

تشخیص خطا با استفاده از ماشین بردار پشتیبان چندکلاسه

علی رنجبر^۱، امیرحسین رحمانی^{۲*}

۱- دانشجوی کارشناسی ارشد، گروه برق، واحد دزفول، دانشگاه آزاد اسلامی، دزفول، ایران، Takay313@gmail.com

*۲- مربی، گروه برق، واحد دزفول، دانشگاه آزاد اسلامی، دزفول، ایران، A_h_rahmani@yahoo.com

تاریخ پذیرش: ۹۸/۹/۱۱

تاریخ دریافت: ۹۸/۶/۳۰

چکیده: در محیط‌های صنعتی، مقدار زیادی از داده‌ها تولید می‌شود که به نوبه خود انبار پایگاه داده و داده‌ها را از همه مناطق مربوطه مانند برنامه‌ریزی، طراحی فرآیند، مواد، مونتاژ، تولید، کیفیت، کنترل فرآیند، برنامه‌ریزی، تشخیص خطا، خاموش کردن، مدیریت ارتباط با مشتری و غیره جمع‌آوری می‌کند. داده‌کاوی به ابزار مورد استفاده برای کسب دانش برای روند صنعتی ساخت آهن و فولاد تبدیل شده است. با توجه به رشد سریع داده‌کاوی، صنایع مختلف از فناوری داده‌کاوی برای جستجوی الگوهای پنهان که ممکن است بیشتر به سیستم با دانش جدیدی که مدل‌های جدید را برای بهبود کیفیت تولید، هزینه مطلوب بهره‌وری و تعمیر و نگهداری و غیره پردازند استفاده کردند. بهبود مستمر تمام فرایند تولید فولاد با توجه به اجتناب از کمبود کیفیت و بهبود تولید مرتبط با آن، یک وظیفه اساسی تولیدکننده فولاد است. بنابراین، استراتژی نقص صفر امروزه محبوب است و برای حفظ آن، چندین تکنیک تضمین کیفیت استفاده می‌شود. در این مقاله سعی می‌شود با استفاده از داده‌کاوی حسگرهای موثر در وضعیت سیستم شناسایی شوند و سپس با استفاده از ماشین بردار پشتیبان مدل مناسبی برای پیش‌بینی وضعیت سیستم به دست آورده شود که در این مقاله با دقت بیش از ۹۵ درصد حالت خطا تشخیص داده شده است.

واژه‌های کلیدی: روش‌های گروهی، تصمیم‌گیری، الگوها، بردار پشتیبان، نزدیکترین همسایگی.

۱-مقدمه

از وضعیت اجزای متحرک سیستم را به واحد کنترل ارسال نموده و باعث تغییر وضعیت عملکرد دستگاه‌ها می‌شوند.

حسگرها جهت تبدیل عوامل فیزیکی مانند حرارت، فشار، نیرو، طول، زاویه چرخش، دبی و غیره به سیگنال‌های الکتریکی بکار برده می‌شوند و به همین منظور حسگرهای مختلفی که قابلیت تبدیل این عوامل را به جریان برق دارا می‌باشند، ساخته شده‌اند. با پیشرفت سریع تکنیک خودکارسازی و پیچیده‌تر شدن پروسه‌های صنعتی و کاربرد روزافزون این شاخه از تکنیک نیاز شدیدی به کاربرد حسگرهای مختلف که اطلاعات مربوط به عملیات تولید را درک و بر اساس این اطلاعات مقتضی صادر گردد، احساس می‌شود. حسگرها به عنوان اعضای حسی یک سیستم، وظیفه جمع‌آوری و با تبدیل اطلاعات را به صورتی که برای یک سیستم کنترل و با اندازه‌گیری قابل تجزیه و تحلیل باشد به عهده دارند. در مهر و موم‌های اخیر حسگرها به صورت یک عنصر قابل تفکیک

حسگر المان، حس‌کننده‌ای است که کمیت‌های فیزیکی مانند فشار، حرارت، رطوبت، دما، و ... را به کمیت‌های الکتریکی پیوسته (آنالوگ) یا غیر پیوسته (دیجیتال) تبدیل می‌کند. درواقع حسگر یک وسیله الکتریکی است که تغییرات فیزیکی یا شیمیایی را اندازه‌گیری می‌کند و آن را به سیگنال الکتریکی تبدیل می‌نماید. حسگرها در انواع دستگاه‌های اندازه‌گیری، دستگاه‌های کنترل آنالوگ و دیجیتال مانند PLC مورد استفاده قرار می‌گیرند. عملکرد حسگرها و قابلیت اتصال آن‌ها به دستگاه‌های مختلف از جمله PLC باعث شده است که حسگر بخشی از اجزای جدانشدنی دستگاه کنترل خودکار و رباتیک باشد. حسگرها اطلاعات مختلف

اطلاعات تولیدشده در طول تولید. داده کاوی فرآیند الگوهای جالب و دانش از مقدار زیادی از داده ها است. منبع داده می تواند شامل پایگاه های داده، انبار داده ها، وب، سایر مخازن اطلاعات، داده هایی است که به صورت خودکار به سیستم منتقل می شوند [۳].

فولاد شامل آلیاژ آهن، کربن و منگنز با مقدار کمی از سیلیکون، فسفر و گوگرد است. مراحل تولید فولاد: حرارت دادن، خنک کردن، ذوب کردن، انجماد. هند سومین گروه بزرگ از این گروه است (از هشتم در سال ۲۰۰۳) و انتظار می رود که دومین تولیدکننده در آینده نزدیک باشد.

بیشتر برنامه های کاربردی صنعتی استخراج داده ها در صنایع فولاد عبارتند از مدل سازی سیستم، نزدیک شدن به فن آوری های تولید جدید و بهبود کیفیت محصولات، خواص ضد خوردگی، فولاد گالوانیزه یک محصول که تقاضای بیشتری را در بخش های مختلف تجربه می کند. مدل سازی غیرخطی شامل تکنیک های مهم کاربرد گسترده ای به دلیل کارایی در سرعت. هدف از تجزیه و تحلیل داده ها، پیش بینی طول ضریب عملکرد، ضریب نفوذ به عنوان توابع تعداد زیادی از متغیرها، از جمله ترکیب شیمیایی، حرارت عملیات، سرعت نوار در فرایند انلینگ بود. همچنین در بهینه سازی و کنترل فرآیند صنعتی استفاده می شود. این مطالعه گسترده ای است که به بررسی تکنیک های موجود در کوره های انفجاری در صنایع فولاد می پردازد. هدف اصلی از کوره انفجار، کاهش فرایند شیمیایی و تبدیل اکسید آهن به آهن مایع است که به نام فلز داغ شناخته می شود. فرایند تولید فلز داغ، حدود ۷۰ درصد کل انرژی تولیدشده از تولید فولاد را به مصرف می رساند. علاوه بر این نیازهای اجتماعی و صنعتی فولاد آهن، قیمت بالای مواد خام و عوامل کاهش دهنده نیز ضرورت مدل سازی این فرآیند پیچیده را برای افزایش بهره وری و کاهش هزینه افزایش داده است [۴]. برای اجرای این، تکنیک های داده کاوی در بسیاری از مراحل عملیات انفجار کوره مورد آزمایش قرار گرفته اند.

با توجه به این که D.Y. و P. Hu, W. Zhang, G.J. Zheng و Shi [۵] در این فرایند شکل گیری گرم، بسیاری از متغیرهای طراحی تأثیری بر نتیجه روش های مختلف و پیچیده مانند ویژگی های هندسه و پارامترهای فرایند تشکیل دارند. دشوار است درک ارتباط بین متغیرهای طراحی و نتایج، که برای هدایت طراحی بسیار مهم است. در این مقاله داده کاوی برای کشف تأثیر ویژگی های هندسی گذشته و پارامترهای شکل گیری گرم بر نتایج گرمایی وسیله مدل B استایل معرفی شد و محدوده پارامتر بهینه تعیین شد. ابتدا چند پارامتر متغیر انتخاب و ۱۰۰ گروه داده تجربی انتخاب شدند

با روش سوپر لاتین تولید شد و سپس نتایج تجزیه و تحلیل متشکل از روش های محدود به دست آمد. تجزیه و تحلیل بعدی و تکامل شبیه سازی با استفاده از الگوریتم های تصمیم گیری درخت

دستگاه های مختلف صنعتی مورد استفاده قرار گرفته و پیشرفت سریعی در جهت جوابگویی به تقاضاهای صنعت در این شاخه از علم الکترونیک انجام پذیرفته است.

نوآوری های این مقاله به شرح زیر است:

۱. در این مقاله هیچ محدودیتی برای سیستم مورد مطالعه در نظر گرفته نمی شود.
۲. با کاهش ابعاد پایگاه داده و انتخاب حسگرهای موثر در گزارش وضعیت سیستم سرعت پردازش بیشتر می شود.
۳. علاوه بر پایش سیستم، روش پیشنهادی میتواند به بهره بردار وضعیت سیستم را تنها با کارکرد چند حسگر گزارش دهد.

همان طور که اشاره شد، حسگرهای صنعتی نمونه های متفاوت و گسترده ای دارند که در صنعت کاربردهای بسیاری دارند. از آنجا که امروزه استفاده از حسگرها ما را با حجم زیادی از اطلاعات روبرو می کند، در این مقاله سعی می شود اطلاعات حسگرها دسته بندی شود و در نهایت بتوان بدون حضور اپراتور متخصص وضعیت سیستم را بررسی و پیش بینی کرد.

ساختار این مقاله به صورت زیر می باشد:

۱. در بخش ۱ اطلاعات اولیه در خصوص انواع حسگرها داده خواهد شد و نیز به مرور کارهای انجام شده در زمینه داده - کاوی و کاربردهای آن در صنعت و روش های پیش بینی وضعیت سیستم می پردازیم.
۲. در بخش ۲ به بیان مدل پیشنهادی و فلوچارت حل مساله می پردازیم.
۳. در بخش ۳ با معرفی سیستم مورد مطالعه به بیان نتایج شبیه سازی روش پیشنهادی بر روی این سیستم می - پردازیم.
۴. در آخرین بخش به تحلیل نتایج به دست آمده و کاربردهای مقاله می پردازیم.

در اکثر بخش های فولادی، [۱] تولید بسیار رقابتی است. برای برطرف کردن چالش های فرا منطقه ای، یک شرکت باید تولید کم هزینه داشته باشد و هنوز کارکنان با مهارت بالا، انعطاف پذیر و کارآمد را حفظ کند که بتوانند محصولات با کیفیت بالا و ارزان قیمت را طراحی و تولید کنند. این می تواند با استفاده از تکنیک های داده کاوی به منظور بهبود تصمیم گیری به دست آید.

داده کاوی به طور کلی می تواند به عنوان یک تکنیک برای پیدا کردن الگوهای (استخراج) یا اطلاعات جالب در مقدار زیادی از داده ها باشد [۲]. این تکنیک به طور گسترده ای در زمینه تحقیقاتی مانند مهندسی، بازاریابی، کسب و کار، آموزش و پرورش و در حال حاضر به ویژه در صنایع مانند آهن، فولاد، لاستیک و غیره استفاده می شود. با این حال، دانش می تواند انواع مختلفی داشته باشد و شناسایی نوع دانش برای استخراج هنگام آزمایش مقدار زیادی

(DT) انجام شد. در نهایت، یک سری از قوانین شکل‌گیری گرم پایه B اصلاح شد، از جمله درجه حرارت اولیه فلز ورق باید بین ۷۲۰ C تا ۸۰۰ C کنترل شود. بنابراین، یک مدل واقعی ستون b برای آزمودن قوانین طراحی شده و نتیجه درست و مؤثر بود.

در مرجع [۶] نویسندگان یک سیستم پشتیبانی تصمیم را ایجاد کردند که می‌تواند خوردگی را تعیین کند و زمان رشد خوردگی را برای نگهداری با استفاده از رویکرد مدل‌سازی احتمالی طراحی کند. در این کار، فرایند کنترل خوردگی بر مبنای داده کاوی معرفی شده است، این فرایند به جای عملیات پیشگیرانه و پیشگیرانه که با استفاده از بازرسی‌های منظم برنامه‌ریزی شده و استفاده از حسگرهای خوردگی به‌طور مداوم برای تسهیل اقدامات پیشگیرانه انجام می‌شود، به جای نگهداری پیشگیرانه و پیش‌بینی شده است. استفاده از فرایند داده کاوی شامل تجزیه و تحلیل داده‌های تاریخی که در طول زمان جمع‌آوری می‌شوند در مورد ضخامت فلز معروف که قبلاً در معرض عوامل محیطی مانند بارندگی، دما، رطوبت و غیره قرار دارد. نتایج تجزیه و تحلیل، تصمیم‌گیری‌های پیشگیرانه و دانش مبتنی بر کنترل خوردگی را برای تصمیم‌گیری فراهم می‌کند. روش طبقه‌بندی مانند روش قضیه بیز، شبکه‌های عصبی برای مدل‌سازی عدم قطعیت در وقوع خوردگی با توجه به عدم اطمینان دانش و عدم اطمینان داده‌ها برای تصمیم‌گیری‌های اطلاعات استفاده شده است. مدل‌سازی احتمالاتی برای توسعه یک روش معادله داده استفاده می‌شود که می‌تواند برای به دست آوردن اطلاعات از یک پایگاه داده بر روی فلزات که قبلاً در معرض محیط جو قرار گرفته‌اند استفاده شود. نویسندگان محمد ساعیه، مهدی مقیمی و ایوب بقیری [۷] پیشنهاد می‌کنند که فرایند انحلال یکی از عملیات مهم تولید ورق‌های فولادی سرد و نورد است که کیفیت نهایی محصولات نورد سرد را تحت تأثیر قرار می‌دهد. در فرایند خود، کوئل‌های سرد نورد گرم به آرامی به دمای موردنظر گرم می‌شوند و سپس خنک می‌شوند. مدل‌سازی فرایند خنک‌سازی (پیش‌بینی زمان گرمایش و خنک‌کننده و پیش‌بینی روند گرمایش سیم‌پیچ کوئل) کار بسیار پیچیده و گران است. مدل‌سازی فرایند انحلال را می‌توان با استفاده از مدل‌های حرارتی انجام داد. در این مقاله مدل‌سازی فرایند انحلال با استفاده از تکنیک‌های داده کاوی به دلیل سرعت بالا در پردازش داده‌ها، نتایج قابل قبول و سادگی آن‌ها برای استفاده از آن پیشنهاد شده است. در این مطالعه، آن‌ها روش مدل‌سازی فرایند انلینگ را با استفاده از تکنیک‌های داده کاوی پیشنهاد کردند. پس از آزمودن تکنیک‌های مختلف داده کاوی، شبکه عصبی پروتکل انتقال خورده پیش‌بینی شده برای پیش‌بینی زمان گرمایش و خنک‌کننده و درجه حرارت هسته کوئل در طی فرایند انحلال انتخاب شده است. نتیجه پیش‌بینی شده شبکه‌های عصبی برای مدل‌سازی فرایند انحلال به اندازه کافی دقیق بود، اما درحالی که ما از دقت بیشتری استفاده می‌کنیم، دقت پیش‌بینی

کننده‌ها بهبود می‌یابد. روش حاضر برای پیش‌بینی رفتار فرایندهایی است که نمی‌توان با هر معادلی تحلیلی یا فیزیکی توصیف کرد. تکنیک‌های دیگر مانند طبقه‌بندی، رگرسیون‌های مختلف و یا روش خوشه‌بندی می‌تواند برای بهینه‌سازی پارامترهای ورودی فرایند انحلال برای به حداکثر رساندن بهره‌وری از عملیات خنثی استفاده شود.

به گفته سید مهران شریفی و حمیدرضا اسماعیلی [۸]، استفاده از روش‌های داده کاوی برای پیش‌بینی نقص در سطح فولاد یک چالش بود. در صنعت فولاد، به‌خصوص فولاد آلیاژی، ایجاد یک محصول متفرقه متفاوت می‌تواند هزینه‌های زیادی را برای تولیدکنندگان فولاد تحمیل کند. یکی از نقایص مشترک در تولید نمرات فولاد کم‌کربن، نقص حفره و پلک است. نقص آن یک اتلاف وقت و هزینه برای از بین بردن این نقص است، ما باید سطح محصول را تمیز کنیم. در بعضی موارد، شدت نقص‌ها ممکن است منجر به قطعه‌قطعه شدن محصول شود. سنگ‌زنی باعث تلف شدن زمان و هزینه تولید خواهد شد. فراوانی نقص‌ها به عوامل متعددی مربوط است

تجزیه و تحلیل مواد و فرایندهای تولید. در این مطالعه، نویسندگان یک مدل برای پیش‌بینی این گسل با روش داده کاوی، از جمله درخت تصمیم‌گیری، شبکه عصبی، ایجاد کردند. آن‌ها اعمال کردند

این تکنیک‌ها به داده‌های جمع‌آوری شده از شرکت آلیاژ فولاد ایران است. مدل ایجاد شده با استفاده از درخت تصمیم‌گیری دارای دقت بالاتر است. مزایای این مدل عبارت‌اند از:

- * کاهش روند پردازش و هزینه انرژی با حذف مرحله سنگ‌زنی.
- * یک ابزار بهینه برای پیش‌بینی نقص و درخواست برای تولید محصولات بدون نقص وجود دارد.

بنا به گفته نویسندگان [۹] این مقاله یک روش داده کاوی برای انتخاب متغیر و استخراج دانش از مجموعه داده‌ها را نشان داد. این رویکرد مبتنی بر رگرسیون نمادین غلط است (هر متغیر موجود در مجموعه داده به عنوان متغیر هدف در رگرسیون چندگانه رفتار می‌شود) و متغیر رگرسیون متغیر جدید برای برنامه‌نویسی ژنتیکی. ارتباط هر متغیر ورودی محاسبه شده و مدل تقریبی متغیر هدف ایجاد شده است. پیکربندی‌های برنامه‌نویسی ژنتیکی با متغیرهای مختلف هدف چندین بار اجرا می‌شود تا اثرات تصادفی را کاهش دهد و نتایج جمع‌آوری شده به عنوان شبکه تعامل متغیر نمایش داده می‌شود. این شبکه تعامل، مؤلفه‌های مهم سیستم و روابط ضمنی بین متغیرها را برجسته می‌کند. کل رویکرد به دلیل پیچیدگی کوره انفجار و بسیاری از ارتباطات بین متغیرها بر روی یک مجموعه داده کوره انفجار آزمایش شده است.

بسیاری از متغیرها در فرایند کوره انفجاری به‌طور ضمنی مرتبط هستند، به علت روابط فیزیکی و یا به دلیل کنترل خارجی

رگرسیون نمادین نسبتاً استاندارد مورد بررسی قرار گرفته است. اثرات ترکیب این سازوکارها بر مجموعه داده‌های واقعی در دو فرآیند تولید فولاد نشان داده شده است. بهترین پیشرفت‌ها از لحاظ کیفیت به علت استفاده از عملکرد تناسب‌انداز $as2R$ به دست آمد.

با توجه به Jong-Hag Jeon [۱۲]، این ارائه نمونه‌ای از کاربرد داده‌کاوی در کارخانه فولاد است. مهندسان مسئول کنترل کیفیت با دو نوع مشکلات در حل مشکلات با استفاده از روش‌های آماری برخورد می‌کنند. کوره‌های انفجاری داغ برای مصرف‌کنندگان انرژی مهم هستند. گرمایش انبساط داغ 10% - 15% از کل انرژی مصرف‌شده توسط کوره‌های انفجاری است. کاهش مصرف انرژی برای گرمایش انفجاری داغ تأثیر مهمی در هزینه تولید گرما دارد.

تکنیک‌های داده‌کاوی در صنایع فولادی نتایج بسیار عالی‌ای را برای اهداف زیر به دست می‌آورند. نظارت بر کیفیت محصول نظارت بر فرآیند نظارت بر استراتژی‌های نگهداری.

همچنین، به نظر می‌رسد روشن است که داده‌کاوی با استفاده از شبیه‌سازی مدل‌سازی شمارشی که داده‌های زمان واقعی در دسترس نیست بهبود می‌یابد. در تمام موارد، تأثیر این فناوری‌ها در فرایند صنعتی مدرن، ضروری است که شرکت‌ها از این امر آگاهی یابند تا دانش را از فرایند استخراج کنند و بتوانند آن‌ها را بهبود ببخشند. همانطور که در مرور مطالعات قبلی گفته شد بیشتر روش‌های موجود سعی در مدل کردن سیستم و پایش برخط سیستم داشتند. در بسیاری از موارد به دلیل بالا بودن هزینه امکان پیاده سازی طرح از نظر اقتصادی فراهم نبوده است. بعضی از مقالات بررسی شده سیستم را در حالات خاص بررسی میکردند و روش پیشنهادی آنها تنها مختص به شرایط کار خاص سیستم می باشد. از طرفی بعضی از روش‌های بالا به دلیل حجم بالای محاسبات توانایی پاسخگویی سریع ندارند. در ادامه سعی می‌شود با روش پیشنهادی خود معایب گفته شده در بالا را تا حد امکان پوشش داد.

۲- مدل سازی روش پیشنهادی

ماشین‌های بردار پشتیبانی (SVMs) روش‌های یادگیری مبتنی بر هسته است که برای حل مشکلات رگرسیون موفقیت‌آمیز است. با این حال، استفاده از آن‌ها در برنامه‌های قابلیت اطمینان به‌طور گسترده‌ای مورد بررسی قرار نگرفته است. در این مقاله، یک تحلیل مقایسه‌ای برای ارزیابی اثربخشی SVM در پیش‌بینی زمان به شکست و قابلیت اطمینان قطعات مهندسی بر اساس داده‌های سری زمانی ارائه شده است. کارایی در مورد مطالعات موردی رگرسیون SVM بر روی دیگر روش‌های پیشرفته یادگیری مانند عملکرد پایه شعاعی، مدل Box-Jenkins یکپارچه متحرک و شبکه‌های محلی عادی مجازی محاسبه می‌شود. مقایسه نشان می‌دهد که در موارد مورد تجزیه و تحلیل SVM بهتر از روش‌های دیگر قابل مقایسه است.

پارامترهای کوره انفجار. مثال‌هایی برای متغیرهایی با رابطه ضمنی با متغیرهای دیگر عبارت‌اند از: شعله یا پارامترهای انفجار داغ. معمولاً، چنین رابطه ضمنی پیشینی در سناریوهای مدل‌سازی مبتنی بر داده شناخته نشده است، اما می‌تواند از اطلاعات مربوط به متغیر جمع‌آوری شده از پرونده‌های متعدد معیار سنجش استخراج شود.

با استفاده از رویکرد داده‌کاوی رگرسیون نامتجانس، چندین مدل شناسایی شده‌اند که مقادیر مشاهده شده در فرآیند کوره انفجاری را تقریباً دقیق نشان می‌دهند. این آزمایش‌ها همچنین منجر به تعداد زیادی از مدل‌های مختلف اجزای کوره انفجار می‌شود. مدل تولید شده برای استخراج اطلاعات در مورد روابط ضمنی در مجموعه داده استفاده می‌شود تا مجموعه‌ای از متغیرهای ورودی مربوط را کاهش داده و متقارن کند.

در مرجع بعدی همچنین استفاده از رگرسیون نمادین در کوره انفجار با استفاده از داده‌های آسیاب خورشیدی [۱۰] توصیف می‌شود. این نشان می‌دهد که استفاده از یک سیستم رگرسیون سمبلیک اقتباس شده در دو مجموعه داده‌های مختلف، شامل اندازه‌گیری از فرآیند کوره انفجاری است که شایع‌ترین تولید فلز گرم (آهن مایع) است. اگرچه فرایند شیمیایی و فیزیکی به خوبی درک شده است، از دست دادن گرما در مناطق خاص کوره به‌طور کامل درک نمی‌شود. دانش در مورد چنین روابط می‌تواند برای بهینه‌سازی فرایند کوره انفجاری استفاده شود و بنابراین مدل‌سازی فرایند کوره انفجار بر اساس داده‌های واقعی دنیای جمع‌آوری شده از اهمیت خاصی برخوردار است. تجزیه و تحلیل رگرسیون یک فرعی از داده‌کاوی است که تلاش می‌کند تا دانش موجود در مجموعه داده را آشکار سازد. در این کار، رگرسیون نمادین با استفاده از یک درخت بر اساس یک سیستم برنامه‌نویسی ژنتیکی برای تکمیل فرمول‌های ریاضی انجام می‌شود. برنامه‌نویسی ژنتیکی [۱۱] یک الگوریتم تکاملی است که برنامه‌هایی را برای حل یک مسئله ارائه می‌دهد. در این روش مدل‌سازی، سه جنبه الگوریتمی گنجانده شده و با یک رویکرد رگرسیون سمبلیک استاندارد مقایسه می‌شود. دقیقاً نویسندگان اثرات انتخاب به‌صورت آفلاین را با استفاده از ضریب همبستگی r به عنوان عملکرد تناسب‌انداز مورد آزمایش قرار دادند و نمونه‌های ارزیابی شده یا موارد تناسب‌انداز را نیز نمونه‌برداری کردند. انتخاب بیرونی یک انتخاب اضافی در الگوریتم ژنتیک و برنامه‌نویسی ژنتیکی است که بعد از نو ترکیب و جهش اعمال می‌شود. تناسب فرد در تجزیه و تحلیل رگرسیون نماد معمولاً به عنوان خطای میانگین میدان بین ارزش پیش‌بینی شده مدل و ارزش مشاهده شده متغیر هدف محاسبه می‌شود. الگوریتم ژنتیک سریع‌تر از الگوریتم ژنتیک بدون انتخاب فرزند است. در این سهم، سه سازگاری الگوریتمی، انتخاب فرزند، ضریب تعیین r به عنوان تابع تناسب و نمونه‌برداری، به یک سیستم

فضای جستجو را در یکی از ابعاد، به دو نیم‌بخش تقسیم می‌کند به گونه‌ای که هر نیم‌بخش تقریباً نیمی از نقاط بخش بزرگتر را در برخواهد گرفت. پرس و جو با پیمایش درخت از ریشه تا برگ صورت می‌پذیرد به گونه‌ای که در هر مرحله، مقدار بردار گره در بُعد مشخصی، با همان بُعد از نقطه پرس و جو، مقایسه می‌شود؛ اگر مقدار نقطه پرس و جو بیشتر بود به زیر درخت سمت راست و در غیر این صورت به زیردرخت چپ حرکت می‌کنیم. همچنین بسته به اینکه فاصله این دو نقطه چقدر است، شاخه‌های کناری نیز ممکن است نیاز به بررسی داشته باشند. برای مسائلی که تعداد بعد فضا ثابت است، میانگین پیچیدگی زمانی این روش از مرتبه $O(\log N)$ است. در مورد نقاط راندم توزیع شده تحلیل پیچیدگی زمانی مربوط به بدترین حالت صورت می‌گیرد. روش دیگر، استفاده از ساختمان داده درخت مستطیلی است که برای استفاده NNS در زمینه‌های پویا بکار گرفته می‌شود. این ساختمان داده دارای الگوریتم‌های مؤثری برای درج و حذف از درخت می‌باشد.

در رابطه با فضاهای متریک کلی، روش انشعاب و تحدید با نام درخت متریک شناخته می‌شود. به عنوان نمونه می‌توان به vp-tree و BK-tree اشاره کرد.

با استفاده از مجموعه‌ای از نقاطی از یک فضای ۳ بعدی و قرار دادن آن‌ها در یک درخت BSP و نیز با داشتن یک نقطه پرس و جو از همان فضا، یک راه حل ممکن برای مسئله پیدا کردن نزدیک‌ترین نقطه در ابری از نقاط، به نقطه پرس و جو، در زیر در قالب یک الگوریتم توضیح داده می‌شود. گفتنی است، ممکن است هیچ نقطه‌ای به عنوان جواب وجود نداشته باشد، زیرا ممکن است این جواب، یکتا نباشد اما در عمل، معمولاً ما تنها دنبال پیدا کردن یکی از زیر مجموعه‌های تمام ابر-نقاط داده شده، که در کوتاه‌ترین فاصله نسبت به نقطه پرس و جو، واقع شده هستیم.

ایده این است که، برای هر شاخه از درخت، حدس می‌زنیم که نزدیک‌ترین نقطه در ابر، در نیم فضای شامل نقطه پرس و جو قرار دارد. اگر چه این حدس ممکن است درست نباشد، اما ابتکار خوبی است. پس از پیمایش بازگشتی نیم فضای حدس زده شده، فاصله محاسبه شده را با کمترین فاصله نقطه پرس و جو تا صفحه پارتیشن‌کننده فضا مقایسه می‌کنیم. این فاصله، کمترین فاصله ایست که بین نقطه پرس و جو و نقاط نیم فضای جستجو نشده می‌تواند وجود داشته باشد. اگر فاصله نزدیک‌ترین نقطه در این نیم-فضا کمتر بود، پاسخ قطعاً در همین نیم-فضاست و نیازی به جستجوی نیم فضای مجاور نیست، در غیر اینصورت باید آن نیم-فضا نیز به صورت بازگشتی جستجو شده و پاسخ آن با کوتاه‌ترین فاصله بدست آمده تاکنون مقایسه گردد و مقدار درست برگردانده شود. عملکرد این الگوریتم زمانی که نقطه پرس و جو در نزدیکی ابر باشد، بیشتر به زمان لگاریتمی نزدیکتر است تا زمان خطی. زیرا هنگامی که فاصله بین نقطه پرس و جو و نزدیک‌ترین نقطه در ابر

در ادامه ابتدا روش انتخاب ویژگی با توجه به نزدیک‌ترین همسایگی معرفی می‌شود و سپس مدل سازی بردار پشتیبان بیان می‌شود.

جستجوی نزدیک‌ترین همسایه (یا NNS)، که همچنین با نام‌های جستجوی مجاورت، جستجوی همسانی یا جستجوی نزدیک‌ترین نقطه شناخته می‌شود، یک مسئله بهینه سازی برای پیدا کردن نزدیک‌ترین نقطه‌ها در فضاهای متریک است. مسئله بدین صورت است که: مجموعه S شامل تعدادی نقطه در یک فضای متریک مانند M و نیز یک نقطه پرس و جو $q \in M$ داده شده‌است، هدف پیدا کردن نزدیک‌ترین نقطه در S به q است. در بسیاری از موارد، فضای M به صورت یک فضای اقلیدسی d-بعدی و فاصله بین نقاط با معیار فاصله اقلیدسی، فاصله منتهن یا دیگر فاصله‌های متریک سنجیده می‌شود. دونالد کنوت در جلد ۳ از هنر برنامه نویسی کامپیوتر (۱۹۷۳) این مسئله را "مسئله دفتر پست" نامید. این نام‌گذاری اشاره یک به برنامه کامپیوتری دارد که برای اختصاص نزدیک‌ترین اداره پست به یک محل اقامت توشته شده بود.

راه حل‌های متعددی برای مسئله NNS پیشنهاد شده‌اند. کیفیت و سودمندی این الگوریتم‌ها بر اساس پیچیدگی زمانی و همچنین پیچیدگی حافظه مصرفی ساختمان داده‌های آنها، ارزیابی می‌شود. مشاهدات غیررسمی که معمولاً آن را با عنوان نفرین ابعاد یاد می‌کنند، بیان می‌دارد که در واقع هیچ راه حل دقیق و همه منظوره‌ای برای مسئله NNS در فضای اقلیدسی با ابعاد بالا وجود ندارد که از پیش پردازش با مرتبه چندجمله‌ای و زمان جستجوی چندلگاریتمی $O((\log n)^k)$ بهره ببرد.

۲-۱- جستجوی خطی

ساده‌ترین راه حل برای مسئله NNS، محاسبه فاصله نقطه پرسش شده تا تمامی نقاط پایگاه داده، همراه با نگهداری "بهترین جواب پیدا شده تا کنون" است. این الگوریتم که به آن الگوریتم سراسر است نیز گفته می‌شود، دارای مرتبه زمانی $O(nd)$ است که در آن N اندازه مجموعه S و d تعداد ابعاد فضای M است. در این روش هیچ داده ساختاری از جستجوها ذخیره نخواهد شد و الگوریتم بجز ذخیره نقاط پایگاه داده، هزینه حافظه‌ای دیگری ندارد. روش ذکر شده، بطور میانگین، از روش‌های پارتیشن‌بندی فضا، در ابعاد بالاتر، بهتر عمل می‌کند.

۲-۲- پارتیشن‌بندی فضا

از سال ۱۹۷۰، روش انشعاب و تحدید بر روی مسئله NNS اعمال شده‌است. در فضای اقلیدسی، این کار با نام شاخص‌های مکانی یا روش دسترسی مکانی شناخته می‌شود. روش‌های متعددی برای پارتیشن بندی فضا به منظور حل مسئله NNS ارائه و توسعه داده شده‌اند. در این بین، ساده‌ترین روش، استفاده از ساختمان داده درخت کی‌دی (به انگلیسی: k-d tree) است که به صورت تکراری

ما مجموعه داده‌های آزمایش D شامل n عضو (نقطه) را در اختیار داریم که به صورت زیر تعریف می‌شود:

$$D = \{(X_i, y_i) \mid X_i \in \mathbb{R}^p, y_i \in \{-1; 1\}\}_{i=1}^n \quad (3)$$

جایی که مقدار y برابر ۱ یا -۱ و هر X یک بردار حقیقی p -بعدی است. هدف پیدا کردن ابرصفحه جداکننده با بیشترین فاصله از نقاط حاشیه‌ای است که نقاط با $y=1$ را از نقاط $y=-1$ جدا کند. هر ابرصفحه می‌تواند به صورت مجموعه‌ای از نقاط $\{x\}$ که شرط زیر را ارضا می‌کنند نوشته شود:

$$W * X - b = 0 \quad (4)$$

بردار نرمال است. که به ابرصفحه عمود است. ما می‌خواهیم w, b را طوری انتخاب کنیم که بیشترین فاصله بین ابرصفحه‌های موازی که داده‌ها را از هم جدا می‌کنند، ایجاد شود. این ابرصفحه‌ها با استفاده از رابطه زیر توصیف می‌شوند.

اگر داده‌های آموزشی جدایی پذیر خطی باشند، ما می‌توانیم دو ابرصفحه در حاشیه نقاط به‌طوری که هیچ نقطه مشترکی نداشته باشند، در نظر بگیریم و سپس سعی کنیم، فاصله آن‌ها را، حداکثر کنیم. با استفاده از هندسه، فاصله این دو صفحه است؛ بنابراین ما باید را مینیمم کنیم. برای اینکه از ورود نقاط به حاشیه جلوگیری کنیم، شرایط زیر را اضافه می‌کنیم: برای هر i

$$\begin{aligned} w \cdot x_i - b &\geq 1, \text{ if } y_i = 1 \\ w \cdot x_i - b &\leq -1, \text{ if } y_i = -1 \end{aligned} \quad (5)$$

با کنار هم قرار دادن این دو یک مسئله بهینه‌سازی به دست می‌آید:

$$\begin{aligned} \min_{(w,b)} \quad & ||w|| \\ \text{s.t.} \quad & \forall 1 \leq i \leq n \\ & y_i(w \cdot x_i - b) \geq 1. \end{aligned} \quad (6)$$

همان‌طور که در بالا توضیح داده شد روش پیشنهادی ما شامل مراحل زیر می‌باشد که باید به ترتیب انجام شوند:

۱. جمع آوری داده‌های سیستم به تعداد زیاد. هرچه پایگاه داده بزرگتر داشته باشیم شرایط و موقعیت‌های بیشتری از سیستم در نظر گرفته می‌شود و نتیجه کامل‌تر خواهد بود.
۲. دسته‌بندی داده‌ها با استفاده از الگوریتم نزدیکترین همسایگی. در این قسمت داده‌ها را به چند دسته دلخواه تقسیم کرده و مراکز دسته‌هایی که بهترین جواب مساله همسایگی را به ما بدهند انتخاب می‌شوند.
۳. با مراکز داده شده در مرحله قبل و استفاده از الگوریتم ژنتیک سعی می‌شود دسته‌بندی مرحله ۲ را با تعداد کمتری از ویژگی‌ها به دست آورد.

به صفر میل کند، الگوریتم تنها با استفاده از جستجوی نقطه پرس و جو به عنوان یک کلید می‌تواند به نتیجه درست برسد.

فاصله دو نقطه p و q اندازه پاره‌خطی است که آنها را به هم متصل می‌کند (pq vector). در مختصات دکارتی اگر:

$$\begin{aligned} p &= (p_1, p_2, \dots, p_n) \\ q &= (q_1, q_2, \dots, q_n) \end{aligned} \quad (1)$$

دو نقطه در فضای اقلیدسی n بعدی باشند، آنگاه فاصله بین آنها به صورت زیر تعریف می‌شود:

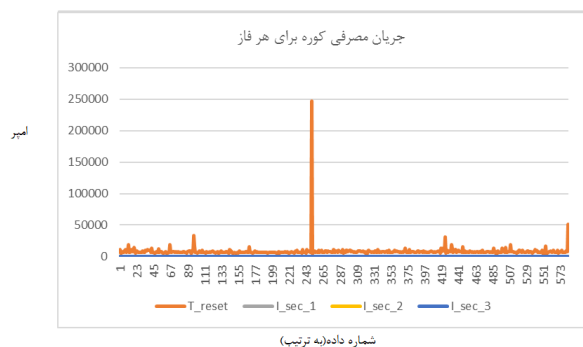
$$\begin{aligned} d(p, q) &= \sqrt{(p_1 - q_1)^2 + (p_2 - q_2)^2 + \dots + (p_n - q_n)^2} \\ &= \sqrt{\sum_{i=1}^n (p_i - q_i)^2} \end{aligned} \quad (2)$$

حال با استفاده از الگوریتم ژنتیک ابتدا بهترین مراکز دسته انتخاب می‌شود. به صورتی که مجموع فواصل همه اعضای هر دسته تا مرکز کمترین مقدار خود باشد.

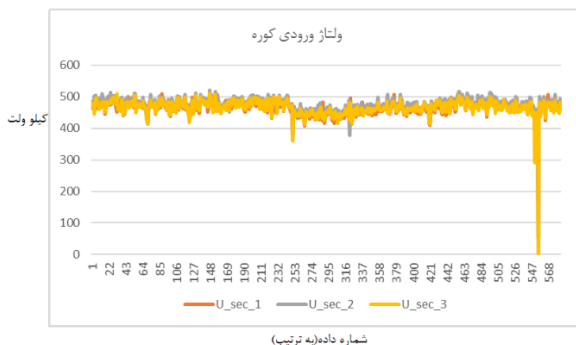
سپس با حذف داده‌ها به صورت تصادفی و تکرار سعی می‌شود حداقل تعداد ویژگی‌ها را به صورتی پیدا کرد که نماینده دسته بندی اولیه باشند.

در این مقاله برای انتخاب ویژگی بیشترین تغییرات در خروجی در نظر گرفته می‌شود. در ادامه روش بردار پشتیبان را به اختصار توضیح می‌دهیم.

ماشین بردار پشتیبانی (SVMs) یکی از روش‌های یادگیری با نظارت است که از آن برای طبقه‌بندی و رگرسیون استفاده می‌کنند. این روش از جمله روش‌های نسبتاً جدیدی است که در سال‌های اخیر کارایی خوبی نسبت به روش‌های قدیمی‌تر برای طبقه‌بندی نشان داده‌است. مبنای کاری دسته‌بندی کننده SVM دسته‌بندی خطی داده‌ها است و در تقسیم خطی داده‌ها سعی می‌کنیم خطی را انتخاب کنیم که حاشیه اطمینان بیشتری داشته باشد. حل معادله پیدا کردن خط بهینه برای داده‌ها به وسیله روش‌های QP که روش‌های شناخته شده‌ای در حل مسائل محدودیت‌دار هستند صورت می‌گیرد. قبل از تقسیم خطی برای اینکه ماشین بتواند داده‌های با پیچیدگی بالا را دسته‌بندی کند داده‌ها را به وسیله تابع ϕ به فضای با ابعاد خیلی بالاتر می‌بریم. برای اینکه بتوانیم مسئله ابعاد خیلی بالا را با استفاده از این روش‌ها حل کنیم از قضیه دوگانگی لاگرانژ برای تبدیل مسئله مینیمم‌سازی مورد نظر به فرم دوگانگی آن که در آن به جای تابع پیچیده ϕ که ما را به فضایی با ابعاد بالا می‌برد، تابع ساده‌تری به نام تابع هسته که ضرب برداری تابع ϕ است ظاهر می‌شود استفاده می‌کنیم. از توابع هسته مختلفی از جمله هسته‌های نمایی، چندجمله‌ای و سیگموئید می‌توان استفاده نمود.



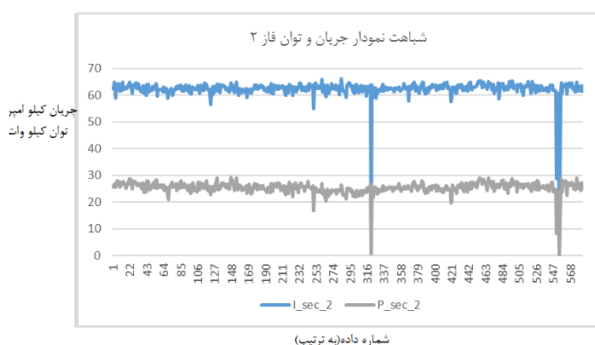
شکل (۳): جریان مصرفی کوره برای هر فاز



شکل (۴): ولتاژ ورودی کوره



شکل (۵): شباهت نمودار جریان و توان راکتیو فاز ۳



شکل (۶): شباهت نمودار جریان و توان فاز ۲

شاید به نظر آید طبق قوانین برقی توان به جریان وابسته است و باید گفت این مطلب درست است ولی در فرایند ویژگی رابطه مستقیم در نظر گرفته می شود و اثبات می شود درحالی که توان اکتیو به ولتاژ و زاویه بار نیز بستگی دارد.

۴. داده های به دست آمده از مرحله ۳ را با استفاده از ماشین بردار پشتیبان دسته بندی کرده و سعی می شود صفحات جداکننده داده ها، را با بیشترین حاشیه به دست آورد. در ادامه بر روی نمونه آزمایشی مراحل گفته شده در بالا را انجام می دهیم و نتایج را بررسی می کنیم.

۳- نتایج شبیه سازی

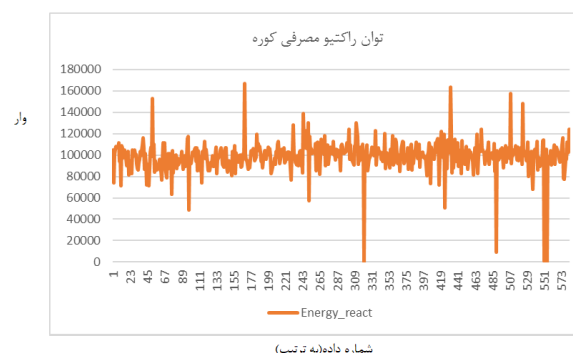
شرکت فولاد خوزستان، دومین تولیدکننده فولاد خام در ایران بعد از شرکت فولاد مبارکه اصفهان می باشد. کارخانه مرکزی این شرکت با وسعت ۳/۸ کیلومتر مربع، در مجاورت شهر اهواز (کیلومتر ۱۰ جاده اهواز بندرامام خمینی (ره)) واقع شده است و دفتر مرکزی آن نیز در اهواز قرار دارد.

شرکت فولاد خوزستان در سال ۱۲ فروردین ۱۳۶۸ به عنوان نخستین مجتمع تولید آهن و فولاد ایران، به روش احیاء مستقیم و کوره قوس الکتریکی راه اندازی شد.

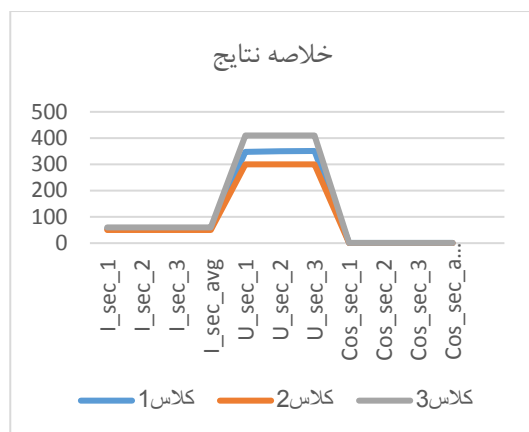
هدف از انجام این مرحله کاهش ابعاد پایگاه داده از ماتریس ۱۵۰ ستونی به ماتریسی با تعداد ستون های کمتر می باشد. اطلاعات استفاده شده در این مقاله از داده های کوره شماره ۲ واحد کارخانه های فولادسازی است. اطلاعات شامل پایش های الکتریکی و حرارتی کوره می باشند. در ادامه قسمتی از این اطلاعات نشان داده می شود.



شکل (۱): توان اکتیو مصرفی کوره



شکل (۲): توان راکتیو مصرفی کوره



شکل (۷): خلاصه نتایج

۴- نتیجه گیری

همان‌طور که می‌بینیم کلاس ۱ که کلاس میانی نیز است کلاس ناحیه خطر می‌باشد که معمولاً بعد از این مرحله احتمال خطای سیستم بالا می‌رود.

نمودار کلاس ۲ که کلاس خطا نیز نامیده می‌شود مرزبندی ناحیه خطا را نشان می‌دهد. با بررسی داده‌ها در صورتی که هر کدام از پارامترها به زیر نمودار کلاس ۲ بیایند می‌توان به این نتیجه رسید که در سیستم خطایی رخ داده است.

نمودار کلاس ۳ که کلاس نرمال نیز نامیده می‌شود مرز بندی کارکرد نرمال را نشان می‌دهد و هر پارامتری که در روی این خط یا بالای آن قرار داشته باشد در حالت نرمال به سر می‌برد.

در بررسی مناطق حاشیه ای ناحیه بین نمودار کلاس ۱ و کلاس ۲ همان‌طور که انتظار می‌رود نه می‌توان به این ناحیه ناحیه خطا نام داد نه ناحیه خطر. ولی هر چه به محدوده ناحیه کلاس ۲ نزدیک می‌شویم احتمال خطا بیشتر می‌شود.

ناحیه بین نمودار کلاس ۳ و کلاس ۲ همان‌طور که انتظار می‌رود نه می‌توان به این ناحیه ناحیه نرمال نام داد نه ناحیه خطر. ولی هر چه به محدوده ناحیه کلاس ۱ نزدیک می‌شویم احتمال خطا بیشتر می‌شود.

از آن‌جا که دستگاه‌های الکتریکی حفاظتی اجازه عبور جریان یا ولتاژ بیشتر از مقدار خاصی را نمی‌دهند می‌توان تمام منطقه بالای نمودار کلاس ۳ را منطقه کارکرد عادی سیستم در نظر گرفت و همان‌طور که دیدیم روش پیشنهادی به خوبی توانست داده‌های حاصل از سال پایش سیستم را به خوبی خلاصه و کلاس بندی کند و با دقتی بیش از ۹۵ درصد حالت خطا را تشخیص دهد. در انجام این مقاله وضعیت سیستم را به ۳ قسمت تقسیم کردیم که می‌توان بنا به تشخیص بهره‌بردار این دست بندی کاهش و افزایش یابد.

همان‌طور که به‌صورت نمونه در بالا دیده می‌شود دمای کوره و جریان ورودی باهم هیچ شباهتی ندارند. در نتیجه نمی‌توان هیچ‌کدام از این دو را از ویژگی‌ها حذف کرد. پس از انجام مرحله انتخاب ویژگی و یا کاهش ابعاد ویژگی‌های زیر انتخاب می‌شوند.

جدول (۱): ویژگی‌های انتخاب شده

شماره	پارامتر	معرفی
۱	I_sec_1	جریان فاز ۱
۲	I_sec_2	جریان فاز ۲
۳	I_sec_3	جریان فاز ۳
۴	I_sec_avg	جریان متوسط
۵	U_sec_1	ولتاژ فاز ۱
۶	U_sec_2	ولتاژ فاز ۲
۷	U_sec_3	ولتاژ فاز ۳
۸	T_reset	دمای کوره
۹	Cos_sec_1	زاویه بین ولتاژ و جریان فاز ۱
۱۰	Cos_sec_2	زاویه بین ولتاژ و جریان فاز ۲
۱۱	Cos_sec_3	زاویه بین ولتاژ و جریان فاز ۳
۱۲	Cos_sec_avg	متوسط زاویه بین ولتاژ و جریان

بردارهای پشتیبان به زبان ساده، مجموعه‌ای از نقاط در فضای n بعدی داده‌ها هستند که مرز دسته‌ها را مشخص می‌کنند و مرزبندی و دسته‌بندی داده‌ها بر اساس آن‌ها انجام می‌شود و با جابجایی یکی از آن‌ها، خروجی دسته‌بندی ممکن است تغییر کند. در این مقاله بعد از اعمال مرحله انتخاب ویژگی به دنبال بردارهای پشتیبان ۱۰ بعدی می‌گردیم.

بعد از انجام الگوریتم ماشین بردار پشتیبان به نتایج زیر می‌رسیم.

جدول (۲): نتایج ویژگی‌های انتخابی

پارامتر	کلاس ۱	کلاس ۲	کلاس ۳
I_sec_1	۵۵	۵۰	۶۰
I_sec_2	۵۵	۵۰	۶۰
I_sec_3	۵۵	۵۰	۶۰
I_sec_avg	۳۴۷	۳۰۰	۴۱۰
U_sec_1	۳۵۰	۳۰۰	۴۱۰
U_sec_2	۳۵۱	۳۰۰	۴۱۰
U_sec_3	۳۵۱	۳۰۰	۴۱۰
T_reset	۵۰۰۰	۱۰۰۰۰	۲۰۰۰۰
Cos_sec_1	۰/۸	۰/۵	۰/۸۸
Cos_sec_2	۰/۸	۰/۵	۰/۸۸
Cos_sec_3	۰/۸	۰/۵	۰/۸۸
Cos_sec_avg	۰/۸	۰/۵	۰/۸۸

مراجع

رزومه



علی رنجبر در اهواز متولد شده است (۱۳۷۲). تحصیلات دانشگاهی خود را در مقطع کارشناسی مهندسی برق - کنترل از دانشگاه آزاد اسلامی واحد ماهشهر (۱۳۹۵)، کارشناسی ارشد مهندسی برق - کنترل از دانشگاه آزاد اسلامی واحد دزفول (۱۳۹۸) سپری کرده است. فعالیت‌های پژوهشی و علاقه‌مندی ایشان در زمینه داده کاوی، یادگیری ماشین، اتوماسیون صنعتی، ابزار دقیق، کنترل بهینه، کنترل مقاوم و کنترل فازی است و در حال حاضر کارشناس ارشد برق - کنترل و مدیر پروژه ریومپ سیستم برق و ابزار دقیق پمپ خانه احیا ۲ شرکت فولاد خوزستان می باشد.



امیرحسین رحمانی در اصفهان متولد شده است (۱۳۴۷). مدارک کارشناسی و کارشناسی ارشد خود را در سال‌های (۱۳۷۰) و (۱۳۷۴) به ترتیب در رشته‌های مهندسی برق و کنترل از دانشگاه تهران اخذ نموده است. ایشان هم اکنون دانشجوی دکترای کنترل دانشگاه آزاد اسلامی واحد علوم و تحقیقات تهران و همچنین عضو هیأت علمی دانشگاه آزاد اسلامی واحد دزفول است. زمینه تحقیقاتی ایشان تحلیل سیستم‌های کنترل و روش‌های بهینه‌سازی هوشمند است.

- [1] Doyen L, Gaudoin O. ۲۰۰۴ Classes of imperfect repair models based on reduction of failure intensity or virtual age. Reliability Engineering & System Safety; ۸۴:۴۵-۵۶
- [2] Fuxing Yu, Yina Suo, Xin Zang, Aidind Yan, Fulong Liu, (۲۰۱۳), Data mining in blast furnace smelting parameter, Applied mechanics and materials, vol. ۳۰۳-۳۰۶, pp. ۱۰۹۳-۱۰۹۶
- [3] Hand D. J, Manila H, & Smyth, (۲۰۰۱): Principles of Data mining, MIT press, Cambridge, Massachusetts. ISBN-۰۲۶۲-۰۸۲۹۰-X
- [4] Jiawei Han, Micheline Kamber, Jian Pei, (۲۰۱۲), Data Mining: Concepts and Techniques, Third Edition, USA, Morgan Kaufmann Publishers.
- [5] John R. Koza, ۱۹۹۲. Consulting Associate Professor in computer science department Stanford university, Genetic programming: On the programming of computers by means of natural selection, the MIT press.
- [6] Jong-Hag Jeon, POSCO, Pohang South Korea, Data mining application of six-sigma project, SUGI ۲۹ solutions, paper. ۱۸۶-۲۹
- [7] Mahamad saraee school of computing, science and Eng., university of Salford, greater Manchester, UK, Mehdi Moghimi, Dept. of Elec. & computer Eng., Islamic Azad university, Najafabad branch, Isfahan, Iran, Ayoub bagheri, Dept. of Elec and computer Eng, Isfahan university of technology, Isfahan Iran, (۲۰۱۱), Modeling Batch Annealing Process using Data Mining Techniques. ACM journal.
- [8] Manisha Verma, Mauly Srivastava, Neha Chack, Abul Kumar Diswar, Nidhi Gupta, (۲۰۱۲), 'Comparative study of various clustering algorithms in data mining'. International journal for engineering research and applications, Vol ۲, Issue ۳, pp. ۱۳۷۹-۱۳۸۴
- [9] Michael Kommenda, Gabriel Kronberger Christoph Feilmayr and Michael Affenzeller, (۲۳ Sep ۲۰۱۳), Data mining using unguided symbolic regression on a blast furnace dataset, arXiv:۱۳۰۹.۵۹۳۱v۱[cs.NE].
- [10] Michael Kommenda, Gabriel Kronberger, Christoph Feilmayr, Leonhard Schickmair, Michael Affenzeller, Stephan Winkler and Stefan Wagner, Application of symbolic regression on blast furnace and temper mill datasets, [n.d].
- [11] Moura MC, Droguett EL. Mathematical formulation and numerical treatment based on transition frequency densities and quadrature methods for non-homogeneous semi-Markov processes. Reliability Engineering & System
- [12] Nine law of Data Mining by Tom Khabaza (<http://www.kdnuggets.com/۱۶/۲۰۱۵/nine-datamining-part-۱.html>).

Fault Diagnosis by Using the Multi-class Support Vector Machine

Ali Ranjbar¹, Amirhossein Rahmani^{2,*}

1- Department of Electrical Engineering, Dezful Branch, Islamic Azad University, Dezful, Iran,
Takay313@gmail.com

*2- Department of Electrical Engineering, Dezful Branch, Islamic Azad University, Dezful, Iran

Abstract: In industrial environments, a large amount of data is generated which in turn stores the database and data from all relevant areas such as planning, process design, materials, assembly, production, quality, process control, scheduling, error detection, shutdown, relationship management. Collects with the customer, etc. Data mining has become the tool used to gain knowledge of the industrial process of iron and steel making. Due to the rapid growth of data mining, various industries have been using data mining technology to search for hidden patterns that may be more relevant to the new windshield system, which will introduce new models to improve production quality, optimal cost of productivity and maintenance, and so on. Continuous improvement of the entire steel production process due to the avoidance of quality deficiencies and associated production improvement is an essential task of the steel producer. Therefore, the zero defect strategy is popular today and several quality assurance techniques are used to maintain it. This article attempts to identify the effective state-of-the-art sensors in the system using data mining and then to obtain a suitable model using a support vector machine to predict the system status in this article, the accuracy of more than 95% of the error states is detected.

Keywords: Group Methods, Decision Making, Patterns, Support Vector, Nearest Neighbor.