer-Verlag, 1980.
- [8] S. Sadrezaad, A. Ferdowsi, and H. Payab, "Mathematical model for a straight grate iron ore pellet induration process of industrial scale," Computational Materials Science, vol. 44, no. 2, pp. 296-302, 2008.
- [9] M. Barati, "Dynamic simulation of pellet induration process in straight-grate system," International journal of mineral processing, vol. 89, no. 1-4, pp. 30-39, 2008.
- [10] H. Nezamabadipour, M. Rezaei, and A. Atighi. Statistical and Intelligent Modeling in Engineering and Industry. Shahid Bahonar University of Kerman Press, 2024. [In Persian].
- [11] T. Umadevi, P. P. Kumar, P. Kumar, N. F. Lobo, and M. Ranjan, "Investigation of factors affecting

بیشترین بار محاسباتی مربوط به آموزش شبکه است که در هر دور آموزشی به‌طور تقریبی $O(E \times N \times P)$ که در آن E تعداد دوره‌های آموزشی و P تعداد پارامترها هستند. در مجموع، از دیدگاه عملی، مدل‌های RFR و MIMO با توجه به تعادل بین دقت بالا و زمان اجرای قابل‌قبول، گزینه‌های مناسبی برای استفاده در محیط‌های واقعی فرآیند گندله‌سازی محسوب می‌شوند.

۵- نتیجه‌گیری

این پژوهش به مدل‌سازی داده‌محور فرآیند کوره گندله‌سازی مجتمع گل‌گهر با استفاده از الگوریتم‌های یادگیری ماشین و شبکه‌های عصبی پرداخت. روش پیشنهادی شامل سه مرحله اصلی پیش‌پردازش داده‌ها، انتخاب ویژگی و مدل‌سازی بود. نتایج انتخاب ویژگی نشان داد که معیار اطلاعات متقابل به دلیل توانایی در شناسایی روابط غیرخطی، عملکرد بهتری نسبت به امتیاز-اف و ضریب همبستگی پیرسون داشته و دقت مدل‌های غیرخطی مانند RFR و KNR را به‌طور محسوسی ارتقا داده است. در مرحله مدل‌سازی نیز مشخص شد که الگوریتم‌های غیرخطی نسبت به رگرسیون خطی چندمتغیره برتری دارند، به‌گونه‌ای که KNR کمترین خطای مطلق و RFR پایداری بیشتری در کاهش خطاهای بزرگ ارائه کرد. مدل‌های MIMO نیز قابلیت پیش‌بینی همزمان فشار و دما را فراهم کردند. به‌طور کلی، انتخاب مدل وابسته به

- pellet strength in straight grate induration machine,” *Ironmaking & steelmaking*, vol. 35, no. 5, pp. 321-326, 2008.
- [12] H. Zhou, Y. He, B. Li, D. Song, Q. Zhu, and Y. Li, “Ironmaking process under artificial intelligence technology: A review,” *Ironmaking & Steelmaking*, 2024.
- [13] A. J. Deo, S. K. Behera, and D. P. Das, “Online monitoring of iron ore pellet size distribution using lightweight convolutional neural network,” *IEEE Transactions on Automation Science and Engineering*, vol. 21, no. 2, pp. 1974-1985, 2023.
- [14] F. Honpeng and G. Runlong, “Practice of Artificial Intelligence for Iron Ore Pelletizing System,” in 2024 5th IEEE Global Conference for Advancement in Technology (GCAT), 2024, pp. 1-4.
- [15] Y.-h. Gong, C.-h. Wang, J. Li, M. N. Mahyuddin, and M. T. A. Seman, “Application of deep learning in iron ore sintering process: a review,” *Journal of Iron and Steel Research International*, vol. 31, no. 5, pp. 1033-1049, 2024.
- [16] Q. Cao, C. Chi, and J. Shan, “Can artificial intelligence technology reduce carbon emissions? A global perspective,” *Energy Economics*, vol. 143, p. 108285, 2025.
- [17] M. M. Carvalho, D. G. Faria, M. G. Pérez, M. Cardoso, and E. K. Vakkilainen, “Review on mathematical models for travelling-grate iron oxide pellet induration furnaces,” *Energy Procedia*, vol. 120, pp. 588-595, 2017.
- [18] C. Cavalcanti, G. Wanderley, D. Braga, R. Brito, L. Vasconcelos, and K. D. Brito, “Rigorous modeling of the Traveling Grate stage in the iron ore pellet induration process,” *Journal of Materials Research and Technology*, vol. 31, pp. 3410-3424, 2024.
- [19] Y. Wang, Y. Xu, X. Song, Q. Sun, J. Zhang, and Z. Liu, “Novel method for temperature prediction in rotary kiln process through machine learning and CFD,” *Powder Technology*, vol. 439, p. 119649, 2024.
- [20] S. Mun and J. Yoo, “Operating Key Factor Analysis of a Rotary Kiln Using a Predictive Model and Shapley Additive Explanations,” *Electronics*, vol. 13, no. 22, p. 4413, 2024.
- [21] J. Mourão, M. Huerta, U. de Medeiros, I. Cameron, K. O’Leary, and C. Howey, “Guidelines for selecting pellet plant technology,” in *Proceedings of the 6th International Congress on the Science and Technology of Ironmaking—ICSTI*, 2012.
- [22] S. Ding, “Feature selection based F-score and ACO algorithm in support vector machine,” in *Second International Symposium on Knowledge Acquisition and Modeling*, 2009, pp. 19-23.
- [23] C.-K. Zhang and H. Hu, “Feature selection using the hybrid of ant colony optimization and mutual information for the forecaster,” in *International Conference on Machine Learning and Cybernetics*, 2005, vol. 3, pp. 1728-1732.
- [24] A. Sheikhi. *Statistical Analysis and Modeling*. Hamrah Elm Publishing, 2018. [In Persian].
- [25] L. Breiman, “Random forests,” *Machine learning*, vol. 45, pp. 5-32, 2001.
- [26] M. Rezaei and H. Nezamabadi-pour, “A Hybridization Method of Prototype Generation and Prototype Selection for K-NN Rule Based on GSA,” *Journal of AI and Data Mining*, vol. 10, no. 2, pp. 257-268, 2022.
- [27] F. Rosenblatt, “The perceptron, a perceiving and recognizing automaton: (Project Para),” *Cornell Aeronautical Laboratory*, 1957.
- [28] M. T. Mitchell, *Machine learning*, New York: McGraw-hill, 2007.
- [29] Y. Bengio, I. Goodfellow, and A. Courville, *Deep learning*, MIT press Cambridge, MA, USA, 2017.
- [30] M. Vazan. *Deep Learning: Principles, Concepts, and Approaches*. Miad Andisheh Publishing, 2020. [In Persian].
- [31] Z. Liu, “A method of SVM with normalization in intrusion detection,” *Procedia Environmental Sciences*, vol. 11, pp. 256-262, 2011.
- [32] K. Zhong, P. Jain, and I. S. Dhillon, “Mixed linear regression with multiple components,” *Advances in neural information processing systems*, vol. 29, 2016.
- [33] P. Cunningham and S. J. Delany, “K-nearest neighbour classifiers-a tutorial,” *ACM computing surveys (CSUR)*, vol. 54, no. 6, pp. 1-25, 2021.
- [34] G. Louppe, *Understanding random forests: From theory to practice*. Universite de Liege (Belgium), 2014.

Intelligent Process Modeling of a Pelletizing Furnace: A Case Study at Golgohar Mining and Industrial Company

Mohadese Rezaei ^{1*}, Hossein Nezamabadi-pour ², Reza Khksar-pour³

¹ PhD Student, Department of Electrical Engineering, Shahid Bahonar University of Kerman, Kerman, Iran

² Professor, Department of Electrical Engineering, Shahid Bahonar University of Kerman, Kerman, Iran

³ MSc, Golgohar Mining and Industrial Company, Sirjan, Kerman, Iran

Article Information

Original Research Paper

Received:

2025 August 22

Accepted:

2025 October 26

Keywords:

Intelligent Modeling, Machine Learning, Feature Selection, multilayer perceptron neural network, pelletizing furnace

Corresponding Author*:

mo.rezaei@eng.uk.ac.ir

Abstract

The pelletizing process is a key stage in the steel production chain, during which iron ore concentrate is transformed into pellets suitable for use in blast furnaces. The induration (firing) stage within the furnace plays a critical role in determining product quality, as the distribution of temperature and pressure directly affects the physical and mechanical properties of the pellets. The objective of this study is to develop a data-driven model to predict the behavior of the pelletizing furnace at Golgohar Mining and Industrial Company based on real operational data. To this end, process input and output variables were collected and preprocessed, and feature selection was performed using three methods: F-score, mutual information, and Pearson correlation coefficient. Subsequently, Multiple Linear Regression, random forest, k-nearest neighbors, and multilayer perceptron neural network models were trained within a multiple-input-multiple-output (MIMO) framework to simultaneously predict 36 temperature variables and 31 pressure variables. The results indicated that the mutual information-based feature selection method improved the performance of nonlinear models, while the multilayer perceptron achieved the highest accuracy, with an average coefficient of determination of 91.50% for temperature and 89.62% for pressure. These findings demonstrate that data-driven learning approaches can accurately model the complex behavior of the furnace and provide a suitable foundation for designing intelligent control systems and optimizing firing conditions.



: 10.22034/ABMIR.2025.23568.1156

E-ISSN: [2821-2037](#)

/© 2025. Published by Yazd University This is an open access article

under the CC BY 4.0 License (<https://creativecommons.org/licenses/by/4.0/>).

