





دانشگاه صنعتی شاهرود
دانشکده شیمی

پایان نامه کارشناسی ارشد شیمی تجزیه

کاربرد روش های کمومتری کس جهت پیش بینی اندیس بازداري برخی ترکیبات آلی

نگارنده:

الهه سلامتی

اساتید راهنما:

دکتر ناصر گودرزی

دکتر داوود شاهسونی

استاد مشاور:

دکتر منصور عرب چم جنگلی

آذر ماه ۱۳۹۷



مدیریت تحصیلات تکمیلی

باسمه تعالی

شماره: ۱۸۷۶ د، س
تاریخ: ۳۰، ۱۰، ۹۷

فرم شماره (۳) صورتجلسه نهایی دفاع از پایان نامه دوره کارشناسی ارشد

با نام و یاد خداوند متعال، ارزیابی جلسه دفاع از پایان نامه کارشناسی ارشد خانم الهه سلامتی با شماره دانشجویی ۹۴۰۹۴۳۴ رشته شیمی گرایش تجزیه تحت عنوان کاربرد روش‌های کمومتریکس جهت پیش‌بینی اندیس بازداري برخی ترکیبات آلی که در تاریخ ۱۳۹۷/۰۹/۲۷ با حضور هیأت محترم داوران در دانشگاه صنعتی شاهرود برگزار گردید به شرح ذیل اعلام می‌گردد:

قبول (با درجه: ☒ قبول) ☐ مردود			
نوع تحقیق: نظری ☒ عملی ☐			
عضو هیأت داوران	نام و نام خانوادگی	مرتبه علمی	امضاء
۱- استاد راهنمای اول	دکتر ناصر گودرزی	دانشیار	
۲- استاد راهنمای دوم	دکتر داوود شاهسونی	دانشیار	
۳- استاد مشاور	دکتر منصور عرب چم جنگلی	استاد	
۴- نماینده تحصیلات تکمیلی	دکتر علی کیوانلو	دانشیار	
۵- استاد ممتحن اول	دکتر فاطمه مصدرا لامور	استادیار	
۶- استاد ممتحن دوم	دکتر قدمعلی باقریان	دانشیار	



نام و نام خانوادگی رئیس دانشکده: دکتر مهدی میرزایی

تاریخ و امضاء و مهر دانشکده:

تصیر: در صورتی که کسی مردود شود حداکثر یکبار دیگر (در مدت مجاز تحصیل) می‌تواند از پایان نامه خود دفاع نماید (دفاع مجدد نباید زودتر از ۴ ماه برگزار شود).

تقدیم به

همسرم

به پاس تعبیر عظیم و انسان‌انوار از کلمه ایثار و از خودگذشتگی به پاس عاطفه سرشار و گرمای
امیدبخش و جودش که در این سردترین روزگار این بهترین پشتیبان است. به پاس قلب بزرگش که
فریاد در سراسر و سرگردانی و ترس در پناهش به شجاعت مگر آید و به پاس صحبت‌ها و برادریش
که هرگز فروکش نمی‌کند.

مادرم

تو روشنائی‌بخش تاریک‌ها هستی و ظلمت‌اندیشه را نور مریخی. چگونه سپاسگویم
مهربانی و لطف تو را که سرشار از عشق و یقین است. چگونه سپاسگویم تأثیر عشق آموز
تو را که چراغ روشن هدایت را بر کلبه رمق و جودم فروزان ساخته‌ای و در مقابل این همه
عظمت و شکوه تو مرا نه توان سپاس است و نه کلام و صفت

تشکر و قدر دانس

تشکر و سپاس بر پايان مضمون خداي است که بشر را آفریده و به او قدرت اندیشیدن داده و تواناها را بالقوه را در وجود انسان قرار داده و او را امر به تلاش و کوشش نموده و راهنمایان را برابر هدایت بشر فرستاده است. پس از ارادت خاضعانه به درگاه خداوند بر همتا لازم است از استیدارجمند جناب دکتر ناصر گودرز و جناب دکتر دلاورد شاسونر و جناب دکتر منصور عرب جم جنگلر به خاطر سع صدر و رهنمودها و دلسوزانه شای که در تهیه این تحقیق مرا مورد لطف خود قرار داده اند و راهنمایان را لازم را نموده اند تشکر و قدر دانس نموده و موفقیت همگان را از درگاه احدیت خواهانم، و همچنین از استیدار داور جناب آقا سردکتر باقریای و سرکار خانم صدر الامور به خاطر مطالعه پایانی نام و شرکت در جلسه دفاع قدر دانس مینمایم

همسر عزیزم که مشکلات زندگی را برایم تسهیل نمود و پیمودن این راه بدون همراهی او ممکن نمیشد، سپاسگذارم

و از سرکار خانم مظفر که مرا در انجام این پروژه یار نمودند کمال تشکر را دارم

از درگاه لایزال احدیت، توفیق روز افزون آنگ بزرگوار را مسئلت داشته و همواره سعادت و

بهر روزیشای را آرزو مندم

تعهدنامه

اینجانب الهه سلامتی دانشجوی دوره کارشناسی ارشد رشته شیمی تجزیه دانشکده شیمی دانشگاه صنعتی شاهرود نویسنده پایان نامه کاربرد روش های کمومتری کس جهت پیش بینی اندیس بازداري برخی ترکیبات آلی. تحت راهنمایی دکتر ناصر گودرزی و دکتر داوود شاهسونی متعهد می شوم .

- تحقیقات در این پایان نامه توسط اینجانب انجام شده است و از صحت و اصالت برخوردار است .
- در استفاده از نتایج پژوهش های محققان دیگر به مرجع مورد استفاده استناد شده است .
- مطالب مندرج در پایان نامه تاکنون توسط خود یا فرد دیگری برای دریافت هیچ نوع مدرک یا امتیازی در هیچ جا ارائه نشده است .
- کلیه حقوق معنوی این اثر متعلق به دانشگاه صنعتی شاهرود می باشد و مقالات مستخرج بانام « دانشگاه صنعتی شاهرود » و یا « Shahrood University of Technology » به چاپ خواهد رسید .
- حقوق معنوی تمام افرادی که در به دست آمدن نتایج اصلی پایان نامه تأثیرگذار بوده اند در مقالات مستخرج از پایان نامه رعایت می گردد.
- در کلیه مراحل انجام این پایان نامه ، در مواردی که از موجود زنده (یا چینی جاهای آنها) استفاده شده است ضوابط و اصول اخلاقی رعایت شده است .
- در کلیه مراحل انجام این پایان نامه، در مواردی که به حوزه اطلاعات شخصی افراد دسترسی یافته یا استفاده شده است اصل رازداری ، ضوابط و اصول اخلاق انسانی رعایت شده است .

تاریخ

امضای دانشجو

مالکیت نتایج و حق نشر

- کلیه حقوق معنوی این اثر و محصولات آن (مقالات مستخرج ، کتاب ، برنامه های رایانه ای ، نرم افزارها و تجهیزات ساخته شده است) متعلق به دانشگاه صنعتی شاهرود می باشد. این مطلب باید به نحو مقتضی در تولیدات علمی مربوطه ذکر شود .
- استفاده از اطلاعات و نتایج موجود در پایان نامه بدون ذکر مرجع مجاز نمی باشد.

*متن این صفحه نیز باید در ابتدای نسخه های چاپ شده پایان نامه وجود داشته باشد .

چکیده

به‌طور کلی میزان گوگرد موجود در نفت خام، بسته به نوع منبع آن بین ۰/۰ تا بعضاً حدود ۷/۸۹ درصد وزنی متغیر است. گوگرد به دو صورت معدنی و آلی در نفت خام وجود دارد. علاوه بر گوگرد عنصری بیش از ۲۰۰ ترکیب آلی گوگرددار در نفت شناسایی شده، که سولفیدها و تیول‌ها یا مرکاپتان‌ها و تیوفن‌ها را شامل می‌شود، تیوفن‌ها و دی بنزوتیوفن‌ها مولکول‌های بسیار پیچیده‌تری دارند. این ترکیبات به عنوان آلوده‌کننده‌های پایدار شناخته می‌شوند علاوه بر این، این ترکیبات سرطان‌زا هستند. روش‌های مختلفی مانند کروماتوگرافی گازی برای جداسازی و شناسایی این ترکیبات پیشنهاد شده است که اندیس بازداری را با استانداردها مقایسه کند. بنابراین اندیس بازداری یک حل‌شونده، پارامتر مهم می‌باشد که می‌تواند برای شناسایی ترکیبات هتروسیکل آروماتیک سولفوردار به کار رود. اما از طرف دیگر این روش‌ها پرهزینه و وقت‌گیر هستند. اما از این رو، برای غلبه بر این مشکلات، می‌توان از روش‌های کمومتریکس برای پیش‌بینی اندیس بازداری این ترکیبات استفاده نمود. با استفاده از این روش‌های کمومتریکس می‌توان اندیس بازداری ترکیبات PASHs را پیش‌بینی نمود که با استفاده از روش‌های جنگل‌های تصادفی (RF) و مارس (MARS) و شبکه عصبی مصنوعی SR-ANN و RF-ANN اقدام به پیش‌بینی اندیس بازداری این ترکیبات شده است. نتایج حاصله نشان می‌دهد که استفاده از روش آماری ذکر شده توانایی خوبی در پیش‌بینی اندیس بازداری تجربی این ترکیبات دارند که R^2 حاصل از روش RF و MARS و SR-ANN و RF-ANN به ترتیب ۰/۹۴۶ و ۰/۹۸۴ و ۰/۹۵۸ و ۰/۹۳۲ می‌باشد.

کلمات کلیدی: کمومتریکس، جنگل تصادفی، MARS، شبکه عصبی مصنوعی اندیس بازداری، ترکیبات هتروسیکل آروماتیک گوگرددار

فهرست مطالب

چکیده ز

فصل اول: مقدمه

۱-۱- اندیس بازداري ۲

۲-۱- معرفي تركيبات گوگرددار و شناسايي و جداسازي آنها ۲

۳-۱- معرفي تركيبات سولفور آروماتيك چند حلقه‌اي ۵

۴-۱- ضرورت انجام تحقيق ۶

۵-۱- مروري بر كارهاي انجام شده ۷

فصل دوم: مطالعات كمومتركس و کاربرد آن در QSPR

۱-۲- كمومتركس ۱۴

۲-۲- مراحل مطالعه ارتباط كمی ساختار - خاصيت (QSPR) ۱۵

۳-۲- ارتباط كمی ساختار- خاصيت ۱۶

۱-۳-۲- فراهم کردن سری داده‌ها ۱۷

۲-۳-۲- بهینه‌سازی ساختار هندسی تركيبات ۱۷

۳-۳-۲- محاسبه توصيف كننده ها ۱۸

۴-۳-۲- انتخاب توصيف كننده های مهم ۲۰

۱-۴-۳-۲- رگرسيون گام به گام ۲۰

۵-۳-۲- ساخت مدل ۲۱

۴-۲- مقدمه‌اي بر شبکه عصبی مصنوعی ۲۱

۱-۴-۲- عملکرد نرون مصنوعی ۲۲

۲-۴-۲- تابع انتقال ۲۴

۳-۴-۲- ساختارهای شبکه ۲۵

۲۶	۴-۴-۲ آموزش شبکه‌های پیش‌خور با تکنیک پس انتشار
۲۷	۵-۲ جنگل‌های تصادفی
۲۹	۱-۵-۲ روش درخت رگرسیونی (تصمیم)
۳۰	۲-۵-۲ الگوریتم تشکیل درخت رگرسیونی
۳۲	۳-۵-۲ اندازه درخت و هرس کردن
۳۳	۴-۵-۲ الگوریتم جنگل تصادفی
۳۶	۵-۵-۲ تعیین اهمیت متغیرها در روش جنگل‌های تصادفی
۳۷	۶-۵-۲ مزیت روش جنگل تصادفی
۳۷	۶-۲ رگرسیون اسپلاین تطبیقی چند متغیره (MARS)
۳۹	۱-۶-۲ MARS یک متغیره
۴۲	۲-۶-۲ MARS دو متغیره
۴۲	۳-۶-۲ گام‌های ساخت MARS
۴۳	۴-۶-۲ مزیت‌های MARS
۴۴	۷-۲ ارزیابی قدرت پیش‌بینی مدل
۴۴	۱-۷-۲ استفاده از پارامترهای آماری
۴۷	۲-۷-۲ استفاده از نمودار خطای باقی مانده
۴۷	۳-۷-۲ استفاده از روش ارزیابی متقاطع
۴۸	۴-۷-۲ آزمون Y-تصادفی

فصل سوم: مطالعه ارتباط کمی ساختار-خاصیت بر روی

۴۹	تعدادی از ترکیبات آروماتیک سولفور چند حلقه ای با استفاده از ANN و RF و MARS
۵۰	۱-۳ مدل سازی جهت پیش بینی اندیس بازداري ترکیبات هتروسیکل چند حلقه ای سولفور
۵۰	۲-۳ نرم افزارهای مورد استفاده
۵۰	۱-۲-۳ نرم افزار Hyperchem
۵۱	۲-۲-۳ نرم افزار Dragon

۵۱	 نرم افزار SPPS ۳-۲-۳
۵۲	 نرم افزار MATLAB ۴-۲-۳
۵۳	 نرم افزار R ۵-۲-۳
۵۳	 سری داده ها ۳-۳
۷۱	 بهینه سازی ساختار هندسی مولکول ها و محاسبه توصیف کننده ها ۴-۳
۷۲	 انتخاب بهترین توصیف کننده ها به روش رگرسیون گام به گام ۱-۴-۳
۷۳	 مدل سازی و بهینه سازی پارامتر های مؤثر بر روش جنگل تصادفی ۵-۳
۷۵	 تعیین اهمیت نسبی توصیف کننده ها با روش جنگل تصادفی ۶-۳
۷۷	 ارزیابی مدل جنگل های تصادفی ۱-۵-۳
۷۷	 ارزیابی مدل با استفاده از داده های سری آزمون ۱-۱-۵-۳
۷۹	 ارزیابی مدل با استفاده از نمودار خطای باقی مانده ۲-۱-۵-۳
۷۹	 مدل سازی شبکه عصبی مصنوعی با توصیف کننده های انتخاب شده توسط روش رگرسیون گام به گام (SR-ANN) ۶-۳
۸۰	 ساختار شبکه عصبی بهینه شده ۱-۶-۳
۸۱	 ارزیابی مدل شبکه عصبی مصنوعی با توصیف کننده های SR ۲-۶-۳
۸۱	 ارزیابی مدل با استفاده از داده های سری ارزیابی ۱-۲-۶-۳
۸۲	 ارزیابی مدل با استفاده از داده های آزمون ۲-۲-۶-۳
۸۴	 ارزیابی مدل توسط نمودار خطای باقی مانده ۳-۲-۶-۳
۸۵	 ارزیابی مدل (SR-ANN) توسط روش ارزیابی متقاطع ۴-۲-۶-۳
۸۸	 مدل سازی شبکه عصبی مصنوعی با توصیف کننده های انتخاب شده توسط RF (RF-ANN) ۷-۳
۸۹	 ساختار شبکه عصبی بهینه شده ۱-۷-۳
۸۹	 ارزیابی مدل شبکه عصبی مصنوعی با توصیف کننده های RF ۲-۷-۳
۸۹	 ارزیابی مدل با استفاده از داده های سری ارزیابی ۱-۲-۷-۳

۹۱	 ۲-۲-۷-۳- ارزیابی مدل با استفاده از داده‌های سری آزمون
۹۳	 ۳-۲-۷-۳- ارزیابی مدل توسط نمودار خطای باقی مانده
۹۳	 ۴-۲-۷-۳- ارزیابی مدل (RF-ANN) توسط روش ارزیابی متقاطع
۹۶	 ۸-۳- مقایسه مدل‌های ارائه شده با استفاده از پارامترهای آماری
۹۷	 ۹-۳- ارزیابی مدل با استفاده از Y -تصادفی
۹۸	 ۱۰-۳- MARS
۱۰۱	 ۱۱-۳- نتیجه گیری مدل‌سازی با مجموع توصیف‌کننده‌های انتخاب‌شده توسط رگرسیون گام به گام
۱۰۲	 ۱۲-۳- بررسی ارتباط بین توصیف‌کننده‌های منتخب و خاصیت مورد نظر
۱۰۲	 ۱-۱۲-۳- توصیف‌کننده‌های 2D frequency finger prints
۱۰۲	 ۲-۱۲-۳- توصیف‌کننده‌های Topological
۱۰۲	 ۳-۱۲-۳- توصیف‌کننده‌های 3D MORSE
۱۰۳	 ۴-۱۲-۳- توصیف‌کننده‌های RDF
۱۰۳	 ۵-۱۲-۳- توصیف‌کننده‌های Charge
۱۰۳	 ۶-۱۲-۳- توصیف‌کننده‌های GATWAY
۱۰۴	 ۷-۱۲-۳- توصیف‌کننده‌های connectivity indices
۱۰۶	 ۱۳-۳- نتیجه گیری:
۱۰۷	 آینده نگری:
۱۰۸	 منابع:

فهرست جداول

جدول ۱-۱- میزان عناصر موجود در یک نمونه نفت خام	۳
جدول ۱-۲- انواع توصیف کننده های محاسبه شده توسط نرم افزار Dragon	۱۹
جدول ۱-۳- نام و اندیس بازداري سری داده ها	۵۴
جدول ۲-۳- توصیف کننده های انتخاب شده با روش رگرسیون گام به گام	۷۲
جدول ۳-۳- ماتریس همبستگی توصیف کننده های انتخاب شده توسط روش رگرسیون گام به گام	۷۳
جدول ۴-۳- کمترین مقادیر MSE همراه با Mtry و ntree متناظر با آن ها	۷۴
جدول ۵-۳- توصیف کننده های انتخاب شده توسط جنگل تصادفی	۷۶
جدول ۶-۳- ماتریکس همبستگی توصیف کننده ها	۷۷
جدول ۷-۳- نتایج حاصل از ارزیابی مدل جنگل تصادفی با استفاده از داده های آزمون	۷۸
جدول ۸-۳- توابع و پارامترهای بهینه شبکه عصبی مصنوعی	۸۰
جدول ۹-۳- مشخصات شبکه بهینه (SR-ANN)	۸۰
جدول ۱۰-۳- نتایج حاصل از مدل شبکه عصبی مصنوعی برای داده های سری ارزیابی (SR-ANN)	۸۱
جدول ۱۱-۳- نتایج حاصل از مدل شبکه عصبی مصنوعی برای داده های سری آزمون (SR-ANN)	۸۳
جدول ۱۲-۳- نتایج حاصل از ارزیابی مدل SR-ANN با استفاده از ارزیابی متقاطع	۸۵
جدول ۱۳-۳- توابع و پارامترهای بهینه شبکه عصبی مصنوعی	۸۹
جدول ۱۴-۳- مشخصات شبکه بهینه (RF-ANN)	۸۹
جدول ۱۵-۳- نتایج حاصل از مدل شبکه عصبی مصنوعی برای داده های سری ارزیابی (RF-ANN)	۹۰
جدول ۱۶-۳- نتایج حاصل از مدل شبکه عصبی مصنوعی برای داده های سری آزمون (RF-ANN)	۹۱
جدول ۱۷-۳- نتایج حاصل از ارزیابی مدل RF-ANN با استفاده از ارزیابی متقاطع	۹۴

جدول ۳-۱۸-مقادیر R^2 ، GCV، RSS، MSE برای داده‌های آزمون و آموزش براساس

مقادیر مختلف nk.....۹۹

جدول ۳-۱۹-محاسبه پارامترهای آماری برای داده‌های آزمون و ارزیابی.....۹۷

جدول ۳-۲۰-مقادیر R^2 برای داده‌های ارزیابی و آزمون بعد از ۱۰ بار تست Y-تصادفی(SR-

ANN).....۹۸

جدول ۳-۲۱-مقادیر R^2 برای داده‌های ارزیابی و آزمون بعد از ۱۰ بار تست Y-تصادفی(RF-

ANN).....۹۸

فهرست اشکال

- شکل ۱-۲- شمای کلی مراحل به کار گرفته شده در یک فرآیند QSPR ۱۷
- شکل ۲-۲- شبکه عصبی مصنوعی ۲۳
- شکل ۳-۲- توابع انتقال ۲۴
- شکل ۴-۲- شکل (الف) افراز قابل قبول شکل (ب) افراز غیرقابل قبول ۳۰
- شکل ۵-۲- نمودار درختی فضای افراز شده در افراز قابل قبول ۳۰
- شکل ۶-۲- نمودار پراکندگی داده های دومتغیره ۳۱
- شکل ۷-۲- سه افراز ممکن در راستای متغیر x_1 ۳۱
- شکل ۸-۲- چهار افراز ممکن در راستای متغیر x_2 ۳۲
- شکل ۹-۲- نمای کلی الگوریتم ۳۵
- شکل ۱۰-۲- نمودار توابع مبنای h_1 و h_2 ۳۸
- شکل ۱-۳- نمودار تغییرات خطای OOB در مقادیر متفاوت Mtry و Ntree ۷۴
- شکل ۲-۳- نمودار اهمیت نسبی ۱۹۲ توصیف کننده ۷۵
- شکل ۳-۳- نمودار میزان اهمیت نسبی توصیف کننده ها ۷۶
- شکل ۴-۳- نمودار تغییرات مقادیر پیش بینی شده در مقابل مقادیر تجربی توسط RF برای داده های سری آزمون ۷۸
- شکل ۵-۳- نمودار خطای باقی مانده برحسب مقادیر تجربی برای داده های سری آزمون با مدل های جنگل تصادفی ۷۹
- شکل ۶-۳- نمودار مقادیر پیش بینی شده برحسب مقادیر تجربی برای داده های ارزیابی توسط روش SR-ANN ۸۲
- شکل ۷-۳- نمودار مقادیر پیش بینی شده اندیس بازداري برحسب مقادیر تجربی برای داده های آزمون توسط روش SR-ANN ۸۴
- شکل ۸-۳- نمودار خطای باقی مانده برحسب مقادیر تجربی برای داده های آزمون (SR-ANN) ... ۸۴

- شکل ۳-۹- نمودار خطای باقی مانده برحسب مقادیر تجربی برای داده‌های سری ارزیابی (SR-ANN) ۸۴
- شکل ۳-۱۰- نمودار مقادیر پیش بینی شده بر حسب مقادیر تجربی برای ارزیابی مدل SR-ANN به روش ارزیابی متقاطع برای کل داده ها ۸۸
- شکل ۳-۱۱- نمودار مقادیر باقی مانده بر حسب مقادیر تجربی برای ارزیابی مدل SR-ANN به روش ارزیابی متقاطع برای کل داده ها ۸۸
- شکل ۳-۱۲- نمودار مقادیر پیش‌بینی شده برحسب مقادیر تجربی برای داده های ارزیابی توسط روش RF-ANN ۹۱
- شکل ۳-۱۳- نمودار مقادیر پیش‌بینی شده اندیس بازداري برحسب مقادیر تجربی برای داده های آزمون توسط روش RF-ANN ۹۲
- شکل ۳-۱۴- نمودار خطای باقی مانده برحسب مقادیر تجربی برای داده‌های آزمون (RF-ANN) ۹۳
- شکل ۳-۱۵- نمودار خطای باقیمانده برحسب مقادیر تجربی برای داده‌های ارزیابی (RF-ANN) ۹۳
- شکل ۳-۱۶- نمودار مقادیر پیش بینی شده بر حسب مقادیر تجربی برای ارزیابی مدل RF-ANN به روش ارزیابی متقاطع برای کل داده ها ۹۶
- شکل ۳-۱۷- نمودار مقادیر باقیمانده بر حسب مقادیر تجربی برای ارزیابی مدل SR-ANN به روش ارزیابی متقاطع برای کل داده ها ۹۶
- شکل ۳-۱۸- نمودار میانگین خطای برحسب توابع پایه (nk) ۹۹
- شکل ۳-۱۹- نمودار مقادیر پاسخ در مقابل برازش داده شده توسط مارس برای سری آزمون ۱۰۰
- شکل ۳-۲۰- نمودار مقادیر پاسخ در مقابل برازش داده شده توسط مارس برای سری آموزش ۱۰۱

فصل اول

مقدمه

۱-۱-اندیس بازداری^۱

اندیس بازداری (RI) در سال ۱۹۵۸ برای اولین بار به عنوان یک پارامتر برای شناسایی ترکیبات شیمیایی با استفاده از کروماتوگرافی معرفی شد. تعیین این اندیس برای هر ترکیب، از کروماتوگرام مربوط به مخلوط آن جسم و دو پارافین نرمال که زمان بازداری ترکیب موردنظر بین زمان بازداری آن دو باشد، امکان پذیر است. در این مورد پارافین های نرمال به عنوان استاندارد در نظر گرفته می شوند. طبق قرارداد مقدار این اندیس برای این پارافین ها صد برابر تعداد کربن های آن ها در نظر گرفته می شوند که مقداری ثابت بوده و مستقل از نوع فاز ساکن، درجه حرارت و سایر شرایط می باشد. درحالی که برای سایر ترکیبات با تغییر شرایط ستون، مقدار این اندیس تغییر می کند. اندیس بازداری ماده مورد تجزیه با استفاده از معادله (۱-۱) محاسبه می گردد که در این رابطه t_{Ra} ، t_{Rn} به ترتیب عبارت اند از زمان های بازداری هیدروکربن سنگین تر بعد از مجهول، زمان بازداری ترکیب مجهول، زمان بازداری ترکیب سبک تر از مجهول و n تعداد کربن های نرمال می باشد [۱].

$$RI = n \times 100 + 100 \frac{\log t_{Ra} - \log t_{Rn}}{\log t_{Rn+1} - \log t_{Rn}} \quad (1-1)$$

۱-۲-معرفی ترکیبات گوگرددار و شناسایی و جداسازی آنها

پس از غذا سوخت مهم ترین منبع انرژی بشر است به طوری که ۸۵ درصد انرژی از سوخت های فسیلی، ۸ درصد از انرژی هسته ای و ۷ درصد از منابع دیگر که به طور کلی آب، برق و... است، تأمین می شود. زغال سنگ، نفت و گاز سه سوخت اصلی هستند که نفت ۴۰ درصد، زغال سنگ حدود ۲۴ درصد و گاز طبیعی حدود ۱۲ درصد از مصرف کل جهان را در بر می گیرد [۲]. بر این اساس، نفت مهم ترین این سوخت هاست.

¹ Retention index

بسیاری از ترکیبات در نفت خام شامل عناصری غیر از هیدروژن و کربن نیز می‌باشند که گاهی درون ساختمان حلقه و یا در گروه‌های استخلافی متصل شده به ساختمان هیدروکربن مشاهده می‌شوند. ترکیبات آلی گوگرددار در نفت خام مولکول‌های هتروسیکل هستند، در مقابل عمده مولکول‌های اکسیژن‌دار به صورت کربوکسیلیک اسیدها هستند. اسیدهای آلیفاتیک سیرشده (اسیدهای نفتیک)، تیوفن آروماتیک‌ها و به مقدار ناچیزی از فنول‌ها و فوران‌ها نیز ممکن است موجود باشد. مولکول‌های بسیار بزرگ و سنگین رزین‌ها و آسفالت‌ها که مخلوطی از ساختمان آروماتیک و هتروسیکلیک هستند نیز وجود دارند.

به‌طور کلی پس از کربن و هیدروژن گوگرد سومین عنصر فراوان در این نوع سوخت‌هاست [۳].

جدول (۱-۱) - میزان عناصر موجود در یک نمونه نفت خام

درصد	عنصر تشکیل‌دهنده
۸۳-۸۷	کربن
۰-۳	گوگرد
۰-۰/۵	اکسیژن
۱۱-۱۴	هیدروژن
۰/۱	نیتروژن
۰-۰/۲	عناصر فلزی

استفاده از سوخت‌های فسیلی به‌عنوان منبع انرژی موجب آلودگی محیطی با ترکیباتی با منشأ غیر زیستی شده است. زیرا این سوخت‌های فسیلی محتوای گوگرد هستند که سوختن مستقیم آن‌ها موجب آزادسازی مقدار زیادی از اکسید گوگرد به اتمسفر می‌شود [۴]. اکسید گوگرد تولیدی مسئول

کیفیت بدهوا، باران اسیدی و غیرفعال شدن کاتالیست‌های شیمیایی و کاهش لایه اوزون است علاوه بر این اکسید گوگرد موجب به خطر افتادن سلامتی انسان می‌شود[۵]. به گونه‌ای که مشکلاتی از قبیل سرطان دستگاه تنفسی، بیماری‌های قلبی را به همراه دارد[۶]. آنالیز این ترکیبات در ماتریکس‌های مختلف، کار دشواری است، زیرا اکثر آن‌ها با غلظت پائینی در بافت‌های پیچیده وجود دارند که سایر گونه‌ها با غلظتی چند برابر، باعث ایجاد مزاحمت در اندازه‌گیری این ترکیبات می‌شود[۹]. همچنین، آنالیز این ترکیبات در هوا پیچیده است زیرا اکسند‌های اتمسفری مقدار قابل اندازه‌گیری این ترکیبات را کاهش می‌دهند. چندین روش از جمله کروماتوگرافی گازی، اسپکتروفتومتری، پلاروگرافی، کولومتری، پتانسیومتری، و..... برای اندازه‌گیری این ترکیبات وجود دارد که از بین این روش‌ها کروماتوگرافی گازی با دتکتور انتخابی گوگرد، به سایر روش‌ها ترجیح داده می‌شود زیرا انتخاب پذیری و حد تشخیص بالایی دارد. اما استفاده از این تکنیک‌ها با مشکلات زیادی همراه است که در زیر به مهمترین آن‌ها اشاره شده است:

۱- استانداردهای بسیاری از این ترکیبات در دسترس نیست.

۲- امکان جذب و از بین رفتن گونه‌ها در سیستم کروماتوگرافی وجود دارد.

۳- در بافت پیچیده، به یک دتکتور حساس و انتخاب پذیر نیاز است.

به دلیل وجود این مشکلات، یک روش تئوری جایگزین خوبی برای روش‌های تجربی است. مطالعه ارتباط کمی ساختار-خاصیت (QSPR)، یک روش توانمند برای پیش بینی اندیس بازداري است.

۱-۳- معرفی ترکیبات هتروسیکل آروماتیک سولفوردار چند حلقه‌ای^۱

تیوفن^۲ با فرمول شیمیایی C_4H_4S یک ترکیب شیمیایی است که شکل ظاهری آن، مایع بی‌رنگ است. تیوفن یک ترکیب هتروسیکل آروماتیک است که شامل چهار اتم کربن و یک اتم گوگرد می‌باشد. آنالوگ‌های تیوفن شامل فوران و پیرول می‌باشند که در آن‌ها اتم S جایگزین اتم‌های O و NH شده است. تیوفن به وسیله ویکتور میر^۳ در سال ۱۸۸۳ به‌عنوان یک جز باقی‌مانده در بنزن کشف شد. بعضی از انواع مهم تیوفن‌ها شامل بنزو تیوفن و دی بنزو تیوفن می‌باشد که به ترتیب دارای یک و دو گروه بنزنی هستند. تیوفن در دمای محیط مایع بی‌رنگ است و بویی مشابه بنزن دارد. تیوفن به‌عنوان ترکیبی آروماتیک شناخته شده است ولی محاسبات تئوری نشان داده که میزان آروماتیک بودن آن از بنزن کمتر است. سهم جفت الکترون گوگرد در اوربیتال مشهود می‌باشد و به خاطر همین خواص آروماتیک تیوفن است که این ترکیب خواصی متفاوت نسبت به تیواتر از خود نشان می‌دهد. به‌عنوان مثال این ترکیب توانایی آلکیل دار شدن توسط متیل یدید را ندارد. همچنین اتم گوگرد در این ترکیب واکنش ناپذیر است. واکنش‌پذیری بالای تیوفن در برابر سولفوناسیون اساس جداسازی بنزن از تیوفن را تشکیل می‌دهد. (از آنجایی که تفاوت نقطه جوش بنزن و تیوفن در دمای محیط حدود ۴ درجه سانتی‌گراد می‌باشد، جداسازی آن‌ها از طریق تقطیر مشکل می‌باشد). بنابراین اضافه کردن اسید سولفوریک به مخلوط بنزن و تیوفن در فرآیند سولفوناسیون تیوفن باعث تشکیل ترکیب محلول در آب تیوفن سولفونیک اسید می‌شود. در بعضی موارد حلقه بنزنی توسط تیوفن جایگزین می‌شود بی‌آنکه در خواص ترکیب تغییر قابل ملاحظه‌ای روی دهد. گوگردزدایی از این ترکیبات توسط رانی نیکل^۴ باعث ایجاد مشتقاتی از بوتان با استخلاف در موقعیت‌های ۱ و ۴ می‌شود. پلیمرهای حاصل از پیوند شدن تیوفن در موقعیت‌های ۱ و ۵ پلی تیوفن نامیده می‌شوند. پلی تیوفن‌ها در اثر اکسایش

¹ Polycyclic aromatic sulfur heterocycles

² Thiophen

³ Victor Meyer

⁴ Raney Nickel

جزئی دارای خاصیت رسانایی می گردند. با در نظر گرفتن پایداری بالای تیوفن ها، این ترکیبات از بسیاری از ترکیبات حاوی گوگرد و وهیدروکربن (به خصوص هیدروکربن های غیراشباع مانند استیلن) به وجود می آیند. درروش سنتز کلاسیک این ترکیبات از واکنش ۱،۴-دی کتون با معرف های سولفور کننده مانند P_4S_{10} تهیه می شوند. تیوفن های خاص می توانند از طریق واکنش تراکمی استرها در حضور گوگرد عنصری تهیه شوند. تیوفن و مشتقات آن معمولاً به همراه مواد نفتی در غلظت های یک تا سه درصد یافت می شوند. انواع مایعات تیوفنی در صنعت معمولاً توسط فرآیند هیدرو دسولفوریزاسیون جداسازی می گردد. در این فرآیند مایع یا گاز از روی بستر کاتالیزوری شامل مولیبدونیم دی سولفید و تحت فشار هیدروژن عبور داده شده و ترکیبات گوگردی ازجمله تیوفن ها جداسازی می شوند. بسیاری از این ترکیبات اثرات مخربی روی ساختار و اعمال ارگانسیم های زنده دارند. آن ها می توانند از طریق تنفس وارد بدن شده و صدمات زیادی به شش ها وارد نمایند آب رودخانه ها و دریاچه ها موجود در نزدیکی مکان هایی که زباله و زوائد کارخانه ها وجود دارند نیز آلوده به این ترکیبات هستند. خوردن آب و غذای آلوده باعث وارد شدن آن ها به خون شده که برای کبد و کلیه مضر هستند و تماس با خاک آلوده، می تواند این ترکیبات را جذب بدن نموده و باعث ایجاد سرطان های پوستی گردد [۱۰].

۱-۴- ضرورت انجام تحقیق

در سال های اخیر پیشرفت های زیادی در جهت آسان تر شدن، دقیق تر شدن و سازگار بودن بیشتر روش های آزمایشگاهی با محیط زیست صورت گرفته است. با این حال این روش ها هنوز محدودیت هایی مانند عدم توانایی در اندازه گیری همه ترکیبات، هزینه بالا و در دسترس نبودن امکانات مورد نظر دارند. استفاده از روش های نظری می تواند در رفع این محدودیت ها مفید واقع شود. به کارگیری این روش ها می تواند علاوه بر پیش بینی خواص مورد نظر، به هدفمند و جهت دارتر ساختن آزمایش

های تجربی و همینطور توضیح پارامترهای موثر بر نتایج این آزمایش‌ها و مکانیسم‌های درگیر کمک کند.

در این تحقیق سعی شده با به‌کارگیری روش‌های QSPR مدلی برای پیش‌بینی اندیس بازداری ترکیبات هتروسیکل آروماتیک سولفوردار چند حلقه‌ای (PASHs) طراحی شود. هنگامی که رابطه معتبری بدست آمد، امکان پیش‌بینی این خواص برای ترکیباتی با ساختار مشابه وجود خواهد داشت. در تحقیق حاضر تلاش شده تا مقایسه‌ای بین پاسخ سری داده‌ها به مدل خطی و غیر خطی نیز صورت پذیرد.

۱-۵- مروری بر کارهای انجام‌شده

در سال ۲۰۰۳ میلر^۱ و برانو^۲، اندیس بازداری کواتس^۳ برای ترکیبات گوگردار را محاسبه نموده و آن‌ها را با مقادیر به‌دست‌آمده از روش تجربی، مقایسه کردند. نتایج نشان داد مقادیر پیش‌بینی شده با مقادیر بدست آمده با پیش‌بینی لگاریتمی مرسوم موافق بوده‌اند. [۱۱].

در سال ۲۰۰۴ دیو^۴ و همکارانش، مدل رگرسیون خطی چندگانه^۵ را برای پیش‌بینی زمان و اندیس بازداری ترکیبات گوگردار به کاربردند. نتایج به‌دست‌آمده توافق خوبی با مقادیر تجربی زمان و اندیس بازداری نشان می‌دهد [۱۲].

در سال ۲۰۰۵ کن^۶ و گروهش مطالعات کمومتریکس^۷ بروی ۸۰ ترکیب PASHs انجام داده‌اند و از روش شبکه عصبی مصنوعی^۸ برای بهینه‌سازی تعداد توصیف‌کننده ورودی^۹ استفاده کردند،

¹ miller

² Bruno

³ kovats

⁴ Du

⁵ Multiple linear regression (MLR)

⁶ Can

⁷ Chemometrics

⁸ Artificial neural network

⁹ Input descriptors

همچنین میزان مشارکت توصیف کننده‌های وارد شده در مدل را نیز مورد تجزیه و تحلیل قرار دادند و توسط روش متقاطع^۱ مدل را مورد ارزیابی قرار دادند [۱۳]. مدل مورد ارزیابی مدل مناسبی برای پیش بینی اندیس بازداری ترکیبات گزارش شد.

در سال ۲۰۰۴ توسط گرکانی نژاد و همکارانش برای مجموعه‌ای از ۸۴۶ ترکیبات آلی مربوط به شیمی تجزیه با استفاده از توصیف کننده‌های استخراج شده از دراگون مدل سازی صورت گرفت. مدل سازی به روش‌های حداقل مربعات جزئی (PLS)^۲ و مدل رگرسیون خطی چند گانه (MLR) انجام شد. انحراف معیار خطاهای پیش بینی شده (SEP) در چهار آزمایش با مجموعه های آموزشی و پیش بینی شده توزیع شده برآورد شد. SEP برای بهترین مدل ها حدود ۸۰ واحد شاخص نگهداری است [۱۴].

لیائو^۳ و همکارانش در سال ۲۰۰۷، روش MLR را برای پیش بینی زمان بازداری ۶۹ ترکیب روغن ضروری موجود در گونه خاصی از گل ها به کار بردند [۱۵]. ضریب همبستگی ۰/۹۴ در این روش برای سری داده ها گزارش شده است.

در سال ۲۰۰۵، جینگ هو و همکارانش اندیس بازداری را برای ترکیبات آروماتیک چند حلقه‌ای نیتروژن با استفاده از ارتباط کمی ساختار خاصیت (QSPR)^۴ پیش بینی کردند [۱۶]. برای نشان دادن تاثیر توصیف کننده‌های ملکولی بر اندیس بازداری سه مدل خطی چند متغیره از سه گروه توصیف کننده ملکولی ساخته شد. علاوه بر این، هر یک از توصیف کننده های مولکولی در این مدل

¹ Leave One Out Cross Validation

² Partial least square (PLS)

³ Liao

⁴ Quantitative structure property relationship

ها به خوبی درک رابطه بین ساختارهای مولکولی و اندیس بازداري آنها مورد بحث قرار گرفت. مدل های پیشنهادی به نتایج زیر رسیدند: با ضرایب تعیین R^2 برای مدل ها با یک، دو و سه توصیف کننده مولکولی به ترتیب ۰,۹۵۷۱، ۰,۹۷۷۶ و ۰,۹۴۸۶ بود.

در سال ۲۰۰۸ ریاحی و همکارانش روش GA-MLR را برای مدل سازی شاخص بازداري ۴۳ ترکیب آلی به کار گرفتند. R^2 گزارش شده برای سری آموزش متشکل از ۳۵ ترکیب، ۰/۹۷۷ و برای بقیه ترکیبات که به عنوان سری تست در نظر گرفته شده اند، ۰/۹۷۸ بود [۱۷].

ریاحی و همکارانش تحقیقاتی در سال ۲۰۰۸، روش ماشین بردار پشتیبان^۱ را برای پیش بینی شاخص بازداري صد ترکیب آلی، با روش های خطی مورد مقایسه قرار دادند و آن را به عنوان روش بهتری نسبت به MLR گزارش کردند [۱۸].

در سال ۲۰۰۸ هیوینگ^۲ و همکارانش اندیس بازداري را برای ۱۱۴ ترکیب PASHs محاسبه نموده و آن ها را با مقادیر به دست آمده از روش تجربی مقایسه کردند و نتایج نشان داد روش تئوری در پیش بینی اندیس بازداري این ترکیبات موفق عمل کرده است [۱۹].

در سال ۲۰۰۵ ژنگ^۳ و همکارانش اندیس بازداري را برای ۲۰۷ ساختار متیل آلکان پیش بینی و ۱۲۷ توصیف کننده توپولوژیکی محاسبه کرده اند. روش GAPLS، که یک روش انتخاب متغیر ترکیبی با الگوریتم های ژنتیک (GA)، حداقل مقادیر جزئی و جزئی (PLS) است، هفت توصیف کننده توپولوژیکی از ۱۲۷ توصیف کننده توپولوژیکی با استفاده از روش GAPLS برای ساخت مدل QSPR با کیفیت رگرسیون بالا انتخاب شده است، ضریب تعیین (R^2) از ۰/۹۹۹۹، انحراف معیار ۲/۸۸. خطای مدل مشابه خطای آزمایشی است [۲۰]. اعتبارسنجی مدل با استفاده از تکنیک

¹ Support vector machine(SVM)

² Hui-Ying

³ Xiang, Zheng

اعتبارسنجی مجدد بررسی می شود. نتیجه اعتبارسنجی نشان می دهد که مدل ساخته شده با کیفیت پیش بینی بالا قابل اعتماد و پایدار است.

در سال ۲۰۰۴ هونگبین^۱ و همکارانش ترکیبات گوگرد شامل مرکاپتان‌ها، سولفیدها و تیوفن‌ها با استفاده از کروماتوگرافی گازی با تشخیص انتشار اتمی شناسایی شدند. زمان و اندیس بازداری آن‌ها با توصیف‌کننده‌های مولکولی تولید شده از ساختار مولکولی آن‌ها مرتبط بود. بهترین مدل رگرسیون خطی چند گانه پارامتر توانایی پیش بینی خوبی را نشان داد. توصیف‌کننده‌هایی که در این مدل‌ها دخیل هستند، خواص هندسی، توپولوژیکی و الکترونیکی مولکول‌ها را نشان می دهند که مربوط به تعامل بین حلال و فاز ثابت است. برای ۳۴ ترکیبات گوگرد (بنزوتیوفن و دیابنسوتیوفن)، دو مدل دیگر برای زمان بازداری و اندیس بازداری با خطاهای استاندارد به ترتیب ۰/۶۱ و ۱/۶۳ به دست آمد. چنین مدل‌هایی برای زمان بازداری و یا اندیس بازداری می‌تواند برای شناسایی قله‌های کروماتوگرافی ناشناخته با تطابق زمان بازداری آن‌ها یا اندیس بازداری با ترکیبات گوگرد از ساختار شناخته شده مولکولی زمانی که استانداردهای شیمیایی مربوطه در دسترس نیستند [۲۱].

در سال ۲۰۱۶ پاشا زانوسی و همکارانش [۲۲] با استفاده از سه روش استخراج به نام‌های ریز استخراج با فاز جامد در فضای فوقانی^۲ استخراج مایکروویو بدون حلال^۳ و تقطیر با آب^۴ مقدار اندیس بازداری را برای ترکیب روغن‌های فرار پیش بینی کردند و مدل خود را توسط QSPR ارزیابی کردند و از روش رگرسیون خطی چند گانه به بالاترین ضریب همبستگی ($R^2 = 0.952$) و $F=104.4$ ^۵ و پایین ترین خطای استاندارد^۶ دست پیدا کردند.

¹ Du, Hongbin

² headspace solid-phase microextraction (HS-SPME)

³ solvent-free microwave extraction (SFME)

⁴ Hydrodistillation (HD)

⁵ Statistic Term

⁶ Standard error(SE)

در سال ۲۰۱۶ گودرزی و همکاران از روش‌های جنگل تصادفی^۱ و شبکه عصبی مصنوعی و رگرسیون خطی چندگانه برای پیش‌بینی اندیس بازداري تعدادی از ترکیبات آلی گوگرددار استفاده کردند و نتایج نشان داد مدل جنگل تصادفی از مدل‌های دیگر برتر و ارجح است به دلیل R^2 بالا با مقدار ۰/۹۹۵ و خطای پیش‌بینی^۲ پائین نسبت به سایر روش‌ها با مقدار ۱/۱۴۶ و F با مقدار ۱۱۹۷ باعث شده این بهتر عمل کند [۲۳].

در سال ۲۰۱۴ گودرزی و دیگر همکارانش بر اساس روش جدید جنگل‌های تصادفی (RF) برای پیش‌بینی شاخص‌های نگهداری (RIs) برخی از ترکیبات هیدروکربن آروماتیک چند حلقه‌ای (PAH)، تحقیق رابطه QSRR کمی انجام دادند. RIs این ترکیبات با استفاده از توصیف‌های نظری تولید شده از ساختار مولکولی آنها محاسبه شد. اثرات پارامترهای مهم موثر بر توانایی قدرت پیش‌بینی RF مانند تعداد درختان (nt) و تعداد متغیرهای تصادفی انتخاب شده برای تقسیم هر گره (m) مورد بررسی قرار گرفت. بهینه‌سازی این پارامترها نشان داد که در نقطه $m = 70$ ، $nt = 460$ ، روش RF می‌تواند بهترین نتایج را بدست آورد. همچنین عملکرد مدل RF با شبکه عصبی مصنوعی (ANN) و تکنیک‌های رگرسیون خطی چندگانه (MLR) مقایسه شد. نتایج بدست آمده نشان دهنده برتری نسبی روش RF در سطح MLR و ANN است [۵۵].

¹ Random forests

² Relative Error Prediction

فصل دوم

مطالعات کمومتریکس و کاربرد آن در QSPR

۲-۱- کمومتریکس^۱

عبارت کمومتریکس توسط اسوانت ولد^۲ و بروس کووالسکی^۳ در اوایل ۱۹۷۰ تعریف شد. قبلاً عباراتی مانند بیومتریکس^۴ و اکونومتریکس^۵ نیز در زمینه‌های علوم زیستی و همین‌طور اقتصاد معرفی شده‌اند. پس از آن بود که انجمن بین‌المللی کمومتریکس تاسیس شد. از آن زمان تا کنون پیشرفت‌های زیادی در این زمان به‌وجود آمده و اکنون کمومتریکس در شاخه‌های مختلف شیمی، به ویژه شیمی تجزیه مورد استفاده قرار می‌گیرد. برخی از کاربردهای آن عبارتند از: کنترل فرآیندها، تجزیه و تحلیل و شناخت الگوها، پردازش علائم، بهینه کردن شرایط.

یکی از تعاریفی که از کمومتریکس وجود دارد، "سازماندهی شیمیایی با استفاده از روش‌های ریاضی و آمار برای (الف)-طراحی و انتخاب بهترین فرآیندهای محاسباتی و آزمایشگاهی و (ب)-فراهم نمودن بیشترین اطلاعات شیمیایی با بررسی داده‌های در دسترس" می‌باشد [۲۴].

هاوری^۶ و هیرش^۷ در اوایل سال ۱۹۸۰ پیشرفت‌های کمومتریکس را در مراحل مختلف رده بندی کردند. اولین مرحله قبل از ۱۹۷۰ است که در آن تعدادی از روش‌های ریاضی توسعه یافتند و در زمینه‌های مختلف ریاضی، علم رفتار و مهندسی به‌کار گرفته شدند. در این محدوده زمانی، شیمیدانان خود را به روش‌های آنالیز داده‌ها، یعنی محاسبات پارامترهای آماری مانند میانگین، انحراف استاندارد و حدود اطمینان محدود کرده بودند. هاوری و هیرش تحقیقات خود را به‌طور ویژه روی همبستگی تعدادی از داده‌های شیمیایی و ویژگی‌های ملکولی مربوط، متمرکز کردند. این اقدامات اولیه اساس یکی از شاخه‌های مهم کمومتریکس یعنی ارتباط کمی ساختار-فعالیت (QSAR)^۸ را فراهم کرد. دومین مرحله تکامل کمومتریکس در سال ۱۹۷۰ واقع شد، هنگامی که کمومتریکس توجه

¹ chemometrics

² Svant Wold

³ Bruce R. Kowalaski

⁴ biometrics

⁵ Econometrics

⁶ Howery

⁷ Hirsch

⁸ Quantitative Structure- Activity Relationship(QSAR)

شیمیدانان، به خصوص شیمیدانان تجزیه را که به این ترتیب روش‌های بیشتری برای آنالیز داده‌های دردسترس و همچنین روش‌های نوینی برای بدست آوردن اطلاعات پیش‌روی خود می‌دیدند، به خود جلب کرد. مرحله بعدی تکامل کمومتریکس که ابتدا توسط هاوری و هیرش و سپس توسط براون^۱ پیش‌بینی شده بود از اوایل ۱۹۸۰ آغاز شد. هم‌اکنون کمومتریکس ابزار قوی و کارآمد برای شیمیدانان به‌شمار رفته و به عنوان یک زیر شاخه رشته شیمی در بسیاری از دانشگاه‌های آمریکا و اروپا و بعضی از دانشگاه‌های چین و سایر کشورها تدریس می‌شود. پیشرفت‌ها در علوم کامپیوتر فناوری اطلاعات آمار و ریاضی کاربردی عناصر جدیدی را در کمومتریکس معرفی کرده است [۲۵،۲۶]. اولین مقاله منتشر شده در زمینه کمومتریکس در ایران به سال ۱۹۹۸ برمی‌گردد [۲۷]. با رشد سریع کمومتریکس در جهان، استفاده شیمی تجزیه‌دانان ایرانی از کمومتریکس در تحقیقاتشان افزایش یافت. تا سال ۲۰۰۵ حدود ۲۰۰ مقاله علمی توسط محققین کمومتریکس در ایران منتشر و بیش از ۸ گروه تحقیقاتی فعال در این زمینه مشغول فعالیت می‌باشند. انجمن کمومتریکس ایران نیز به عنوان یکی از شاخه‌های فعال انجمن شیمی ایران در سال ۲۰۰۰ تأسیس شد و همچنین سمینارهای داخلی و بین‌المللی در زمینه کمومتریکس در ایران برگزار می‌گردد. زمینه‌هایی که تاکنون بیشتر در ایران مورد توجه قرار گرفته‌اند مطالعات (QSPR^۲/QSAR)، کالبراسیون چند متغیره^۳ و هوش مصنوعی^۴ هستند [۲۸].

۲-۲- مراحل مطالعه ارتباط کمی ساختار – خاصیت (QSPR)

در مطالعات QSPR برای بیان ارتباط ساختار ملکول با خاصیت مورد نظر باید مراحل زیر به ترتیب طی شود.

۱- فراهم کردن سری داده‌ها

^۱ Brown

^۲ Quantitative Structure- Property Relationship(QSPR)

^۳ Multivariate Calibration(MVC)

^۴ Artificial Intelligence(AI)

۲-رسم و بهینه کردن ساختار هندسی ترکیبات

۳-محاسبه توصیف کننده‌ها

۴-انتخاب بهترین توصیف کننده‌ها

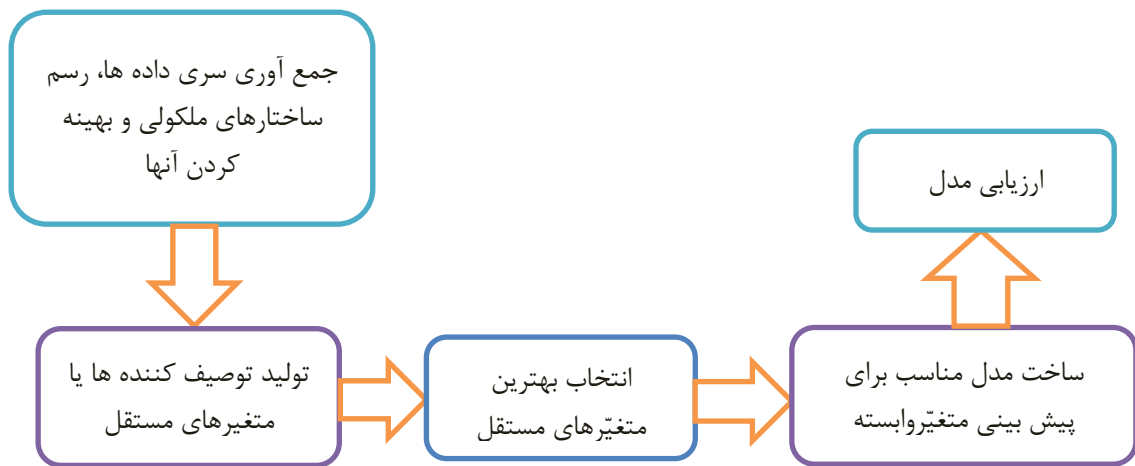
۵-مدل سازی

۶-ارزیابی قدرت پیش‌بینی مدل [۲۹].

۲-۳-ارتباط کمی ساختار-خاصیت^۱

QSPR در لغت به معنی برقراری ارتباط کمی بین ساختار و خواص مولکول می‌باشد. درواقع با استفاده از آنچه قبلاً به صورت تجربی انجام شده، این ارتباط برقرار می‌شود و سپس از آن برای پیش‌بینی فعالیت یا خواص ترکیبات جدید استفاده می‌گردد. زمانی که نمونه‌های استاندارد از لحاظ فعالیت در دسترس نباشند، آزمایش‌ها وقت‌گیر، پیچیده، پرهزینه و یا با خطر همراه باشند. روش QSPR روش مناسبی برای حل مشکل خواهد بود. درروش QSPR فرض بر آن است که بین خواص فیزیکی-شیمیایی و ساختار مولکولی ترکیبات، یک رابطه خطی یا غیرخطی وجود دارد و در همین راستا سعی بر آن دارد که با برقراری روابط ساده ریاضی، فعالیت/خاصیت مجموعه‌ای از ترکیبات را پیش‌بینی نماید. بدین منظور، ابتدا بایستی توصیف‌کننده‌های مولکولی مناسب که بتواند معرف ترکیبات مختلف باشند، انتخاب شوند و سپس با استفاده از الگوریتم‌هایی نظیر رگرسیون خطی چندگانه و یا شبکه‌های عصبی مصنوعی، مدل QSPR مربوطه ساخته شود. مهم‌ترین مسئله در توسعه مدل‌های QSPR انتخاب توصیف‌کننده‌هایی است که بستگی بسیار زیادی به خاصیت/فعالیت موردنظر داشته باشد و معمولاً این انتخاب باید از میان تعداد زیادی توصیف‌کننده‌های مولکولی صورت پذیرد.

¹ Quantitative structure-property relationship



شکل ۱-۲- شمای کلی مراحل به کار گرفته شده در یک فرآیند QSPR

۲-۳-۱- فراهم کردن سری داده‌ها

برای دستیابی به مدل مناسب، اولین مرحله انتخاب یک دسته از مولکول‌هاست که مقادیر تجربی یک خصوصیت فیزیکی یا شیمیایی آن‌ها مانند زمان بازداری، نقطه جوش و... در دسترس باشد. در این خصوص لازم است شرایط اندازه‌گیری خاصیت موردنظر برای همه ترکیبات یکسان باشد. بدیهی است که هرچقدر اطلاعات تجربی قابل دسترس برای طراحی مدل بیشتر باشد بدون تردید مدل کارایی بهتری خواهد داشت. فراهم کردن سری داده‌های مناسب از طریق جستجو در مقالات صورت می‌گیرد.

۲-۳-۲- بهینه‌سازی ساختار هندسی ترکیبات

کمومتریکس از روش‌های مفیدی که قادر به استخراج اطلاعات موجود در داده هستند، استفاده می‌کند. به این منظور از شیمی محاسباتی استفاده می‌شود. شیمی محاسباتی ساختارهای ملکولی را به صورت پارامترهای عددی معرفی و رفتار آن‌ها را با معادلات کوانتومی و فیزیک کلاسیک شبیه‌سازی می‌نماید. این امر به دانشمندان امکان می‌دهد که بتوانند از این طریق به اطلاعات ملکول از جمله ساختار هندسی، انرژی‌ها و خواص الکترونیکی و اثرات حلال دست پیدا کنند. در سال‌های اخیر

روش‌های محاسباتی در بین شیمی‌دانان رواج یافته است. در این روش‌ها می‌توان به راحتی محاسبات را انجام داد، بدون اینکه از اصول اولیه و روش محاسبه آگاهی دقیق داشت. روش‌های محاسباتی در شیمی به چندین دسته تقسیم می‌شوند. روشی که در این تحقیق به کار گرفته شده، روش $AM1^1$ است که جز روش‌های محاسباتی نیمه تجربی^۲ می‌باشد. در روش‌های نیمه تجربی، که در برنامه‌هایی نظیر Hyperchem وارد شده‌اند، محاسبات براساس مکانیک کوانتومی صورت می‌گیرد. Hyperchem یک نرم‌افزار قدرتمند مدل‌سازی مولکولی است که تمام ابزارهای موردنیاز برای ترسیم ساختارهای دوبعدی و سه‌بعدی مولکولی را در محیطی ساده و انعطاف‌پذیر در خود گردآوری کرده است که برای بهینه‌سازی ساختار هندسی مولکول‌ها از این نرم‌افزار استفاده می‌شود که توسط روش نیمه تجربی $AM1$ با استفاده از الگوریتم Polak-Rabiere تا ریشه میانگین مربع گرادیان 0.001 انجام شد. ایجاد توصیف‌کننده‌های هندسی و هیبریدی، بر مبنای ساختار و هندسه دقیق مولکولی استوار است و اگر ساختارها به شکل صورت‌بندی^۳ با حداقل انرژی نباشند مقادیر غیر صحیحی برای این توصیف‌کننده‌ها ایجاد می‌شود [۳۰]. به کمک این نرم‌افزار می‌توان طول پیوند، زاویه پیوندی و زوایای پیچشی را در مولکول تعیین کرد. داده‌های حاصل از این نرم‌افزار را می‌توان به عنوان ورودی به سایر نرم‌افزارها معرفی نمود.

۲-۳-۳ - محاسبه توصیف‌کننده‌ها

توصیف‌کننده‌ها کمیت‌هایی هستند که به ساختار مولکول ارتباط دارند و برای هر مولکولی مقادیر خاصی را به خود اختصاص می‌دهند. درواقع هر توصیف‌کننده بیان‌گر خصوصیت ویژه‌ای از مولکول است که ممکن است بر فعالیت و خاصیت موردنظر مؤثر باشد و هر توصیف‌کننده اطلاعات خاصی از مولکول را که بر کمیت مورد مدل‌سازی اثر می‌گذارد را در اختیار می‌دهد. محاسبه توصیف‌کننده‌ها توسط نرم‌افزار Dragon انجام می‌شود. از این نرم‌افزار جهت استخراج توصیف‌کننده‌های

¹ Austin methods

² Semi-empirical methods

³ Conformation

مولکولی استفاده می‌شود، Dragon قادر به محاسبه بیش از ۳۲۲۴ توصیف‌کننده می‌باشد که به ۲۲ دسته اصلی تقسیم‌بندی می‌شوند که در جدول (۱-۲) نام این ۲۲ دسته اصلی آورده شده است. علاوه بر محاسبه ساده‌ترین نوع اتم‌ها و گروه‌های عاملی و شمارش اجزا می‌توان تعداد زیادی از توصیف‌کننده‌های توپولوژیکی و هندسی را نیز توسط این نرم‌افزار محاسبه نمود.

جدول (۱-۲)- انواع توصیف‌کننده‌های محاسبه‌شده توسط نرم‌افزار Dragon

۱- توصیف‌کننده‌های زیر ساختاری ^۱	۲- توصیف‌کننده‌های توپولوژیکی ^۲
۳- شمارنده‌های مولکولی مورس ^۳	۴- توصیف‌کننده‌های BCUT ^۴
۵- شاخص‌های بار ^۵	۶- خود ارتباطی‌های دوبعدی ^۶
۷- توصیف‌کننده‌های بار ^۷	۸- شاخص‌های آروماتیسیت ^۸
۹- پروفایل‌های مولکولی راندیک ^۹	۱۰- توصیف‌کننده‌های هندسی ^{۱۰}
۱۱- توصیف‌کننده‌های عملکرد توزیع شعاعی ^{۱۱}	۱۲- توصیف‌کننده سه‌بعدی پراش الکترونی ^{۱۲}
۱۳- توصیف‌کننده‌های ملکولی مقیاس وزنی ^{۱۳}	۱۴- توصیف‌کننده‌های GETAWAY ^{۱۴}
۱۵- گروه‌های عاملی ^{۱۵}	۱۶- اجزای میان اتمی
۱۷- توصیف‌کننده‌های تجربی ^{۱۶}	۱۸- خصوصیات مولکولی
۱۹- شاخص‌های اتصال ^{۱۷}	۲۰- اثرانگشت فرکانسی ^{۱۸}
۲۱- شاخص‌های اطلاعاتی ^{۱۹}	۲۲- شاخص‌های مجاورت ^{۲۰}

¹ Substructural Descriptors

² Topological Descriptors

³ Morse Molecular Counts

⁴ BCUT Descriptors

⁵ Charge Indices

⁶ 2D Autocorrelations

⁷ Charge Descriptors

⁸ Aromatic Indices

⁹ Rancid Molecular Profiles

¹⁰ Geometrical Descriptors

¹¹ Radial Distribution Function Descriptors(RDF)

¹² 3D Molecule Representation of Structure Backed on Electron Diffraction Descriptors

¹³ Weighted Holistic Invariant Molecular Descriptors(WHIM)

¹⁴ Geometry, Topology and Atom Weights Assembly

¹⁵ Functional Groups

¹⁶ Empirical Descriptors

¹⁷ Connectivity indices

¹⁸ fingerprint Frequency

¹⁹ Information indices

²⁰ Adge adjacency indices

۲-۳-۴- انتخاب توصیف‌کننده‌های مهم

یکی از مهم‌ترین مراحل QSPR، انتخاب توصیف‌کننده‌های مناسب می‌باشد زیرا توصیف‌کننده‌های نامناسب کار برازش و مدل‌سازی را طولانی می‌کنند و تأثیری در بهبود نتایج نخواهند داشت. به این منظور باید در انتخاب توصیف‌کننده‌ها دقت لازم به عمل آید و توصیف‌کننده‌هایی که بیشترین و نزدیک‌ترین ارتباط را به پارامتر موردنظر دارند، انتخاب شوند (توصیف‌کننده‌هایی که بیشترین وابستگی را با خاصیت موردنظر دارند). ازجمله روش‌های انتخاب توصیف‌کننده^۱ می‌توان به رگرسیون گام‌به‌گام^۲، سهم گروه^۳، الگوریتم ژنتیک^۴ و روش جایگزینی^۵ اشاره کرد. که در این پایان‌نامه از روش رگرسیون گام‌به‌گام در نرم‌افزار SPSS^۶ برای انتخاب بهترین توصیف‌کننده‌ها استفاده شده است.

۲-۳-۴-۱- رگرسیون گام‌به‌گام

این روش با محاسبه ضرایب همبستگی که معیاری برای سنجش همبستگی بین یک توصیف‌کننده و متغیر پاسخ است، انجام می‌شود. مقدار این ضریب بین -۱ و ۱ است. در این روش ابتدا توصیف‌کننده‌ای که بیشترین همبستگی با متغیر وابسته (خاصیت موردنظر) را دارد، وارد مدل شده و با هر توصیف‌کننده جدید، کلیه‌ی توصیف‌کننده‌های موجود در مدل بررسی شده و اگر هر کدام از آن‌ها سطح معنادار خود را از دست داده باشند، قبل از ورود توصیف‌کننده جدید، از مدل خارج می‌شوند [۳۱].

¹ descriptor

² Step by Step regression (stepwise)

³ Group contribution

⁴ Genetic algorithm (GA)

⁵ Replacement method

⁶ Statistical Package for the social science (SPSS)

۲-۳-۵- ساخت مدل

مدل درواقع یک رابطه ریاضی است که بیان کننده رابطه بین متغیر وابسته و متغیرهای مستقل می باشد و به کمک آن می توان با داشتن مقادیر متغیرهای مستقل، متغیر وابسته را ارزیابی کرد. پس از انتخاب مناسب ترین توصیف کننده ها با استفاده از روش های آماری مختلف به جستجوی مدل مناسبی پرداخته می شود که بتواند ارتباط بین توصیف کننده های انتخابی و پارامترهای مورد مدل سازی را به درستی بیان کند. برای مدل سازی از روش های گوناگون خطی و غیرخطی می توان استفاده کرد. روش های آماری خطی شامل: رگرسیون خطی چندگانه (MLR)، حداقل مربعات جزئی^۱، رگرسیون اجزا اصلی و جنگل های تصادفی و روش های آماری غیرخطی شامل: شبکه عصبی مصنوعی (ANN) و رگرسیون اسپلاین تطبیقی چند متغیره^۲ (MARS) می باشد.

در این پایان نامه از روش خطی RF و روش های غیرخطی ANN و MARS برای ساخت مدل استفاده شده است.

۲-۴- مقدمه ای بر شبکه عصبی مصنوعی

شبکه عصبی مصنوعی یک الگوی پردازش اطلاعات است که از سیستم های عصبی بیولوژیکی مانند مغز الهام می گیرد. عنصر کلیدی این ایده، ایجاد ساختارهایی جدید برای سامانه پردازش اطلاعات است. این سیستم از شمار زیادی عناصر پردازشی فوق العاده به هم پیوسته بانام نرون^۳ تشکیل شده که برای حل یک مسئله با هم هماهنگ عمل می کنند و توسط سیناپس^۴ ها (ارتباطات الکترومغناطیسی) اطلاعات را منتقل می کنند. در این شبکه ها اگر یک سلول آسیب ببیند بقیه

^۱ Partial least square (PLS)

^۲ Multivariate Adaptive Regression Spline (MARS)

^۳ neuron

^۴ synapse

سلول‌ها می‌توانند نبود آن را جبران کرده و نیز در بازسازی آن سهیم باشند. این شبکه‌ها قادر به یادگیری‌اند، یادگیری در این سیستم‌ها به صورت تطبیقی صورت می‌گیرد، یعنی با استفاده از مثال‌ها وزن سیناپس‌ها به گونه‌ای تغییر می‌کند که در صورت دادن ورودی‌های جدید، سیستم پاسخ درستی تولید کند. به طور کلی یک شبکه عصبی مصنوعی از سه لایه ورودی، خروجی و پردازش تشکیل می‌شود. هر لایه شامل گروهی از سلول‌های عصبی (نرون) است که عموماً با کلیه نرون‌های لایه دیگر در ارتباط هستند، مگر اینکه کاربرد ارتباط بین نرون‌ها را محدود کند ولی نرون‌های هر لایه با سایر نرون‌های همان لایه، ارتباطی ندارد. نرون کوچک‌ترین واحد پردازشگر اطلاعات است که اساس عملکرد شبکه‌های عصبی را تشکیل می‌دهد. یک شبکه عصبی مجموعه‌ای از نرون‌هاست که با قرار گرفتن در لایه‌های مختلف، معماری خاصی را بر مبنای ارتباطات بین نرون‌ها در لایه‌های مختلف تشکیل می‌دهند. نرون می‌تواند یک تابع ریاضی غیرخطی باشد، در نتیجه یک شبکه عصبی که از اجتماع این نرون‌ها تشکیل می‌شود، نیز می‌تواند یک سامانه کاملاً پیچیده و غیرخطی باشد. در شبکه عصبی هر نرون به طور مستقل عمل می‌کند و رفتار کلی شبکه، برآیند رفتار نرون‌های متعدد است. به عبارت دیگر، نرون‌ها در یک روند همکاری، یکدیگر را تصحیح می‌کنند.

۲-۴-۱- عملکرد نرون مصنوعی

یک نرون با برچسب j که ورودی $P_j(t)$ را از نرون‌های ماقبل خود دریافت می‌کند از اجزا زیر تشکیل می‌شود:

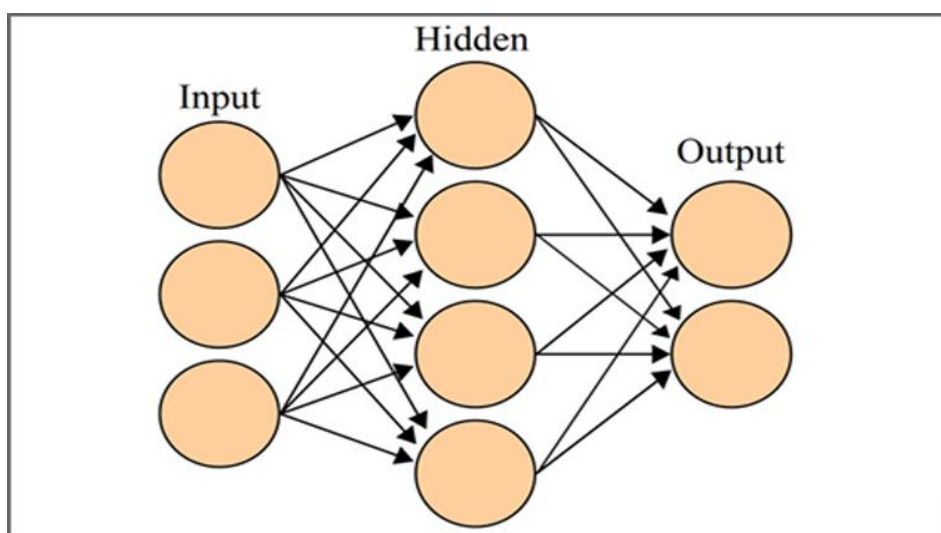
- یک فعال‌ساز که به یک پارامتر گسسته زمان وابسته است.
- احتمالاً یک آستانه θ_j ، ثابت است مگر آنکه توسط تابع یادگیرنده تغییر کند.
- یک تابع فعال‌سازی f که فعال‌ساز جدید را در زمان داده شده $t+1$ از $\theta_j(t)$ ، ورودی خالص P_j حساب می‌کند و رابطه (۱-۲) را بدست می‌آورد:

$$a_j(t+1) = f(a_j(t), P_j(t), \theta_j) \quad (1-2)$$

- و یک تابع خروجی که خروجی را از فعال ساز حساب می کند.

$$O_j(t) = f_{out}(a_j(t)) \quad (2-2)$$

تابع خروجی اغلب تابع همانی است. نرون ورودی نرون ماقبل ندارد اما برای کل شبکه نقش رابط ورودی را ایفا می کند. به طور مشابه نرون خروجی تالی ندارد و لذا به عنوان رابط خروجی کل شبکه ایفای نقش می کند.



شکل ۲-۲- شبکه عصبی مصنوعی

شکل (۲-۲) یک شبکه عصبی مصنوعی می باشد که شبیه شبکه بزرگ نرون ها در مغز است. در اینجا هر رأس دایره نمایشگر یک نرون مصنوعی است و پیکان نمایشگر یک اتصال از خروجی یک نرون مصنوعی به ورودی نرون دیگر است

۲-۴-۲- تابع انتقال^۱

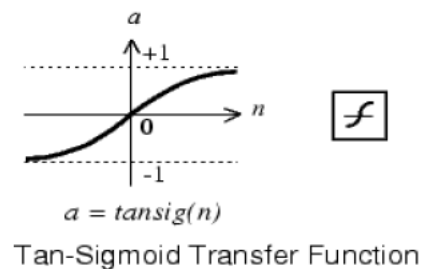
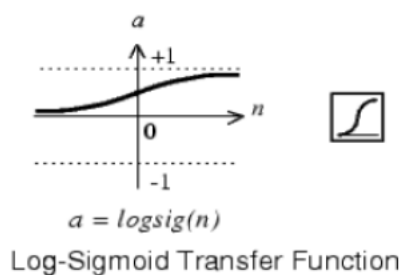
زمانی که ورودی‌ها با توجه به اهمیت آن‌ها با یکدیگر جمع جبری شوند توسط تابع انتقال به نرون بعدی منتقل می‌شوند. درواقع تابع انتقال یکی از اجزای شبکه عصبی می‌باشد که در این پایان‌نامه برای بهینه‌سازی شبکه به کار گرفته شده عبارت‌اند از:

➤ تابع انتقال لگاریتم سیگموئید^۲ (log sig)

از این تابع انتقال در شبکه‌های پس انتشار^۳ استفاده می‌شود. این تابع انتقال مقادیر ورودی را در محدوده $-\infty$ و $+\infty$ دریافت کرده و خروجی بین ۰ و ۱ تولید می‌نماید.

➤ تابع انتقال تانژانت سیگموئید^۴ (tan sig)

این تابع انتقال مقادیر ورودی را در محدوده $-\infty$ و $+\infty$ دریافت کرده و خروجی بین -۱ و ۱ تولید می‌نماید.



شکل ۲-۳- توابع انتقال

¹ Transfer function

² Logarithm sigmoid transfer function

³ Back propagation

⁴ Hyperbolic tangent transfer function

۲-۴-۳- ساختارهای شبکه

بر اساس ساختار اتصال نرون‌ها، شبکه عصبی مصنوعی را می‌توان در دودسته کلی طبقه‌بندی کرد:

➤ شبکه‌های پیش‌خور^۱

➤ شبکه‌های برگشتی^۲

شبکه‌های پیش‌خور: شبکه‌هایی هستند که مسیر پاسخ در آن‌ها همواره روبه‌جلو پردازش می‌شود و به نرون‌های لایه قبل باز نمی‌گردد. در این نوع شبکه‌ها به سیگنال^۳ اجازه می‌دهد که در مسیر یک‌طرفه عبور کند یعنی از ورودی تا خروجی، بنابراین بازخوردی وجود ندارد بدین معنی که خروجی هر لایه تأثیری بر همان لایه و همچنین لایه‌های قبلی ندارد.

شبکه‌های برگشتی: که در گراف آن‌ها به دلیل وجود بازخورد در ساختار شبکه، حلقه به وجود می‌آید. در معمول‌ترین خانواده‌ی شبکه پیش‌خور که پرسپترون چندلایه نامیده می‌شوند، نرون‌ها در لایه‌هایی قرار می‌گیرند و اتصال بین آن‌ها یک‌طرفه است. اتصالات مختلف سبب رفتارهای متفاوت شبکه‌ها می‌شود. به‌طور کلی می‌توان گفت شبکه‌های پیش‌خور استاتیک هستند، به این معنی که از ورودی داده‌شده تنها یک دسته مقدار خروجی تولید می‌کنند نه یک دنباله از مقدار خروجی. این شبکه‌ها بی‌حافظه هستند و پاسخ آن‌ها به یک ورودی مستقل از وضعیت قبلی شبکه است، از طرف دیگر شبکه‌های برگشتی، سیستم‌های دینامیک هستند و زمانی که یک دنباله ورودی جدید به آن‌ها داده شود، خروجی نرون‌ها محاسبه می‌شود.

^۱ Feed forward

^۲ Recurrent networks

^۳ Signal

۲-۴-۴- آموزش شبکه‌های پیش‌خور با تکنیک پس انتشار^۱

شبکه پس انتشار نوعی شبکه عصبی چندلایه با تابع انتقال غیرخطی می‌باشد. از بردار ورودی و هدف در راستای آموزش این نوع شبکه برای تقریب زدن یک تابع، یافتن رابطه بین ورودی و خروجی و دسته‌بندی ورودی‌ها استفاده می‌شود. یک شبکه BP با دارا بودن بایاس^۲، یک‌لایه سیگموئید و یک‌لایه خروجی خطی، توانایی تخمین هر تابعی با تعداد نقاط ناپیوستگی محدود را داراست [۳۲]. BP یک الگوریتم استاندارد با کاهش شیب می‌باشد که در آن وزن‌های شبکه در جهت خلاف شیب تابع کارایی^۳ حرکت می‌کنند. لغت پس انتشار به رفتار شبکه BP در محاسبه شیب در شبکه‌های غیرخطی چندلایه اشاره دارد. پس انتشار خطا یک روش متداول آموزش با ناظر برای شبکه‌های پیش‌خور است یعنی برای به دست آوردن ارتباط بین متغیرهای ورودی و خروجی در یادگیری به الگوی آموزشی نیاز است؛ بنابراین باید داده‌های آموزشی به شبکه ارائه شود. دو روش مختلف برای ارائه الگوهای آموزشی و پیاده‌سازی الگوریتم وجود دارد:

➤ آموزش گام‌به‌گام^۴

➤ آموزش دسته‌ای^۵

در آموزش گام‌به‌گام وزن‌ها و بایاس‌ها بعد از اعمال هر ورودی به‌روز می‌شوند درحالی‌که روش دسته‌ای پس از اعمال تمام ورودی‌ها (اعضای مجموعه آموزشی) عملیات به‌روزرسانی وزن‌ها انجام می‌شود. بدین‌صورت که شیب‌های محاسبه‌شده برای هر ورودی باهم جمع می‌شوند تا درنهایت وزن‌ها و بایاس‌ها از طریق آن به‌روز شود [۳۳].

به‌طور کلی آموزش به کمک تکنیک پس انتشار بر طبق مراحل زیر انجام می‌شود [۳۴]:

¹ Back propagation

² bias

³ Performance function

⁴ Incremental training

⁵ Batch training

۱-انتشار ورودی‌ها از نرون‌های ورودی به سمت نرون‌های خروجی

۲-اختصاص ماتریس وزن‌های تصادفی به هر یک از اتصالات

۳-مقایسه خروجی‌های شبکه با مقادیر واقعی (مقادیر هدف) و محاسبه خطای شبکه

۴-پس انتشار خطا از نرون‌های خروجی به سمت نرون‌های ورودی و اصلاح وزن‌ها

۵-ارزیابی عملکرد شبکه با توجه به تابع کارایی تعیین شده

مراحل فوق تا زمانی تکرار می‌شود که به حداکثر تکرار^۱ مجاز رسیده باشد یا مقدار تابع کارایی از مقداری که تعیین شده کمتر باشد. به دلیل فرآیندی که در طی این تکنیک رخ می‌دهد، این روش انتشار به عقب نامیده می‌شود زیرا ابتدا وزن‌های ارتباطی بین لایه خروجی و لایه پنهان و لایه ورودی اصلاح می‌گردند.

۲-۵- جنگل‌های تصادفی^۲

جنگل تصادفی یک روش مدرن از روش‌های درخت-مبنا^۳ هستند که شامل انبوهی از درخت‌های رده‌بندی و رگرسیونی^۴ می‌باشد [۳۵]. علاوه بر عملکرد بسیار بالای جنگل‌های تصادفی در پیشگویی متغیر پاسخ، این توانایی را نیز دارند که میزان اهمیت متغیرهای توضیحی را در پیشگویی متغیر پاسخ تعیین کنند. یک جنگل تصادفی مجموعه‌ای از درخت‌های هرس نشده است که هر درخت با الگوریتم

¹Epoch

² Random forests

³ Tree-based

⁴ Classification and Regression trees

افرازهای بازگشتی^۱ [۳۶] بدست می‌آید. الگوریتم ساخت یک جنگل تصادفی با T درخت از یک مجموعه داده با n مشاهده و p متغیر بدین صورت است:

(i) یک نمونه تصادفی با جایگذاری n تایی از مشاهدات انتخاب می‌شود.

(ii) برای نمونه انتخاب شده یک درخت کلاس‌بندی با استفاده از الگوریتم افرازهای بازگشتی، رشد می‌کند. در هر گره افراز بر اساس یک نمونه تصادفی m تایی از p متغیر پیشگو انجام می‌شود.

(iii) الگوریتم افرازهای بازگشتی آن قدر ادامه می‌یابد تا درخت به بزرگ‌ترین اندازه خود برسد بدون آن که درخت هرس شود.

(iv) مراحل (i) تا (iii) T بار تکرار می‌شود تا یک جنگل تصادفی ساخته شود [۳۷]. انتخاب‌های رایج برای T ، ۱۰۰۰ درخت و برای m ، $\sqrt{p}$ و $\log P$ هستند [۳۸]. یک جنگل تصادفی آن قدر بزرگ است که تفسیر آن کار بسیار دشواری است، لذا نیازمند خلاصه کردن اطلاعات آن با استفاده از شاخص‌های کمی هستیم. یکی از این شاخص‌ها اهمیت متغیر^۲ (VI) است، VI شاخصی برای رتبه‌بندی متغیرها برحسب اهمیت آن‌ها در اثرگذاری رویه پاسخ^۳ است. معروف‌ترین شاخص‌های اهمیت شاخص اهمیت جینی^۴ و شاخص اهمیت جایگشتی^۵ می‌باشد.

شاخص اهمیت جایگشتی: برای محاسبه این شاخص، الگوریتم جنگل‌های تصادفی از تمام مشاهدات نمونه برای ساخت درخت استفاده نمی‌کند بلکه یک نمونه تصادفی با جایگذاری به حجم n_1 (معمولاً برابر $2n/3$) از مشاهدات انتخاب می‌شود. به مشاهدات انتخاب شده نمونه آموزشی^۶ (LS) به

بقیه آن‌ها نمونه خارج کیسه^۷ گفته می‌شود. درخت‌ها با مشاهدات LS ساخته می‌شوند و از OOB

¹ Recurrent partitioning

² Variable importance

³ Response surface

⁴ Gini importance index

⁵ Permutaine importance mdex

⁶ Learning sample

⁷ Out of bag

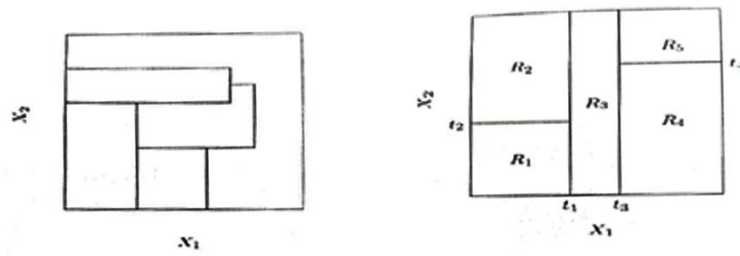
برای اندازه‌گیری ناخالصی درخت استفاده می‌شود. در هر درخت ابتدا اندازه ناخالصی روی مشاهدات OOB محاسبه می‌شود، سپس مقادیر متغیر X_i مشاهدات OOB به‌طور تصادفی جابه‌جا^۱ می‌شوند و اندازه ناخالصی درخت روی مقادیر جابه‌جاشده محاسبه می‌شود. اندازه اهمیت متغیر X_i در هر درخت، اختلاف بین این دو اندازه ناخالصی است و میانگین این مقادیر شاخص اهمیت جایگشتی است. انگیزه این روش این است که اگر X_i متغیر مهم می‌باشد جابه‌جا شدن مقادیر آن به‌طور تصادفی منجر به افزایش ناخالصی درخت می‌شود درحالی‌که اگر متغیر تأثیرگذاری نباشد، تغییری ایجاد نمی‌شود. کلیه تحلیل‌ها با استفاده از نرم‌افزار Matlab انجام شد.

۲-۵-۱- روش درخت رگرسیونی^۲ (تصمیم)

در روش رگرسیون خطی یک مدل واحد برای کل فضای داده‌ها در نظر گرفته می‌شود. اما ممکن است به دلیل تفاوت رفتار متغیر پاسخ در نواحی مختلف فضای داده‌ها، برقراری یک مدل واحد، کارایی لازم را نداشته باشد. لذا یک روش جایگزین، تقسیم‌بندی فضای داده‌ها به بخش‌های کوچک‌تر است تا بتوان رفتار متغیر پاسخ را به‌طور موضعی مدل‌سازی نمود. در روش درخت رگرسیونی هدف این است که مقادیر متغیرهای توضیحی در هر ناحیه از ابر مکعب p بعدی، به‌گونه‌ای تقسیم‌بندی می‌شوند که داده‌های واقع در هر ناحیه تا حد ممکن همگون باشند به‌طوری‌که بتوان آن‌ها را توسط ساده‌ترین مدل پیش‌بینی نمود. به‌طوری‌که افراز فضای متغیرهای توضیحی را می‌توان به هر شکلی انجام داد. اما در روش درخت رگرسیونی، تنها افرازی قابل قبول است که تمامی نواحی ساخته‌شده به شکل مربع یا مستطیل باشند. در شکل (۲-۴) نمونه‌ای از افراز قابل قبول و غیرقابل قبول در فضای دومتغیره با روش درخت رگرسیونی نشان داده شده است.

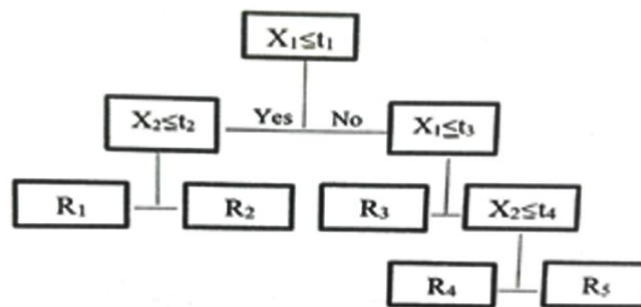
^۱ Permute

^۲ Regression tree



شکل ۲-۴- (الف) افزار قابل قبول شکل (ب) افزار غیر قابل قبول

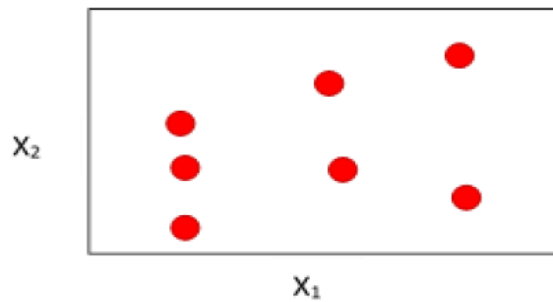
در شکل (۲-۵)، نوع دیگری از نمایش این افراز نشان داده شده است. دلیل نام گذاری درخت رگرسیونی را می توان به نمایش درختی افراز فضای متغیرهای ورودی نسبت داد که در آن هر مشاهده از نقطه بالای نمودار درختی وارد شده و در یکی از نواحی افراز قرار می گیرد.



شکل ۲-۵- نمودار درختی فضای افراز شده در افراز قابل قبول

۲-۵-۲- الگوریتم تشکیل درخت رگرسیونی

برای ساده کردن موضوع، ابتدا یک مسئله با دو توصیف کننده را در نظر می گیریم. فرض کنید نمودار پراکندگی داده ها به صورت شکل (۲-۶) باشد:

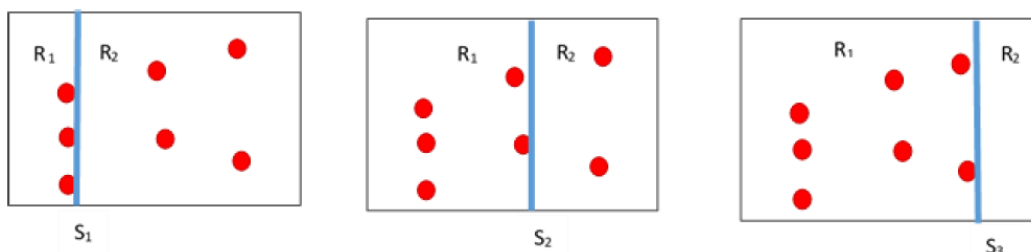


شکل (۶-۲)- نمودار پراکندگی داده‌های دومتغیره

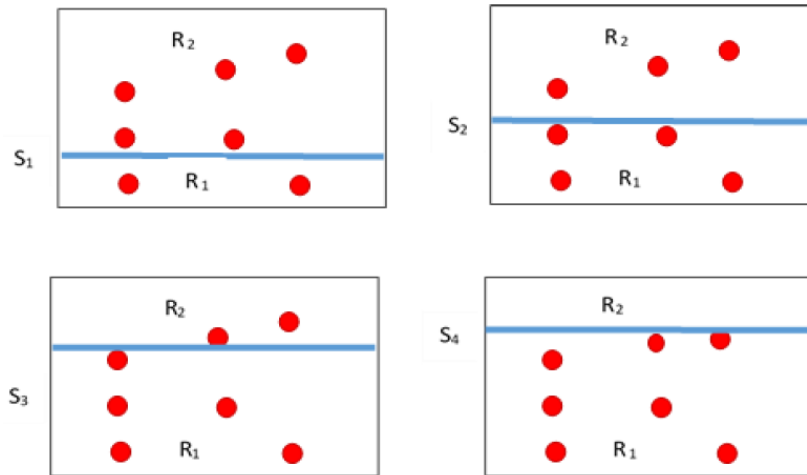
برای اولین افراز، انتخاب‌های مختلفی وجود دارد. شکل‌های (۷-۲) و (۸-۲)، برخی از افرازهای ممکن در راستای دو توصیف‌کننده را نشان می‌دهند. برای انتخاب بهترین تقسیم‌بندی، حاصل جمع مجموع مربعات خطا (A) در دو قسمت ایجاد شده به صورت زیر محاسبه می‌شود:

$$A = \sum_{i=1}^n (y_{i1} - \bar{y}_{i1})^2 + \sum_{i=1}^m (y_{i2} - \bar{y}_{i2})^2 \quad (۳-۲)$$

در این رابطه n تعداد داده‌های قرار گرفته در قسمت اول و m تعداد داده‌های قرار گرفته در قسمت دوم، y_{i1} و y_{i2} متغیر پاسخ داده‌های قرار گرفته در قسمت اول و دوم $\bar{y}_i$ میانگین مقادیر متغیر پاسخ در هر قسمت است. افراز بهینه افرازی است که در آن، مقدار عبارت فوق کمینه شود. پس از انتخاب بهترین افراز و به دست آمدن دو ناحیه مجزا، فرآیند فوق در هریک از نواحی تولید شده تکرار می‌شود.



شکل (۷-۲)- سه افراز ممکن در راستای متغیر x_1



شکل (۲-۸) - چهار افراز ممکن در راستای متغیر x_2

در حالت چند متغیره باید تمام افرازهای ممکن در راستای هر توصیف‌کننده در نظر گرفته شده و عبارت متناظر (۳-۲) محاسبه شود. به عبارت دیگر، روش $CART^1$ در هر مرحله افراز فضای توصیف-توصیف‌کننده‌ها باید هم‌زمان جستجو کند که کمترین مقدار عبارت (۳-۲) به ازای افراز در چه راستایی و از کدام نقطه (گره) حاصل می‌شود [۳۹].

۲-۵-۳- اندازه درخت و هرس کردن

یک مطلب مهم در ساخت درخت رگرسیونی این است که هر درخت تا کجا می‌تواند رشد کند. اگر درخت به اندازه کافی رشد نکند و یا به عبارت دیگر تفکیک فضای توصیف‌کننده‌ها در مراحل اولیه متوقف شود، ممکن است مدل مناسبی ارائه نشود و اگر هم درخت بیش از حد رشد کند، احتمال رخ دادن بیش‌برازشی^۲ وجود دارد. برای انتخاب اندازه درخت نیاز به یک پارامتر در ساختار مدل است

¹ Classification and regression trees

² Over fitting

به طوری که این پارامتر مقدار بهینه اندازه درخت را به دست آورد. در هر مرحله از رشد درخت، مدل حاصل از روش درخت رگرسیونی دقیق تر شده و مجموع توان های دوم خطا کاهش می یابد. یک راه توقف برای رشد درخت می تواند براساس میزان کاهش مجموع توان های دوم خطا باشد. بدین ترتیب که اگر در یک مرحله از رشد درخت، مجموع توان های دوم خطا دچار کاهش چندانی نشود، می توان از رشد درخت در آن مرحله صرف نظر کرد. این میزان کاهش در مجموع توان های دوم خطا می تواند در اختیار کاربر باشد [۳۹].

۲-۵-۴- الگوریتم جنگل تصادفی

اساس روش جنگل های تصادفی وابسته به ماهیت درخت رگرسیونی است و متشکل از ترکیب مجموعه ای از درخت های رگرسیونی می باشد [۴۰]. هر درخت یک مدل تولید می کند و از برآیند یا ترکیب این مدل ها، مدل نهایی به دست می آید. برخلاف روش درخت رگرسیونی، در این روش برای ساخت هر درخت، تنها بخشی از داده ها مورد استفاده قرار می گیرند، زیرا انتخاب داده های آموزش با جایگذاری^۱ انجام می شود. همچنین در جنگل های تصادفی برای افزایش فضای توصیف کننده ها و ساخت هر گره، به جای استفاده از تمام توصیف کننده ها، از زیر مجموعه ای از توصیف کننده ها استفاده می شود. فرض کنید $i = 1, \dots, N$ و (X_i, Y_i) مجموعه ای از داده های آموزش باشد که در آن $X_i = (x_{i1}, \dots, x_{iM})$ برداری از M توصیف کننده و Y_i متغیر پاسخ متناظر آن هاست. اگر تعداد کل درخت های مدل با n_{tree} نشان داده شود، مراحل زیر بیان گر ساخت هر درخت است:

الف- یک نمونه N تایی از داده های آموزش به روش جایگذاری برای ساخت مدل انتخاب می شوند. حدود یک سوم از داده های اصلی که در ساخت درخت شرکت نمی کنند، داده های خارج از کیسه

¹ With replacement

(OOB) نام دارند که برای هر درخت متفاوت هستند و نقش داده‌های ارزیابی را برای آن درخت ایفا می‌کنند.

ب- داده‌های انتخاب شده در مرحله اول و انتخاب تصادفی m توصیف‌کننده که زیر مجموعه‌ای از تعداد کل توصیف‌کننده‌ها است، جستجو برای یافتن بهترین توصیف‌کننده، فقط در بین m توصیف-کننده مذکور انجام شده و فضای توصیف‌کننده‌ها به دو قسمت تقسیم می‌شود. در این مرحله پیشنهاد بریمن این است که در مدل رگرسیونی می‌توان این تعداد را برابر با $M/3$ ، $M/6$ و یا $2M/3$ در نظر گرفت.

ج- برای هریک از دو ناحیه تولید شده، با استفاده از همان نمونه آموزش انتخاب شده در مرحله (الف)، مجدداً مرحله (ب) انجام می‌شود.

چ- مرحله (ج) برای تمامی نواحی افزار شده تا زمانی تکرار می‌شود که تعداد مشاهدات در تمامی نواحی کمتر از $2n_r$ باشد. پارامتر n_r در اختیار کاربر بوده و بیانگر حداقل تعداد مشاهدات موجود در هر ناحیه است. مقدار پیش فرض این پارامتر برای مسائل رگرسیون برابر با ۵ است.

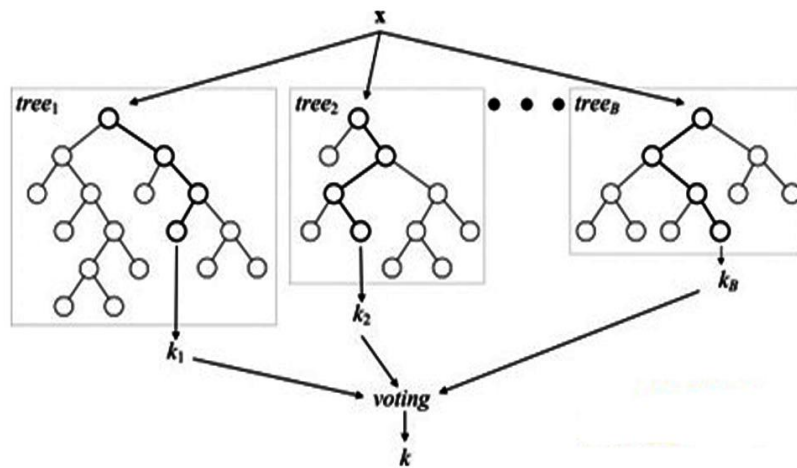
ح- مدل درخت رگرسیونی آم ساخته شده که فضای توصیف‌کننده‌ها را به ناحیه‌ی $R_{i1}, R_{i2}, \dots, R_{ir}$ تقسیم کرده است، به صورت زیر می‌باشد:

$$\hat{f}_i(x) = \sum_{j=1}^{r_i} \hat{C}_j I_{R_{ij}}(x) \quad (4-2)$$

که در آن، $\hat{C}_j = \bar{y}_j$ و $I_{R_{ij}} = \begin{cases} 1, & x \in R_{ij} \\ 0, & x \notin R_{ij} \end{cases}$ می‌باشد.

از آنجایی که در روش جنگل‌های تصادفی تعداد n_{tree} درخت رگرسیونی وجود دارد، می‌توان گفت که تعداد n_{tree} مدل به صورت معادله (۲-۴) خواهیم داشت. اگر برای هر داده، مدل خروجی درخت i ام را به صورت $\hat{f}_i(X)$ نشان دهیم، برآورد متغیر پاسخ برای این داده، با میانگین‌گیری از مدل خروجی تمام درخت‌ها یعنی از رابطه زیر بدست می‌آید [۴۱]:

$$\hat{y}(x) = \frac{1}{n_{tree}} \sum_{i=1}^{n_{tree}} \hat{f}_i(x) \quad (2-5)$$



شکل ۲-۹- نمای کلی الگوریتم

شکل ۲-۹ نمای کلی از الگوریتم جنگل تصادفی است که یک متد یادگیری تجمعی برای رده‌بندی و رگرسیون است. مدل‌های پایه، درختان رگرسیونی CART هستند و خروجی K ، بر اساس رأی‌گیری خروجی‌های B عدد درخت تعیین می‌شود. در مورد رده‌بندی، رأی نهایی بر مبنای مُد خروجی‌ها و در مورد رگرسیون بر اساس میانگین‌گیری تعیین می‌شود.

بکارگیری این استراتژی باعث عملکرد مناسب و کارا در مقایسه با دیگر الگوریتم‌ها همچون آنالیزهای جداسازی، بردارهای پشتیبان و شبکه عصبی مصنوعی می‌شود. همچنین این تکنیک به

تنظیم دو پارامتر (۱) تعداد متغیرهای قرارگرفته در زیرمجموعه متغیرها به تصادف انتخاب شده (۲) تعداد درختان مورد استفاده، نیاز دارد و علاوه بر آن، حساسیت شدیدی نسبت به مقادیر این پارامترها ندارد.

۲-۵-۵- تعیین اهمیت متغیرها در روش جنگل‌های تصادفی

در ساختار روش RF امکانی وجود دارد که می‌توان میزان اهمیت متغیرها^۱ را در مدل، تعیین نموده و متغیرهایی که دارای نقش بیشتری در هر درخت و در مدل نهایی دارند شناسایی کرد. همان‌طور که در الگوریتم روش RF اشاره شد، برای تشکیل هر درخت، یک نمونه با جایگذاری از داده‌های اصلی مورد استفاده قرار می‌گیرد. همچنین داده‌هایی که در این نمونه حضور ندارند را OOB نامیدیم که به نوعی نقش داده‌های آزمایشی را برای ارزیابی آن درخت ایفا می‌کند. فرض کنید خطای پیش‌بینی Y برای داده‌های OOB در درخت λ ام، با نماد $EOOB_i$ نشان داده شود. برای تعیین اهمیت متغیر λ ام، مقادیر این متغیر را به‌طور تصادفی n_{perm} مرتبه (n_{perm} پارامتری است که در اختیار کاربر می‌باشد) جابه‌جا کرده و مجدداً خطا، به ازای مجموعه جدید محاسبه می‌شود که مقدار این خطا با نماد $\widetilde{EOOB}_i^j$ نشان داده می‌شود. میزان اهمیت متغیر λ ام در مدل درخت λ ام با $(VI_i(x^j))$ و به‌صورت زیر تعریف می‌شود:

$$VI_i(x^j) = \widetilde{EOOB}_i^j - EOOB_i \quad (۶-۲)$$

سرانجام، میزان اهمیت متغیر λ ام در مدل نهایی $(VI_i(x^j))$ به‌صورت زیر است:

$$VI_i(x^j) = \frac{1}{ntree} \sum_{i=1}^{ntree} (VI_i(x^j)) \quad (۷-۲)$$

¹ Descriptor importance

برای دقت بیشتر می‌توان جابه‌جا نمودن تصادفی مقادیر متغیر λ را به تعداد بیشتر از یک‌بار انجام داد و در هر بار خطای برآورد را محاسبه نمود و سر انجام مقدار $\widehat{EOOB}_t^J$ را به‌صورت میانگین خطاهای حاصله در نظر گرفت [۴۲،۴۳]. قابل‌ذکر است که اندازه اهمیت یک متغیر، به‌تنهایی قابل تفسیر نبوده و فقط برای رتبه‌بندی متغیرها بر اساس اهمیت آن‌ها در مدل به کار می‌رود.

۲-۵-۶- مزیت روش جنگل تصادفی

۱- در این روش در مجموعه داده‌هایی که تعداد توصیف‌کننده زیادی دارد از کارایی بالایی برخوردار است. زیرا در هر مرحله از مدل‌سازی تنها زیرمجموعه‌ای از توصیف‌کننده‌ها انتخاب می‌شود.

۲- در این روش بیش برآزشی رخ نمی‌دهد.

۳- در صورت کم بودن تعداد داده‌ها، به دلیل وجود داده‌های OOB، نیاز به در نظر گرفتن تعدادی از داده‌ها به‌عنوان سری ارزیابی وجود ندارد.

۴- این روش توانایی انتخاب متغیر را داراست و برخلاف سایر روش‌های یادگیری ماشین، نیازمند روشی دیگر برای انتخاب متغیر نیست [۴۴].

۲-۶- رگرسیون اسپلاین تطبیقی چند متغیره (MARS)

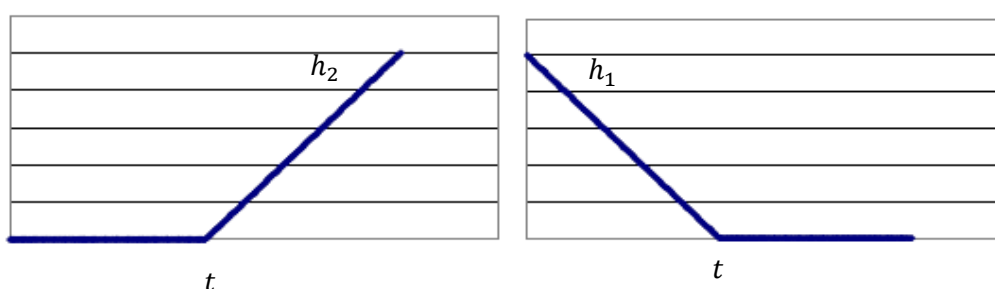
الگوی رگرسیون اسپلاین تطبیقی چند متغیره توسط فریدمن^۱ در سال ۱۹۹۱ معرفی شد که تعمیمی از رگرسیون خطی گام به گام یا شکل اصلاح شده درخت رگرسیون است [۴۵]. در این روش توابع زیر موسوم به توابع مبنا، مدل رگرسیونی را مشخص می‌کنند:

¹ Friedman

$$h_1(x) = (x - t)_+ = \begin{cases} x - t & x > t \\ 0 & x \leq t \end{cases} \quad (۸-۲)$$

$$h_2(x) = (t - x)_+ = \begin{cases} t - x & x < t \\ 0 & x \geq t \end{cases} \quad (۹-۲)$$

که در آن ثابت t گره نامیده شده و نماد "+" به معنای بخش مثبت است. گره t در عمل یکی از مشاهدات متغیر x است. توابع $h_1(x)$ و $h_2(x)$ را چوب هاکی و یا زوج منعکس یافته^۱ نیز می نامند. شکل (۱۰-۲) نمودار این توابع را نشان می دهد.



شکل (۱۰-۲) - نمودار توابع مبنای h_1 و h_2

مجموعه C (رابطه (۱۰-۲)) را تمامی زوج توابع منعکس یافته در نظر می گیریم که $\{x_1, x_2, x_3, \dots, x_n\}$ مقادیر مشاهده ای x هستند.

$$C = \{(x - t)_+, (t - x)_+\} \quad \forall t \in \{x_1, x_2, x_3, \dots, x_n\} \dots \dots \dots (۱۰-۲)$$

ساختار کلی الگوی رگرسیون یک متغیره مارس برای برآورد Y مطابق رابطه (۱۱-۲) است که h_i یکی از توابع مجموعه C است.

$$Y = f_{2m}(x) = b_0 + \sum_{i=1}^{2m} b_i h_i(x) \quad (۱۱-۲)$$

¹ Reflected pair

ضرائب $\{b_0, b_1, \dots, b_{2m}\}$ با کمینه‌سازی مجموع مربعات خطا برآورد می‌شوند. روش MARS سه مرحله دارد (۱) آرایش (ایجاد الگوی اولیه به روش پیشرو) (۲) پیرایش (اصلاح الگوی اولیه) و (۳) گزینش (انتخاب الگوی بهینه) [۴۵].

الگوهای مختلفی در مراحل آرایش و پیرایش براساس توابع پایه بر داده‌ها برازش و مجموع مربعات خطا^۱ اندازه‌گیری می‌شود. الگوی نهایی در مرحله گزینش و با کمترین مقدار معیار اعتبار متقابل تعمیم یافته^۲، $(1 - \frac{vm_j}{n})$ ، $GCV_j = SSE_j \div$ ، $(j=1,2,\dots,2n-2)$ انتخاب می‌شود. m_j تعداد توابع پایه در الگوی j ، SSE_j مجموع مربعات خطای الگوی j ، v پارامتر هموارسازی و n حجم نمونه است.

۲-۶-۱ MARS یک متغیره

در ابتدا مدل مارس را فقط با یک متغیر ورودی X بیان کرده، سپس در بخش‌های بعد آن را به ورودی‌های چندگانه تعمیم می‌دهیم. در این حالت، داده‌های مشاهده شده را به صورت $((X_1, Y_1), \dots, (X_n, Y_n))$ که (X_i, Y_i) معرف i امین مشاهده است، در نظر می‌گیریم.

➤ مرحله پیشرو

در ابتدا مدل در نظر گرفته شده فقط شامل عرض از مبدا است، یعنی

$$\hat{f}_1(x) = \hat{\beta}, \quad (12-2)$$

که در آن $\bar{y} = \hat{\beta}$ این مدل، مدل ۱ نامیده می‌شود. سوال این است که کدامیک از زوج توابع پایه مجموعه C باید وارد مدل شوند، تمام مدل‌های ممکن که با هریک از زوج توابع پایه ساخته می‌شوند (به اندازه K زوج توابع پایه تعداد $2K$ مدل داریم) را در نظر می‌گیریم و سپس برای ساختن مدل ۲،

¹ Sum of square error (SSE)

² Generalized cross validation

از بین این مدل‌ها، مدلی را که دارای کمترین مقدار مجموع توان‌های دوم خطا نسبت به سایر مدل‌ها است، انتخاب می‌کنیم.

$$\hat{f}_{21}(X) = \hat{\beta}_0 + \hat{\beta}_1(X - x_1)_+ + \hat{\beta}_2(x_1 - X)_+ \quad (13-2)$$

$$\hat{f}_{22}(X) = \hat{\beta}_0 + \hat{\beta}_3(X - x_2)_+ + \hat{\beta}_4(x_2 - X)_+ \quad (14-2)$$

$$\hat{f}_{2n}(X) = \hat{\beta}_0 + \hat{\beta}_{2n-1}(X - x_n)_+ + \hat{\beta}_{2n}(x_n - X)_+ \quad (15-2)$$

فرض کنید مدل $\hat{f}_{22}$ دارای کمترین مقدار SSE است لذا دومین مدل عبارت است از :

$$\hat{f}_2(X) = \hat{\beta}_0 + \hat{\beta}_3(X - x_2)_+ + \hat{\beta}_4(x_2 - X)_+ \quad (16-2)$$

که به صورت مختصرتر می‌توان آن را به صورت زیر نمایش داد و آن را مدل ۲ نامید.

$$\hat{f}_2(X) = \hat{\beta}_0 + \hat{\beta}_3 h_3(X)_+ + \hat{\beta}_4 h_4(X)_+. \quad (17-2)$$

اکنون مدل ۳ را با افزودن سایر توابع پایه به مدل ۲ که افزودن آن توابع، بزرگترین کاهش را در مقدار مجموع توان دوم‌های خطا دارد، ساخته می‌شود. بنابراین مدل ۳ از بین مدل‌های زیر انتخاب خواهد شد.

$$\hat{f}_{31}(X) = \hat{\beta}_0 + \hat{\beta}_1 * h_3(X) + \hat{\beta}_2 * h_4(X) + \hat{\beta}_3 * h_1(X) + \hat{\beta}_4 * h_2(X). \quad (18-2)$$

$$\hat{f}_{31}(X) = \dots + \hat{\beta}_3 * h_5(X) + \hat{\beta}_4 * h_6(X), \quad (19-2)$$

$$\hat{f}_{3(n-1)}(X) = \dots + \hat{\beta}_3 * h_{2n-1}(X) + \hat{\beta}_4 * h_{2n}(X) \quad (20-2)$$

همانطور که گفته شد مدلی از بین مدل‌های فوق انتخاب می‌شود که بیشترین کاهش در معادله (۲۱-۲) داشته باشد. شایان ذکر است که در هر بار ساختن مدل جدید ضرایب β موجود در مدل

جدید توسط مینیمم کردن مجموع توان‌های دوم مجدداً برآورد می‌شوند. این فرآیند تا زمانی که k تابع پایه به مدل اضافه شود، ادامه می‌یابد. توجه شود $k \leq 2n$ یا توسط کاربر مشخص می‌شود.

$$SSE = \sum_{i=1}^n (y_i - \hat{f}(x_i))^2 \quad (21-2)$$

➤ مرحله پس‌رو

در پایان مرحله پیش‌رو، مدل تولید شده مدلی است که اگر چه برای داده‌های مدل‌ساز ممکن است مناسب باشد اما موجب بیش‌برازش داده‌های آزمون است. برای رفع این مشکل، روش مارس وارد مرحله دوم یعنی "حذف توابع پایه از مدل" می‌شود. بنابراین اکنون به حذف ممکن هریک از توابع مدل f_k که دارای یک عرض از مبدأ و دو $2(k-1)$ تابع پایه است، توجه می‌شود. توابه پایه‌ای از مدل حذف خواهند شد که حذف آن‌ها کمترین افزایش را در مقدار مجموع توان‌های دوم خطا داشته باشد. این فرآیند تا زمانی که همه توابع از مدل، به جز عرض از مبدأ حذف شوند، ادامه می‌یابد.

وقتی که فرآیند حذف توابع پایه به‌طور کامل انجام شد، تعداد $2k-2$ مدل ساخته می‌شود (در هر مرحله یک مدل) که هر یک از آن‌ها، نامزد برای مدل نهایی هستند. برای هر کدام از این مدل‌ها، معیار اعتبار سنجی متقابل تعمیم یافته (GCV)، محاسبه می‌شود. اگر GCV_l مقدار شاخص GCV برای l -امین مدل در فرآیند پس‌رو به ازای $l=1, \dots, 2K-2$ باشد داریم [۴۳].

$$GCV_l = \frac{SSE_l}{1 - \left(\frac{vm_l + 1}{n} \right)} \quad (22-2)$$

که در آن m_l و SSE_l به ترتیب تعداد توابع پایه و مجموع توان‌های دوم خطا در مدل l ام و v جریمه هر تابع پایه است. در واقع می‌توان گفت که v به عنوان یک پارامتر هموارسازی که کاربر تعریف می‌کند، عمل می‌نماید. در حقیقت v مبادله بین مدل‌های پیچیده و ساده را کنترل می‌نماید.

مدلی که در بین 2k-2 مدل دارای کمترین مقدار GCV به عنوان مدل نهایی و برآورد مارس انتخاب می‌شود.

۲-۶-۲ MARS- دو متغیره

اگر بیش از یک متغیر پیش بین وجود داشته باشد، مارس این توانایی را دارد که جز به توابع پایه، حاصلضرب آن‌ها یعنی اثر متقابل‌شان را در مدل لحاظ کند [۴۶].

$$h_m(X) = h_i(X)(X_j - t)_+, h_{m+1}(X) = h_i(X)(t - X_j)_+ \quad (2-23)$$

که در آن $\{(X_j - t)_+, (t - X_j)_+\}$ جفت‌های منعکس شده از مجموعه C به ازای $J=1, \dots, P$ می‌باشند و $h_i(X)$ $1 \leq i \leq m-1$ توابع پایه‌ای هستند که از قبل در مدل وجود داشته‌اند. انتخاب جفت توابع پایه $\{h_{m+1}(X), h_m(X)\}$ از طریق جستجوی حاصلضرب‌هایی از توابع پایه $h_i(X)$ از مدل و جفت‌های منعکس شده در مجموعه C انجام می‌شود، به طوری که $h_i(X)$ و جفت‌های منعکس شده هیچ‌یک در ساختارشان متغیر یکسان X_j را نداشته و وقتی به مدل اضافه می‌شوند، بیشترین کاهش را در مجموعه توان‌های دوم خطا داشته باشند [۵۶].

۲-۶-۳ گام‌های ساخت MARS

مارس روش رگرسیون غیر خطی و ناپارامتریک است که پاسخ‌های غیر خطی را بین ورودی‌ها و خروجی‌های یک سیستم به وسیله مجموعه‌ای از قطعه‌های خطی تکه‌ای (کثیرالجمله‌های چند قطعه‌ای) با گرادیان‌های متفاوت مدل‌سازی می‌کند. فرضی ثابت در رابطه تابعی اساسی بین متغیرهای ورودی و خروجی لازم نیست. نقاط انتهایی این قطعه‌ها گره نامیده می‌شوند. گره انتهایی یک ناحیه از داده‌ها و ابتدای ناحیه‌ای دیگر از داده‌ها را مشخص می‌کند. منحنی‌های تکه‌ای منتج شناخته شده به-

عنوان توابع پایه‌ای انعطاف پذیری بیشتری به مدل می‌دهند و انحناها، آستانه‌ها و دیگر انحراف‌های حاصل از توابع خطی را در نظر می‌گیرند. مارس توابع پایه‌ای را با جستجو به روش مرحله‌ای ایجاد می‌کند. الگوریتم رگرسیون انطباقی برای انتخاب موقعیت گره به کار می‌رود. مدل‌های مارس به روش‌های دو مرحله‌ای ایجاد می‌شوند، مرحله مقدم توابع را جمع می‌بندد و گره‌های احتمالی را برای بهبود عملکرد می‌یابد که به مدلی با برازش کامل می‌انجامد. مرحله دوم دربرگیرنده حذف کمترین جمله‌های حقیقی است. مدل‌سازی مارس فرآیند پدیدآمده از داده‌هاست برای برازش مدل ابتدا روش انتخابی مقدم روی داده‌های آموزشی انجام می‌شود و سپس مدل تنها با جفت پایه ایجاد می‌شود که بزرگترین کاهش را در خطای آموزشی به وجود می‌آورد که با توجه به مدل جاری با توابع پایه M جفت بعدی به مدل اضافه می‌شود. برای کاهش تعداد جمله‌ها از رشته حذفی مؤخر پیروی می‌شود. هدف این روش پیدا کردن مدل نزدیک به حد مطلوب با حذف متغیرهای غیر اصلی است، مسیر این روش به این صورت است که از مدل با حذف توابع پایه همراه با کمترین سهم نسبت به مدل می‌زند تا اینکه بهترین مدل را بیابد، بنابراین توابع پایه حفظ شده در مدل بهینه نهایی از مجموعه تمام توابع پایه انتخاب می‌شود که در مرحله انتخابی مقدم مورد استفاده قرار می‌گیرد. زیر مجموعه‌های مدل با استفاده از روش ارزیابی متقابل تعمیم یافته (GCV) که به لحاظ محاسباتی کم هزینه است محاسبه می‌شوند. GCV نه تنها تعداد توابع پایه‌ای مدل بلکه تعداد گره‌ها را نیز برآورد می‌کند [۴۷].

۲-۶-۴- مزیت‌های MARS

مارس به لحاظ محاسباتی در پیدا کردن مدل بهینه کارآمدتر است چرا که در اصل با برازش مدل‌های قابل انعطافی می‌سازد و مدل را با تقسیم به شیب‌های جداگانه در فاصله‌های مشخص متغیرهای ورودی انتخاب می‌کند. متغیرها و محل‌های گره برای هر متغیر از طریق روش جستجو

سریع اما متمرکز تعیین می‌شود همچنین در مرحله پس‌رو و پیش‌رو تضمین می‌کند که مدل بهینه می‌تواند تشخیص داده شود.

۲-۷- ارزیابی قدرت پیش‌بینی مدل

۲-۷-۱- استفاده از پارامترهای آماری

برای اطمینان از اینکه مدل به‌دست‌آمده، مدل مناسبی است که توانایی پیشگویی نمونه‌های مختلفی از یک جمعیت را داراست، باید مدل را ارزیابی کرد. این ارزیابی از طریق شاخص‌های کمی است که به‌وسیله آن‌ها صحت نتایج ارائه‌شده توسط مدل موردسنجش قرار می‌گیرند. برخی از این شاخص‌های کمی عبارت‌اند از:

ضریب همبستگی^۱: ساده‌ترین راه برای بررسی میزان همبستگی دو یا چند متغیر، محاسبه آماری ضریب همبستگی آن‌هاست. ضریب همبستگی دو متغیر X, Y با رابطه زیر تعریف می‌شود. مقدار این آماره بین -1 و 1 است مقدار بزرگ‌تر آن نشان‌دهنده این است که ارتباط خطی بیشتری میان متغیر وابسته و متغیرهای مستقل وجود دارد.

$$R = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum (x_i - \bar{x})^2 \sum (y_i - \bar{y})^2}} \quad (2-24)$$

ضریب تعیین^۲: شاخصی برای بیان دقت خط رگرسیون برآورد شده است و مقدار آن بین صفر و یک تغییر می‌کند و مبین واریانس مشترک بین متغیر وابسته و مستقل است یعنی نسبتی از واریانس متغیر وابسته که توسط متغیرهای مستقل به‌حساب آمده معین می‌شود و توسط رابطه زیر تعریف می‌شود.

$$R^2 = \frac{SSR}{SST} \quad (2-25)$$

¹ Correlation coefficient

² Coefficient of determination

که SSR در این رابطه بیانگر مجموع مربعات انحراف مقادیر پیش‌بینی‌شده متغیر وابسته از میانگین مقادیر آن است.

$$SSR = \sum_{i=1}^n (\hat{y}_i - \bar{y})^2 \quad (26-2)$$

^۱ SST بیانگر مجموع مربعات انحراف مقادیر واقعی متغیر وابسته از میانگین مقادیر آن است.

$$SST = \sum_{i=1}^n (y_i - \bar{y})^2 \quad (27-2)$$

که در این روابط $\hat{y}_i$ مقدار پیش‌بینی‌شده متغیر وابسته، y_i مقدار واقعی متغیر وابسته و $\bar{y}$ در هر رابطه، میانگین مقادیر متغیر وابسته است. بنابراین با توجه به روابط فوق می‌توان نوشت:

$$R^2 = \frac{\sum_{i=1}^n (\hat{y}_i - \bar{y})^2}{\sum_{i=1}^n (y_i - \bar{y})^2} \quad (28-2)$$

طبق این رابطه اگر تمام مشاهدات بر روی خط برازش شده قرار گرفته باشند، یعنی به ازای تمام نقاط $y_i = \hat{y}_i$ باشد، مقدار R^2 برابر یک می‌شود و هرگونه انحرافی از این حالت باعث می‌شود که مقدار R^2 از یک کوچک‌تر شود.

خطای استاندارد پیش‌بینی (SEP):^۲

$$SEP = \sqrt{\frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{n}} \quad (29-2)$$

مجموع مربعات باقی‌مانده پیش‌بینی (PRESS):^۳

$$PRESS = \sum_{i=1}^n (y_i - \hat{y}_i)^2 \quad (30-2)$$

^۱ Sum square of total

^۲ Standard error of prediction

^۳ Predictive residual

میانگین مربع خطاها (MSE):^۱

$$MSE = \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{n} \quad (31-2)$$

میانگین خطای مطلق (MAE):^۲

$$MAE = \frac{\sum_{i=1}^n |y_i - \hat{y}_i|}{n} \quad (32-2)$$

میانگین خطای نسبی (MRE):^۳

$$MRE = \frac{\sum_{i=1}^n \left| \frac{y_i - \hat{y}_i}{y_i} \right|}{n} \quad (33-2)$$

خطای نسبی پیش‌بینی (REP):^۴

$$REP\% = \frac{100}{\bar{y}} \times \sqrt{\frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{n}} \quad (34-2)$$

ضریب تعیین تعدیل یافته:^۵

$$R_{adj}^2 = 1 - \frac{n-1}{n-p-1} \times \frac{SSE}{SST} = 1 - (1 - R^2) \frac{n-1}{n-p-1} \quad (35-2)$$

آمار F^6 : یا آزمون فیشر درواقع آزمون معنی‌دار بودن آماری در تحلیل رگرسیون ساده و چند

متغیره است و برابر با نسبت میانگین رگرسیون (MSR)^۷ به میانگین مربعات باقی‌مانده (MSE)

است.

¹Mean Square Error

²Mean Square Regression

³Mean relative error

⁴Relative error of Prediction Coefficient

⁵Adjusted determination coefficient

⁶Statistic term

⁷Statistic term

$$F = \frac{MSR}{MSE} = \frac{SSR/p}{SSE/(n-p-1)} \quad (2-36)$$

¹**FIT**: این پارامتر نیز به عنوان یک معیار برای مقایسه مدل های مختلف با تعداد متفاوت توصیف کننده به کار می رود.

$$FIT = \frac{R^2(n-p-1)}{(n+p^2) \times (1-R^2)} \quad (2-37)$$

در این روابط n تعداد ترکیبات مربوط به مدل و p تعداد توصیف کننده های مدل است، هرچه مقدار F و FIT بیشتر باشد (FIT به صورت مقایسه ای است و هر مدلی که FIT بالاتری داشته باشد مدل مناسب تری است) بهتر است. [۴۷، ۴۸].

۲-۷-۲- استفاده از نمودار خطای باقی مانده

منظور از عبارت خطای باقی مانده، اختلاف بین مقادیر پیش بینی شده و مقادیر تجربی است. اگر پراکندگی مقادیر در دو طرف خط صفر یکسان باشد، این امر نشان دهنده تصادفی بودن خطاست، ولی اگر عمده نقاط در این نمودار، در یک طرف خط صفر باشند، این بدان معناست که خطای جهت داری رخ داده است.

۲-۷-۳- استفاده از روش ارزیابی متقاطع

در این روش هم برای تعیین اعتبار مدل های پیش بینی کننده هر بار یک ترکیب از سری داده ها کنار گذاشته می شود و در شرایط بهینه به دست آمده برای سری ارزیابی یا سری آموزش با ترکیبات باقی مانده مدل سازی صورت می گیرد، سپس مدل به دست آمده برای پیش بینی فعالیت/خاصیت^۲

¹ Fitness function

² Activity/property

ترکیب کنار گذاشته شده به کار گرفته می شود و این فرآیند برای تمام اعضای سری داده ها تکرار می گردد.

۲-۷-۴-آزمون Y-تصادفی^۱

این تکنیک برای مطالعه همبستگی تصادفی مدل طراحی شده است. وقتی در مدلی، همبستگی متغیرهای مستقل با متغیرهای پاسخ به صورت تصادفی باشد، دارای همبستگی تصادفی هستند. در این آزمون، مقادیر پاسخ در محدوده پاسخ ها به صورت تصادفی تغییر داده می شود. سپس همبستگی متغیرهای مستقل با متغیرهای پاسخ با استفاده از یکی از شاخص های آماری مورد بررسی قرار می گیرد تفاوت قابل ملاحظه بین مقادیر ضریب تعیین مدل های اصلی و مدل های جدید بیانگر عدم وجود ارتباط تصادفی در مدل ها است. این فرآیند را چندین بار باید تکرار کرد [۴۸].

¹ Y-randomization test

فصل سوم

مطالعهٔ ارتباط کمی ساختار-خاصیت بر روی
تعدادی از ترکیبات هتروسیکل آروماتیک
سولفوردار چند حلقه‌ای با استفاده از
ANN و RF و MAR

۳-۱- مدل سازی جهت پیش بینی اندیس بازداری ترکیبات سولفور آروماتیک

چند حلقه ای سولفوردار

در این فصل به منظور مطالعه ارتباط کمی ساختار-خاصیت بروی ترکیبات سولفور چند حلقه ای، از جنگل های تصادفی، شبکه عصبی مصنوعی و رگرسیون اسپلاین تطبیقی چند متغیره به عنوان روش های مدل سازی برای پیش بینی اندیس بازداری این ترکیبات استفاده شده و توانایی مدل های به دست آمده از هر سه روش مورد ارزیابی قرار گرفته است. به طور کلی بخش تجربی این تحقیق شامل معرفی سری داده ها، بهینه سازی ساختار مولکول ها، محاسبه توصیف کننده های مولکولی، انتخاب توصیف کننده های مولکولی مناسب با استفاده از روش رگرسیون گام به گام و همچنین مدل سازی و ارزیابی مدل های برتر است.

۳-۲- نرم افزارهای مورد استفاده

در این تحقیق از بسته نرم افزاری متنوعی در مراحل مدل سازی استفاده شده است که هر یک به اختصار توضیح داده می شوند.

۳-۲-۱- نرم افزار Hyperchem

این نرم افزار جهت رسم و بهینه سازی اولیه ساختار ترکیبات شیمیایی مورد استفاده قرار می گیرد. در این نرم افزار، ساختار ترکیبات می تواند با چهار روش آغازین^۱، نیمه تجربی^۲، مکانیک مولکولی^۳ و تابع دانسیته^۴ بهینه سازی شود. در این پروژه، بهینه سازی تا زمانی ادامه یافت که گرادیان انرژی به ۰/۰۰۱ کیلوکالری بر مول رسید. به کمک این برنامه می توان طول پیوند، زاویه پیوندی و زوایای پیچشی را در مولکول تعیین کرد. داده های حاصل از این نرم افزار را می توان به عنوان ورودی به سایر

¹ Ab initio methods

² Semi-empirical methods

³ Molecular mechanicse

⁴ Density functional theory(DFT)

نرم افزارها معرفی نمود. همچنین به کمک این نرم افزار می توان تعدادی از توصیف کننده ها از جمله حجم مولی و قطبش پذیری را محاسبه کرد.

۳-۲-۲- نرم افزار Dragon

از این نرم افزار جهت استخراج توصیف کننده های مولکولی استفاده می شود. نرم افزار Dragon توسط گروه QSPR و کمومتریکس دانشگاه میلانو طراحی شده است. این نرم افزار قادر به محاسبه ۳۲۲۴ توصیف کننده می باشد. علاوه بر محاسبه ساده ترین نوع اتم ها و گروه های عاملی و شمارش اجزا می توان تعداد زیادی از توصیف کننده های توپولوژیکی و هندسی را نیز توسط این نرم افزار محاسبه نمود. برای اجرای این نرم افزار به فایل های ساختار مولکولی ایجاد شده به وسیله نرم افزارهای مدل سازی مولکولی، مانند Hyperchem نیاز می باشد. اکثر فایل های مولکولی با فرمت های معمول در این نرم افزار قابل قبول هستند. برای استفاده بهتر از محاسبات این نرم افزار باید بهینه سازی سه بعدی ساختارها با هیدروژن هایشان به کار رود. لازم به ذکر است که Dragon به عنوان یک نرم افزار QSPR/QSAR طراحی نشده است یعنی فقط توصیف کننده های مولکولی را نمایش می دهد و هیچ گونه آنالیز QSPR/QSAR را انجام نمی دهد. با این حال Dragon یک فایل خروجی کامل را که به سادگی توسط هر نرم افزار آنالیز همبستگی قابل کاربرد است فراهم می سازد تا به کمک آن بتوان مطالعات QSPR/QSAR را انجام داد و همبستگی های خاصی را می توان توسط آن حذف کرد.

۳-۲-۳- نرم افزار SPSS

SPSS یک نرم افزار آماری است که توانایی تحلیل کلی اطلاعات را دارا بوده و عمومی ترین نرم افزار آماری می باشد. اولین نسخه آن در سال ۱۹۷۰ توسط جمعی از فارغ التحصیلان دانشگاه استنفورد آمریکا ارائه شده است. پس از ۲۸ جولای ۲۰۰۹ که شرکت سازنده این نرم افزار توسط

IBM خریداری شد این نرم افزار با نام PASW^۱ منتشر گردید. SPSS از جمله نرم افزارهایی است که برای تحلیل های آماری در علوم اجتماعی، به صورت بسیار گسترده ای استفاده می شود. افزون بر تحلیل های آماری، مدیریت داده ها و مستندسازی داده ها نیز از ویژگی های این نرم افزار هستند. از جمله توانایی های نرم افزار SPSS عبارتند از:

※تهیه خلاصه های آماری مانند گراف ها، جداول، آمارها

※محاسبه انواع توابع ریاضی مانند قدر مطلق، تابع علامت، لگاریتم، توابع مثلثاتی

※تهیه انواع جداول سفارشی مانند جداول فراوانی، فراوانی تجمعی، درصد فراوانی

※انواع توزیع های آماری شامل توزیع های گسسته و پیوسته

※تهیه انواع طرح های آماری

※انجام آنالیز واریانس یک طرفه، دو طرفه، چند طرفه و آنالیز کوواریانس

※محاسبه میانگین ساده برای داده ها

※تکنیک های تجزیه و تحلیل سری های زمانی

※محاسبه انواع آمارهای توصیفی

※ایجاد داده های تصادفی و پیوسته

※انجام رگرسیون تک متغیره و چند متغیره

۳-۲-۴- نرم افزار MATLAB^۲

نرم افزار MATLAB برنامه ای کامپیوتری است که برای کسانی که با محاسبات عددی سر و کار دارند تهیه شده است. هدف اولیه این نرم افزار قادر ساختن مهندسين و دانشمندان به حل مسائل، شامل عملیات ماتریسی بدون نیاز به نوشتن برنامه در زبان های فورترن و بود. کاربردهای نوعی آن شامل ریاضیات، الگوریتم محاسباتی، شبیه سازی، مدل سازی و آنالیز اطلاعات می باشد. یکی از

^۱ Predictive Analytics Software

^۲ Matrix Laboratory

کاربردهای مهم آن که مورد استقبال شیمیدانان قرار گرفته امکان استفاده از آن برای آنالیز آماری و مدل سازی داده های شیمیایی است. مدل سازی به روش شبکه ی عصبی مصنوعی از جمله کاربردهای مهم دیگر آن است که در تحقیق حاضر مورد استفاده قرار گرفته است. اطلاعات شیمیایی ترکیبات به صورت ماتریس به عنوان ورودی به نرم افزار ارائه می شوند. سپس در محیط MATLAB با ورژن ۲۰۱۵ و با استفاده از آرایه ها و فرامین موجود امکان مدل سازی غیر خطی فعالیت شیمیایی ترکیبات امکان پذیر می گردد. سایر روش های کمومتری کس شامل آنالیز (رگرسیون) اجزای اصلی، روش کمترین مربعات جزئی، الگوریتم ژنتیک و غیره نیز توسط این نرم افزار قابل اجرا است. برای رسم نمودار سه بعدی در این تحقیق از نرم افزار sigmaplot با ورژن ۱۲ استفاده شده است.

۳-۲-۵- نرم افزار R

R یک زبان برنامه نویسی و محیط نرم افزاری برای محاسبات آماری و علم داده هاست و اکثر زمینه های آمار کاربردی مانند تحلیل سری های زمانی، رگرسیون خطی و غیر خطی، آزمون فرض های کلاسیک، کدگذاری، خوشه بندی و را پوشش می دهد. R یک نرم افزار متن باز^۱ و رایگان است. R نرم افزار قدرتمندی برای ایجاد اشکال گرافیکی و نمودارهاست [۴۹].

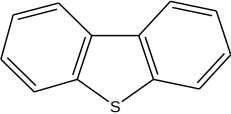
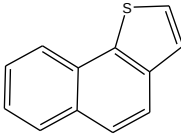
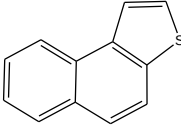
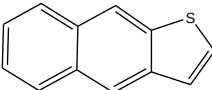
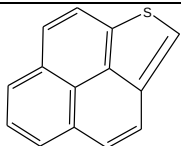
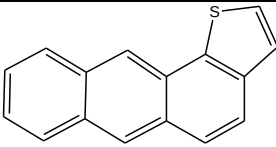
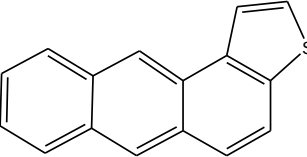
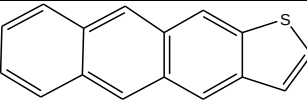
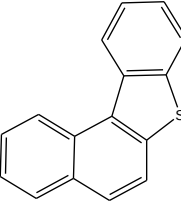
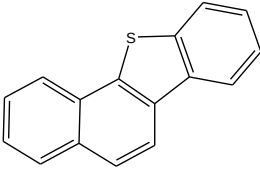
۳-۳- سری داده ها

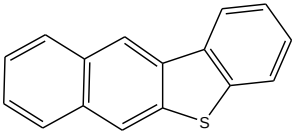
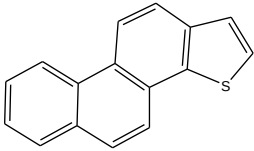
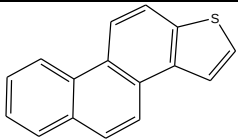
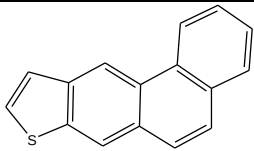
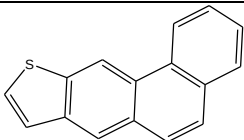
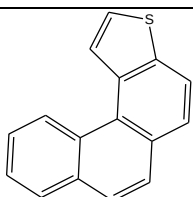
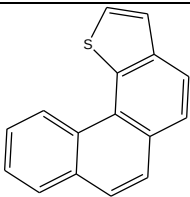
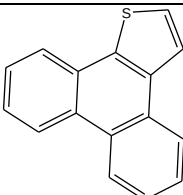
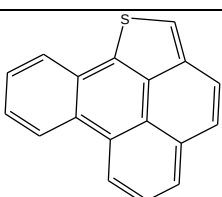
سری داده ها شامل نتایج اندیس بازداری ۱۵۲ ترکیب آروماتیک سولفور چند حلقه ای می باشد که توسط والتر ویلسون^۲ و همکارانش گزارش شده است [۵۱، ۵۰]. نام، ساختار و اندیس بازداری (RI) تجربی این ترکیبات در جدول (۳-۱) نشان داده شده است. شایان ذکر است که اندیس بازداری به عنوان خاصیت ترکیبات ذکر شده، متغیر وابسته محسوب می شود.

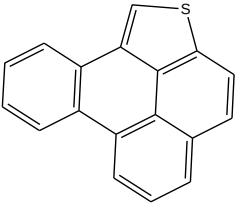
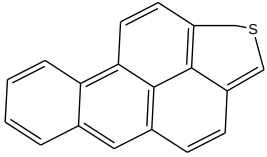
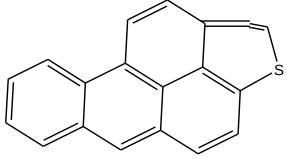
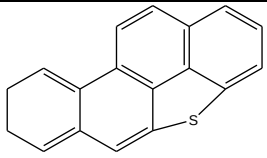
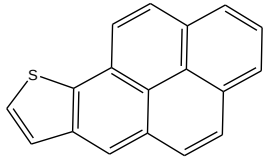
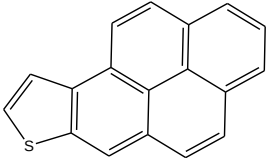
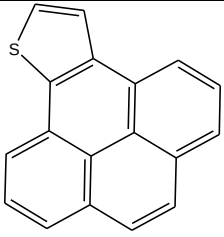
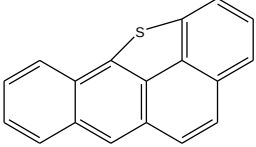
¹ Open source

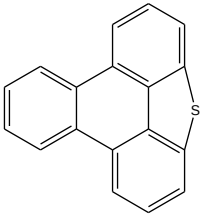
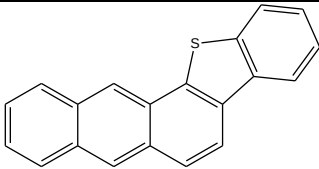
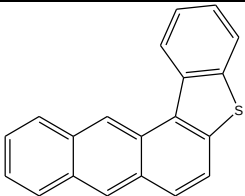
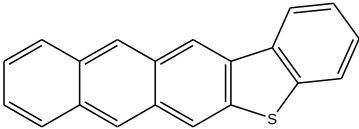
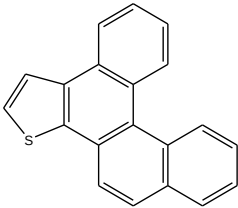
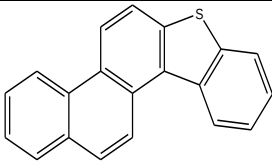
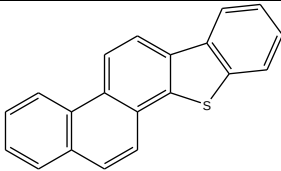
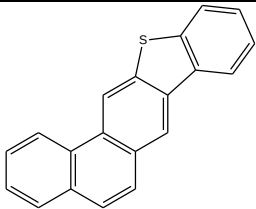
² Walter B . wilson

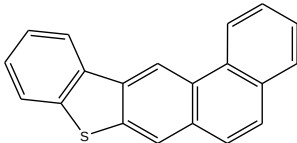
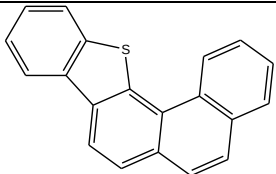
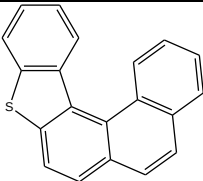
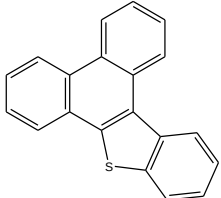
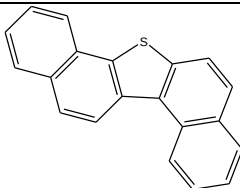
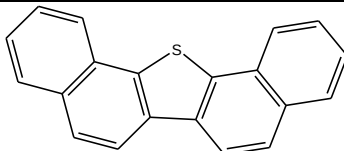
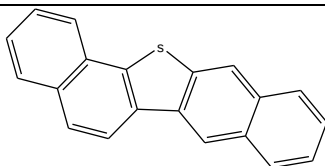
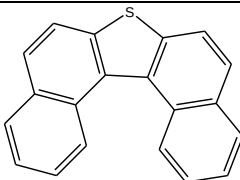
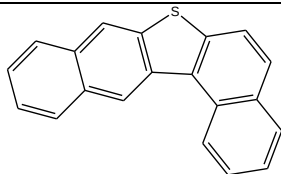
جدول (۳-۱) نام و اندیس بازداري سری داده‌ها

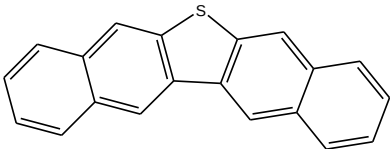
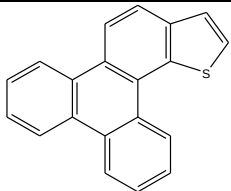
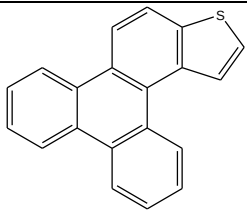
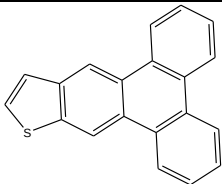
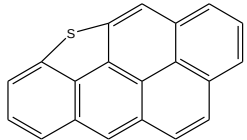
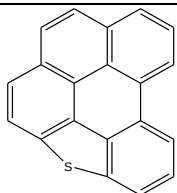
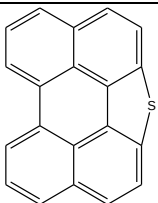
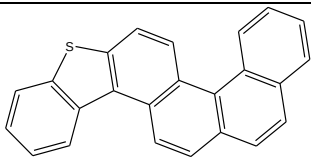
شماره	نام	ساختار	RI
۱	Dibenzothiophene		۳/۲۸
۲	Naphtho[1,2-b]thiophene		۲/۸۴
۳	Naphtho[2,1-b]thiophene		۲/۷۸
۴	Naphtho[2,3-b]thiophene		۲/۹۱
۵	Phenaleno[1,9-bc]thiophene		۳/۴۷
۶	Anthra[1,2-b]thiophene		۳/۹۲
۷	Anthra[2,1-b]thiophene		۳/۷۶
۸	Anthra[2,3-b]thiophene		۳/۹۰
۹	Benzo[b]thiophene		۴/۰۰
۱۰	Benzo[b]thiophene		۴/۱۳

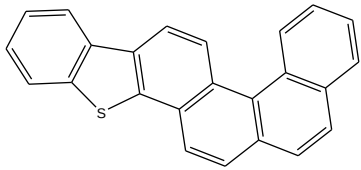
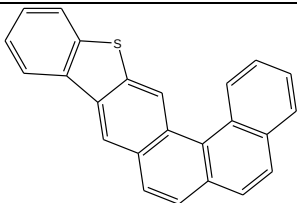
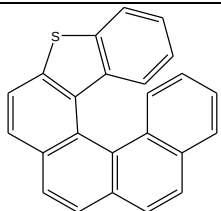
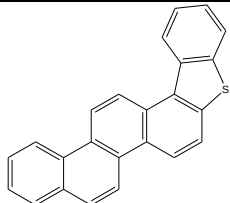
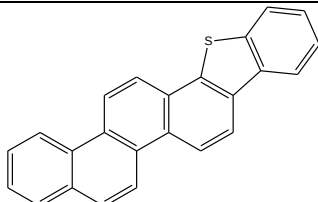
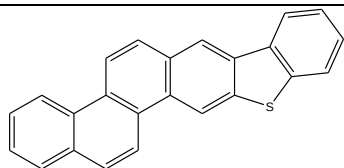
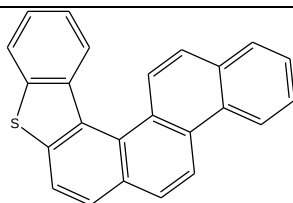
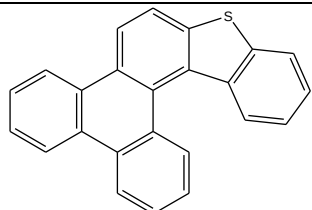
شماره	نام	ساختار	RI
۱۱	Benzo[b]thiophene		۴/۰۵
۱۲	Phenanthro[1,2-b]thiophene		۳/۹۰
۱۳	Phenanthro[2,1-b]thiophene		۳/۷۰
۱۴	Phenanthro[2,3-b]thiophene		۳/۷۰
۱۵	Phenanthro[3,2-b]thiophene		۳/۶۶
۱۶	Phenanthro[3,4-b]thiophene		۳/۷۳
۱۷	Phenanthro[4,3-b]thiophene		۳/۸۶
۱۸	Phenanthro[9,10-b]thiophene		۳/۸۸
۱۹	Benzo[1,2]phenaleno[3,4-bc]thiophene		۴/۶۳

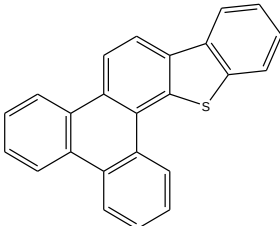
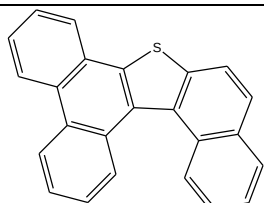
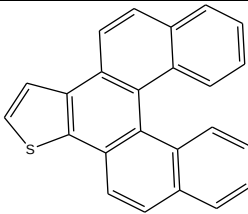
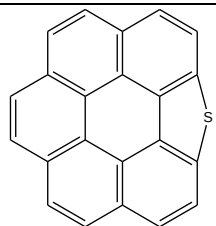
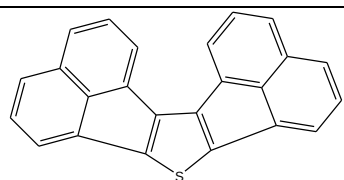
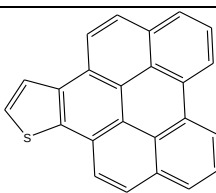
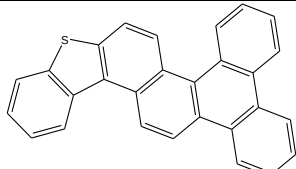
شماره	نام	ساختار	RI
۲۰	Benzo[1,2]phenaleno[4,3-bc]thiophene		۵/۰۲
۲۱	Benzo[4,5]phenaleno[1,9-bc]thiophene		۴/۴۷
۲۲	Benzo[4,5]phenaleno[9,1-bc]thiophene		۴/۴۷
۲۳	Chryseno[4,5-bcd]thiophene		۴/۷۶
۲۴	Pyreno[1,2-b]thiophene		۴/۵۹
۲۵	Pyreno[2,1-b]thiophene		۴/۴۳
۲۶	Pyreno[4,5-b]thiophene		۴/۵۴
۲۷	Tetrapheno[1,12-bcd]thiophene		۴/۹۷

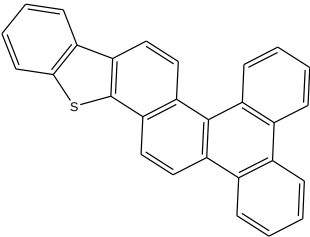
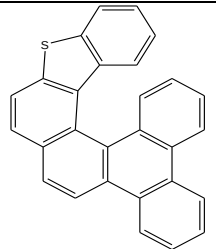
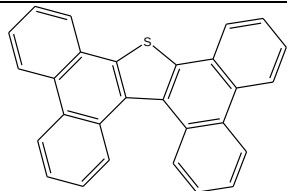
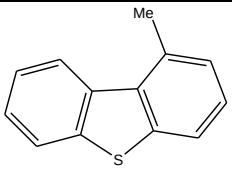
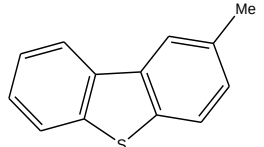
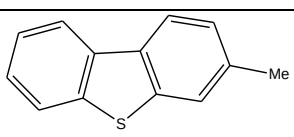
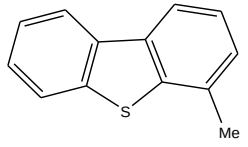
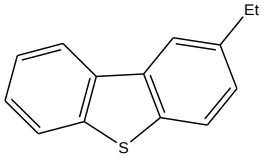
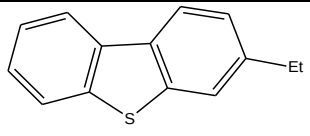
شماره	نام	ساختار	RI
۲۸	Triphenyleno[4,5-bcd]thiophene		۴/۶۰
۲۹	Anthra[1,2-b]benzo[d]thiophene		۵/۳۷
۳۰	Anthra[2,1-b]benzo[d]thiophene		۴/۹۹
۳۱	Anthra[2,3-b]benzo[d]thiophene		۵/۰۶
۳۲	Benzo[3,4]phenathro[1,2-b]thiophene		۴/۸۹
۳۳	Benzo[b]phenathro[1,2-b]thiophene		۵/۲۲
۳۴	Benzo[b]phenathro[2,1-b]thiophene		۵/۳۶
۳۵	Benzo[b]phenathro[2,3-b]thiophene		۴/۵۱

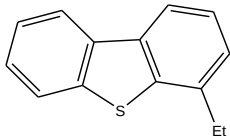
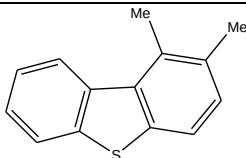
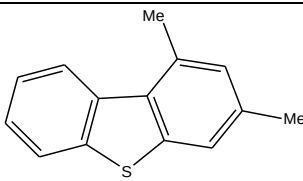
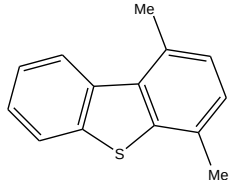
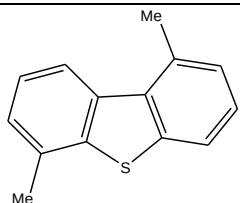
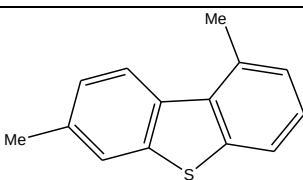
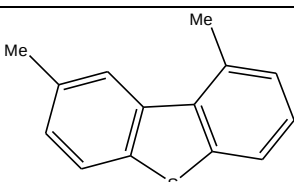
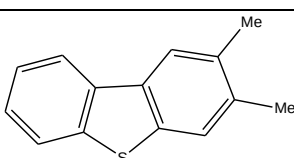
شماره	نام	ساختار	RI
۳۶	Benzo[b]phenathro[3,2-b]thiophene		۴/۹۱
۳۷	Benzo[b]phenathro[3,4-b]thiophene		۵/۳۲
۳۸	Benzo[b]phenathro[4,3-b]thiophene		۴/۹۱
۳۹	Benzo[b]phenathro[9,10-b]thiophene		۵/۳۰
۴۰	Dinaphtho[1,2:b-1,2-d]thiophene		۵/۴۰
۴۱	Dinaphtho[1,2:b-2,1-d]thiophene		۵/۶۶
۴۲	Dinaphtho[1,2:b-2,3-d]thiophene		۵/۴۷
۴۳	Dinaphtho[2,1:b-1,2-d]thiophene		۴/۹۷
۴۴	Dinaphtho[2,1:b-2,3-d]thiophene		۴/۹۹

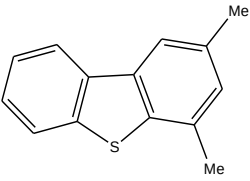
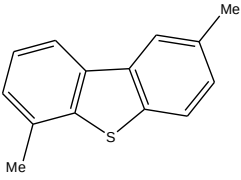
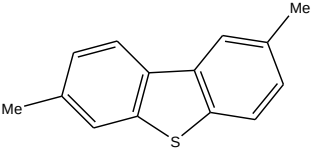
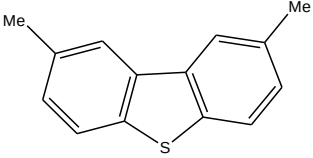
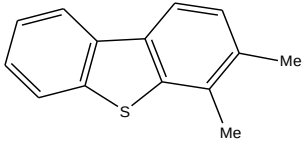
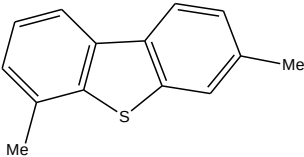
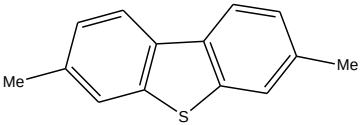
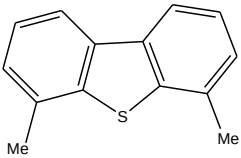
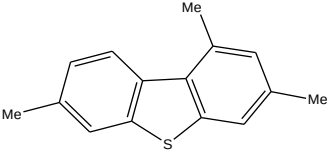
شماره	نام	ساختار	RI
۴۵	Dinaphtho[2,3-b-2,3-d]thiophene		۵/۱۹
۴۶	Triphenyleno[1,2-b]thiophene		۴/۷۶
۴۷	Triphenyleno[2,1-b]thiophene		۴/۶۴
۴۸	Triphenyleno[2,3-b]thiophene		۴/۵۱
۴۹	Benzo[10,11]chryseno[4,5-bcd]thiophene		۵/۸۳
۵۰	Benzo[4,5]triphenyleno[1,12-bcd]thiophene		۵/۵۷
۵۱	Peryleno[1,12-bcd]thiophene		۵/۰۵
۵۲	Benzo[b]benzo[5,6]phenathro[1,2-d]thiophene		۶/۰۵

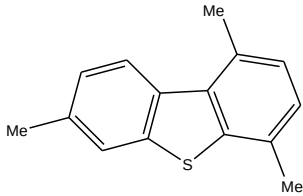
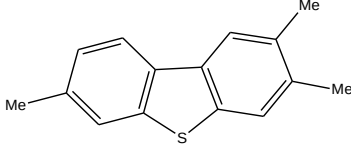
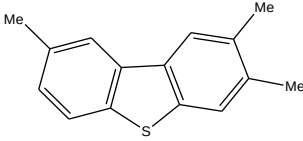
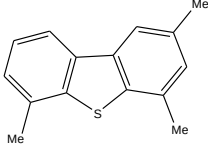
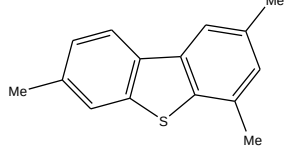
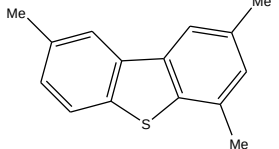
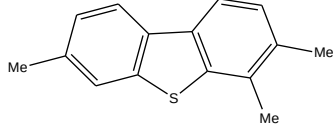
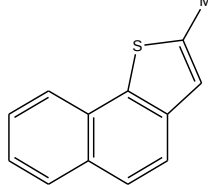
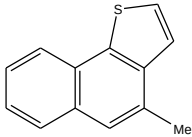
شماره	نام	ساختار	RI
۵۳	Benzo[b]benzo[5,6]phenathro[2,1-d]thiophene		۶/۳۳
۵۴	Benzo[b]benzo[5,6]phenathro[2,3-d]thiophene		۵/۸۲
۵۵	Benzo[b]benzo[5,6]phenathro[4,3-d]thiophene		۵/۱۴
۵۶	Benzo[b]chryseno[1,2-d]thiophene		۶/۱۶
۵۷	Benzo[b]chryseno[2,1-d]thiophene		۵/۷۰
۵۸	Benzo[b]chryseno[2,3-d]thiophene		۵/۰۰
۵۹	Benzo[b]chryseno[4,3-d]thiophene		۵/۸۹
۶۰	Benzo[b]triphenyleno[1,2-d]thiophene		۵/۷۸

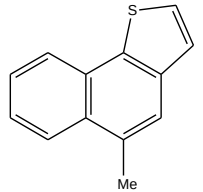
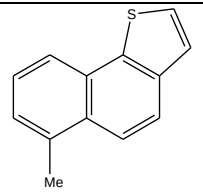
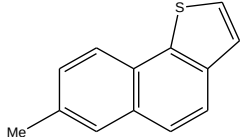
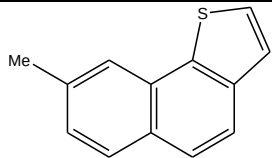
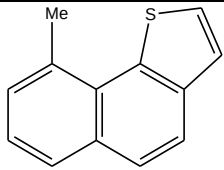
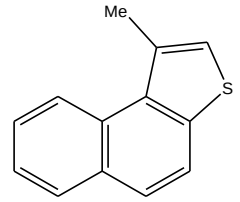
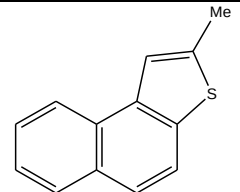
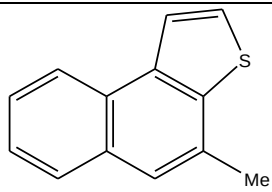
شماره	نام	ساختار	RI
۶۱	Benzo[b]triphenylene[2,1-d]thiophene		۶/۲۷
۶۲	Naphtho[2,1-d]phenanthro[9,10-d]thiophene		۶/۱۹
۶۳	Naphtho[2,1:3,4-d]phenanthro[1,2-b]thiophene		۵/۵۵
۶۴	Benzo[6,7]perylene[1,12-bcd]thiophene		۶/۴۱
۶۵	Diacenaphtho[1,2-b:1,2-d]thiophene		۶/۲۵
۶۶	N[218:345]PY12T		۵/۴۳
۶۷	Benzo[b]benzo[5,6]chryseno[1,2-d]thiophene		۶/۷۰

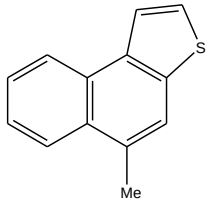
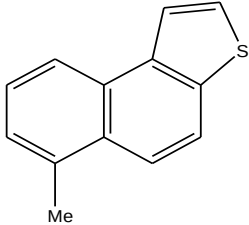
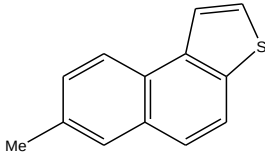
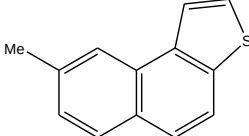
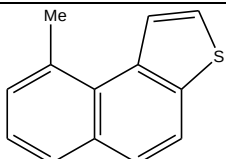
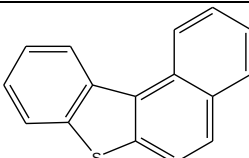
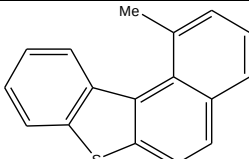
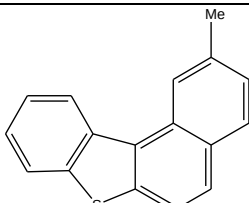
شماره	نام	ساختار	RI
۶۸	Benzo[b]benzo[5,6]chryseno[2,1-d]thiophene		۷/۰۳
۶۹	Benzo[b]benzo[5,6]chryseno[3,4-d]thiophene		۶/۴۳
۷۰	Diphenanthro[9,10-b:9,10-d]thiophene		۷/۳۷
۷۱	1-Methyl Dibenzothiophene		۳/۵۴
۷۲	2-Methyl Dibenzothiophene		۳/۵۷
۷۳	3-Methyl Dibenzothiophene		۳/۶۴
۷۴	4-Methyl Dibenzothiophene		۳/۶۸
۷۵	2-Ethyl Dibenzothiophene		۳/۹۸
۷۶	3-Ethyl Dibenzothiophene		۴/۰۰

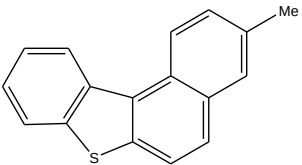
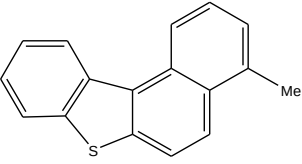
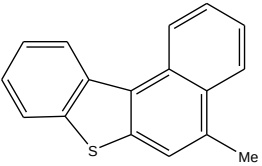
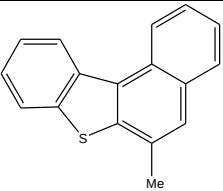
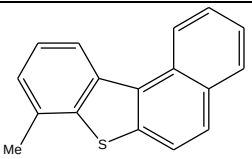
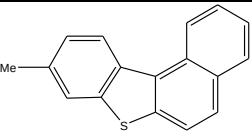
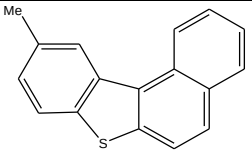
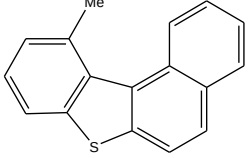
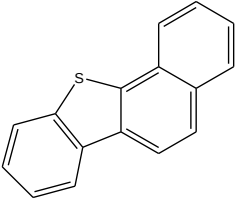
شماره	نام	ساختار	RI
۷۷	4-Ethyl Dibenzothiophene		۴/۰۸
۷۸	1,2-Dimethyl dibenzothiophene		۴/۰۱
۷۹	1,3-Dimethyl dibenzothiophene		۴/۱۱
۸۰	1,4-Dimethyl dibenzothiophene		۴/۱۶
۸۱	1,6-Dimethyl dibenzothiophene		۴/۱۶
۸۲	1,7-Dimethyl dibenzothiophene		۴/۱۰
۸۳	1,8-Methyl Dibenzothiophene		۴/۰۰
۸۴	2,3-Dimethyl dibenzothiophene		۴/۰۱

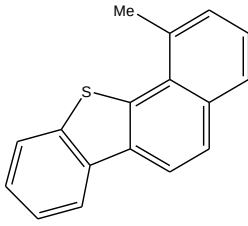
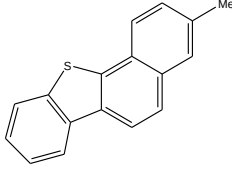
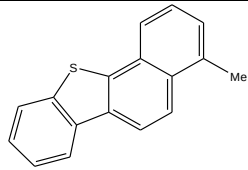
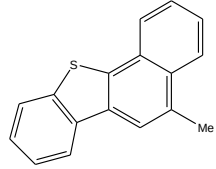
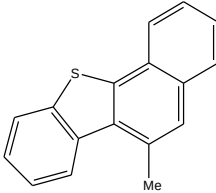
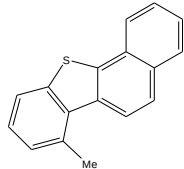
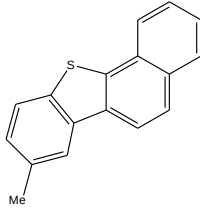
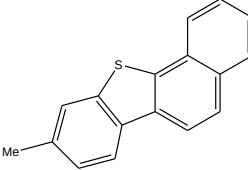
شماره	نام	ساختار	RI
۸۵	2,4-Dimethyl dibenzothiophene		۴/۱۷
۸۶	2,6-Dimethyl dibenzothiophene		۴/۱۷
۸۷	2,7-Dimethyl dibenzothiophene		۴/۱۲
۸۸	2,8-Dimethyl dibenzothiophene		۴/۰۵
۸۹	3,4-Methyl Dibenzothiophene		۴/۱۰
۹۰	3,6-Dimethyl dibenzothiophene		۴/۲۴
۹۱	3,7-Dimethyl dibenzothiophene		۴/۱۹
۹۲	4,6-Dimethyl dibenzothiophene		۴/۲۵
۹۳	1,3,7-Trimethyl dibenzothiophene		۴/۶۶

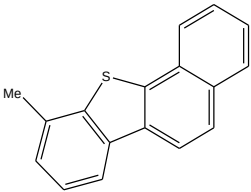
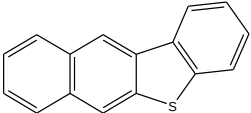
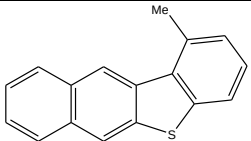
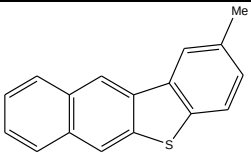
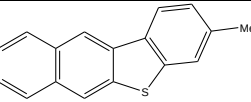
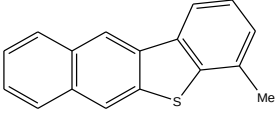
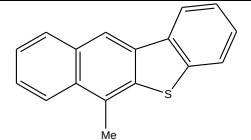
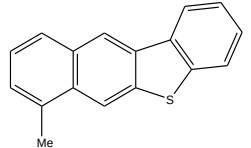
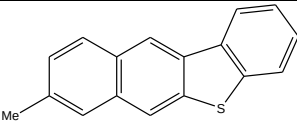
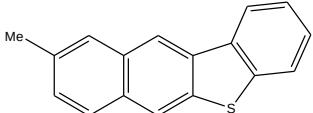
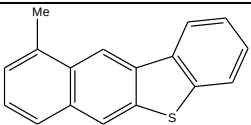
شماره	نام	ساختار	RI
۹۴	1,4,7-Trimethyl dibenzothiophene		۴/۷۱
۹۵	2,3,7-Trimethyl dibenzothiophene		۴/۵۶
۹۶	2,3,8-Trimethyl dibenzothiophene		۴/۴۹
۹۷	2,4,6-Trimethyl dibenzothiophene		۴/۷۵
۹۸	2,4,7-Trimethyl dibenzothiophene		۴/۷۳
۹۹	2,4,8-Trimethyl dibenzothiophene		۴/۶۵
۱۰۰	3,4,7-Trimethyl dibenzothiophene		۴/۶۶
۱۰۱	2-Methyl naphtha[1,2-b]thiophene		۳/۵۹
۱۰۲	4-Methyl naphtha[1,2-b]thiophene		۳/۴۸

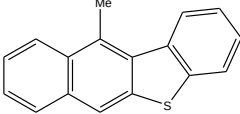
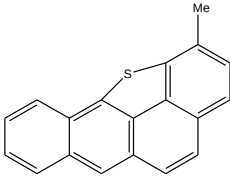
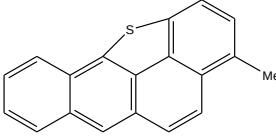
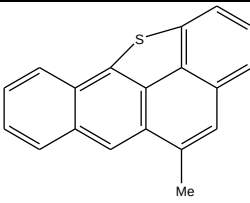
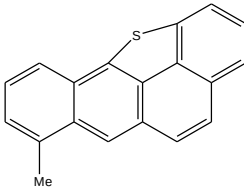
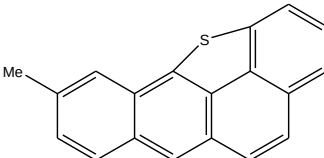
شماره	نام	ساختار	RI
۱۰۳	5-Methyl naphtha[1,2-b]thiophene		۳/۴۷
۱۰۴	6-Methyl naphtha[1,2-b]thiophene		۳/۴۵
۱۰۵	7-Methyl naphtha[1,2-b]thiophene		۳/۵۱
۱۰۶	8-Methyl naphtha[1,2-b]thiophene		۳/۴۲
۱۰۷	9-Methyl naphtha[1,2-b]thiophene		۴/۰۲
۱۰۸	1-Methyl naphtha[2,1-b]thiophene		۳/۳۰
۱۰۹	2-Methyl naphtha[2,1-b]thiophene		۳/۳۹
۱۱۰	4-Methyl naphtha[2,1-b]thiophene		۳/۳۸

شماره	نام	ساختار	RI
۱۱۱	5-Methyl naphtha[2,1-b]thiophene		۳/۳۳
۱۱۲	6-Methyl naphtha[2,1-b]thiophene		۳/۲۹
۱۱۳	7-Methyl naphtha[2,1-b]thiophene		۳/۳۴
۱۱۴	8-Methyl naphtha[2,1-b]thiophene		۳/۲۷
۱۱۵	9-Methyl naphtha[2,1-b]thiophene		۳/۳۱
۱۱۶	Benzo[b]naphtha[1,2-d]thiophene		۴/۰۰
۱۱۷	1-Methyl Benzo[b]naphtha[1,2-d]thiophene		۴/۲۳
۱۱۸	2-Methyl Benzo[b]naphtha[1,2-d]thiophene		۴/۴۴

شماره	نام	ساختار	RI
۱۱۹	3-Methyl Benzo[b]naphtha[1,2-d]thiophene		۴/۵۴
۱۱۲۰	4-Methyl Benzo[b]naphtha[1,2-d]thiophene		۴/۵۲
۱۲۱	5-Methyl Benzo[b]naphtha[1,2-d]thiophene		۴/۵۶
۱۲۲	6-Methyl Benzo[b]naphtha[1,2-d]thiophene		۴/۶۸
۱۲۳	8-Methyl Benzo[b]naphtha[1,2-d]thiophene		۴/۷۰
۱۲۴	9-Methyl Benzo[b]naphtha[1,2-d]thiophene		۴/۶۴
۱۲۵	10-Methyl Benzo[b]naphtha[1,2-d]thiophene		۴/۴۹
۱۲۶	11-Methyl Benzo[b]naphtha[1,2-d]thiophene		۴/۳۷
۱۲۷	Benzo[b]naphtha[2,1-d]thiophene		۴/۱۳

شماره	نام	ساختار	RI
۱۲۸	1-Methyl Benzo[b]naphtha[2,1-d]thiophene		۴/۷۷
۱۲۹	3-Methyl Benzo[b]naphtha[2,1-d]thiophene		۴/۷۹
۱۳۰	4-Methyl Benzo[b]naphtha[2,1-d]thiophene		۴/۷۵
۱۳۱	5-Methyl Benzo[b]naphtha[2,1-d]thiophene		۴/۶۸
۱۳۲	6-Methyl Benzo[b]naphtha[2,1-d]thiophene		۴/۷۴
۱۳۳	7-Methyl Benzo[b]naphtha[2,1-d]thiophene		۴/۷۴
۱۳۴	8-Methyl Benzo[b]naphtha[2,1-d]thiophene		۴/۷۵
۱۳۵	9-Methyl Benzo[b]naphtha[2,1-d]thiophene		۴/۸۲

شماره	نام	ساختار	RI
۱۳۶	10-Methyl Benzo[b]naphtha[2,1-d]thiophene		۴/۸۴
۱۳۷	Benzo[b]naphtha[2,3-d]thiophene		۴/۰۵
۱۳۸	1-Methyl Benzo[b]naphtha[2,3-d]thiophene		۴/۴۸
۱۳۹	2-Methyl Benzo[b]naphtha[2,3-d]thiophene		۴/۷۵
۱۴۰	3-Methyl Benzo[b]naphtha[2,3-d]thiophene		۴/۶۳
۱۴۱	4-Methyl Benzo[b]naphtha[2,3-d]thiophene		۴/۷۱
۱۴۲	6-Methyl Benzo[b]naphtha[2,3-d]thiophene		۴/۵۸
۱۴۳	7-Methyl Benzo[b]naphtha[2,3-d]thiophene		۴/۴۹
۱۴۴	8-Methyl Benzo[b]naphtha[2,3-d]thiophene		۴/۶۲
۱۴۵	9-Methyl Benzo[b]naphtha[2,3-d]thiophene		۴/۵۹
۱۴۶	10-Methyl Benzo[b]naphtha[2,3-d]thiophene		۴/۴۴

شماره	نام	ساختار	RI
۱۴۷	11-Methyl Benzo[b]naphtha[2,3-d]thiophene		۴/۵۱
۱۴۸	1-Methyl Tetrapheno[1,12-bcd]thiophene		۵/۷۹
۱۴۹	3-Methyl Tetrapheno[1,12-bcd]thiophene		۵/۶۷
۱۵۰	5-Methyl Tetrapheno[1,12-bcd]thiophene		۵/۶۰
۱۵۱	7-Methyl Tetrapheno[1,12-bcd]thiophene		۵/۵۸
۱۵۲	9-Methyl Tetrapheno[1,12-bcd]thiophene		۵/۲۸

۳-۴- بهینه‌سازی ساختار هندسی مولکول‌ها و محاسبه توصیف‌کننده‌ها

ابتدا ساختار مولکول‌ها با استفاده از نرم‌افزار Hyperchem رسم و با روش نیمه تجربی AM1 برای رسیدن جذر میانگین مربعات گرادیان انرژی به ۰/۰۰۱ کیلوکالری بر مول بهینه شد. سپس به‌وسیله نرم‌افزار Dragon، ۳۲۲۴ توصیف‌کننده با استفاده از این ساختارهای بهینه‌شده محاسبه شد. و روش رگرسیون گام‌به‌گام برای انتخاب بهترین توصیف‌کننده‌ها استفاده شد که به آن اشاره می‌شود.

۳-۴-۱- انتخاب بهترین توصیف‌کننده‌ها به روش رگرسیون گام‌به‌گام

با توجه به اینکه تعداد زیاد توصیف‌کننده باعث پیچیدگی محاسبات شده و تعدادی از توصیف‌کننده‌ها حاوی اطلاعات یکسانی هستند، تعداد توصیف‌کننده‌ها باید به نوعی کاهش یابد. به همین منظور با استفاده از correlation توصیف‌کننده‌هایی که همبستگی بالای ۰/۹ داشتند حذف شدند و توصیف‌کننده‌هایی که برای تمام ملکول‌ها مقادیر ثابت یا تقریباً ثابت داشتند، حذف شدند و در نهایت از ۳۲۲۴ توصیف‌کننده، ۱۹۲ توصیف‌کننده باقی ماند. با به کارگیری نرم‌افزار SPSS و با اجرای رگرسیون مرحله‌ای برای ۱۹۲ توصیف‌کننده یاد شده به عنوان متغیرهای مستقل، و اندیس بازداري به عنوان متغیر وابسته، تعداد ۱۰ توصیف‌کننده انتخاب شدند که نام و طبقه‌بندی آن‌ها در جدول (۳-۲) نشان داده شده است. همچنین جدول (۳-۳) ماتریس همبستگی بین این توصیف‌کننده‌ها را ارائه می‌دهد، نتایج این جدول نشان می‌دهد که بین این توصیف‌کننده‌ها همبستگی معناداری وجود ندارد.

جدول (۳-۲)-توصیف‌کننده‌های انتخاب‌شده با روش رگرسیون گام‌به‌گام

NO.	Symbol	Class	Meaning
۱	X3v	Connectivity indices	Valance connectivity index chi-3
۲	PCWTe	Charge descriptors	Partial charge weighted topological electronic descriptor
۳	F04CS	2D-frequency fingerprints	Frequency of c-s at topological distance 04
۴	RDF115e	RDF	Radial Distribution Function -11.5/ weighted by aromatic Sanderson electronegativities
۵	Mor19m	3D-MoRSE descriptors	3D-MoRSE-Signal 19/ weighted by aromatic masses
۶	PW3	Topological descriptors	Path/walk3-Randic shape index
۷	H6m	GETAWAY descriptors	H autocorrelation of lag 6/ weighted by aromatic masses
۸	R6u	GETAWAY descriptors	R autocorrelation of lag 6/unweighted
۹	R2m	GETAWAY descriptors	R autocorrelation of lag 2/ weighted by aromatic masses
۱۰	Mor29m	3D-MoRSE descriptors	3D-MoRSE-Signal 29/ weighted by aromatic masses

جدول (۳-۳) - ماتریس همبستگی توصیف‌کننده‌های انتخاب‌شده توسط روش رگرسیون گام‌به‌گام

	x3v	PCWTe	F04CS	Mor19m	PW3	H6m	RDF115e	R6u	R2m	Mor29m
x3v	۱									
PCWTe	۰/۵۷۵	۱								
F04CS	۰/۴۹۹	۰/۳۷۴	۱							
Mor19m	۰/۳۸۹	۰/۵۷	۰/۳۳۵	۱						
PW3	۰/۷۰۶	۰/۰۹۹	۰/۴۴۷	۰/۵۶۹	۱					
H6m	-۰/۵۴۲	۰/۴۲۹	۰/۱۸۴	۰/۰۳۷	۰/۲۱۶	۱				
RDF115e	۰/۶۲۱	۰/۴۷۰	۰/۲۳۷	۰/۰۸۳	۰/۲۲۶	۰/۷۵۲	۱			
R6u	۰/۲۲۹	-۰/۰۱۷	-۰/۱۳۴	-۰/۳۳۴	-۰/۱۳۲	-۰/۱۸۴	-۰/۱۴۰	۱		
R2m	-۰/۳۳۵	-۰/۵۴۱	-۰/۵۰۱	-۰/۲۶۱	-۰/۱۹۷	-۰/۳۰۳	-۰/۲۷۵	۰/۲۶۴	۱	
Mor29m	-۰/۱۰۳	-۰/۰۲۵	۰/۰۱۳	۰/۴۱	-۰/۱۶۹	-۰/۱۷۴	۰/۲۳۵	-۰/۰۱	-۰/۲۷۶	۱

۳-۵- مدل‌سازی و بهینه‌سازی پارامترهای مؤثر بر روش جنگل تصادفی

برای مدل‌سازی توسط جنگل‌های تصادفی، ابتدا باید پارامترهای مؤثر بهینه شوند. برای انجام این کار

ابتدا ترکیبات به دو سری آموزش (۱۲۲ ترکیب) و سری آزمون (۳۰ ترکیب) تقسیم شدند. مقادیر

مربوط به ۱۹۲ توصیف‌کننده به عنوان متغیر مستقل و مقادیر اندیس بازداري به عنوان متغیر وابسته

منسوب شدند. داده‌های سری OOB نیز برای بهینه‌سازی تعداد درختان (ntree)، تعداد

توصیف‌کننده‌های انتخاب‌شده در هر مرحله افراز (Mtry) و تعداد مشاهدات باقی‌مانده در هر گره

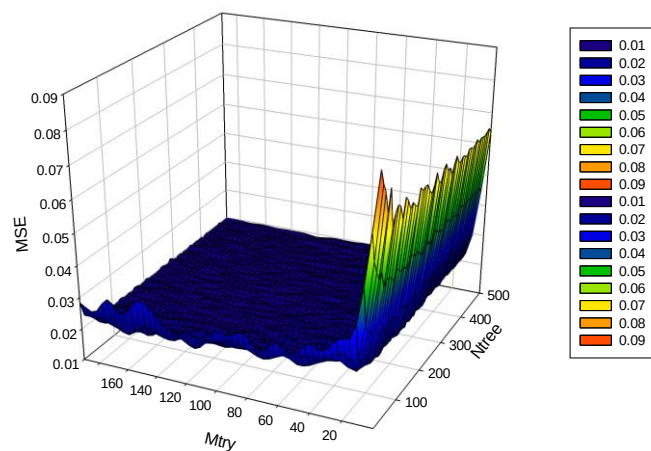
(Node Size) در نظر گرفته شدند. برای بهینه نمودن، تعداد درخت از ۱۰ تا ۵۰۰ با گام ده، تعداد

توصیف‌کننده‌های انتخاب‌شده در هر مرحله افراز (Mtry) از ۱ تا ۱۹۲ با گام ده و تعداد مشاهدات

باقی‌مانده در هر گره از ۱ تا ۵ تغییر داده شد و در هر مرحله مقدار خطای مربوط به مجموعه OOB

محاسبه گردید. نتایج بهینه‌سازی پارامترهای تعداد درختان و تعداد توصیف‌کننده‌های وارد شده در

هر مرحله افزار در گره بهینه در شکل (۳-۱) آمده است. شکل نشان می‌دهد که با بالا رفتن $mtry$ خطا کمتر شده تا جایی که مقدار خطا ثابت می‌شود که در قسمت آبی رنگ قابل مشهود است و با پائین آمدن $mtry$ خطا بیشتر شده به طوری که به ماکسیمم مقدار خود می‌رسد. همچنین برای متغیرهای وارد شده جدول (۳-۴) چندین حالت که کمترین MSE را برای OOB دارد، نشان می‌دهد.



شکل (۳-۱)-نمودار تغییرات خطای OOB در مقادیر متفاوت $Mtry$ و $Ntree$

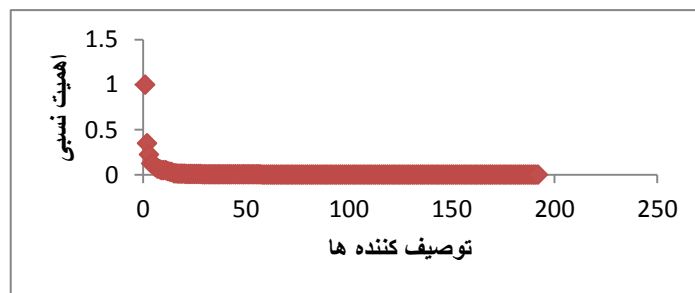
جدول (۳-۴)-کمترین مقادیر MSE همراه با $Mtry$ و $ntree$ متناظر با آن‌ها

$ntree$	$Mtry$	Node Size	MSE Of OOB
۱۷۰	۱۱۱	۱	۰/۰۱۶۳
۱۷۰	۱۱۱	۲	۰/۰۱۶۲
۱۵۰	۹۱	۳	۰/۰۱۶۹
۱۱۰	۱۷۱	۴	۰/۰۱۷۵
۱۹۰	۱۴۱	۵	۰/۰۱۸۷

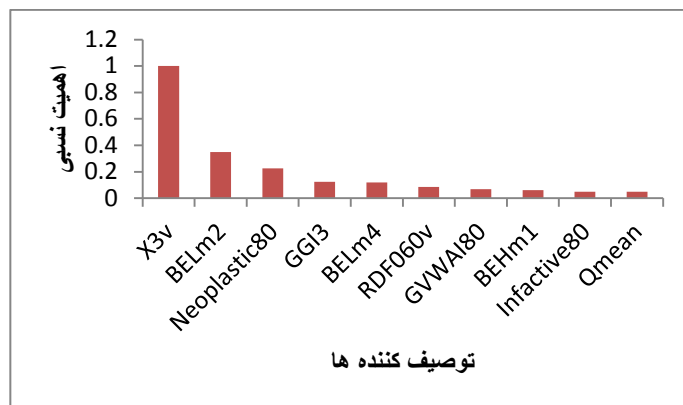
با توجه به جدول (۳-۴) و بر اساس مدل‌سازی جنگل‌های تصادفی با ۱۹۲ توصیف‌کننده ورودی، شرایط بهینه $N_{tree}=170$ و $M_{try}=111$ و $Nodesize=2$ به دست آمد و همانطور که می‌شود در این حالت مقدار MSE مینیمم است. این مدل با توجه به ارزیابی‌های پیش‌رو می‌تواند با مدل‌های برتر در رقابت باشد.

۳-۶- تعیین اهمیت نسبی توصیف‌کننده‌ها با روش جنگل تصادفی

همان‌طور که در بخش قبل توضیح داده شد، روش جنگل تصادفی علاوه بر اینکه به‌عنوان یک روش مدل‌سازی به کار می‌رود، می‌تواند به‌عنوان یک روش انتخاب متغیر نیز به‌کاربرده شود. پس از ساخت مدل جنگل‌های تصادفی و بهینه کردن پارامترهای مؤثر بر جنگل‌های تصادفی، با وارد کردن ۱۹۲ توصیف‌کننده به عنوان توصیف‌کننده ورودی در جنگل تصادفی طبق شرایط بهینه منتخب که در بخش ۳-۵ بدست آمد اهمیت نسبی توصیف‌کننده‌های ورودی محاسبه گردید. این توصیف‌کننده‌های ورودی براساس اهمیت چیده شدند. با توجه به شکل (۳-۲) از توصیف‌کننده ۱۰ به بعد مقدار اهمیت نسبی ثابت می‌شود به همین خاطر پراهمیت‌ترین آن‌ها از توصیف‌کننده اول تا دهم انتخاب می‌شود و به همین منظور ۱۰ توصیف‌کننده برتر به عنوان ورودی شبکه عصبی انتخاب می‌شوند. شکل (۳-۳) میزان اهمیت نسبی ۱۰ توصیف‌کننده را نمایش می‌دهد. همان‌طور که مشاهده می‌شود بسیاری از توصیف‌کننده‌ها دارای اهمیت نسبی بالایی هستند و توصیف‌کننده X3v دارای بالاترین اهمیت نسبی نسبت به سایر توصیف‌کننده‌ها است. ۱۰ توصیف‌کننده مورد نظر همراه با مقدار اهمیت و نوع طبقه آن‌ها در جدول (۳-۵) نشان داده شده‌اند. همچنین ماتریس همبستگی این توصیف‌کننده‌ها در جدول (۳-۶) آمده است که نشان‌دهنده عدم همبستگی معنی‌دار بین آن‌ها می‌باشد.



شکل (۳-۲)- نمودار اهمیت نسبی ۱۹۲ توصیف‌کننده



شکل (۳-۳)- نمودار میزان اهمیت نسبی توصیف کننده‌ها

جدول (۳-۵)- توصیف کننده‌های انتخاب شده توسط جنگل تصادفی

NO.	Symbol	class	Meaning	Importance	Relative Importance
۱	X3v	Connectivity indices	Valance connectivity index chi-3	۲۴/۶۷	۱
۲	BELm2	Burden eigen values	Lowest eigenvalue n.2 of Burden matrix/weighted by atomic masses	۸/۶۲	۰/۳۴۹
۳	Neoplastic80	Molecular properties	Ghose-viswanadhan-wendoloski antineoplastic-like- index at 80%	۵/۵۶۰	۰/۲۲۵
۴	GGI3	Topological charge indices	Topological charge index of order 3	۳/۰۴۷	۰/۱۲۳
۵	BELm4	Burden eigen values	highest eigenvalue n.1 of Burden matrix/weighted by atomic masses	۲/۹۵۴	۰/۱۱۹
۶	RDF060v	RDF	Radial Distribution Function -6.0/ weighted by aromatic van der waals volumes	۲/۱۳	۰/۰۸۶۴
۷	GVWAI80	Molecular properties	Ghose-viswanadhan-wendoloski drug-like- index at 80%	۱/۶۷۱	۰/۰۶۷۷
۸	BEHm1	Burden eigen values	highest eigenvalue n.1 of Burden matrix/weighted by atomic masses	۱/۵۱۸	۰/۰۶۱۵
۹	infective80	Molecular properties	Ghose-viswanadhan-wendoloski antiinfective-like index at 80%	۱/۲۳۳	۰/۰۴۹۹
۱۰	Qmean	Charge descriptors	Mean absolute charge	۱/۲۲۴	۰/۰۴۹۶

جدول (۳-۶) - ماتریکس همبستگی توصیف‌کننده‌ها

	X3v	BELm2	Neoplastic80	GGI3	BELm4	RDF060v	GVWAI80	BEHm1	infective80	Qmean
X3v	۱									
BELm2	۰/۷۷۵	۱								
Neoplastic80	-۰/۵۹۱	-۰/۷۶۴	۱							
GGI3	۰/۸۳۴	۰/۷۶۶	-۰/۷۰۷	۱						
BELm4	۰/۸۴۸	۰/۶۴۳	-۰/۵۱۳	۰/۷۸۲	۱					
RDF060v	۰/۸۱۸	۰/۷۳۱	-۰/۵۶۰	۰/۷۶۲	۰/۷۶۵	۱				
GVWAI80	-۰/۷۴۳	-۰/۵۶۸	۰/۲۹۴	-۰/۵۵۸	-۰/۶۵۲	-۰/۶۲۷	۱			
BEHm1	۰/۷۴۵	۰/۸۵۳	-۰/۷۶۷	۰/۸۰۸	۰/۵۳۴	۰/۷۴۸	-۰/۴۴۹	۱		
infective80	-۰/۶۷۵	-۰/۶۸۲	۰/۵۹۷	-۰/۶۵۸	-۰/۶۴۰	-۰/۶۱۶	۰/۴۹۲	-۰/۶۲۳	۱	
Qmean	-۰/۷۶۳	-۰/۸۴۳	۰/۸۱۸	-۰/۷۹۳	-۰/۶۵۹	-۰/۶۶۰	۰/۴۵۰	-۰/۸۲۴	۰/۷۰۲	۱

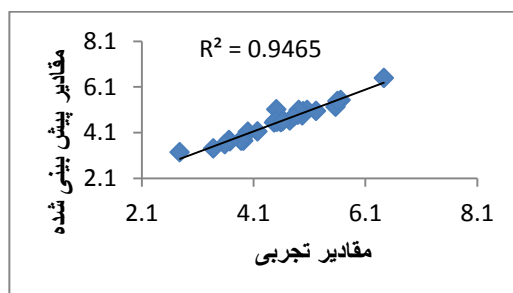
۳-۵-۱- ارزیابی مدل جنگل‌های تصادفی

۳-۵-۱-۱- ارزیابی مدل با استفاده از داده‌های سری آزمون

اهمیت مدل پیش‌بینی وقتی مشخص می‌گردد که مدل حاصله بتواند خاصیت مولکول‌هایی را که در مدل‌سازی به کار نرفته‌اند پیش‌بینی کند. بدین منظور مدل منتخب، برای پیش‌بینی اندیس بازداری داده‌های سری آزمون به کار برده شد. جدول (۳-۷) نتایج پیش‌بینی حاصل از ارزیابی مدل‌های جنگل تصادفی (RF) با استفاده از داده‌های سری آزمون را نشان می‌دهد. نتایج بیانگر این است که مدل ارائه شده خطای پیش‌بینی کمی دارد. شکل (۳-۴) نمودار مقادیر پیش‌بینی‌شده در مقابل مقادیر تجربی را برای داده‌های سری آزمون نشان می‌دهد. نتایج بیانگر این است که مدل ارائه شده خطای پیش‌بینی کمی دارند. ضریب تعیین بالای مدل با ۱۹۲ توصیف‌کننده ورودی، برتری قابل ملاحظه‌ای مدل را در پیش‌بینی اندیس بازداری مربوط به داده‌های سری آزمون نشان می‌دهد.

جدول (۷-۳) نتایج حاصل از ارزیابی مدل جنگل تصادفی با استفاده از داده‌های آزمون

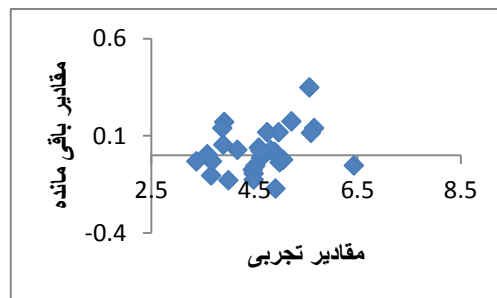
درصد خطا	مقدار پیش‌بینی‌شده	مقدار تجربی	شماره ترکیب
۱۶/۴۱	۳/۲۳	۲/۷۸	۳
-۴/۳۳	۳/۷۴	۳/۹۲	۶
۳/۲۲	۴/۱۲	۴	۹
-۱/۳۹	۳/۸۴	۳/۹	۱۲
۲/۸۵	۳/۷۶	۳/۶۶	۱۵
-۳/۵۷	۳/۷۴	۳/۸۸	۱۸
۱/۶۵	۴/۵۴	۴/۴۷	۲۱
-۰/۸۴	۴/۵۵	۴/۵۹	۲۴
-۲/۳۸	۴/۸۵	۴/۹۷	۲۷
۰/۷۵	۵/۰۲	۴/۹۹	۳۰
۰/۴۵	۵/۰۸	۵/۰۶	۳۱
-۰/۳۴	۴/۸۷	۴/۸۹	۳۲
-۳/۳۱	۵/۰۴	۵/۲۲	۳۳
۱۳/۴۵	۵/۱۱	۴/۵۱	۳۵
۳/۴۹	۵/۰۸	۴/۹۱	۳۶
-۲/۴۶	۵/۵۲	۵/۶۶	۴۱
-۶/۲۴	۵/۲۲	۵/۵۷	۵۰
۰/۸۲	۶/۴۸	۶/۴۳	۶۹
۰/۸۶	۳/۷۱	۳/۶۸	۷۴
-۰/۶۴	۴/۱۴	۴/۱۷	۸۵
۲/۱۰	۴/۵۸	۴/۴۹	۹۶
-۲/۴۶	۴/۶۳	۴/۷۵	۹۷
-۰/۱۲	۳/۵۸	۳/۵۹	۱۰۱
۰/۹۲	۳/۴۱	۳/۳۸	۱۱۰
۱/۲۰	۴/۵۷	۴/۵۲	۱۲۰
-۰/۵۴	۴/۷۹	۴/۸۲	۱۳۰
۰/۲۹	۴/۶۴	۴/۶۳	۱۴۰
۲/۷۸	۴/۶۱	۴/۴۹	۱۴۳
-۰/۵۰	۴/۵۹	۴/۶۲	۱۴۴
-۲/۰۵	۵/۴۸	۵/۶	۱۵۰



شکل (۴-۳) - نمودار تغییرات مقادیر پیش‌بینی شده در مقابل مقادیر تجربی توسط مدل RF برای داده‌های سری آزمون

۳-۵-۱-۲- ارزیابی مدل با استفاده از نمودار خطای باقی مانده

اختلاف مقادیر پیش‌بینی شده و مقادیر تجربی، باقی مانده نامیده می‌شود. توزیع متقارن داده‌ها حول محور افقی (خطای صفر) حاکی از عدم وجود خطای سیستماتیک است. نمودار خطای باقی مانده برحسب مقادیر تجربی، برای مدل‌های ذکر شده در شکل (۳-۵) نشان داده شده است.



شکل (۳-۵) - نمودار خطای باقی مانده برحسب مقادیر تجربی با استفاده از داده‌های سری آزمون توسط روش جنگل تصادفی

۳-۶- مدل سازی شبکه عصبی مصنوعی با توصیف کننده‌های انتخاب شده توسط

روش رگرسیون گام به گام (SR-ANN)

ابتدا ۱۰ توصیف کننده انتخاب شده باروش رگرسیون گام به گام به عنوان متغیرهای ورودی و خاصیت متناظر آن‌ها به عنوان متغیر وابسته در نظر گرفته شد تا پاسخ شبکه با آن‌ها سنجیده شود. برای به دست آوردن بهترین مدل با کمترین خطا، پارامترهای مؤثر تابع انتقال، تابع آموزش، تعداد متغیرهای ورودی شبکه، تعداد گره‌ها در لایه پنهان و تعداد دوره‌های آموزشی به طور همزمان بهینه شدند. تابع کارایی شبکه نیز خطای مطلق میانگین (MAE) در نظر گرفته شد.

در فرآیند بهینه سازی پارامترهای شبکه، سری داده‌ها به سه بخش سری آزمون (۳۰ ترکیب) و سری ارزیابی (۳۰ ترکیب) و سری آموزش (۹۲ ترکیب) تقسیم شدند. سری ارزیابی برای محاسبه گرادینان و به روز کردن اوزان و بایاس‌های شبکه استفاده شد. برای این منظور هر شبکه با ورودی از ۲

تا ۱۰ با گام یک به عنوان توصیف‌کننده ورودی و با دو الگوریتم آموزشی لونبرگ-مارکوات^۱ و تنظیم بایزن^۲ و دو تابع انتقال لگاریتم سیگموئید و تانژانت سیگموئید و تعداد گره از ۲ تا ۱۰ با گام یک و تعداد دوره‌های آموزش از ۵ تا ۵۰ با گام ۵، به‌طور هم‌زمان آموزش داده شد. در روند بهینه‌سازی فوق، حداقل میانگین مربعات سری ارزیابی به‌عنوان معیار انتخاب شد. نتیجه بهینه‌سازی این پارامترها در جدول (۸-۳) گردآوری شده است.

جدول (۸-۳)- توابع و پارامترهای بهینه شبکه عصبی مصنوعی

R	تعداد دوره آموزش	MSE	تعداد گره	تابع انتقال	تابع آموزش	تعداد توصیف‌کننده
۰/۹۸۴	۱۰	۰/۰۲۳۱	۷	لگاریتم-سیگموئید	لونبرگ-مارکوات	۶
۰/۹۸۶	۱۵	۰/۰۲۱۹	۸	تانژانت-سیگموئید	لونبرگ-مارکوات	۶
۰/۹۸۷	۱۵	۰/۰۲۵	۴	لگاریتم-سیگموئید	تنظیم بایزن	۷
۰/۹۸۶	۲۵	۰/۰۲۲۶	۴	تانژانت-سیگموئید	تنظیم بایزن	۶

با توجه به نتایج به‌دست آمده، الگوریتم آموزشی لونبرگ-مارکوات و تابع انتقال تانژانت سیگموئید، ۶ توصیف‌کننده ورودی و ۸ گره در لایه پنهان با تعداد دور آموزشی ۱۵ کمترین MSE را نشان می‌دهد. بنابراین این شبکه برای مدل‌سازی در نظر گرفته شد.

۳-۶-۱- ساختار شبکه عصبی بهینه شده

در جدول (۹-۳) مشخصات شبکه بهینه به دست آمده نشان داده شده است.

جدول (۹-۳)- مشخصات شبکه بهینه (SR-ANN)

تابع آموزش	لونبرگ-مارکوات
تابع انتقال لایه پنهان	تانژانت سیگموئید
تابع انتقال لایه خروجی	Purelin
تعداد گره	۸
تعداد توصیف‌کننده ورودی	۶
تعداد دور آموزشی	۱۵
پارامتر MSE	۰/۲۱۹

¹ Levenberg-Marquart

² bayesian regularization

۳-۶-۲- ارزیابی مدل شبکه عصبی مصنوعی SR-ANN

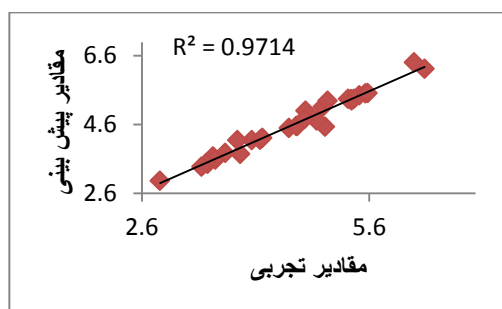
۳-۶-۲-۱- ارزیابی مدل با استفاده از داده‌های سری ارزیابی

بدین منظور، شبکه عصبی با ۹۲ مولکول سری آموزش، آموزش داده شد. معیار بهینه‌سازی شبکه، کمترین خطای مطلق میانگین برای سری ارزیابی است. نتایج حاصل از شبکه عصبی مصنوعی در جدول (۳-۱۰) و ترسیمی از مقادیر پیش‌بینی‌شده برحسب مقادیر تجربی برای سری ارزیابی در شکل (۳-۶) آورده شده است. درصد خطای مشاهده شده و ضریب تعیین به دست آمده نشان از برتری مدل به دست آمده دارد.

جدول (۳-۱۰)- نتایج حاصل از مدل شبکه عصبی مصنوعی برای داده‌های سری ارزیابی (SR-ANN)

درصد خطا	مقدار پیش‌بینی‌شده	مقدار واقعی	شماره ترکیب
۳/۹۱	۲/۹۵	۲/۸۴	۲
-۰/۵۸	۳/۴۴	۳/۴۷	۵
-۴/۰۶	۳/۷۴	۳/۹	۸
۲/۲۵	۴/۱۴	۴/۰۵	۱۱
۱/۸۴	۳/۷۶	۳/۷	۱۴
۷/۱۶	۴/۱۳	۳/۸۶	۱۷
-۹/۸۱	۴/۵۲	۵/۰۲	۲۰
۵/۰۲	۴/۹۹	۴/۷۶	۲۳
-۰/۹۷	۴/۴۹	۴/۵۴	۲۶
-۱/۴۲	۵/۲۹	۵/۳۷	۲۹
-۰/۴۶	۵/۳۳	۵/۳۶	۳۴
۰/۳۶	۵/۳۳	۵/۳۲	۳۷
-۴/۱۹	۴/۷۰	۴/۹۱	۳۸
-۱/۱۲	۵/۳۳	۵/۴۰	۴۰
-۰/۴۳	۵/۴۴	۵/۴۷	۴۲
۲/۳۳	۵/۱۰	۴/۹۹	۴۴
۴/۷۱	۵/۲۸	۵/۰۵	۵۱
-۱/۸۰	۶/۲۱	۶/۳۳	۵۳
۳/۴۹	۶/۴۰	۶/۱۹	۶۲
-۰/۹۸	۵/۴۹	۵/۵۵	۶۳
۳/۵۹	۳/۶۶	۳/۵۴	۷۱
-۰/۰۳	۳/۵	۳/۵۷	۷۲

۸۰	۴/۱۶	۴/۱۴	-۰/۳۸
۹۱	۴/۱۹	۴/۲۰	۰/۲۹
۹۹	۴/۶۵	۴/۵۴	-۱/۲
۱۰۴	۳/۴۵	۳/۴۶	۰/۴۶
۱۰۹	۳/۳۹	۳/۳۶	-۰/۷۳
۱۲۴	۴/۶۴	۴/۵۴	-۱/۹۸
۱۳۰	۴/۷۵	۴/۷۵	۰/۰۷۰
۱۵۱	۵/۵۸	۵/۵۱	-۱/۲۹



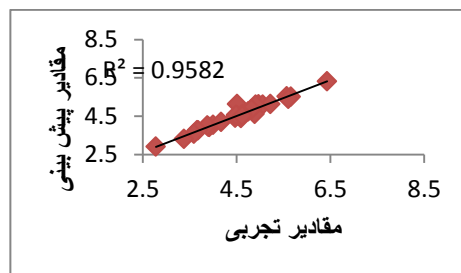
شکل (۳-۶) - نمودار مقادیر پیش‌بینی شده برحسب مقادیر تجربی برای داده‌های سری ارزیابی توسط روش ANN-SR

۳-۶-۲-۲- ارزیابی مدل با استفاده از داده‌های سری آزمون

بعد از انتخاب شبکه عصبی بهینه با استفاده از سری ارزیابی، قدرت پیش‌بینی شبکه به وسیله ۳۰ داده که در سری آزمون قرار دارند مورد بررسی قرار گرفت. نتایج حاصل در جدول (۳-۱۱) نشان داده شده است. همچنین شکل (۳-۷) مقادیر پیش‌بینی شده برحسب مقادیر تجربی برای سری آزمون نشان داده است. با توجه به میانگین خطای بدست آمده که مقدار آن معادل ۰/۸۴۰ است و مقدار ضریب تعیین در SR-ANN برای داده‌های سری ارزیابی نسبت به داده‌های سری آزمون بیشتر بوده که انتظار نیز همین است همچنین میانگین خطای ارزیابی ۰/۱۰۳ است و ضریب تعیین بالاتری دارد.

جدول (۳-۱۱)-نتایج حاصل از مدل شبکه عصبی مصنوعی برای داده‌های سری آزمون (SR-ANN)

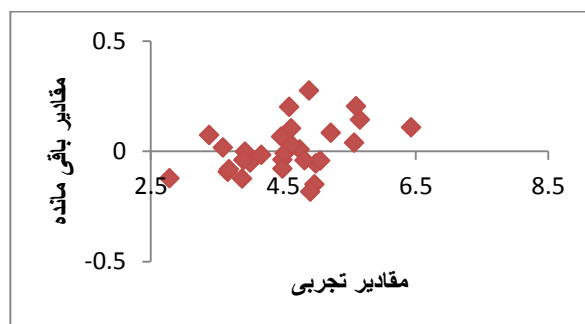
درصد خطا	مقدار پیش‌بینی شده	مقدار تجربی	شماره ترکیب
۴/۴۱	۲/۹۰	۲/۷۸	۳
۰/۰۴۰	۳/۹۲	۳/۹۲	۶
۱/۲۸۰	۴/۰۵۱	۴	۹
۱/۰۷	۳/۹۳	۳/۹	۱۲
۲/۵۳	۳/۷۵	۳/۶۶	۱۵
۳/۱۹	۴/۰۰	۳/۸۸	۱۸
-۱/۴۸	۴/۴۰	۴/۴۷	۲۱
-۴/۳۷	۴/۳۸	۴/۵۹	۲۴
۳/۰۲	۵/۱۲	۴/۹۷	۲۷
۱/۱۱	۵/۰۴	۴/۹۹	۳۰
۰/۸۵	۵/۱۰	۵/۰۶	۳۱
-۵/۶۲	۴/۶۱	۴/۸۹	۳۲
-۱/۶۰	۵/۱۳	۵/۲۲	۳۳
۱۳/۷۲	۵/۱۲	۴/۵۱	۳۵
۳/۷۴	۵/۰۹	۴/۹۱	۳۶
-۲/۵۲	۵/۵۱	۵/۶۶	۴۱
-۰/۶۸	۵/۵۳	۵/۵۷	۵۰
-۱/۶۸	۶/۳۲	۶/۴۳	۶۹
۲/۲۳	۳/۷۶	۳/۶۸	۷۴
۰/۴۰۷	۴/۱۸	۴/۱۷	۸۵
۰/۸۶	۴/۵۳	۴/۴۹	۹۶
-۰/۱۶	۴/۷۴	۴/۷۵	۹۷
-۰/۴۵	۳/۵۷	۳/۵۹	۱۰۱
-۲/۱۸	۳/۳۰	۳/۳۸	۱۱۰
۰/۲۰	۴/۵۲	۴/۵۲	۱۲۰
۰/۸۴	۴/۸۶	۴/۸۲	۱۳۰
-۰/۶۱	۴/۶۰	۴/۶۳	۱۴۰
۱/۷۶	۴/۵۶	۴/۴۹	۱۴۳
-۲/۲۵	۴/۵۱	۴/۶۲	۱۴۴
-۳/۶۳	۵/۳۹	۵/۶	۱۵۰



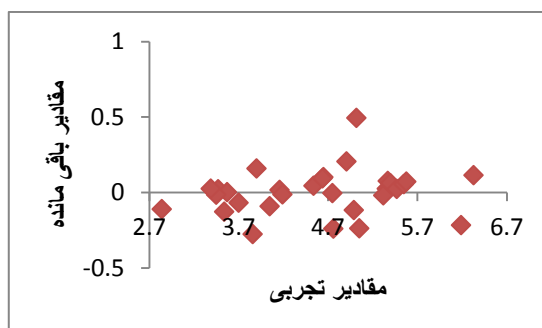
شکل (۷-۳)-نمودار مقادیر پیش‌بینی شده اندیس بازداري برحسب مقادیر تجربی برای داده‌های سری آزمون توسط روش SR-ANN

۳-۶-۲-۳- ارزیابی مدل SR-ANN توسط نمودار خطای باقی‌مانده

شکل‌های (۸-۳) و (۹-۳) نمودارهای خطای باقی‌مانده برحسب مقادیر تجربی سری‌های آزمون و ارزیابی برای مدل SR-ANN را نشان می‌دهد. توزیع متقارن داده‌ها حول محور افقی حاکی از عدم وجود خطای سیستماتیک است.



شکل (۸-۳)-نمودار خطای باقی‌مانده برحسب مقادیر تجربی برای داده‌های آزمون (SR-ANN)



شکل (۹-۳)- نمودار خطای باقی‌مانده برحسب مقادیر تجربی برای داده‌های سری ارزیابی (SR-ANN)

۳-۶-۲-۴- ارزیابی مدل SR-ANN توسط روش ارزیابی متقاطع

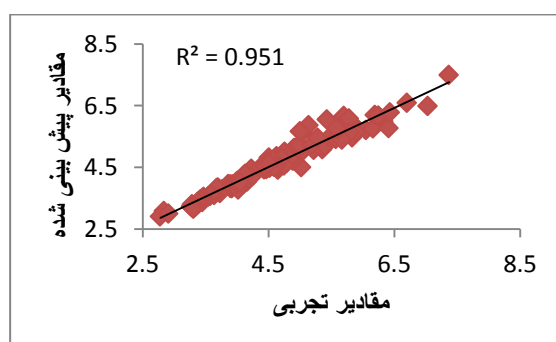
در این روش هر بار یک ترکیب از ۱۵۲ ترکیب به عنوان داده‌های آزمون کنار گذاشته شد و در شرایط بهینه به دست آمده برای سری آموزش با ترکیبات باقی مانده مدل سازی صورت گرفت. سپس مدل به دست آمده برای پیش بینی اندیس بازداری ترکیب کنار گذاشته شده به کار گرفته شد و این فرآیند برای تک تک ترکیبات تکرار گردید. نتایج حاصل از این روش در جدول (۳-۱۲) ارائه شده است. این نتایج تعمیم پذیری بالای مدل طراحی شده را بیان می کند. همچنین ضریب همبستگی مشاهده شده بین مقادیر پیش بینی شده ی اندیس بازداری و مقادیر تجربی در شکل (۳-۱۰)، بر نزدیکی مقادیر پیش بینی شده توسط مدل نسبت به مقادیر تجربی آن ها دلالت دارد. درصد خطاهای مشاهده شده برای ترکیبات سولفور چند حلقه ای زیر ۱۵ درصد می باشد که نشان می دهد مدل مورد نظر صحت خوب و قابل قبولی را داراست، همچنین در شکل (۳-۱۱) مقادیر باقی مانده بر حسب مقدار تجربی اندیس بازداری ترکیبات مورد بحث ترسیم شده است که تقارن پراکندگی نقاط در دو طرف محور افقی عدم وجود خطای سیستماتیک را برای بیشتر ترکیبات نشان می دهد. طبق تکنیک ارزیابی متقاطع این مدل می تواند به عنوان یک مدل مناسب معرفی شود.

جدول (۳-۱۲)- نتایج حاصل از ارزیابی مدل SR-ANN با استفاده از ارزیابی متقاطع

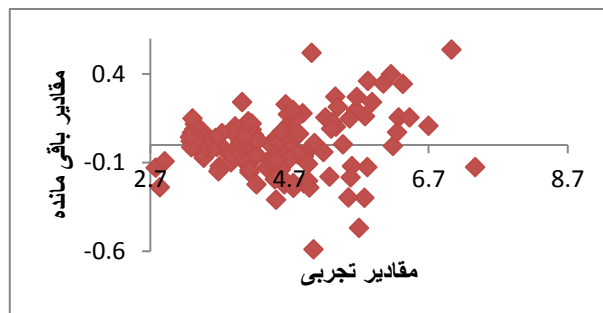
NO.	مقدار تجربی	مقدار پیش بینی	درصد خطا	NO.	مقدار تجربی	مقدار پیش بینی	درصد خطا
۱	۳/۲۸	۳/۲۶	-۰/۵۸	۸۱	۴/۱۶	۴/۰۴	-۲/۸۷
۲	۲/۸۴	۳/۰۸	۸/۴۶	۸۲	۱/۴	۴/۲۱	۲/۷۸
۳	۲/۷۸	۲/۹۱	۴/۶۵	۸۳	۴	۳/۹۸	-۰/۴۰
۴	۲/۹۱	۳/۰۰	۳/۱۹	۸۴	۴/۰۱	۴/۰۶	۱/۲۷
۵	۳/۴۷	۳/۵۴	۲/۱۱	۸۵	۴/۱۷	۴/۲۲	۱/۳۵
۶	۳/۹۲	۳/۸۱	-۲/۶۲	۸۶	۴/۱۷	۴/۱۵	-۰/۲۵
۷	۳/۷۶	۳/۷۸	۰/۶۲	۸۷	۴/۱۲	۴/۱۸	۱/۶۵
۸	۳/۹	۳/۸۳	-۱/۷۱	۸۸	۴/۰۵	۴/۰۶	۰/۴۸
۹	۴	۴/۰۰	۰	۸۹	۱/۴	۴/۱۹	۲/۲۱
۱۰	۴/۱۳	۴/۲۸	۳/۷۶	۹۰	۴/۲۴	۴/۲۹	۱/۳۹
۱۱	۴/۰۵	۴/۰۰	-۱/۰۷	۹۱	۴/۱۹	۴/۱۳	-۱/۲۸
۱۲	۳/۹	۳/۹۴	۱/۱۲	۹۲	۴/۲۵	۴/۲۲	-۰/۵۹

۱۳	۳/۷	۳/۸۳	۳/۷۱	۹۳	۴/۶۶	۴/۶۴	-۰/۲۲
۱۴	۳/۷	۳/۸۲	۳/۲۶	۹۴	۴/۷۱	۴/۶۲	-۱/۸۴
۱۵	۳/۶۶	۳/۶۷	-۰/۳۳	۹۵	۴/۵۶	۴/۵۵	-۰/۱۲
۱۶	۳/۷۳	۳/۶۶	-۱/۷۴	۹۶	۴/۴۹	۴/۵۴	۱/۲۵
۱۷	۳/۸۶	۳/۹۵	۲/۵۷	۹۷	۴/۷۵	۴/۶۵	-۱/۹۴
۱۸	۳/۸۸	۳/۹۳	۱/۵۲	۹۸	۴/۷۳	۴/۶۶	-۱/۵۱
۱۹	۴/۶۳	۴/۸۵	۴/۸۲	۹۹	۴/۶۵	۴/۴۲	-۴/۸۸
۲۰	۵/۰۲	۴/۵۰	-۱۰/۳۲	۱۰۰	۴/۶۶	۴/۷۱	۱/۰۶
۲۱	۴/۴۷	۴/۴۷	-۰/۳۵	۱۰۱	۳/۵۹	۳/۶۱	-۰/۶۲
۲۲	۴/۴۷	۴/۴۴	-۰/۵۴۱	۱۰۲	۳/۴۸	۳/۴۵	-۰/۷۵
۲۳	۴/۷۶	۵/۰۰	۵/۱۰	۱۰۳	۳/۴۷	۳/۴۹	-۰/۸۱
۲۴	۴/۵۹	۴/۵۵	-۰/۶۶	۱۰۴	۳/۴۵	۳/۴۸	۱/۱۵
۲۵	۴/۴۳	۴/۴۲	-۰/۰۰۷	۱۰۵	۳/۴۵	۳/۳۷	-۲/۲۲
۲۶	۴/۵۴	۴/۶۲	۱/۸۲	۱۰۶	۳/۴۲	۳/۴۶	۱/۲۲
۲۷	۴/۹۷	۵/۱۷	۴/۰۹	۱۰۷	۴/۰۲	۳/۷۷	-۵/۹۹
۲۸	۴/۶	۴/۷۹	۴/۱۳	۱۰۸	۳/۳	۳/۲۳	-۲/۰۰
۲۹	۵/۳۷	۵/۲۶	-۱/۹۰	۱۰۹	۳/۳۹	۳/۳۲	-۲/۰۶
۳۰	۴/۹۹	۵/۰۱	-۰/۵۰	۱۱۰	۳/۳۸	۳/۳۸	-۰/۰۶۰
۳۱	۵/۰۶	۵/۰۵	-۰/۱۹	۱۱۱	۳/۳۳	۳/۲۴	-۲/۵۹
۳۲	۴/۸۹	۴/۷۱	-۳/۶۴	۱۱۲	۳/۲۹	۳/۳۰	-۰/۴۲۰
۳۳	۵/۲۲	۵/۰۶	-۲/۹۶	۱۱۳	۳/۳۴	۳/۲۲	-۳/۵۴
۳۴	۵/۳۶	۵/۰۸	-۵/۰۴	۱۱۴	۳/۲۷	۳/۲۲	-۱/۳۴
۳۵	۴/۵۱	۴/۶۹	۴/۱۷	۱۱۵	۳/۳۱	۳/۱۶	-۴/۵۰
۳۶	۴/۹۱	۵/۱۲	۴/۴۵	۱۱۶	۴	۳/۹۸	-۰/۴۹
۳۷	۵/۳۲	۵/۱۹	-۲/۳۶	۱۱۷	۴/۲۳	۴/۴۵	۵/۳۱
۳۸	۴/۹۱	۴/۹۷	۱/۳۷	۱۱۸	۴/۴۴	۴/۵۶	۲/۸۶
۳۹	۵/۳	۵/۲۰	-۱/۷۲	۱۱۹	۴/۵۴	۴/۵۷	-۰/۶۸
۴۰	۵/۴	۵/۱۸	-۳/۹۳	۱۲۰	۴/۵۲	۴/۴۸	-۰/۷۱
۴۱	۵/۶۶	۵/۴۶	-۳/۴۵	۱۲۱	۴/۵۶	۴/۵۴	-۰/۳۰۰
۴۲	۵/۴۷	۵/۴۷	-۰/۰۲۸	۱۲۲	۴/۶۸	۴/۵۱	-۳/۶۵
۴۳	۴/۹۷	۵/۰۷	۲/۱۸	۱۲۳	۴/۷	۴/۵۸	-۲/۵۱
۴۴	۴/۹۹	۵/۲۳	۴/۸۵	۱۲۴	۴/۶۴	۴/۵۶	-۱/۶۰
۴۵	۵/۱۹	۵/۲۳	-۰/۸۰	۱۲۵	۴/۴۹	۴/۵۴	۱/۲۱
۴۶	۴/۷۶	۴/۹۷	۴/۵۳	۱۲۶	۴/۳۷	۴/۴۹	۲/۹۶
۴۷	۴/۶۴	۴/۶۲	-۰/۳۹	۱۲۷	۴/۱۳	۴/۲۷	۳/۴۵
۴۸	۴/۵۱	۴/۸۲	۶/۹۰	۱۲۸	۴/۷۷	۴/۷۳	-۰/۷۹
۴۹	۵/۸۳	۵/۴۶	-۶/۱۷	۱۲۹	۴/۷۹	۴/۶۴	-۳/۰۴
۵۰	۵/۵۷	۵/۴۲	-۲/۶۵	۱۳۰	۴/۷۵	۴/۶۷	-۱/۶۶
۵۱	۵/۰۵	۵/۶۴	۱۱/۶۹	۱۳۱	۴/۶۸	۴/۷۲	۱/۰۳
۵۲	۶/۰۵	۵/۷۰	-۵/۷۰	۱۳۲	۴/۷۴	۴/۷۱	-۰/۵۱

۵۳	۶/۳۳	۵/۹۸	-۵/۴۱	۱۳۳	۴/۷۴	۴/۷۷	۰/۷۲
۵۴	۵/۸۲	۵/۹۴	۲/۱۵	۱۳۴	۴/۷۵	۴/۷۶	۰/۲۵
۵۵	۵/۱۴	۵/۸۶	۱۴/۱۵	۱۳۵	۴/۸۲	۴/۸۷	۱/۱۲
۵۶	۶/۱۶	۵/۷۶	-۶/۴۵	۱۳۶	۴/۸۴	۴/۷۷	-۱/۲۸
۵۷	۵/۷	۶/۱۷	۸/۲۴	۱۳۷	۴/۰۵	۴/۰۰	-۱/۰۷
۵۸	۵	۵/۶۷	۱۳/۴۴	۱۳۸	۴/۴۸	۴/۵۸	۲/۳۹
۵۹	۵/۸۹	۵/۶۴	-۴/۰۸	۱۳۹	۴/۷۵	۴/۵۵	-۴/۱۴
۶۰	۵/۷۸	۶/۰۷	۵/۱۸	۱۴۰	۴/۶۳	۴/۶۹	۱/۴۴
۶۱	۶/۲۷	۶/۱۱	-۲/۴۶	۱۴۱	۴/۷۱	۴/۷۸	۱/۶۴
۶۲	۶/۱۹	۶/۱۹	۰/۱۰۵	۱۴۲	۴/۵۸	۴/۷۲	۳/۰۷
۶۳	۵/۵۵	۵/۸۴	۵/۳۶	۱۴۳	۴/۴۹	۴/۶۸	۴/۳۴
۶۴	۶/۴۱	۵/۷۷	-۹/۹۶	۱۴۴	۴/۶۲	۴/۵۳	-۱/۷۵
۶۵	۶/۲۵	۶/۱۸	-۱/۱۱	۱۴۵	۴/۵۹	۴/۵۱	-۱/۵۹
۶۶	۵/۴۳	۶/۰۵	۱۱/۵۱	۱۴۶	۴/۴۴	۴/۵۳	۲/۰۶
۶۷	۶/۷	۶/۵۹	-۱/۶۰	۱۴۷	۴/۵۱	۴/۵۹	۱/۸۶
۶۸	۷/۰۳	۶/۴۹	-۷/۶۴	۱۴۸	۵/۷۹	۵/۶۲	-۲/۷۸
۶۹	۶/۴۳	۶/۲۷	-۲/۴۰	۱۴۹	۵/۶۷	۵/۳۹	-۴/۷۶
۷۰	۷/۳۷	۷/۴۹	۱/۷۱	۱۵۰	۵/۶	۵/۷۲	۲/۱۶
۷۱	۳/۵۴	۳/۵۵	۰/۳۸	۱۵۱	۵/۵۸	۵/۷۶	۳/۳۱
۷۲	۳/۵۷	۳/۵۷	۰/۰۴۲	۱۵۲	۵/۲۸	۵/۴۶	۳/۴۱
۷۳	۳/۶۴	۳/۶۰	-۱/۰۶				
۷۴	۳/۶۸	۳/۸۲	۴/۰۶				
۷۵	۳/۹۳	۳/۹۴	۰/۴۸				
۷۶	۴	۴/۰۳	۰/۸۶				
۷۷	۴/۰۸	۴/۰۰	-۱/۸۴				
۷۸	۴/۰۱	۴/۰۰	-۰/۱۱				
۷۹	۴/۱۱	۳/۹۷	-۳/۲۲				
۸۰	۴/۱۶	۴/۰۷	-۲/۰۲				



شکل (۳-۱۰) - نمودار مقادیر پیش‌بینی شده بر حسب مقادیر تجربی برای ارزیابی مدل SR-ANN به روش ارزیابی متقاطع برای کل داده‌ها



شکل (۳-۱۱) - نمودار مقادیر باقی‌مانده بر حسب مقادیر تجربی برای ارزیابی مدل SR-ANN به روش ارزیابی متقاطع برای کل داده‌ها

۳-۷- مدل‌سازی شبکه عصبی مصنوعی با توصیف‌کننده‌های انتخاب‌شده

توسط روش RF (RF-ANN)

در این بخش، از شبکه عصبی با یک لایه پنهان و لایه خروجی و توصیف‌کننده‌های منتخب روش جنگل‌های تصادفی استفاده گردید. برای یافتن مقادیر بهینه پارامترهای موثر بر کارایی شبکه، تعداد ورودی‌ها از ۲ تا ۱۰ و با دو الگوریتم آموزشی لونیبرگ-مارکوات (train lm) و تنظیم بایزن (trainbr) انتقال لگاریتم سیگموئید (logsig) و تانژانت سیگموئید (tansig) و تعداد گره از ۲ تا ۱۰ و تعداد دور آموزش از ۵ تا ۵۰ با گام ۵ داده شد. در این مرحله، تنظیمات مختلفی برای همه پارامترها در نظر گرفته شد و خطای مطلق میانگین به عنوان معیار انتخاب شد و با به‌دست آمدن تنظیمات مدل بهینه کمترین مقدار میانگین مربعات خطا محاسبه گردید. جدول (۳-۱۳) بهترین شبکه‌های به دست آمده در حین بهینه‌سازی را نمایش می‌دهد، که در آن شبکه عصبی با ۶ توصیف‌کننده ورودی، ۷ نرون در لایه پنهان، تعداد دورهای آموزش ۱۰، تابع آموزشی لونیبرگ-مارکوات و تابع انتقال تانژانت سیگموئید، کمترین MSE را نسبت به سایر شبکه‌ها داراست و لذا شبکه‌ای با این ساختار می‌تواند به عنوان شبکه بهینه برای مدل‌سازی اندیس بازدارندگی ترکیبات در نظر گرفته شود.

جدول (۳-۱۳)-توابع و پارامترهای بهینه شبکه عصبی مصنوعی

R	MSE	تعداد دور آموزش	تعداد گره	تابع انتقال	تابع آموزش	تعداد توصیف کننده
۰/۹۸۰	۰/۰۴۰۶	۵	۶	لگاریتم-سیگموئید	لونبرگ-مارکوات	۵
۰/۹۷۶	۰/۰۳۴۹	۱۰	۷	تانژانت-سیگموئید	لونبرگ-مارکوات	۶
۰/۹۷۶	۰/۰۴۰۵	۲۰	۹	تانژانت-سیگموئید	تنظیم بایزن	۵
۰/۹۷۷	۰/۰۳۷۲	۳۰	۹	لگاریتم-سیگموئید	تنظیم بایزن	۵

۳-۷-۱-ساختار شبکه عصبی بهینه شده

با توجه به نتایج به دست آمده در مرحله قبل، مشخصات شبکه عصبی مصنوعی در جدول (۳-۱۴) نشان داده شده است.

جدول (۳-۱۴)-مشخصات شبکه بهینه (RF-ANN)

تابع آموزش	لونبرگ-مارکوات
تابع انتقال لایه پنهان	تانژانت سیگموئید
تابع انتقال لایه خروجی	Purelin
تعداد گره	۷
تعداد توصیف کننده ورودی	۶
تعداد دور آموزشی	۱۰
پارامتر MSE	۰/۳۴۹

۳-۷-۲-ارزیابی مدل شبکه عصبی مصنوعی با توصیف کننده های RF

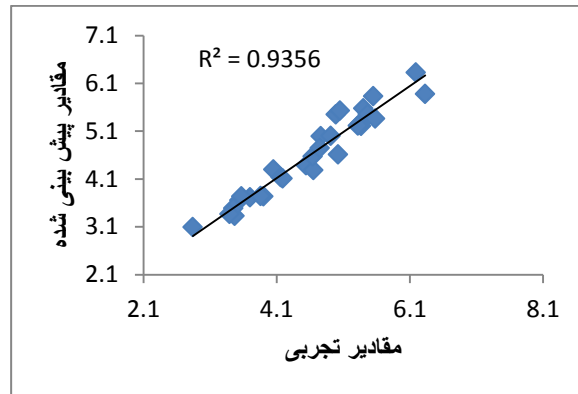
۳-۷-۲-۱-ارزیابی مدل با استفاده از داده های سری ارزیابی

جدول (۳-۱۵) نتایج پیش بینی حاصل از ارزیابی مدل RF-ANN در پیش بینی اندیس بازداري با استفاده از داده های سری ارزیابی و شکل (۳-۱۲) نیز نمودار تغییرات مقادیر پیش بینی شده در مقابل مقادیر تجربی را برای داده های سری ارزیابی نشان می دهد. براساس مشاهدات و ارزیابی های انجام

شده می‌توان استناد کرد که روش SR-ANN برای داده‌های سری ارزیابی دارای ضریب تعیین بالاتر و میانگین درصد خطای کمتری به نسبت RF-ANN است و بنابراین مدل برتر است.

جدول (۳-۱۵)- نتایج حاصل از مدل شبکه عصبی مصنوعی با استفاده از داده‌های سری ارزیابی (RF-ANN)

درصد خطا	مقدار پیش‌بینی شده	مقدار واقعی	شماره ترکیب
۸/۸۱۶	۳/۰۹۰	۲/۸۴	۲
-۴/۲۴۵	۳/۳۲۲	۳/۴۷	۵
-۴/۲۷۶	۳/۷۳۳	۳/۹	۸
۶/۰۸۹	۴/۲۹۶	۴/۰۵	۱۱
۰/۲۹۱	۳/۷۱۰	۳/۷	۱۴
-۳/۰۵۳	۳/۷۴۲	۳/۸۶	۱۷
-۸/۱۴۴	۴/۶۱۱	۵/۰۲	۲۰
۴/۷۹۸	۴/۹۸۸	۴/۷۶	۲۳
-۳/۴۸۵	۴/۳۸۱	۴/۵۴	۲۶
-۳/۰۵۹	۵/۲۰۵	۵/۳۷	۲۹
-۱/۹۰۴	۵/۲۵۷	۵/۳۶	۳۴
-۲/۰۳۰	۵/۲۱۱	۵/۳۲	۳۷
۱/۸۴۶	۵/۰۰۰	۴/۹۱	۳۸
۳/۲۵۶	۵/۵۷۵	۵/۴۰	۴۰
-۰/۳۱۵	۵/۴۵۲	۵/۴۷	۴۲
۹/۱۰۲	۵/۴۴۴	۴/۹۹	۴۴
۹/۴۸۴	۵/۵۲۸	۵/۰۵	۵۱
-۷/۲۱۰	۵/۸۷۳	۶/۳۳	۵۳
۲/۱۲۳	۶/۳۲۱	۶/۱۹	۶۲
۵/۰۱۱	۵/۸۲۸	۵/۵۵	۶۳
۳/۱۵۸	۳/۶۵۱	۳/۵۴	۷۱
۴/۵۹۳	۳/۷۳۳	۳/۵۷	۷۲
-۰/۶۲۸	۴/۱۳۳	۴/۱۶	۸۰
-۱/۹۲۶	۴/۱۰۹	۴/۱۹	۹۱
-۷/۹۱۴	۴/۲۸۱	۴/۶۵	۹۹
۱/۰۷۲	۳/۴۸۷	۳/۴۵	۱۰۴
-۰/۷۹۲	۳/۳۶۳	۳/۳۹	۱۰۹
-۱/۴۸۵	۴/۵۷۱	۴/۶۴	۱۲۴
-۰/۱۲۷	۴/۷۴۳	۴/۷۵	۱۳۰
-۳/۹۹۰	۵/۳۵۷	۵/۵۸	۱۵۱



شکل (۳-۱۲)- نمودار مقادیر پیش‌بینی شده برحسب مقادیر تجربی برای داده‌های ارزیابی توسط روش RF-ANN

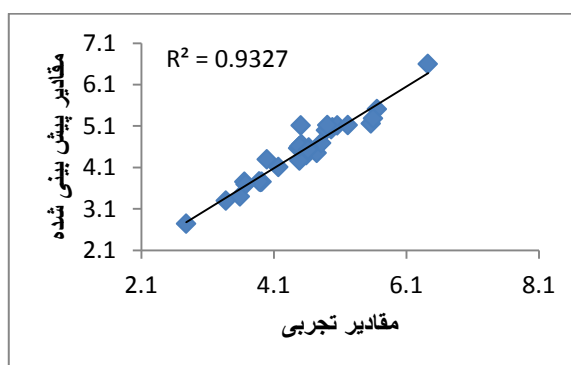
۳-۷-۲-۲- ارزیابی مدل با استفاده از داده‌های سری آزمون

اعتبار و اهمیت مدل با پیش‌بینی اندیس بازداري داده‌هایی که در مدل‌سازی حضور نداشته‌اند مشخص می‌شود. بدین منظور مدل‌های انتخاب شده براساس توصیف‌کننده‌های انتخاب شده توسط RF جهت پیش‌بینی اندیس بازداري ۳۰ ترکیب سری آزمون که در فرآیند مدل‌سازی استفاده نشده‌اند، به کار گرفته شد و با استفاده از شبکه عصبی بهینه شده مقادیر اندیس بازداري این ترکیبات پیش‌بینی شد. جدول (۳-۱۶) نتایج حاصل از ارزیابی مدل شبکه عصبی مصنوعی در پیش‌بینی اندیس بازداري داده‌های سری آزمون را نشان می‌دهد. خطای کم در پیش‌بینی اندیس بازداري ترکیبات مورد بررسی، قدرت پیش‌بینی مدل را نشان می‌دهد. شکل (۳-۱۳) نمودار تغییرات مقادیر پیش‌بینی شده در مقابل مقادیر تجربی را برای داده‌های سری آزمون نشان می‌دهد. برتری مدل SR-ANN نسبت به مدل RF-ANN مشهود است.

جدول (۳-۱۶)- نتایج حاصل از ارزیابی مدل شبکه عصبی مصنوعی با استفاده از داده‌های سری آزمون (RF-ANN)

درصد خطا	مقدار پیش‌بینی شده	مقدار تجربی	شماره ترکیب
-۱/۳۲۴	۲/۷۴۳	۲/۷۸	۳
-۴/۱۸۹	۳/۷۵۵	۳/۹۲	۶
۷/۳۰۰	۴/۲۹۲	۴	۹
-۳/۹۰۸	۳/۷۴۷	۳/۹	۱۲

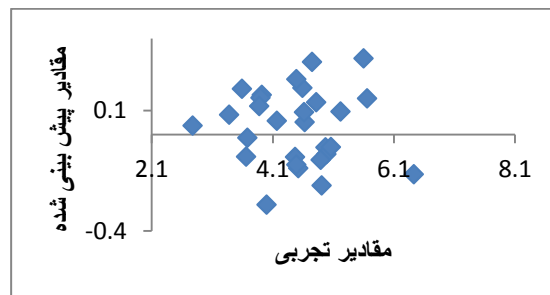
۱۵	۳/۶۶	۳/۷۵۲	۲/۵۲۹
۱۸	۳/۸۸	۳/۷۶۱	-۳/۰۴۲
۲۱	۴/۴۷	۴/۵۶۳	۲/۰۹۹
۲۴	۴/۵۹	۴/۳۹۶	-۴/۲۱۵
۲۷	۴/۹۷	۵/۰۲۳	۱/۰۸۵
۳۰	۴/۹۹	۵/۰۷۰	۱/۶۱۸
۳۱	۵/۰۶	۵/۱۱۱	۱/۰۲۵
۳۲	۴/۸۹	۴/۹۹۵	۲/۱۵۲
۳۳	۵/۲۲	۵/۱۲۵	-۱/۸۱۹
۳۵	۴/۵۱	۵/۱۱۳	۱۳/۳۸۰
۳۶	۴/۹۱	۵/۱۲۲	۴/۳۲۱
۴۱	۵/۶۶	۵/۵۱۰	-۲/۶۴۸
۵۰	۵/۵۷	۵/۱۶۰	-۷/۳۵۳
۶۹	۶/۴۳	۶/۵۹۵	۲/۵۶۷
۷۴	۳/۶۸	۳/۶۹۲	۰/۳۵۱
۸۵	۴/۱۷	۴/۱۱۲	-۱/۳۷۵
۹۶	۴/۴۹	۴/۲۶۰	-۵/۱۰۴
۹۷	۴/۷۵	۴/۴۴۸	-۶/۳۴۱
۱۰۱	۳/۵۹	۳/۳۹۹	-۵/۳۰۳
۱۱۰	۳/۳۸	۳/۲۹۷	-۲/۴۴۲
۱۲۰	۴/۵۲	۴/۶۵۹	۳/۰۹۴
۱۳۰	۴/۸۲	۴/۶۸۵	-۲/۷۸۷
۱۴۰	۴/۶۳	۴/۵۷۹	-۱/۱۰۱
۱۴۳	۴/۴۹	۴/۶۱۵	۲/۷۹۵
۱۴۴	۴/۶۲	۴/۵۲۷	-۲/۰۰۷
۱۵۰	۵/۶	۵/۲۸۴	-۵/۶۴۰



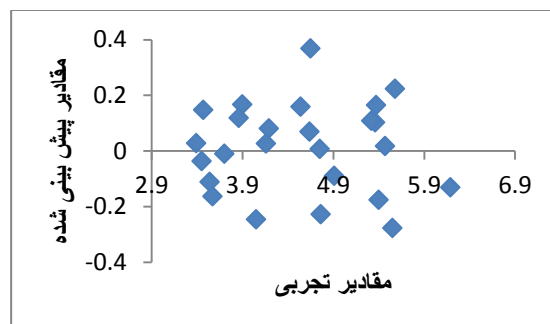
شکل (۳-۱۳)-نمودار مقادیر پیش‌بینی شده‌ی اندیس بازاری برحسب مقادیر تجربی با استفاده از سری داده‌های
آزمون توسط روش RF-ANN

۳-۷-۲-۳- ارزیابی مدل توسط نمودار خطای باقی مانده

شکل‌های (۳-۱۴) و (۳-۱۵) نمودارهای خطای باقی مانده بر حسب مقادیر تجربی، برای سری‌های آزمون و ارزیابی را در مدل RF-ANN نشان می‌دهد. توزیع متقارن داده‌ها حول محور افقی حاکی از عدم وجود خطای سیستماتیک است. مشاهده می‌شود.



شکل (۳-۱۴) نمودار خطای باقی مانده بر حسب مقادیر تجربی برای داده‌های سری آزمون (RF-ANN)



شکل (۳-۱۵) نمودار خطای باقی مانده بر حسب مقادیر تجربی برای داده‌های سری ارزیابی (RF-ANN)

۳-۷-۲-۴- ارزیابی مدل RF-ANN توسط روش ارزیابی متقاطع

از دیگر روش‌هایی که برای ارزیابی مدل ارائه شده توسط ANN به کار رفته است، تکنیک ارزیابی متقاطع می‌باشد که نتایج حاصل از این روش در جدول (۳-۱۷) و شکل (۳-۱۶) ارائه شده است. این نتایج تعمیم‌پذیری بالای مدل طراحی شده را بیان می‌کند، اما در مقایسه با ارزیابی متقاطع به کار

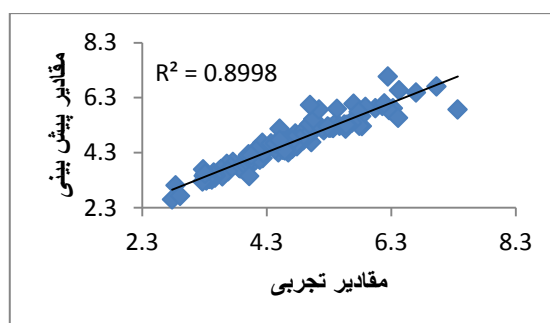
گرفته شده در SR-ANN درصد خطا بیشتر و ضریب تعیین پایین تری را داراست. در شکل (۳-۱۷) مقادیر باقی مانده بر حسب مقادیر تجربی مورد بررسی ترسیم شده است که تمرکز نقاط به سمت محور صفر، نشان دهنده عدم وجود خطای سیستماتیک در مدل است.

جدول (۳-۱۷)-نتایج حاصل از ارزیابی مدل RF-ANN با استفاده از ارزیابی مقاطع

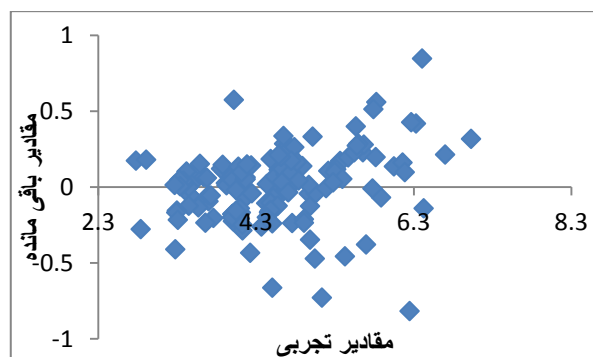
NO.	مقدار تجربی	مقدار پیش بینی	درصد خطا	NO.	مقدار تجربی	مقدار پیش بینی	درصد خطا
۱	۳/۲۸	۳/۶۹۰	۱۲/۵۱۰	۸۱	۴/۱۶	۴/۲۳۵	۱/۸۰۵
۲	۲/۸۴	۳/۱۱۷	۹/۷۶۲	۸۲	۱/۴	۴/۲۴۰	۳/۴۳۵
۳	۲/۷۸	۲/۶۰۵	-۶/۲۷۹	۸۳	۴	۴/۱۷۹	۴/۴۸۲
۴	۲/۹۱	۲/۷۳۰	-۶/۱۷۰	۸۴	۴/۰۱	۴/۱۸۵	۴/۳۸۰
۵	۳/۴۷	۳/۳۸۸	-۲/۳۴۸	۸۵	۴/۱۷	۴/۱۰۵	-۱/۵۳۶
۶	۳/۹۲	۳/۸۹۲	-۰/۶۹۷	۸۶	۴/۱۷	۴/۱۱۸	-۱/۲۴۳
۷	۳/۷۶	۳/۹۶۳	۵/۴۱۹	۸۷	۴/۱۲	۴/۰۲۲	-۲/۳۶۴
۸	۳/۹	۳/۷۹۶	-۲/۶۴۱	۸۸	۴/۰۵	۴/۰۷۰	۰/۵۰۱
۹	۴	۴/۱۹۸	۴/۹۶۵	۸۹	۱/۴	۴/۳۰۷	۵/۰۶۵
۱۰	۴/۱۳	۴/۴۱۵	۶/۹۰۷	۹۰	۴/۲۴	۴/۰۹۶	-۳/۳۷۵
۱۱	۴/۰۵	۴/۲۰۷	۳/۸۹۳	۹۱	۴/۱۹	۴/۰۴۱	-۳/۵۵۰
۱۲	۳/۹	۳/۷۹۲	-۲/۷۴۳	۹۲	۴/۲۵	۴/۲۹۱	۰/۹۷۵
۱۳	۳/۷	۳/۷۹۳	۲/۵۱۹	۹۳	۴/۶۶	۴/۳۷۳	-۶/۱۳۸
۱۴	۳/۷	۳/۷۵۳	۱/۴۴۳	۹۴	۴/۷۱	۴/۴۳۲	-۵/۸۹۳
۱۵	۳/۶۶	۳/۸۹۴	۶/۴۰۴	۹۵	۴/۵۶	۴/۴۰۵	-۳/۳۹۵
۱۶	۳/۷۳	۳/۸۷۶	۱/۵۱۵	۹۶	۴/۴۹	۴/۳۰۵	-۴/۱۰۹
۱۷	۳/۸۶	۳/۷۳۷	-۳/۱۶۲	۹۷	۴/۷۵	۴/۶۱۷	-۲/۷۹۲
۱۸	۳/۸۸	۳/۷۳۲	-۳/۷۹۵	۹۸	۴/۷۳	۴/۵۱۳	-۴/۵۷۶
۱۹	۴/۶۳	۴/۶۶۸	۰/۸۴۱	۹۹	۴/۶۵	۴/۳۱۳	-۷/۲۲۵
۲۰	۵/۰۲	۴/۶۸۷	-۶/۶۲۵	۱۰۰	۴/۶۶	۴/۴۴۶	-۴/۵۷۸
۲۱	۴/۴۷	۴/۶۳۲	۳/۶۳۷	۱۰۱	۳/۵۹	۳/۴۳۵	-۴/۲۸۹
۲۲	۴/۴۷	۴/۶۳۷	۳/۷۵۶	۱۰۲	۳/۴۸	۳/۴۶۲	-۰/۵۱۰
۲۳	۴/۷۶	۴/۹۹۷	۴/۹۸۱	۱۰۳	۳/۴۷	۳/۵۲۵	۱/۵۹۶
۲۴	۴/۵۹	۴/۴۷۴	-۲/۵۲۳	۱۰۴	۳/۴۵	۳/۵۷۲	۳/۵۶۰
۲۵	۴/۴۳	۴/۵۳۶	۲/۳۹۷	۱۰۵	۳/۴۵	۳/۴۴۹	-۰/۰۲۷۵
۲۶	۴/۵۴	۴/۵۳۷	-۰/۰۶۱۴	۱۰۶	۳/۴۲	۳/۳۱۷	-۳/۰۰۹
۲۷	۴/۹۷	۵/۰۱۴	۰/۸۸۹	۱۰۷	۴/۰۲	۳/۴۴۴	-۱۴/۳۲۳
۲۸	۴/۶	۴/۵۴۳	-۱/۲۳۸	۱۰۸	۳/۳	۳/۴۵۴	۴/۶۷۶
۲۹	۵/۳۷	۵/۱۹۸	-۳/۱۸۶	۱۰۹	۳/۳۹	۳/۳۶۸	-۰/۶۴۷
۳۰	۴/۹۹	۵/۱۱۵	۲/۵۲۴	۱۱۰	۳/۳۸	۳/۳۳۸	-۱/۲۲۴
۳۱	۵/۰۶	۵/۱۰۱	۰/۸۱۲	۱۱۱	۳/۳۳	۳/۲۸۳	-۱/۳۹۷
۳۲	۴/۸۹	۴/۷۵۱	-۲/۸۲۷	۱۱۲	۳/۲۹	۳/۴۵۷	۵/۰۸۳

33	Δ/22	Δ/113	-2/038	113	3/34	3/335	-0/131
34	Δ/36	Δ/202	-2/939	114	3/27	3/256	-0/423
35	4/Δ1	Δ/172	14/697	115	3/31	3/Δ26	6/ΔΔ0
36	4/91	Δ/122	4/319	116	4	4/199	4/977
37	Δ/32	Δ/201	-2/218	117	4/23	4/664	10/272
38	4/91	Δ/145	4/789	118	4/44	4/420	-0/437
39	Δ/3	Δ/201	-1/8Δ6	119	4/Δ4	4/470	-1/Δ3Δ
40	Δ/4	Δ/346	-0/994	120	4/Δ2	4/Δ91	1/Δ86
41	Δ/66	Δ/431	-4/044	121	4/Δ6	4/689	2/834
42	Δ/47	Δ/274	-3/Δ70	122	4/68	4/Δ99	-1/717
43	4/97	4/9Δ6	-0/268	123	4/7	4/728	0/60Δ
44	4/99	Δ/336	6/947	124	4/64	4/Δ19	-2/Δ90
45	Δ/19	Δ/196	0/129	125	4/49	4/Δ94	2/318
46	4/76	4/68Δ	-1/ΔΔ4	126	4/37	4/630	Δ/963
47	4/64	4/614	-0/Δ47	127	4/13	4/377	Δ/999
48	4/Δ1	4/642	2/948	128	4/77	4/746	-0/Δ
49	Δ/83	Δ/269	-9/610	129	4/79	4/Δ26	-Δ/Δ03
Δ0	Δ/Δ7	Δ/170	-7/166	130	4/7Δ	4/723	-0/ΔΔ4
Δ1	Δ/0Δ	Δ/Δ22	9/348	131	4/68	4/690	0/234
Δ2	6/0Δ	Δ/914	-2/239	132	4/74	4/660	-1/683
Δ3	6/33	Δ/910	-6/621	133	4/74	4/72Δ	0/301
Δ4	Δ/82	Δ/623	-3/381	134	4/7Δ	4/6Δ4	-2/020
ΔΔ	Δ/14	Δ/868	14/174	13Δ	4/82	4/6Δ4	-3/437
Δ6	6/16	Δ/999	-2/606	136	4/84	4/79Δ	-0/91Δ
Δ7	Δ/7	6/079	6/662	137	4/0Δ	4/207	3/893
Δ8	Δ	6/034	20/699	138	4/48	4/634	3/447
Δ9	Δ/89	Δ/9Δ8	1/167	139	4/7Δ	4/Δ76	-3/647
60	Δ/78	Δ/789	0/1Δ6	140	4/63	4/Δ97	-0/706
61	6/27	Δ/843	-6/803	141	4/71	4/744	0/732
62	6/19	6/090	-1/600	142	4/Δ8	4/703	2/689
63	Δ/ΔΔ	Δ/320	-4/139	143	4/49	4/642	3/390
64	6/41	Δ/Δ61	-13/231	144	4/62	4/419	-4/338
6Δ	6/2Δ	7/067	13/072	14Δ	4/Δ9	4/376	-4/643
66	Δ/43	Δ/886	8/410	146	4/44	4/639	4/49Δ
67	6/7	6/486	-3/192	147	4/Δ1	4/749	Δ/308
68	7/03	6/711	-4/Δ23	148	Δ/79	Δ/278	-8/894
69	6/43	6/Δ70	2/177	149	Δ/67	Δ/390	-4/923
70	7/37	Δ/867	-20/381	1Δ0	Δ/6	Δ/314	-Δ/099
71	3/Δ4	3/634	2/679	1Δ1	Δ/Δ8	Δ/30Δ	-4/912
72	3/Δ7	3/704	3/767	1Δ2	Δ/28	Δ/248	-0/Δ87

۷۳	۳/۶۴	۳/۵۷۳	-۱/۸۲۷				
۷۴	۳/۶۸	۳/۶۲۰	-۱/۶۲۷				
۷۵	۳/۹۳	۳/۹۱۹	-۰/۲۷۳				
۷۶	۴	۳/۹۵۳	-۱/۱۵۰				
۷۷	۴/۰۸	۳/۹۴۶	-۳/۲۷۴				
۷۸	۴/۰۱	۴/۲۳۴	۵/۶۰۷				
۷۹	۴/۱۱	۴/۲۷۸	۴/۰۹۱				
۸۰	۴/۱۶	۴/۱۰۹	-۱/۲۰۸				



شکل (۳-۱۶)- نمودار مقادیر پیش‌بینی شده بر حسب مقادیر تجربی برای ارزیابی مدل RF-ANN به روش ارزیابی متقاطع برای کل داده‌ها



شکل (۳-۱۷)- نمودار مقادیر باقی‌مانده بر حسب مقادیر تجربی برای ارزیابی مدل RF-ANN به روش ارزیابی متقاطع برای کل داده‌ها

۳-۸- مقایسه مدل‌های ارائه شده با استفاده از پارامترهای آماری

همانطور که در مباحث قبل بیان شد، پارامترهای آماری مقادیر تعیین کننده‌ای برای ارزیابی مدل‌ها هستند. مطابق جدول (۳-۱۹) هفت پارامتر آماری عنوان شده جهت ارزیابی توانایی پیشگویی مدل‌های ساخته شده به روش‌های SR-ANN و RF-ANN و جنگل تصادفی به کاربرده شد. باتوجه به

نتایج موجود در جدول (۳-۱۹)، ضرایب تعیین بزرگتر و خطای استاندارد کوچکتر سری آزمون و ارزیابی برای SR-ANN، برتری روش شبکه عصبی مصنوعی در پیش‌بینی اندیس بازدارنده نشان می‌دهد. این پارامترها طبق روابط توضیح داده‌شده در بخش ۲-۳-۶ محاسبه شده‌اند.

جدول (۳-۱۹)-محاسبه پارامترهای آماری برای داده‌های آزمون و ارزیابی

پارامتر آماری		R2	PRESS	MSE	SEP	MAE	REP%	MRE	F	FIT
SR-ANN	سری ارزیابی	۰/۹۷۱	۰/۶۵۷	۰/۰۲۱۹	۰/۱۴۸	۰/۱۰۴	۳/۲۱۸	۲/۲۶۷	۶۵۱/۶۸۰	۷/۴۰
	سری تست	۰/۹۵۸	۰/۷۲۴	۰/۰۲۴۱	۰/۱۵۵	۰/۱۰۳۴	۳/۴۲۳	۲/۲۸۵		
RF-ANN	سری ارزیابی	۰/۹۳۵	۱/۴۹۷	۰/۰۴۹	۰/۲۲۳	۰/۱۷۶	۴/۸۵۶	۳/۸۰۷	۲۲۴/۸۷۵	۴/۴۰
	سری تست	۰/۹۳۲	۱/۲۱۸	۰/۰۴۰۶	۰/۲۰۱	۰/۱۶۰	۴/۴۴۱	۳/۴۹۷		
RF		۰/۹۴۶	۰/۹۳۹	۰/۰۳۱۳	۰/۱۷۷	۰/۱۱۷	۳/۸۹۹	۲/۷۵۱	۳۶۱/۳۲۲	۷/۰۲

۳-۹- ارزیابی مدل با استفاده از Y -تصادفی

این تکنیک ارزیابی مدل با هدف بررسی هر گونه ارتباط تصادفی بین داده‌ها انجام می‌شود. در این آزمون مقادیر تصادفی از متغیر وابسته تولید گردید. اگر مدل اصلی هیچگونه ارتباط تصادفی نداشته باشند، تفاوت قابل توجهی بین مقدار ضریب تعیین مدل اصلی و مدل QSPR با پاسخ‌های تصادفی توسعه یافته وجود دارد. نتایج حاصل از چندین بار اجرای آزمون Y -تصادفی برای سری آزمون و ارزیابی در روش SR-ANN و نیز آزمون Y -تصادفی برای داده‌های سری ارزیابی و آزمون در روش RF-ANN به ترتیب در جدول‌های (۳-۲۰) و (۳-۲۱) آورده شده است. مقادیر کوچک ضریب تعیین (R^2) بیانگر عدم وجود ارتباط یا وابستگی ساختاری در مدل توسعه یافته توسط شبکه می‌باشد، همانطور که مشاهده می‌شود در روش SR-ANN مقادیر R^2 کوچکتر است و تاییدی بر بهتر بودن این روش نسبت به RF-ANN می‌باشد.

جدول (۳-۲۰)-مقادیر R^2 برای داده‌های سری ارزیابی و آزمون بعد از ۱۰ بار اجرا Y-تصادفی در روش (SR-ANN)

تکرار	سری آزمون	سری ارزیابی
۱	۰/۰۰۶۴	۰/۰۳۸۳
۲	۰/۰۰۵۵	۰/۰۷۴۳
۳	۰/۰۶۸۳	۰/۰۰۱۶
۴	۰/۰۰۷۷	۰/۰۱۹۱
۵	۰/۰۱۱۷	۰/۰۰۰۰۲
۶	۰/۰۰۰۷	۰/۰۲۸۱
۷	۰/۰۰۲۵	۰/۰۰۰۲
۸	۰/۰۰۵۳	۰/۰۳۹۹
۹	۰/۰۰۱۶	۰/۰۵۸۴
۱۰	۰/۰۱۵۵	۰/۰۳۱

جدول (۳-۲۱)-مقادیر R^2 برای داده‌های سری ارزیابی و آزمون بعد از ۱۰ بار اجرا Y-تصادفی در روش (RF-ANN)

تکرار	سری آزمون	سری ارزیابی
۱	۰/۰۰۱۴	۰/۰۳۳۰
۲	۰/۰۰۵۶	۰/۰۰۰۰۶
۳	۰/۱۲۲	۰/۰۲۲۱
۴	۰/۰۷۲۰	۰/۰۰۶۲
۵	۰/۰۳۲۰	۰/۱۳۸
۶	۰/۰۰۶۶	۰/۰۰۰۸
۷	۰/۰۴۲۱	۰/۰۵۲۹
۸	۰/۰۰۱۲۷	۰/۰۲۷۸
۹	۰/۰۰۰۳۲	۰/۰۲۲۹
۱۰	۰/۰۲۰۱	۰/۰۶۰۹

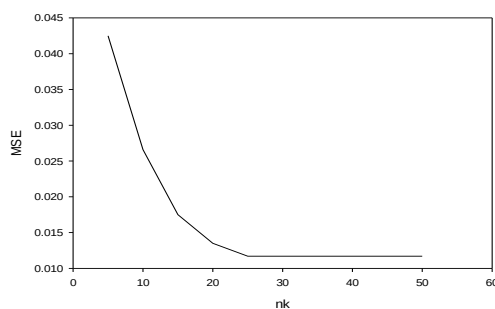
MARS-۱۰-۳

برای مدل‌سازی توسط مارس ابتدا باید پارامترهای موثر بهینه شوند، برای انجام این کار ابتدا کل داده‌ها به دو سری آزمون (۳۰ ترکیب) و سری آموزش (۱۲۲ ترکیب) تقسیم بندی شده‌اند و مقادیر مربوط به ۱۹۲ توصیف‌کننده به عنوان متغیر مستقل و مقادیر اندیس بازداري به عنوان متغیر وابسته منسوب شدند. برای برازش مدل مارس، همانطور که ذکر شد، پارامتری کلیدی این روش یعنی تعداد توابع پایه در گام پیش‌رو یعنی nk باید بهینه شود. با در نظر گرفتن nk از ۵ تا ۵۰ نتایج جدول (۳-۱۸) به دست آمده است. با توجه به جدول (۳-۱۸)، با تغییر nk مقادیر MSE و R-Square همگام با nk تغییر می‌کنند، در ابتدا این تغییر قابل محسوس است اما با افزایش بیشتر nk تغییری در آن‌ها

احساس نمی‌شود تا جایی که ثابت می‌شوند (با توجه به شکل (۳-۱۸)). به همین ترتیب مدل بهینه انتخاب شده مدلی است که دارای بیشترین مقدار ضریب تعیین (R-square) و کمترین مقدار اعتبار سنجی (GCV) با تنظیمات $nk=25$ حاصل می‌شود. این مدل براساس اینکه یک مدل بهینه در مجموعه داده‌های مدل‌ساز می‌باشد، مشاهده می‌شود بالاترین R^2 را برای داده‌های آموزش ایجاد کرده و به عنوان مدل برتر شناخته می‌شود، اما برای داده‌های سری آزمون با مقدار ضریب تعیین ۰/۹۱۴ و با توجه به شکل (۳-۱۹) که مقدار پیش‌بینی شده با مقدار واقعی تطابق زیادی ندارند روش MARS برازش خوبی برای داده‌های آزمون ایجاد نکرده است.

جدول (۳-۱۸)-مقادیر R-square, GCV, RSS, MSE برای داده‌های آزمون و آموزش براساس مقادیر مختلف nk

Training set				Testset	
nk	R-squre	MSE	GCV	R-squre	MSE
۵	۰/۹۴۰	۰/۰۴۵۲	۰/۰۵۰۹	۰/۹۳۰	۰/۰۵۴۱
۱۰	۰/۹۶۴	۰/۰۲۶۶	۰/۰۳۴۵	۰/۹۴۴	۰/۰۴۲۵
۱۵	۰/۹۷۶	۰/۰۱۷۵	۰/۰۲۶۶	۰/۹۱۸	۰/۰۵۶۹
۲۰	۰/۹۸۲	۰/۰۱۳۵	۰/۰۲۲۳	۰/۹۰۶	۰/۰۶۵۱
۲۵	۰/۹۸۴	۰/۰۱۱۷	۰/۰۲۱۹	۰/۹۱۴	۰/۰۵۸۷
۳۰	۰/۹۸۴	۰/۰۱۱۷	۰/۰۲۱۹	۰/۹۱۴	۰/۰۵۸۷
۳۵	۰/۹۸۴	۰/۰۱۱۷	۰/۰۲۱۹	۰/۹۱۴	۰/۰۵۸۷
۴۰	۰/۹۸۴	۰/۰۱۱۷	۰/۰۲۱۹	۰/۹۱۴	۰/۰۵۸۷
۴۵	۰/۹۸۴	۰/۰۱۱۷	۰/۰۲۱۹	۰/۹۱۴	۰/۰۵۸۷
۵۰	۰/۹۸۴	۰/۰۱۱۷	۰/۰۲۱۹	۰/۹۱۴	۰/۰۵۸۷



شکل (۳-۱۸)- نمودار میانگین خطای برحسب توابع پایه (nk)

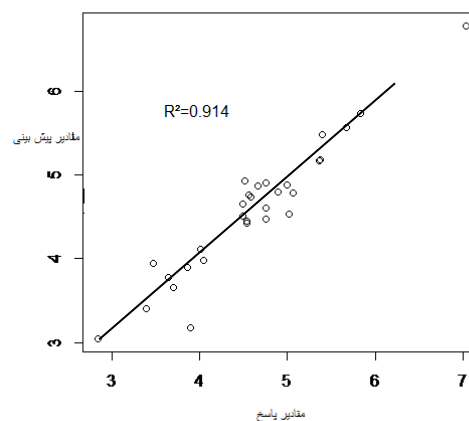
مدل نهایی مارس بعد از طی مراحل پیش‌رو و پس‌رو به صورت زیر است:

$$Y = 4.78 + 0.53h(X3V - 3.93) - 0.82h(PJI3 - 0.79) - 0.045h(4.87 - LBW) - 0.157h(LBW - 4.87) + 0.28h(0.822 - AROM) - 13.94h(AROM - 0.822) - 0.24h(4.169 - RDF040V) - 0.076h(RDF040V - 4.16) - 0.12h(1.554 - RDF115e) - 0.58h(RDF11e - 1.55) - 2.66h(0.029 - Mor19m) - 0.46h(Mor19m - 0.029) + 3.3h(0.154 - Mor11p) - 0.14h(15.06 - PCWTe) - 0.11h(4 - F04CS) + 0.14h(F04CS - 4)$$

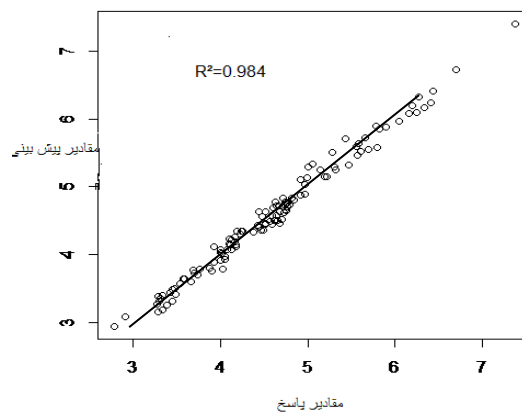
که در آن

$$h(u) = u_+ = \begin{cases} u & \text{اگر } u > 0 \\ 0 & \text{در غیر این صورت} \end{cases}$$

عملکرد مدل MARS را می‌توان در شکل‌های (۳-۱۹) و (۳-۲۰) به وسیله نمودارهای مقادیر پیش بین‌ی شده در مقابل مقادیر پاسخ دید. که می‌توان گفت مدل MARS برازش بهتری را به داده‌ها انجام داده است



شکل (۳-۱۹)-نمودار مقادیر پاسخ در مقابل مقادیر برازش داده شده توسط روش مارس برای داده‌های سری آزمون



شکل (۳-۲۰)- نمودار مقادیر پاسخ در مقابل مقادیر برازش داده شده توسط روش مارس برای داده‌های سری آموزش

۳-۱۱- نتیجه گیری مدل سازی با مجموع توصیف کننده‌های انتخاب شده توسط

رگرسیون گام به گام

ارزیابی عملکرد مدل‌های مختلف به کمک فاکتور R^2 امکان پذیر است. این فاکتور زمانی می‌تواند ارزش رجحانی مدل‌ها را تعیین نماید که بین مقادیر مشاهده شده و مقادیر پیش‌بینی شده محاسبه شوند. توجه به مقادیر R^2 مزیت نسبی مدل‌ها را نشان می‌دهد. بیشترین ضریب تعیین معنی‌دار بهترین و کارآمدترین روش را معرفی می‌نماید بنابراین، ارزش رجحانی روش‌ها به ترتیب MARS و SR-ANN و RF و RF-ANN با میزان ضرائب تعیین ۰/۹۸۴، ۰/۹۷۱، ۰/۹۴۴، ۰/۹۳۸ انتخاب می‌گردد. شکل‌های (۳-۳) و (۵-۳) و (۶-۳) و (۱۱-۳) و (۱۲-۳) و (۱۸-۳) و (۱۹-۳) ارتباط بین مقادیر مشاهده شده و مقادیر پیش‌بینی در مدل‌های چهارگانه تحقیق براساس ضریب تعیین را نشان می‌دهد. در مقایسه کارایی این مدل‌ها هرچه نتایج پیش‌بینی بیشتر و منطبق بر نتایج مشاهده شده باشد مدل از کارایی مطلوب‌تری برخوردار است و همچنین مقدار MSE بدست آمده برای این مدل‌ها به این شرح است. روش‌ها به ترتیب MARS و SR-ANN و RF و RF-ANN با میزان میانگین مربعات خطا ۰/۰۱۱۷، ۰/۰۲۴۱، ۰/۰۳۱۳، ۰/۰۴۰۶ انتخاب می‌گردد.

۳-۱۲- بررسی ارتباط بین توصیف‌کننده‌های منتخب و خاصیت مورد نظر

در این بخش بعضی از توصیف‌کننده‌های وارد شده در مدل MARS که به عنوان توصیف‌کننده ورودی در SR-ANN نیز محسوب می‌شود آورده شده است، که به بررسی اجمالی روی اثرات این متغیرها بر خاصیت ترکیبات مورد مطالعه پرداخته شده است. هرکدام از این توصیف‌کننده‌ها بیانگر خصوصیت دو بعدی، هندسی و سه بعدی ترکیبات مورد بررسی می‌باشند.

۳-۱۲-۱- توصیف‌کننده‌های 2D frequency finger prints

این توصیف‌کننده‌ها اندازه مولکول‌ها را بررسی می‌کند. توصیف‌کننده اثرانگشت فرکانس که به بررسی مجموع فرکانس دو اتم در فاصله هندسی آن اتم‌ها می‌پردازد، اندازه اتم‌ها را به‌طور مستقیم بررسی می‌کند، که از میان توصیف‌کننده‌های اثرانگشت F04CS که فرکانس S-C را در فاصله ای هندسی نشان می‌دهد [۵۲].

۳-۱۲-۲- توصیف‌کننده‌های Topological

این توصیف‌کننده بر اساس نمایش گراف مولکول می‌باشد. در این گراف‌ها هر نقطه نشان دهنده یک اتم بوده و خطوط بین نقاط نیز نشان دهنده پیوند شیمیایی بین اتم‌ها می‌باشد. معمولاً در گراف‌های مولکولی اتم هیدروژن را نشان نمی‌دهد. این توصیف‌کننده اطلاعاتی راجع به ساختمان، اندازه، شکل، تقارن، شاخه دار شدن، نحوه اتصال اتم‌ها و نوع اتم‌های موجود در یک مولکول در اختیار ما قرار می‌دهند. محاسبه این توصیف‌کننده به سادگی از روی ساختمان دو بعدی مولکول امکان پذیر می‌باشد. از میان این توصیف‌کننده‌ها PW3 است [۵۳].

۳-۱۲-۳- توصیف‌کننده‌های 3D MORSE

این توصیف‌کننده ساختار ۳ بعدی مولکول را بر اساس تفرق الکترون نشان می‌دهند. که توصیف‌کننده‌ها Mor19m و Mor29m را شامل می‌شود [۵۳].

۳-۱۲-۴- توصیف‌کننده‌های RDF^1

این توصیف‌کننده می‌تواند تابع توزیع شعاعی یک ترکیب شامل N اتم را به احتمال یافتن یک اتم در یک حجم کروی با شعاع r ارتباط دهد. اهمیت این توصیف‌کننده‌ها به دلیل اختلاف در توزیع اتم‌ها در مولکول‌ها می‌باشد. از میان این توصیف‌کننده‌های $RDF115e$ ، دیده شده است [۵۳].

۳-۱۲-۵- توصیف‌کننده‌های Charge

با توجه به تئوری کلاسیک، برهم‌کنش‌های بین مواد از نظر ماهیت، الکتروستاتیکی یا کووالانسی است. تغییرات بار الکتریکی در ملکول به عنوان نیروی پیش‌برنده برهم‌کنش‌های الکتروستاتیک است. از جمله این توصیف‌کننده $PCWTe$ ، شاخص دو قطبی محلی می‌باشد. علامت منفی آن در مدل، اهمیت برهم‌کنش‌های الکتروستاتیک بین بازدارنده و ملکول هدف را نشان می‌دهد و به عنوان یک شاخص برای اندازه‌گیری میزان این برهم‌کنش‌های الکتروستاتیک به کار می‌رود [۵۴].

۳-۱۲-۶- توصیف‌کننده‌های $GATWAY^2$

این توصیف‌کننده‌ها با توجه به مختصات فضایی اتم‌ها در یک ملکول به راحتی محاسبه هستند. این توصیف‌کننده‌ها براساس ماتریس قدرت نفوذ^۳ تعریف می‌شوند و بیان‌کننده ویژگی‌های هندسی و توپولوژیک ملکول‌ها هستند. توصیف‌کننده‌های $GATWAY$ به دو زیر مجموعه H و R تقسیم می‌شوند.

ماتریس نفوذ یا ماتریس تاثیر ملکول (MIM)^۴ با رابطه زیر بیان می‌شود.

¹ Radial distribution function

² Geometry, Topology, and Atom-weights Assembly

³ Leverage matrix

⁴ Molecular influence matrix

$$H = M \cdot (M^T \cdot M)^{-1} \cdot M^T \quad \text{رابطه (۱-۳)}$$

که M ماتریس مختصات اتمی، T به معنای ماتریس ترانپازه^۱ و H ماتریس قدرت نفوذ است. H یک ماتریس متقارن A^*A است که A تعداد اتم‌هاست. عناصر قطری ماتریس H لوریج‌ها نام دارند که هر یک بیانگر ماتریس تاثیر-فاصله به صورت زیر تعریف می شود

$$[R]_{ij} = \left[\frac{\sqrt{h_{ii} \cdot h_{jj}}}{r_{ij}} \right]_{ij} \quad i \neq j \quad \text{رابطه (۲-۳)}$$

h_{ii} و h_{jj} لوریج‌های دو اتم i و j و r_{ij} فاصله آن دو است. عناصر قطری این ماتریس صفر است. سه توصیف‌کننده $R6u$ و $R2m$ و $H6m$ از این کلاس در مدل‌سازی به کار گرفته شده است. u در این توصیف‌کننده‌ها نشان‌دهنده آن است که آن‌ها با ویژگی خاصی از اتم‌ها محاسبه نشده‌اند [۵۴].

۳-۱۲-۷- توصیف‌کننده‌های connectivity indices

$X3v$ توصیف‌کننده‌ای از گروه اندیس‌های ارتباط مولکولی است. این اندیس‌ها دسته‌ای از توصیف‌کننده‌های توپولوژیکی هستند که اطلاعاتی در مورد ساختمان، اندازه و میزان شاخه‌ای شدن مولکول، نحوه‌ی اتصال اتم‌ها و نوع اتم‌های موجود در یک مولکول را در اختیار قرار می‌دهند. جهت محاسبه‌ی این توصیف‌کننده‌ها ابتدا باید ساختار مولکول را با استفاده از تئوری گراف شیمیایی رسم نمود.

این اندیس‌ها اولین بار به‌وسیله راندیک^۲ در سال ۱۹۷۵ ارائه شدند و سپس در سال ۱۹۷۶ توسط هال^۳ اصلاح گردیدند. این توصیف‌کننده‌ها با رابطه‌ی (۳-۳) قابل محاسبه هستند:

^۱ Transpose

^۲ - Randic

^۳ - Hall

$${}^m\chi^v = \sum_{i=1}^{N_s} \prod_{k=1}^{m+1} \left(\frac{1}{\delta_k^v} \right)^{1/2} \quad (3-3)$$

که در این رابطه δ_k^v ، درجه رأس ظرفیت^۱ اتم k ام در گراف مولکولی است و به شکل رابطه‌ی زیر تعریف می‌شود:

$$\delta_k^v = \frac{(Z_K^v - h_k)}{(Z - Z_k^v - 1)} \quad (4-3)$$

که در آن Z عدد اتمی اتم k ، Z_k^v تعداد الکترون‌های لایه‌ی ظرفیت اتم k و h_k تعداد اتم‌های هیدروژن متصل به اتم غیر هیدروژنی k است.

^۱ - Valence vertex degree

۳-۱۳- نتیجه گیری:

در کار ارائه شده مطالعه ارتباط کمی ساختار-خاصیت جهت پیش بینی اندیس بازداری دسته وسیعی از ترکیبات آلی انجام شد. در ابتدا توصیف کننده هایی ایجاد و مناسب ترین آنها که بیشترین ارتباط را با اندیس بازداری داشتند به کمک روش رگرسیون گام به گام انتخاب گردیدند و به روش هایی از جمله جنگل تصادفی و شبکه عصبی مصنوعی و مارس مدل مناسب ایجاد شد. مدل ایجاد شده دارای توصیف کننده می باشد که این توصیف کننده ها از دسته های گوناگون از جمله توپولوژیکی، هندسی، ساختاری و.... می باشند. با استفاده از این توصیف کننده ها مدلی ایجاد شد که قدرت پیش بینی اندیس بازداری ترکیبات مشابه را داشت. بعضی از این مدل ها با استفاده از روش های مختلف از جمله ارزیابی متقاطع مورد ارزیابی قرار گرفت و همچنین پارامترهای آماری مختلف نشان داد که مدل های ارائه شده می تواند جهت پیش بینی اندیس بازداری سایر ترکیبات قرار گیرد. در ضمن از این روش می توان جهت مدل سازی و پیش بینی سایر خواص فیزیکوشیمیایی استفاده نمود.

آینده نگری:

- در این تحقیق نتایج نشان می‌دهد که مدل‌های خطی و غیر خطی توسعه یافته می‌توانند به‌عنوان روش‌های موفق برای مدل‌سازی و پیش‌بینی خاصیت مورد نظر (اندیس بازداري) ترکیبات مورد مطالعه به کار روند.
- روش جنگل‌های تصادفی را می‌توان برای پیش‌بینی اندیس بازداري یا سایر خواص ترکیبات دیگر که هنوز سنتز نشده‌اند استفاده نمود
- علاوه بر روش رگرسیون گام‌به‌گام که در این تحقیق به‌عنوان روش انتخاب توصیف‌کننده استفاده شده است می‌توان از سایر روش‌های انتخاب متغیر مانند اجتماع مورچگان و راهبرد CDFS و طرح‌ریزی پیایی استفاده کرد.
- می‌توان از روش‌های خطی دیگر مانند حداقل مربعات جزئی و آنالیز اجزای اصلی را جایگزین روش‌های جنگل‌های تصادفی کرد و نتایج به‌دست آمده را باهم مقایسه کرد.

منابع:

- [1]-راجر، ام اسمیت، (۱۳۸۳)، "کروماتوگرافی گاز و مایع در شیمی تجزیه"، ارباب زوارم، سرفراز یزدی ع ، انتشارات دانشگاه فردوسی مشهد، ۴۸۶، ص ۱۸۵-۱۸۴
- [2]-Gupta N., Roychoudhury Pk., Deb Jk. Biotechnology of desulfurization on diesel; Prospects and challenges. Applied Environmental Microbiology 2005; 66: 356-66
- [3]-Baldi F., Pepi M ., Fara F. Growth of Rhodosporidium. Toruloides strain DBVPG 6662 on Dibenzothiophene Crystals and Orimulsion. Applied Environmental Microbiology 2003; 69(8): 4689-96.
- [4]-Monticello Dj. Biodesulfurization and the upgrading of petroleum distillates. Current opinion Microbiology 2000;11:540-6.
- [5]-Schemidt M., Siebert W., Bangall Kw. The chemistry of Sulphur., Selenium., tellurium and polonium pergamon Texin Inorganic chemistry. 115, New York: Pergamon Press, oxford; 1973
- [6]-Prayuen Yongp. Coal biodesulfurization Process. Science thechnology 2012; 24: 493-507
- [7]-Kilbane JJ. Microbial bi Octalyst developments toupgrade fossil fuel. Energy Biotechnology 2006;17: 305-14.
- [8]-Mei H., Mei1 B.w. and Yent.F.,” A new method for obtaining ultra-low sulfur diesel fuel via ultra-sound as-sisted oxidative desulfurization” Fuel 82,PP. 405-414,2003.
- [9]-Lin L., Hong L., Janhua Q., Jin Juan X. Progress in the technology for desulfurization of crude Oil. China petroleum processing Petrochemical technology 2010;12(4):1-6.
- [10]-WWW.fa.wikipedia.org/wiki/تسوفن
- [11]-Miller K.E,Bruno TJ,(2003),” Iso thermal Kovats retention indieces of sulfur compounds on poly(5% diphenyl-95% dimethyl siloxane) Stationary phase”, J Chromatogr. A , 1007,PP 117-125

- [12]-DUH, Ring Z, Briker Y, Arboled P,(2004),”Prediction of gas chromatographic retention times and indices of sulfur compound in light cycle Oil” *Catal. Today*, 98, PP 214-225
- [13]-Can H, Diamoglo A, Kovalishyn V, (2005), "Application of artificial neural networks for the prediction of sulfur polycyclic aromatic compounds retention indices" *J.Mol.Struct.* 723. PP 183-188
- [14]-Garkani-Nejad, Z., et al. “Prediction of gas chromatographic retention indices of a diverse set of toxicologically relevant compounds.” *Journal of Chromatography A* 1028.2 (2004): 287-295.
- [15]-Liao L., Mei H., Li J., Li Z., (2008), “Estimation and prediction on retention times of components from essential oil of paulownia tomentosa flowers by molecular electronegativity-distance vector (MEDV)”, *J Mol Struct*, 850, pp 1-8
- [16]-Hu, Rong-Jing, et al. “QSPR prediction of GC retention indices for nitrogen-containing polycyclic aromatic compounds from heuristically computed molecular descriptors.” *Talanta* 68.1 (2005): 31-39.
- [17]-Riahi S., Ganjali M.R., Pourbasheer S., Norouzi P., (2008), “QSPR Study of GC Retention Indices of Organic Compounds by Multiple Linear Regression with a Genetic Algorithm”, *Chromatographia*, 67, pp 917-922
- [18]-Riahi S., Gaanjali M., Pourbasheer S., Norouzi P.,(2009), “Investigation of different linear and nonlinear chemometric methods for modeling of retention index of organic components: Concerns to support vector machine” *Journal of Hazardous Materials*, 166, pp 853-859
- [19]-Xu, Hui-Yuing et al., "Quantitative structure-chromatographic sulfur heterocycles" *Journal chromatography A*, 1198(2008):202-207.
- [20]-Xiang, Zheng, Yizeng Liang, and Qiannan Hu. “Relativity study of the topological index of methylalkane structures and chromatographic retention index.” *Se pu= Chinese journal of chromatography* 23.2 (2005): 117-122.
- [21]-Du, Hongbin, et al. “Prediction of gas chromatographic retention times and indices of sulfur compounds in light cycle oil.” *Catalysis today* 98.1-2 (2004): 217-225.

[22]-Zanousi mohammad bagher, Mehdi nekoi and majid mohammad Hosseini,(2016),"composition of the essential oils and volatile fractions of Artemisia absinthium by three different extraction methods: Hydrodistillation, solvent-free microwave extraction and headspace solid-phase micro extraction combined with a novel QSRR Evaluation", Journal of essential oil bearing plants, 19:7(2016),1561-1581.

[23]-Goudarzi, Nasser, et al."Quantitative structure-property relationships of retention indices of some sulfure organic compounds using random forest technique as a variable selection and modeling method". Journal of sparation science 39.19(2016): 3835-3842

[24]-Otto M., (2007), "Chemometrics", WILEY-VCH Verlag GmbH & Co. KgaA, Weinheim, pp 1-11

[25]-Howery D.G., Hirsch R.F., (1983), "Chemometrics in the chemistry curriculum", J. Chem. Ed. 60, pp 656-659

[26]-Wold S., "Chemometrics: what do we mean with it, and what do we want from it?" And Brown, S. D., "Has the chemometrics revolution ended? Some views on the past, present and future of chemometrics", papers on chemometrics:Philosoiphy, History and Direction of In CINC'94.

[27]-Moosavi-Movahedi A.A., Gharanfoli M., Nazeri K., Shamsipur M., Chamani J., Hemmateenejad B., Alavi M., Shokrollahi A., Habib-Rezaei M., Sorrenson C., Sheibani N., (2005), Colloids Suurf., 43, pp 150-157

[28]-Hemmateenejad B., (2005) "Chemometrics in Iran" Chemometrics and intelligent Laboratory System, 81, pp 202-208

[29]-Anderson, G.G. Kaufmann, P.(2000),"Development of a generalized neural network",chemometr. Intel. Lab. Syst.50 PP.101-105

[30]-Arab cham Jangali M, (2009), Modelling of Cytotoxicity data of anti- Hiv 1-[(5-chloropheny) sulfpnyl]-1 H-pyrole derrivatives using calculated molecular descriptors and Levenbery-Marquadt Artificial neural network", chem. Bio. Des, 73, PP 456-465.

[۳۱]-فرشاد فرع-(۱۳۸۰)،" اصول و روش های پیشرفته آماری در تجزیه رگرسیون " چاپ دوم، انتشارات طاق بستان.

[32]-Mohammad Asradel, "Prediction of combustion Dynamics in An Experimental Turbulent Swir stabilized combustor with secondary Fuel Injection", university of Tehran, 2015.

[۳۳]-کیا.م، (۱۳۸۷)، "شبکه های عصبی در MATLAB چاپ دوم، انتشارات کیان رایانه سبز

[۳۴]-فردوسی.م.ع، (۱۳۸۹)، پایان نامه کارشناسی ارشد، "پیش بینی ثابت های هنری بعضی از ترکیبات آلی با استفاده از روش های خطی و غیر خطی QSPR"، دانشگاه صنعتی شاهرود

[35]-Braiman L, Random Forests, Machine Learn 2001; 45:P 15-32

[36]-Hastie T TRFJ, The elements of statistical learning data mining inference and prediction, In Springer series in statistics New York, Springer 2001; Xvi:P.533

[37]-Braiman L, Classification and regression trees CA, WadsWorth Intenational Groups 1984.

[۳۸]-جندقی. ف (۱۳۹۲)، پایان نامه کارشناسی ارشد، "کاربرد روش جنگل های تصادفی به عنوان ابزاری برای انتخاب متغیر و پیش بینی اندیس بازداري تعدادی از ترکیبات آلی به عنوان آلاینده ی محیط زیست"، دانشگاه صنعتی شاهرود

[39]-Hastie T, Tibshirani R, Friedman J, (2009), "The Elements of statistical learning Data Mining Inferences, and Prediction", 2nd Ed, Springer, New York

[40]-Braiman L, Random Forests, Machine Learn 2001; 45:P.5-32.

[41]-Genuer R, Poggi L M, Tueleau-Malotch, 2010, "Variable selection using random forests", Pattern Recoyn Lett, 31, PP.2225-2236.

[42]-Friedman JH (1991) Multivariate adaptive regression splines, J. Annals of statistics, 19:1-141

[۴۳]-تکتم ولی زاده، (۱۳۹۱)، پایان نامه کارشناسی ارشد، "اسپلاین های تطبیقی و غیر تطبیقی در مدل های رگرسیونی نیمه پارامتری"، دانشگاه صنعتی شاهرود

[۴۴]-زینب افتخاری، (۱۳۹۰)، پایان‌نامه کارشناسی ارشد، "مطالعه QSPR جهت پیشگویی حلالیت حشره کش‌ها در آب با استفاده از روش MARS-ANFIS و اندازه‌گیری همزمان سرب و کادمیم به روش ولتامتری عاری سازی پالس تفاضلی با استفاده از روش حداقل مربعات جزئی (PLS)"، دانشگاه دامغان.

[45]-Mercader A. G, Pomilio A.B,(2010), "QSAR Study Of Flavonoids and bioflavonoids as in fluensa HINI virus neuramaini dase inhibitors", Eur.J. Med.Chem,45, PP 1724-1730

[46]-Todeschini R, Consonni V, (2000), "Hundbook of molecular descriptors", Wiley-VCH, Wein heim, Germany

[47]-<http://WWW.downloadsoftware.ir>

[48]-Walter B. Wilson., Lane C Sander., "Retention behavior of alkyl-sustituted Polycyclic aromatic sulfur heterocycles in reversed-phase liquid chromatography" J. Chromatography A, 1461(2016), 120-130

[49]-Walter B. Wilson., Lane C Sander., "Retention behavior of isomeric polycyclic aromatic sulfur heterocycles in reversed-phase liquid chromatography" J. Chromatography A, 1461(2016) 107-119

[50]-Crisan L., Iliescu S., (2016), "Structure-flammability relationship study of phosphoester dimers by MLR and PLS", Polimeros, (AHEAD), 0-0.

[51]-Cosonni V., Todeschini R., Pavan M, (2000), "Structure/Response Correlations and similarity/Diversity Analysis by GETAWAY Descriptors. Theory of the novel 3D Molecular Descriptors", J. Chem. Info. Comput. Sci, 42, PP 682-692

[52]-Karelson M., Lobanov V.S., Katritzky A.R. (1996), "Quantum-chemical descriptors in QSAR/QSPR Studies", J. Chem Rev., 96, PP 1027-1043. Charge

[۵۳]-اشرفی. م. (۱۳۸۹)، پایان‌نامه کارشناسی ارشد "مطالعه ارتباط کمی ساختار-فعالیت مشتقات کربامات به عنوان دسته‌ای جدید از بازدارنده‌های غیر نوکلئوزیدی HIV"، دانشگاه صنعتی شاهرود

[۵۴]-عجم. ز. (۱۳۹۳)، پایان‌نامه کارشناسی ارشد "پیش‌بینی فعالیت ضد ایدز مشتقات غیر نوکلئوزیدی تیوکربامات به روش RF"، دانشگاه صنعتی شاهرود

[55]-Goudarzi, N., et al. "Application of random forests method to predict the retention indices of some polycyclic aromatic hydrocarbons." *Journal of Chromatography A* 1333 (2014): 25-31.

[56]- Kao, Szu-Pyng, Yao-Chung Chen, and Fang-Shii Ning. "A MARS-based method for estimating regional 2-D ionospheric VTEC and receiver differential code bias." *Advances in Space Research* 53.2 (2014): 190-200.

Abstract

Generally, the amount of sulfur in crude oil, depending on its source, varies from 0.25 to about 7.89 of Weight percent. Generally sulfur is observed both organic and inorganic forms in crude oil. In addition to elemental sulfur more than 200 sulfur organic compounds are identified in the oil, which contains sulfides, Thiols or mercaptans, thiophene, which theophany, and di-benzoate molecules are much more complex. These compounds are known as persistent contaminants. In addition, these compounds are carcinogenic. Various methods, such as gas chromatography, have been proposed to segregate and identify these compounds, which inhibition index compares with standards. Therefore, the inhibition index of a soluble is an important parameter that can be used to identify sulfur heterocyclic aromatic compounds. But on the other hand, these methods are costly and time-consuming. By applying chemo metrics methods, prediction of inhibitory index of compounds is plausible (PASHs) and by using random forest methods (RF) and MARS methods and the artificial neural network SR-ANN and RF-ANN we predicted the inhibitory effect of these compounds. The results indicate that the use of the mentioned statistical method has a good potential to predict the empirical inhibitory index of these compounds, which R^2 results from RF and MARS and SR-ANN and RF-ANN methods are 0/944 and 0/984, 0/971 and 0/932 respectively.

Keywords: Chemometrics, Random forest, MARS, Artificial Neural Network, Inhibition Indices (index), Aromatic heterocyclic sulfur compounds



University of Shahrood

Faculty of chemistry

M.Sc Thesis in analytical chemistry

**Application of Chemometric methods to Predict the Retention
index of some organic Compounds**

Elahe Salamati

Supervisors:

Dr. Naser goodarzi

Dr. Davood Shamsavani

Advisor:

Dr. Mansour ArabCham jangali

December 2018