

بِسْمِ اللَّهِ الرَّحْمَنِ الرَّحِيمِ



دانشکده شیمی

پایان نامه کارشناسی ارشد شیمی تجزیه

مطالعه ارتباط کمی ساختار - فعالیت برخی از بازدارنده‌های c-Met سولفون‌دار جدید با فعالیت
ضد سرطان

نگارنده:

محدثه لطفی

استاد راهنما:

دکتر منصور عرب چم جنگلی

اسفند ۱۴۰۰

اگر این مختصر را ارزشی است؛

آن ارزش تقدیم به:

دستان آسمانی و ستارگان تبارم

دل دریایی و چشمان پر امیدم

قلب مهربان و پر عطفه خواهر و برادرم

و

تمامی پویندگان علم و دانش

پاسکزاری

این جانب از استاد محترم جناب آقای دکتر منصور عرب جم جلی، شکر می‌نمایم که مراد انجام این پروژه یاری نمودند و از خداوند متعال پیروزی و موفقیت روز افزون ایشان را خواستارم. از اساتید گرام جناب آقای دکتر قدوسی باقریان، دینی و جناب آقای دکتر ناصر کوردزی که زحمات دایمی پیمان نامه‌ی حاضر را داشته و از محضر ایشان در کلاس های خویش بهره‌ی وافر برده‌ام شکر می‌نمایم.

از خانواده‌ی عزیزم و همچنین افسانه‌ی عزیزم به خاطر تمام زحماتشان شکر می‌نمایم. از دوست عزیزم سرکار خانم مریم مفتری برای تمام یاری‌رسانی‌هایی که در پیمان رسانیدن این پیمان نامه داشتند شکر می‌نمایم. همچنین از سایر اساتید محترم دانشکده‌ی شی که در طول دوران تحصیلاتم افتخار ساگردیشان را داشتم، کارمندان محترم دانشکده و تمامی کسانی که مرا به نحوی یاری رساندند نیات شکر و قدر دانی را به جا آورده و از حضرت تعالی سلامت و توفیق روز افزون آن‌ها را مسئلت می‌نمایم.

تعهد نامه

اینجانب **محدثه لطفی** دانشجوی دوره کارشناسی ارشد شیمی تجزیه دانشکده شیمی دانشگاه صنعتی شاهرود، نویسنده پایان نامه **مطالعه ارتباط کمی ساختار - فعالیت بازدارنده‌های c-Met سولفون دار جدید با فعالیت ضد سرطان تحت راهنمایی آقای دکتر منصور عرب چم جنگلی** متعهد می‌شوم:

- تحقیقات در این پایان نامه توسط اینجانب انجام شده است و از صحت و اصالت برخوردار است.
- در استفاده از نتایج پژوهش‌های 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 می‌باشد و مقالات مستخرج با نام «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 بوده‌اند در مقالات مستخرج از پایان نامه رعایت می‌گردد.
- در کلیه مراحل انجام این پایان نامه، در مواردی که از موجود زنده (یا بافت‌های 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، برنامه‌های رایانه‌ای، نرم افزارها و تجهیزات ساخته شده است) متعلق به دانشگاه صنعتی شاهرود می‌باشد. 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 نمی‌باشد.

چکیده

هدف از این مطالعه استفاده از تکنیک غربالگری داده‌ها برای کاهش ابعاد ماتریس داده است. برای این منظور از روش جدید پیش‌پردازش جریمه‌ای Ridge برای کاهش ابعاد داده‌ها استفاده شد. سپس از روش انحراف قدر مطلق به‌طور هموار کوتاه شده (SCAD) به‌عنوان روش انتخاب متغیر جریمه‌ای برای انتخاب مؤثرترین توصیف‌کننده‌ها و مدل‌سازی استفاده شد. توصیف‌کننده‌های انتخاب شده به‌عنوان ورودی برای ساخت مدل QSAR با استفاده از روش SCAD استفاده شد. برای این منظور، مدل QSAR با استفاده از اجرای چندگانه Ridge به‌عنوان یک تکنیک پیش‌پردازش برای غربالگری داده‌ها توسعه داده شد. توصیف‌کننده‌های باقی‌مانده به مدل جریمه شده SCAD وارد شدند. پس از اعمال SCAD به داده‌های مجموعه آموزش، از برخی توصیف‌کننده‌های مؤثر (۹) در ساخت مدل استفاده شد. مدل بهینه هدف (Ridge-SCAD) با استفاده از محاسبه پارامترهای آماری، دامنه کاربرد و آزمون Y-تصادفی مورد ارزیابی قرار گرفت. آزمون‌های مختلف ذکر شده نتایج رضایت بخشی را نشان می‌دهند. عملکرد مدل Ridge-SCAD برای پیشنهاد برخی از مهارکننده‌های قوی c-Met استفاده شد. کارایی ترکیبات پیشنهادی با استفاده از داکینگ مولکولی و محاسبه قانون پنج لیپینسکی بررسی شد. نتایج حاصل قدرت ترکیبات جدید را اثبات می‌کند. برای هدف مقایسه، مدل QSAR بدون استفاده از روش Ridge به‌عنوان یک تکنیک پیش‌پردازش توسعه داده شد و همه توصیف‌کننده‌های محاسبه‌شده به‌عنوان ورودی‌های SCAD تعریف شدند. مدل SCAD با استفاده از داده‌های مجموعه آموزش ساخته شد و ۱۲ توصیف‌کننده به‌عنوان توصیف‌کننده معنی‌دار به‌دست آمد. فعالیت دارویی داده‌های آزمون با استفاده از مدل SCAD بهینه پیش‌بینی شد. پارامترهای آماری محاسبه شده مدل SCAD با مدل Ridge-SCAD مقایسه شد. نتایج نشان دهنده قدرت پیش‌بینی بالای مدل Ridge-SCAD در برابر مدل SCAD است.

کلمات کلیدی: QSAR، c-Met، پیش‌پردازش، Ridge، SCAD، داکینگ مولکولی

مقاله مستخرج از این پایان نامه تحت عنوان زیر در ژورنال ISI تحت داوری قرار

گرفته است:

M. Lotfi, M. Arab Chamjangali and Z. Mozafari. Ridge regression coupled with a new uninformative variable elimination algorithm as a new descriptor screening method: Application of data reduction in QSAR study of some sulfonated derivatives as c-Met inhibitors. Chemometrics and Intelligent Laboratory Systems, under review manuscript on 13 May 2022 .

فهرست مطالب

فصل اول: مقدمه	۱
۱-۱ کمومتریکس	۳
۱-۲ ارتباط کمی ساختار- فعالیت (QSAR)	۴
۱-۲-۱ تعریف	۴
۱-۳ اصول مدل سازی QSAR	۴
۱-۳-۱ جمع آوری داده ها	۵
۱-۳-۲ رسم و بهینه سازی ساختار ترکیبات	۶
۱-۳-۳ محاسبه توصیف کننده ها	۷
۱-۳-۴ پیش پردازش	۷
۱-۳-۴-۱ پیش پردازش اولیه	۸
۱-۳-۴-۲ استفاده از روش رگرسیون Ridge برای غربالگری توصیف کننده ها	۸
۱-۳-۵ انتخاب مؤثرترین توصیف کننده ها و مدل سازی	۹
۱-۳-۵-۱ روش مدلسازی	۱۰
۱-۳-۵-۲ روش کمترین توان های دوم معمولی	۱۰
۱-۳-۵-۳ روش های رگرسیون مبتنی بر جریمه	۱۱
۱-۳-۵-۴ رگرسیون Ridge	۱۱
۱-۳-۵-۵ روش SCAD	۱۴
۱-۴ ارزیابی مدل	۱۵
۱-۴-۱ محاسبه پارامترهای آماری	۱۵
۱-۴-۲ ارزیابی مدل با استفاده از نتایج پیش بینی شده برای مجموعه آزمون	۲۱

- ۱- ۴- ۳ ارزیابی مدل با استفاده از تکنیک رد مرحله ای تک تک (LOO) ۲۱
- ۱- ۴- ۴ استفاده از نمودار خطای باقی مانده ها ۲۲
- ۱- ۴- ۵ استفاده از دامنه کاربرد ۲۲
- ۱- ۴- ۶ استفاده از آزمون Y-تصادفی ۲۳
- ۱- ۴- ۷ بررسی نحوه تأثیر متغیرهای مستقل بر متغیر وابسته ۲۳
- ۱- ۵- ۵ شبیه سازی داکینگ مولکولی ۲۴
- ۱- ۵- ۱ فرآیند اعتبار سنجی ۲۵
- ۱- ۵- ۲ آماده سازی پروتئین ۲۶
- ۱- ۵- ۳ ساختن لیگاند ۲۶
- ۱- ۵- ۴ تنظیم کردن جعبه شبکه ای ۲۶
- ۱- ۵- ۵ گزینه داکینگ ۲۷
- ۱- ۵- ۶ انجام محاسبه داکینگ ۲۷
- ۱- ۵- ۷ آنالیز و تحلیل نتایج ۲۸
- ۱- ۵- ۸ کاربردهای داکینگ مولکولی ۲۸
- ۱- ۶- ۶ سرطان ۲۹
- ۱- ۷- ۷ مروری بر کارهای انجام شده و ضرورت تحقیق ۳۰
- ۱- ۸- ۸ نوآوری تحقیق ۳۲
- فصل دوم: نتایج تجربی ۳۵
- ۲- ۱ نرم افزارهای مورد استفاده ۳۶
- ۲- ۱- ۱ نرم افزار هایپرکم ۳۶

۳۶	۲-۱-۲ نرم افزار دراگون
۳۷	۲-۱-۳ نرم افزار SPSS
۳۷	۲-۱-۴ نرم افزار R و R-Studio
۳۷	۲-۱-۵ نرم افزار متلب
۳۸	۲-۱-۶ نرم افزار اتوداک
۳۸	۲-۱-۷ نرم افزار ویورلایت
۳۸	۲-۱-۸ نرم افزار Discovery Studio
۳۹	۲-۲ مدل سازی فعالیت دارویی مشتقات سولفون دار به عنوان بازدارنده های c-Met
۴۱	۲-۲-۱ معرفی مجموعه داده های مورد استفاده در این مطالعه
۴۵	۲-۲-۲ رسم و بهینه سازی ساختار هندسی ترکیبات
۴۵	۲-۲-۳ محاسبه توصیف کننده ها
۴۵	۲-۲-۴ پیش پردازش اولیه
۴۶	۲-۲-۵ استفاده از روش رگرسیون Ridge برای غربالگری توصیف کننده ها
۵۱	۲-۲-۶ انتخاب متغیر و مدل سازی SCAD
۵۵	۲-۲-۷ ارزیابی مدل Ridge-SCAD
۵۵	۲-۲-۸ ارزیابی مدل Ridge-SCAD و مدل SCAD با استفاده از پیش بینی داده های مجموعه آزمون
۵۸	۲-۲-۹ ارزیابی مدل Ridge-SCAD با استفاده از تکنیک رد مرحله ای تک تک
۶۲	۲-۲-۱۰ ارزیابی مدل Ridge-SCAD و مدل SCAD با استفاده از پارامترهای آماری
۶۴	۲-۲-۱۱ ارزیابی مدل با استفاده از دامنه کاربرد
۶۶	۲-۲-۱۲ ارزیابی مدل Ridge-SCAD با استفاده از Y-تصادفی

فصل سوم: نتیجه گیری و آینده نگری	۶۹
۳-۱ تجزیه و تحلیل توصیف کننده های مدل Ridge-SCAD و تأثیر آن ها بر فعالیت دارویی با استفاده از ضرایب رگرسیون استاندارد شده	۷۰
۳-۱-۱ بررسی میزان و چگونگی تأثیر توصیف کننده ها بر فعالیت دارویی مجموعه بازدارنده های سرطان و کاربرد مدل ارائه شده در پیشنهاد ترکیبات جدید	۷۰
۳-۱-۱-۱ محاسبه سهم مشارکت هر توصیف کننده در مدل Ridge-SCAD	۷۰
۳-۱-۱-۲ توصیف کننده دو بعدی همبسته	۷۲
۳-۱-۱-۳ توصیف کننده مکانی	۷۳
۳-۱-۱-۴ توصیف کننده نمایش سه بعدی ساختار مولکول براساس تفرق الکترون	۷۳
۳-۲ پیشنهاد ترکیبات جدید دارای فعالیت مناسب ضد سرطان با استفاده از مدل Ridge-SCAD و مطالعه داکینگ مولکولی ترکیبات پیشنهادی	۷۴
۳-۲-۱ مطالعه داکینگ مولکولی برای بازدارنده های جدید پیشنهادی	۷۵
۳-۳ نتیجه گیری نهایی	۸۸
۳-۴ آینده نگری	۸۹
۳-۵ فهرست منابع	۹۰

فهرست شکل‌ها:

- شکل ۱-۱ شمای کلی از مراحل QSAR ۵
- شکل ۲-۱ عملکرد رگرسیون Ridge ۱۳
- شکل ۳-۱ اتصال لیگاند به مولکول هدف ۲۵
- شکل ۱-۲ اسکلت اصلی ساختار ترکیبات مورد مطالعه مربوط به ترکیبات ۱ تا ۳۹ ۴۲
- شکل ۲-۲ اسکلت اصلی ساختار ترکیبات مورد مطالعه مربوط به ترکیبات ۴۰ تا ۶۳ ۴۴
- شکل ۳-۲ فلوجارت نمایش روند غربالگری Ridge (زیروند Original مربوط به توصیف کننده های مرحله پیش پردازش، زیروند old مربوط به توصیف کننده های اولین مرحله غربالگری Ridge و زیروند new مربوط به توصیف کننده های حاصل از هر مرحله اجرای روند غربالگری Ridge است.) ۵۰
- شکل ۴-۲ نمودار تغییرات مقادیر پیش بینی شده (pIC_{50}) مدل Ridge-SCAD در مقابل مقادیر تجربی برای مجموعه آزمون ۵۷
- شکل ۵-۲ نمودار مقادیر باقی مانده های پیش بینی فعالیت دارویی با استفاده از مدل Ridge-SCAD و مقادیر تجربی بر حسب مقادیر تجربی ۵۷
- شکل ۶-۲ نمودار تغییرات مقادیر پیش بینی شده (pIC_{50}) مدل SCAD در مقابل مقادیر تجربی برای مجموعه آزمون ۵۸
- شکل ۷-۲ نمودار تغییرات مقادیر پیش بینی شده همه داده ها بر اساس تکنیک LOO در مقابل مقادیر تجربی .. ۶۱
- شکل ۸-۲ نمودار باقی مانده های حاصل از پیش بینی فعالیت دارویی ترکیبات با استفاده از تکنیک LOO و مقادیر تجربی بر حسب مقادیر تجربی ۶۱
- شکل ۹-۲ نمودار ویلیام (دامنه کاربرد) مدل Ridge-SCAD، خطوط نقطه چین افقی و عمودی در دو انتها نمودار ۶۵
- شکل ۱۰-۲ نمودار مقادیر R^2 به دست آمده در آزمون Y-تصادفی بر حسب تعداد اجرا، برای ۱۰۰ بار اجرای آزمون Y-تصادفی و پیش بینی فعالیت ترکیبات مجموعه آزمون به وسیله مدل Ridge-SCAD توسعه یافته با استفاده از پاسخ تصادفی شده ۶۷

- شکل ۱-۳ نمودار سهم مشارکت توصیف کننده ها در مدل Ridge-SCAD ۷۲
- شکل ۲-۳ ساختار کریستالوگرافی ۳LQ8 (منطقه انتخاب شده نشان دهنده لیگاند کریستالوگرافی و مابقی زنجیره های اسید آمینه ای است) ۷۶
- شکل ۳-۳ بر همکنش ترکیب ۶۳ با اسید آمینه های کلیدی ۷۸
- شکل ۴-۳ بر همکنش ترکیب ۵ با اسید آمینه های کلیدی ۷۹
- شکل ۵-۳ بر همکنش ترکیب جدید ۱ با اسید آمینه های کلیدی ۸۰
- شکل ۶-۳ بر همکنش ترکیب جدید ۲ با اسید آمینه های کلیدی ۸۰
- شکل ۷-۳ بر همکنش ترکیب جدید ۳ با اسید آمینه های کلیدی ۸۱

فهرست جدول‌ها:

جدول ۱-۲ جزئیات ساختار ۱ تا ۳۹ به همراه مقادیر λ pIC ترکیبات مورد مطالعه.....	۴۲
جدول ۲-۲ جزئیات ساختار ۴۰ تا ۶۳ به همراه مقادیر λ pIC ترکیبات مورد مطالعه.....	۴۴
جدول ۳-۲ توصیف کننده های انتخاب شده توسط روش Ridge-SCAD.....	۵۳
جدول ۴-۲ ماتریس همبستگی و مقادیر عامل افزایش واریانس توصیف کننده های انتخاب شده مدل Ridge-SCAD.....	۵۳
جدول ۵-۲ توصیف کننده های انتخاب شده توسط روش SCAD.....	۵۴
جدول ۶-۲ نتایج حاصل از پیش بینی فعالیت دارویی داده های مجموعه آزمون توسط مدل Ridge-SCAD.....	۵۶
جدول ۷-۲ نتایج حاصل از ارزیابی مدل Ridge-SCAD به روش رد مرحله ای تک تک برای کل داده ها.....	۵۹
جدول ۸-۲ پارامترهای آماری برای مجموعه آموزش و آزمون داده های پیش بینی شده مدل.....	۶۳
جدول ۹-۲ مقایسه پارامترهای آماری محاسبه شده برای مجموعه آزمون داده های پیش بینی شده مدل Ridge-SCAD و SCAD.....	۶۴
جدول ۱۰-۳ اثر متوسط توصیف کننده های انتخاب شده توسط روش Ridge-SCAD.....	۷۰
جدول ۲-۳ پارامترهای محاسبه شده برای ترکیبات مورد مطالعه و ترکیبات پیشنهادی.....	۸۲
جدول ۳-۳ پارامترهای محاسبه شده قاعده ۵ لیپینسکی.....	۸۷

فهرست رابطه‌ها:

۱۰.....	رابطه ۱-۱.....
۱۰.....	رابطه ۲-۱.....
۱۰.....	رابطه ۳-۱.....
۱۰.....	رابطه ۴-۱.....
۱۱.....	رابطه ۵-۱.....
۱۲.....	رابطه ۶-۱.....
۱۲.....	رابطه ۷-۱.....
۱۲.....	رابطه ۸-۱.....
۱۲.....	رابطه ۹-۱.....
۱۲.....	رابطه ۱۰-۱.....
۱۲.....	رابطه ۱۱-۱.....
۱۲.....	رابطه ۱۲-۱.....
۱۲.....	رابطه ۱۳-۱.....
۱۲.....	رابطه ۱۴-۱.....
۱۳.....	رابطه ۱۵-۱.....
۱۴.....	رابطه ۱۶-۱.....
۱۴.....	رابطه ۱۷-۱.....
۱۴.....	رابطه ۱۸-۱.....
۱۵.....	رابطه ۱۹-۱.....
۱۶.....	رابطه ۲۰-۱.....
۱۶.....	رابطه ۲۱-۱.....
۱۶.....	رابطه ۲۲-۱.....

١٦.....	رابطه ٢٣-١.....
١٦.....	رابطه ٢٤-١.....
١٧.....	رابطه ٢٥-١.....
١٧.....	رابطه ٢٦-١.....
١٧.....	رابطه ٢٧-١.....
١٨.....	رابطه ٢٨-١.....
١٨.....	رابطه ٢٩-١.....
١٨.....	رابطه ٣٠-١.....
١٨.....	رابطه ٣١-١.....
١٨.....	رابطه ٣٢-١.....
١٩.....	رابطه ٣٣-١.....
١٩.....	رابطه ٣٤-١.....
٢٠.....	رابطه ٣٥-١.....
٢٠.....	رابطه ٣٦-١.....
٢٠.....	رابطه ٣٧-١.....
٢٠.....	رابطه ٣٨-١.....
٢٠.....	رابطه ٣٩-١.....
٢٢.....	رابطه ٤٠-١.....
٢٣.....	رابطه ٤١-١.....
٢٣.....	رابطه ٤٢-١.....
٢٤.....	رابطه ٤٣-١.....
٢٤.....	رابطه ٤٤-١.....
٤٧.....	رابطه ١-٢.....
٤٧.....	رابطه ٢-٢.....

۴۷.....	رابطه ۳-۲.....
۵۱.....	رابطه ۴-۲.....
۵۲.....	رابطه ۵-۲.....
۵۲.....	رابطه ۶-۲.....
۵۴.....	رابطه ۷-۲.....
۷۱.....	رابطه ۱-۳.....
۷۲.....	رابطه ۲-۳.....
۷۳.....	رابطه ۳-۳.....

فصل اول

مقدمه

مقدمه

از بزرگ‌ترین مشکلاتی که همواره در مسیر توسعه جوامع بشری قرار داشته، مربوط به سلامت و مبارزه با انواع بیماری‌ها از قبیل سرطان، ایدز، دیابت، آلزایمر و ... بوده است. از این‌رو بایستی به‌دنبال داروها و روش‌هایی برای از بین بردن مشکلات بود. سرطان شاید یکی از پیچیده‌ترین مشکلات بشر در زمینه سلامتی باشد که هر ساله هزینه‌های مالی و عاطفی فراوانی به جوامع تحمیل می‌کند. علی‌رغم پیشرفت در پیشگیری، تشخیص زود هنگام، مدیریت جراحی و پیشرفت در شیمی درمانی، توانایی درمان بیماران سرطانی همچنان هدف مبهمی است و نیاز شدیدی به توسعه درمان‌های مؤثرتر دارد. بنابراین، جستجوی گسترده عوامل ضد سرطانی جدید با آشکار سازی اهداف مولکولی نوین به‌دنبال کشف ترکیباتی است که با چنین اهدافی تعامل دارند. از آنجا که طراحی دارو در آزمایشگاه زمان‌بر و پرهزینه بوده، طراحی محاسباتی دارو کمک شایانی، به صرفه‌جویی در زمان، هزینه‌های آزمایشگاهی و پژوهشی کرده است. از این‌رو نیاز به استفاده از روش‌های تئوری و محاسباتی بدون انجام آزمایش که بتواند ویژگی یا فعالیت ترکیبات را پیش‌بینی کند، ضروری به‌نظر می‌رسد. ظهور علم کمومتریکس توانسته راه حلی برای رفع این محدودیت‌ها باشد. بدین صورت که، با به‌دست آوردن رابطه کمی بین ساختار شیمیایی ترکیبات و فعالیت دارویی آن‌ها می‌توان فعالیت دارویی ترکیبات مشابه را پیش‌بینی کرد.

۱-۱ کمومتریکس^۱

اصطلاح کمومتریکس اولین بار توسط شیمیدان آلی اسوانت ولد^۲ سال ۱۹۷۲ در یک مجله سوئدی ذکر شد. ولد با همکاری بروس آر. کووالسکی^۳ که در دانشگاه واشنگتن بر روی الگوشناسی در شیمی تجزیه مطالعه می‌کرد، انجمن بین‌المللی کمومتریکس^۴ (ICS) را در سال ۱۹۷۴ تأسیس کردند [۱]. کمومتریکس، توسط ICS، به‌صورت زیر تعریف شده است:

"کمومتریکس استفاده از روش‌های ریاضی، آماری برای به‌وجود آوردن ارتباط بین سنجش‌های انجام شده روی یک سیستم یا فرآیند شیمیایی برای فهم بهتر اطلاعات شیمیایی است." اطلاعات شیمیایی که اساس بسیاری از فرآیندهای مهم را تشکیل می‌دهد، به سه ویژگی مهم فرآیند اندازه‌گیری بستگی دارد، از جمله:

۱. خواص شیمیایی: شامل استوکیومتری، تعادل جرم، تعادل شیمیایی، سینتیک و ...
 ۲. خواص فیزیکی: شامل دما، انتقال انرژی، انتقال فاز و ...
 ۳. خواص آماری: شامل منابع خطا در فرآیند اندازه‌گیری، کنترل عوامل تداخل، کالیبراسیون سیگنال‌های پاسخ، مدل‌سازی سیگنال‌های پیچیده چند متغیره و ...
- اگر یکی از این سه ویژگی کامل نباشد، سیستم اندازه‌گیری نمی‌تواند نتایج قابل اطمینان را ارائه دهد، که گاهی اوقات پیامدهای فاجعه‌باری به‌همراه خواهد داشت. نقش آمار و کمومتریکس در رسیدگی به سومین ویژگی، بحرانی است. این نقش اساسی، انگیزه اولیه را برای پیشرفت در زمینه کمومتریکس فراهم می‌کند [۲].

¹ Chemometrics

² Swante Wold

³ Bruce R. Kowaski

⁴ International Chemometrics Society

از بین کاربردهای مختلف کمومتریکس مهم‌ترین و شاخص‌ترین این کاربردها طراحی دارو است، که ارتباط کمی ساختار- فعالیت^۱ (QSAR) را بررسی می‌کند و خواص مولکول‌ها را به ویژگی‌های ساختاری آنها نسبت می‌دهد [۳]. در ادامه به شرح QSAR پرداخته شد.

۱-۲ ارتباط کمی ساختار- فعالیت (QSAR)

۱-۲-۱ تعریف

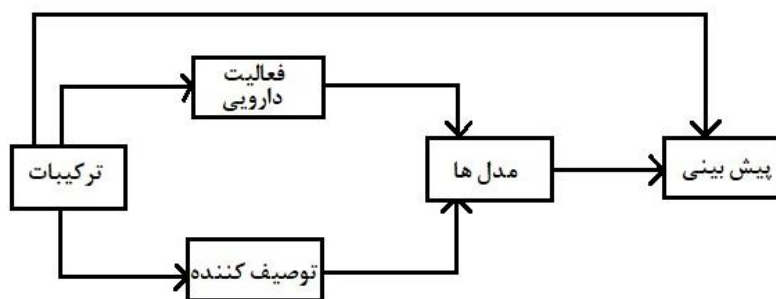
ارتباط کمی ساختار- فعالیت (QSAR)، در ساده‌ترین تعریف، روشی جهت ساختن مدل‌های ریاضی یا محاسباتی است که می‌تواند ارتباط معنا دار آماری را بین ساختار و عملکرد ایجاد کند. تلاش برای برقراری ارتباط بین ساختار شیمیایی و فعالیت زیستی موضوع QSAR است [۴].

۱-۳ اصول مدل سازی QSAR

در مطالعات QSAR ارتباط ساختار مولکول‌ها با فعالیت دارویی آنها طی مراحل زیر بیان می‌شود [۵]:

۱. جمع‌آوری داده‌ها
۲. رسم و بهینه‌سازی ساختار ترکیبات
۳. محاسبه و غربالگری توصیف‌کننده‌ها
۴. انتخاب مؤثرترین توصیف‌کننده‌ها
۵. ایجاد مدل‌های آماری
۶. ارزیابی مدل

¹ Quantitative structure- activity relationship



شکل ۱-۱ شمای کلی از مراحل QSAR

بنابراین همان‌طور که شکل ۱-۱ نشان می‌دهد، اساس رابطه کمی ساختار - فعالیت (QSAR) ایجاد رابطه بین فعالیت دارویی مشاهده شده و ویژگی‌های ساختاری یک مولکول است. با در نظر گرفتن مجموعه‌ای از مولکول‌ها، یک مدل پیش‌بینی ایجاد می‌شود که می‌تواند برای پیش‌بینی فعالیت دارویی سایر مولکول‌ها مورد استفاده قرار گیرد [۶].

۱-۳-۱ جمع‌آوری داده‌ها

اولین مرحله مدل‌سازی جمع‌آوری داده‌ها است که مقادیر تجربی فعالیت آنها باید با صحت قابل قبول اندازه‌گیری شده باشد. داده‌های تجربی می‌توانند در آزمایشگاه تولید شده یا از مقادیر منتشر در مقالات یا پایگاه داده‌ها استخراج شوند. در مدل‌سازی QSAR مجموعه داده‌ها با استفاده از آنالیزهای متفاوت تقسیم‌بندی از جمله کنار-استون^۱ (KS) انجام می‌شود [۷، ۸]. تقسیم‌بندی KS بر اساس محاسبه فاصله اقلیدسی^۲ است. بنابراین مجموعه داده‌ها به دو دسته مجموعه آموزش^۳ و مجموعه آزمون^۴ تقسیم می‌شوند. الگوریتم KS داده‌ها را طوری تقسیم می‌کند که مجموعه‌های تقسیم‌بندی شده مانند مجموعه آزمون، نماینده مناسبی از مجموعه آموزش باشد. از این‌رو داده‌های مجموعه آموزش و آزمون به‌ترتیب برای ساخت

^۱ Kennard-Stone (KS)

^۲ Euclidean distance

^۳ Training set

^۴ Test set

مدل و جهت بررسی قدرت پیش‌بینی و اعتبار مدل استفاده می‌شود [۹]. لازم به ذکر است که مجموعه آزمون در هیچ یک از مراحل انتخاب متغیر و ساخت مدل‌های QSAR شرکت نخواهد داشت.

۱-۳-۲ رسم و بهینه سازی ساختار ترکیبات

بعد از جمع آوری مجموعه داده‌ها ساختار شیمیایی ترکیبات با استفاده از نرم افزار هایپرکم^۱ رسم و بهینه‌سازی می‌شود. ساختار ترکیبات به کمک روش‌های شیمی محاسباتی^۲ بهینه می‌شود تا مختصاتی از ساختار به دست آید که در آن، انرژی مولکول حداقل باشد. که این روش‌ها عبارتند از:

۱. روش‌های مکانیک کوانتومی^۳ (شامل روش‌های آغازین^۴، روش‌های نیمه تجربی^۵، روش‌های تابع

چگالی^۶)

۲. روش مکانیک مولکولی^۷

۳. روش‌های دینامیک مولکولی^۸

در نرم افزارهایی نظیر هایپرکم اغلب از روش نیمه تجربی (AM1) استفاده می‌شود. بهینه‌سازی ساختار ترکیبات شیمیایی با استفاده از این روش انجام می‌شود و پایدارترین حالت با کم‌ترین انرژی برای هر ساختار به دست می‌آید. ساختارهای بهینه شده به عنوان ورودی نرم‌افزارهای محاسبه‌ای توصیف کننده-های مولکولی به کار گرفته می‌شوند.

¹ Hyperchem

² Computational Chemistry

³ Quantum Mechanic

⁴ Ab Initio Methods

⁵ Semi-empirical Methods

⁶ Density Functional Theory

⁷ Molecular Mechanic

⁸ Molecular dynamic

۱-۳-۳ محاسبه توصیف کننده‌ها

توصیف کننده‌ها به عنوان اعداد کمی، نشان دهنده ویژگی مواد شیمیایی هستند و به ایجاد ارتباط ریاضی کمک می‌کنند. بنابراین، انواع مختلف توصیف کننده نقش مهمی در شناسایی و همچنین تجزیه و تحلیل ترکیبات ایفا می‌کنند. اگرچه طیف گسترده‌ای از توصیف کننده‌ها برای استفاده در دسترس هستند، اما، هدف استفاده یا توسعه توصیف کننده‌هایی است که به راحتی محاسبه شوند و مقدار مشخصی از اطلاعات شیمیایی را ارائه دهند [۱۰].

۱-۳-۴ پیش پردازش^۱

بسیاری از برنامه‌ها قادر به تولید صدها یا هزاران توصیف کننده مولکولی مختلف هستند. به طوری که این برنامه‌ها می‌توانند برای تعداد کمی از ترکیبات (n)، هزاران توصیف کننده (p)، محاسبه کنند. به عبارتی $n < p$ ، که به عنوان داده‌های با ابعاد بالا^۲ نامیده می‌شوند. ابعاد داده‌ها باید کاهش یابد تا علاوه بر این که مدل قدرت پیش‌بینی بالایی داشته باشد، قدرت تفسیر پذیری روشی هم داشته باشد. مدل‌هایی که شامل مجموعه بزرگی از توصیف کننده‌های مولکولی هستند، تمایل به داشتن واریانس بزرگ دارند که اغلب منجر به عملکرد پیش‌بینی ضعیف می‌شود. بنابراین انتخاب تنها توصیف کننده‌های مولکولی مهم می‌تواند دقت پیش‌بینی را بهبود بخشد [۱۱]. از طرفی دیگر در داده‌هایی با ابعاد بالا مشکل هم‌خطی بین توصیف کننده‌ها وجود دارد، به طوری که وابستگی خطی منجر به ایجاد پدیده هم‌خطی شدید شده و صحت پیش‌بینی مدل را کم می‌کند. به همین دلیل بهتر است ابعاد داده‌ها کاهش یابد. بنابراین انتخاب توصیف کننده‌های مؤثری که بتواند ارتباط معنا دار بین متغیر وابسته و متغیرهای مستقل را نشان دهد، باعث بهبود صحت

^۱ preprocessing

^۲ High Dimensional Data

پیش‌بینی می‌شود. با توجه به این توضیحات استفاده از روش غربالگری^۱ قبل از توسعه مدل حائز اهمیت است.

۱-۳-۴-۱ پیش پردازش اولیه

به‌منظور پیدا کردن توصیف‌کننده‌های مؤثر، در مرحله‌ی اول، توصیف‌کننده‌هایی با مقادیر ثابت و نسبتاً ثابت (انحراف استاندارد کمتر از ۰/۰۰۱) [۱۲، ۱۳] از مجموعه داده‌ها حذف می‌شود. در مرحله بعد از بین دو توصیف‌کننده با همبستگی بالای ۰/۹ ($R > 0.9$) آن توصیف‌کننده‌ای که ارتباط بیشتری با متغیر وابسته (فعالیت دارویی) دارد، نگه داشته می‌شود و توصیف‌کننده دیگر حذف می‌گردد تا مجموعه‌ای تنک با توصیف‌کننده‌های مؤثرتر توسعه یابد.

۱-۳-۴-۲ استفاده از روش رگرسیون Ridge برای غربالگری توصیف‌کننده‌ها

در این پروژه به‌منظور غربالگری بیش‌تر توصیف‌کننده‌ها و حذف توصیف‌کننده‌های فاقد اطلاعات مفید، از روش رگرسیون انقباضی Ridge [۱۴] به‌عنوان یک روش غربالگری جدید استفاده گردید. روش رگرسیون Ridge با اعمال محدودیت مجموع مجذور پارامترها کمتر از یک مقدار ثابت بر روی تابع کم‌ترین توان‌های دوم موجب انقباض پارامترها و کاهش هم‌خطی و همچنین کاهش واریانس می‌شود و در نهایت دقت پیش‌بینی مطلوب‌تری نسبت به روش برآورد کم‌ترین توان‌های دوم ارائه می‌دهد. اما برآوردگر Ridge نیز در مواقعی که تعداد متغیرها بیش از تعداد ترکیبات است، همانند روش کم‌ترین توان‌های دوم قابل محاسبه نیست. در برآوردگر Ridge با تغییر کوچکی در داده‌ها تغییر زیادی در مدل نهایی صورت نمی‌پذیرد، بنابراین مدل پایداری را ارائه می‌دهد. اما رگرسیون Ridge هیچ یک از پارامترها را حذف نمی‌کند و به این دلیل برای مدل‌هایی با تعداد متغیرهای زیاد مدل ساده‌ای را نخواهیم داشت. با تلفیق یک روش جدید از این روش، برای غربالگری و حذف متغیرهایی با اهمیت کمتر استفاده شده است. این روش با استفاده از

¹ screening

تکنیک رد مرحله‌ای تک تک^۱ روی توصیف کننده‌های باقی‌مانده اجرا می‌شود. به این منظور برای اجرای روش Ridge به‌عنوان یک روش غربالگری به‌ترتیب سه ماتریس ساخته شد. ماتریس زائد^۲ (R) شامل توصیف کننده‌های با کم‌ترین اطلاعات مفید است که ارتباط معنا داری با پاسخ ندارند. ماتریس اصلی (X) متشکل از توصیف کننده‌هایی است، که ارتباط معنا داری با پاسخ دارند و ماتریس XR به‌ترتیب شامل توصیف کننده‌های ماتریس‌های X و R است.

۱-۳-۵ انتخاب مؤثرترین توصیف کننده ها و مدل سازی

انتخاب متغیر و مدل‌سازی به روش خطی و غیرخطی انجام می‌شود، که روش‌های خطی به روش‌های کلاسیک و انقباضی یا جریمه‌ای^۳ تقسیم می‌شوند. از بین روش‌های کلاسیک می‌توان به روش‌های کمترین توان‌های دوم معمولی (OLS) که شامل انتخاب متغیر پیش رونده^۴، انتخاب متغیر حذفی پس رونده^۵ و انتخاب متغیر گام به گام^۶ اشاره کرد [۱۵]. هر چند این روش‌ها، زیرمجموعه مناسبی را از بین تمام زیرمجموعه‌های ممکن انتخاب می‌نمایند، اما زمانی که ابعاد داده‌ها بزرگ باشد تعداد زیر مجموعه‌های منتخب افزایش می‌یابد و با افزایش تعداد متغیرهای مدل، مقدار ضریب تعیین به‌طور کاذب بالا بوده و قدرت پیش‌بینی مدل پایین می‌آید. علاوه بر این، روش‌های انتخاب متغیر کلاسیک دارای مشکل بی‌ثباتی هستند. به‌طوری‌که با تغییری کوچک در داده‌ها، مدل‌های خیلی متفاوتی به‌وجود می‌آید. بنابراین استفاده از روش‌های انتخاب متغیر خطی چندگانه به‌دلیل داشتن معایبی هم‌چون ناپایداری، بایاس بالا، ناکارآمدی در حضور هم‌خطی، برای انتخاب توصیف کننده‌های منتخب توصیه نمی‌شود [۱۶]. روش‌های مبتنی بر جریمه شامل

¹ Leave One Out

² Redundant

³ Shrinkage and Penalized methods

⁴ Forward

⁵ Backward

⁶ Stepwise

روش Ridge [۱۴]، روش انحراف قدر مطلق به‌طور هموار کوتاه شده^۱ (SCAD) [۱۷] و ... می‌باشد. در ادامه به شرح روش‌های کلاسیک و روش‌های جریمه‌ای پرداخته شد.

۱-۳-۵ روش مدل‌سازی

۱-۳-۵-۲ روش کم‌ترین توان‌های دوم معمولی

روش کم‌ترین توان‌های دوم معمولی (OLS) در سال ۱۸۰۹ توسط گوس^۲ معرفی شد، که به‌طور رایج برای تخمین ضرایب رگرسیون استفاده می‌شود.

مدل رگرسیون خطی چندگانه^۳ زیر را در نظر بگیرید رابطه ۱-۱ :

$$y = X\beta + \epsilon \quad \text{رابطه ۱-۱}$$

که در آن $y = (y_1, \dots, y_n)^T$ بردار ستونی پاسخ‌ها، $X = (x_1, \dots, x_n)^T$ ماتریس طرح در اندازه $n \times p$ با رتبه کامل ستونی p شامل متغیرهای مستقل (پیشگو)، $\beta = (\beta_1, \dots, \beta_p)^T$ بردار ضرایب رگرسیونی و $\epsilon = (\epsilon_1, \dots, \epsilon_n)^T$ بردار مؤلفه‌های خطای تصادفی می‌باشد.

تخمین کم‌ترین توان‌های دوم معمولی (OLS) با مینیمم کردن مجموع توان‌های دوم خطاهای مدل یعنی $\epsilon^T \epsilon = (y - X\beta)^T (y - X\beta)$ به‌وجود می‌آید. در این صورت برآوردگر OLS بردار پارامترها به‌صورت رابطه ۲-۱ محاسبه می‌شود:

$$\hat{\beta} = \underset{\beta}{\operatorname{argmin}} \epsilon^T \quad \text{رابطه ۲-۱}$$

با توجه به این‌که

$$\frac{\partial \epsilon^T \epsilon}{\partial \beta} = -2X^T y + 2X^T X \beta = 0 \quad \text{رابطه ۳-۱}$$

می‌توان نتیجه گرفت که :

$$\hat{\beta} = (X^T X)^{-1} X^T y \quad \text{رابطه ۴-۱}$$

^۱ Smoothly Clipped Absolute Deviation

^۲ Gauss

^۳ Multiple Linear Regression (MLR)

برای بهبود مدل‌های خطی، یک پارامتر جریمه^۱ به روش برآورد OLS اضافه شده که روش‌های رگرسیون‌های مبتنی بر جریمه^۲ نامیده می‌شود، که در ادامه به شرح آن‌ها پرداخته شد.

۱-۳-۵ روش‌های رگرسیون مبتنی بر جریمه

روش‌های رگرسیون مبتنی بر جریمه، همه متغیرهای پیش‌بینی کننده را در مدل نگه می‌دارند، اما، با منقبض کردن ضرایب رگرسیون به سمت صفر و صفر کردن برخی از آن‌ها عمل برآوردیابی و انتخاب متغیر را انجام می‌دهند. این انقباض معمولاً از طریق ایجاد محدودیت روی دامنه تغییرات ضرایب در هنگام برآوردیابی صورت می‌گیرد. هدف از این کار وارد کردن اریبی در مدل با هدف کاهش واریانس و در نهایت کاهش مجموع مربعات خطا و بالا بردن دقت پیش‌بینی است. این روش‌ها باعث افزایش دقت پیش‌بینی، افزایش پایداری و بهبود تفسیر مدل‌های توسعه یافته می‌شوند.

۱-۳-۵ Ridge رگرسیون

برای اولین بار رگرسیون Ridge در سال ۱۹۷۰ توسط هورل و کنارد^۳ معرفی شد. پایه و اساس روش برآوردگر OLS یک رگرسیون خطی، این است که $(x^T x)^{-1}$ وجود داشته باشد و به دو دلیل این معکوس وجود ندارد: ۱. ماتریس طرح پر رتبه ستونی نباشد ۲. مشکل هم‌خطی وجود داشته باشد. یکی از روش‌های حل این مشکل استفاده از روش رگرسیون Ridge می‌باشد.

راهی آسان برای تضمین معکوس پذیری، اضافه کردن ماتریس قطری λI به $x^T x$ است که در آن λ یک مقدار ثابت مثبت و I ماتریس همانی می‌باشد. از این‌رو برآوردگر رگرسیون Ridge پارامتر β به صورت رابطه ۵-۱ است:

$$\hat{\beta} = (x^T x + \lambda I)^{-1} x^T \quad \text{رابطه ۵-۱}$$

^۱ Penalty

^۲ penalized regression methods

^۳ Hoerl and Kennard

که $\lambda > 0$ است. همان گونه که با مینیم کردن $\epsilon^T \epsilon$ نسبت به β ، برآوردگر OLS به دست آمد، این برآوردگر را نیز می توان با مینیم کردن عبارت مجموع توان های دوم خطا $\epsilon^T \epsilon$ نسبت به β تحت شرط $\beta^T \beta = \sum_{j=1}^p \beta_j^2 \leq t$ محاسبه کرد. در این جا $t \geq 0$ یک پارامتر تنظیم کننده است که میزان انقباض ضرایب را کنترل می کند. به بیانی هر چه مقدار t کوچک تر باشد ضرایب بزرگ تر بیش تر جریمه می شوند. برای حل مسئله بهینه سازی بالا از ضرایب لاگرانژ می توان استفاده کرد. به بیانی دیگر برآوردگر Ridge به صورت رابطه ۶-۱ به دست می آید:

$$\hat{\beta}_R = \operatorname{argmin}_{\beta} (\epsilon^T \epsilon + \lambda \beta^T \beta) \quad \text{رابطه ۶-۱}$$

که در آن λ ضریب لاگرانژ یا همان پارامتر انقباض Ridge می باشد. از آنجایی که

$$\frac{\partial (\epsilon^T \epsilon + \lambda \beta^T \beta)}{\partial \beta} = -2x^T y + 2x^T x \hat{\beta} + 2\lambda \beta = 0 \quad \text{رابطه ۷-۱}$$

داریم:

$$(x^T x + \lambda I_p) \beta = x^T y \quad \text{رابطه ۸-۱}$$

و بنابراین

$$\hat{\beta}_R = (x^T x + \lambda I_p)^{-1} x^T y \quad \text{رابطه ۹-۱}$$

$\hat{\beta}_R$ یک برآوردگر اریب با میانگین و واریانس زیر است:

$$E(\hat{\beta}_R) = \beta + \lambda (X^T X + \lambda I)^{-1} \beta \quad \text{رابطه ۱۰-۱}$$

$$\operatorname{Var}(\hat{\beta}_R) = (X^T X + \lambda I)^{-1} X^T X (X^T X + \lambda I)^{-1} \sigma^2 \quad \text{رابطه ۱۱-۱}$$

در صورتی که در خصوص برآوردگر OLS داریم:

$$\epsilon \hat{\beta} = \beta \quad \text{رابطه ۱۲-۱}$$

$$\operatorname{Var}(\hat{\beta}) = (X^T X)^{-1} \sigma^2 \quad \text{رابطه ۱۳-۱}$$

هورل و کنارد اثبات کرده اند که اگر $\beta^T \beta$ کراندار باشد می توان $\lambda > 0$ -یی را پیدا کرد به گونه ای که:

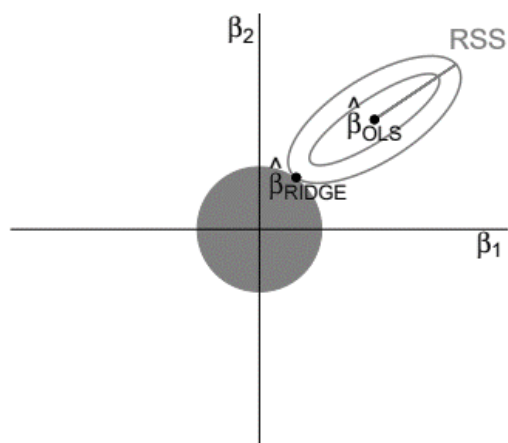
$$\operatorname{MSE}(\hat{\beta}_R) < \operatorname{MSE}(\hat{\beta}) \quad \text{رابطه ۱۴-۱}$$

بنابراین رگرسیون Ridge می تواند برآورد OLS را بهبود ببخشد [۱۴].

برای یافتن درک بصری روش Ridge، با فرض وجود تنها دو متغیر در مدل رگرسیونی ناحیه محدودیت به صورت رابطه ۱-۱۵ :

$$\beta_1^2 + \beta_2^2 \leq t \quad \text{رابطه ۱-۱۵}$$

همان طور که در شکل ۱-۲ نشان داده شده دایره ناحیه تاوان است. محل برخورد مجموع توان های دوم باقی مانده ها^۱ (RSS) روش کمترین توان های دوم با دایره، مختصات برآوردگر Ridge را نمایش می دهد. با عمود کردن این نقطه بر محور افقی، $\hat{\beta}_1$ و با عمود کردن آن بر محور عمودی، $\hat{\beta}_2$ حاصل می شود. روش Ridge فرآیندی پیوسته است که ضرایب تولیدی آن پایدار هستند اما برآوردهایی اریب ارائه می دهد که پذیرفتن مقداری اریبی در مقابل یافتن پاسخی قابل اطمینان تر معقول به نظر می رسد. روش Ridge ضرایب را به سمت صفر منقبض می کند، هرچه مقدار پارامتر Ridge بیش تر باشد ضرایب برآورد شده انقباض بیش تری داشته و کوچک تر می شوند ولی هیچ گاه حتی در بی نهایت نیز به صفر نمی رسند این امر به سبب خاصیت دایره (ناحیه تاوان روش Ridge) رخ می دهد. روش Ridge یک روش مدل سازی است، به دلیل این که همه متغیرها را در مدل نگه می دارد برای انتخاب متغیر استفاده نمی شود [۱۶].



شکل ۱-۲ عملکرد رگرسیون Ridge

^۱ Residual Sum of Squares

۱-۳-۵-۵ روش SCAD

در سال ۲۰۰۱، فن و لی^۱ تابع جریمه انحراف قدر مطلق به طور هموار کوتاه شده (SCAD) را به

معادله OLS اضافه کردند که ضرایب رگرسیون را تخمین می‌زند:

$$\hat{\beta}^{ols} = \arg \min_{\beta \in \mathbb{R}^{p+1}} \|Y - X\beta\|^2 = (X^T X)^{-1} X^T Y \quad \text{رابطه ۱۶-۱}$$

$$\hat{\beta}^{SCAD} = \arg \min_{\beta \in \mathbb{R}^{p+1}} \{\|Y - X\beta\|^2 + \lambda \sum_{i=1}^p p(|\beta_i|); \alpha, \lambda\} \quad \text{رابطه ۱۷-۱}$$

که $p(\cdot; \alpha, \lambda)$ تابع جریمه SCAD هست و سه آرگومان زیر را دنبال می‌کند:

$$P(t; \alpha, \lambda) = \begin{cases} \lambda t & 0 \leq t \leq \lambda \\ \frac{2\alpha\lambda t - t^2 - \lambda^2}{2(\alpha-1)} & \lambda \leq t < \alpha\lambda \\ \frac{(\alpha+1)\lambda^2}{2} & t \geq \alpha\lambda \end{cases} \quad \text{رابطه ۱۸-۱}$$

جایی که $\alpha > 2$ و $\lambda > 0$ پارامترهای جریمه هستند و بزرگی ضرایب را کنترل می‌کنند. فن و لی

مقدار بهینه α را به طور تجربی $3/7$ بیان کردند. با یک مطالعه نشان دادند که مقدار $3/7$ برای α می‌تواند

نتایج خوبی داشته باشد.

فن و لی ویژگی‌های یک تابع توان خوب را به صورت زیر بیان کردند:

۱. نا اریبی: برآوردگر به دست آمده برای ضرایب رگرسیونی که در واقع مقادیر آن‌ها بزرگ

است، تقریباً نا اریب باشد. این ویژگی اریبی مدل را کاهش می‌دهد.

۲. تنگی^۲: برآوردگر نتیجه شده باید به طور خودکار ضرایب برآورد شده‌ای که مقدار کوچکی

دارند را برابر صفر قرار دهد تا متغیرهای مناسب را انتخاب کند. این کار پیچیدگی مدل را

کاهش می‌دهد.

۳. پیوستگی: برآوردگر نتیجه شده دارای تغییرات پیوسته باشد و ناپایداری در پیش‌بینی مدل

را کاهش دهد [۱۷].

^۱ Fan & Li

^۲ Sparsity

رگرسیون Ridge از ویژگی تنکی بی بهره است، زیرا توانایی خارج کردن ضرایب بی اهمیت از مدل را ندارد. تابع جریمه SCAD، دارای این سه ویژگی می باشد. در این تحقیق از روش SCAD برای انتخاب متغیر و مدل سازی استفاده شده است.

۱-۴ ارزیابی مدل

از مرحله های اساسی ساخت و توسعه مدل QSAR، ارزیابی مدل است. در این پروژه، قدرت پیش بینی و تعمیم پذیری مدل با استفاده از پیش بینی فعالیت دارویی داده ها مجموعه آزمون، مجموعه پیش بینی شده به وسیله تکنیک رد مرحله ای تک تک (LOO)، محاسبه پارامترهای آماری، آزمون های دامنه کاربرد^۱ و Y-تصادفی^۲ مورد ارزیابی قرار گرفت.

۱-۴-۱ محاسبه پارامترهای آماری

بعد از اینکه مدل ساخته شد توسط یک مجموعه پارامتر آماری ارزیابی می شود تا از توانایی پیش بینی مدل برای نمونه های مشابه اطمینان حاصل شود. فرض کنید $x_1, x_2, \dots, x_n$ نشان دهنده متغیرهای مستقل (پیشگو) و $y_1, y_2, \dots, y_n$ نشان دهنده متغیرهای وابسته در مدل مورد بررسی باشند. برخی از این پارامترها عبارتند از :

ضریب همبستگی^۳: یکی از معیارهای مورد استفاده در تعیین همبستگی دو متغیر است. این ضریب بین ۱ تا -۱ است. مقدار بزرگتر آن نشان دهنده این است که ارتباط خطی بیش تری میان متغیر وابسته و متغیرهای مستقل وجود دارد، و در صورت عدم وجود رابطه بین دو متغیر، برابر صفر است. همبستگی بین دو متغیر x, y به صورت رابطه ۱-۱۹ تعریف می شود:

$$R = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum (x_i - \bar{x})^2 \sum (y_i - \bar{y})^2}} \quad \text{رابطه ۱-۱۹}$$

^۱ Applicability Domain

^۲ Y- Randomization

^۳ Correlation coefficient

ضریب تعیین^۱: معیاری برای بیان دقت خط رگرسیون برآورد شده، می‌باشد. نسبتی از واریانس

برحسب متغیر وابسته است که از متغیرهای مستقل قابل پیش‌بینی باشد. این ضریب به صورت رابطه ۲۰-۱

تعریف می‌شود:

$$R^2 = \frac{SSR}{SST} = 1 - \frac{SSE}{SST} \quad \text{رابطه ۲۰-۱}$$

اگر $R^2 = ۱$ باشد آن‌گاه $SSR=SST$ یا به عبارتی $SSE = ۰$ یعنی زمانی که از متغیرهای مستقل

استفاده کنید هیچ خطای وجود ندارد که این بهترین حالت ممکن است. اگر $R^2 = ۰$ باشد آن‌گاه $SSR=۰$

یا به عبارتی $SSE = SSR$ یعنی استفاده از متغیرهای مستقل هیچ تأثیری بر برآورد خط رگرسیونی ندارد.

• مجموع توان‌های دوم رگرسیون^۲ (SSR): بیانگر مجموع توان‌های دوم انحراف مقادیر پیش‌بینی

شده متغیر وابسته از میانگین مقادیر آن است. طبق رابطه ۲۱-۱ است:

$$SSR = \sum_{i=1}^n (\hat{y}_i - \bar{y})^2 \quad \text{رابطه ۲۱-۱}$$

• مجموع توان‌های دوم کامل^۳ (SST): طبق رابطه ۲۲-۱ نشانگر مجموع توان‌های دوم انحراف مقادیر

واقعی متغیر وابسته از میانگین مقادیر آن است.

$$SST = \sum_{i=1}^n (y_i - \bar{y})^2 \quad \text{رابطه ۲۲-۱}$$

• مجموع توان‌های دوم خطا^۴ (SSE): مجموع توان‌های دوم انحراف مقادیر واقعی متغیر وابسته از

مقادیر پیش‌بینی شده برای آن است رابطه ۲۳-۱ :

$$SSE = \sum_{i=1}^n (y_i - \hat{y}_i)^2 \quad \text{رابطه ۲۳-۱}$$

بنابراین با توجه به روابط فوق ضریب تعیین را می‌توان به صورت رابطه ۲۴-۱ تعریف کرد:

$$R^2 = 1 - \frac{\sum_{i=1}^n (y_i - \hat{y}_i)^2}{\sum_{i=1}^n (y_i - \bar{y})^2} \quad \text{رابطه ۲۴-۱}$$

^۱ Determination coefficient

^۲ Sum Square Regression

^۳ Sum Square Total

^۴ Sum Square Error

طبق رابطه ۲۴-۱ اگر به ازای تمام نقاط $y_i = \hat{y}_i$ باشد، یعنی تمام مشاهدات بر روی خط برازش شده قرار گرفته باشند، مقدار R^2 برابر یک می‌شود و هرگونه انحرافی از این حالت باعث می‌شود که مقدار R^2 از یک کوچک‌تر شود.

ضریب تعیین تصحیح شده^۱: از آن جا که R^2 (ضریب تعیین)، تعداد پارامترهای موجود در مدل را به حساب نمی‌آورد، در ارزیابی و مقایسه مدل‌های مختلف با تعداد متفاوت متغیرهای مستقل ضریب تعیین تصحیح شده استفاده می‌گردد، که طبق رابطه ۲۵-۱ بیان می‌شود:

$$R_{adj}^2 = 1 - \frac{n-1}{n-p-1} \cdot \frac{SSE}{SST} = 1 - (1 - R^2) \quad \text{رابطه ۲۵-۱}$$

در این رابطه p نشان‌دهنده تعداد متغیرهای مستقل و n نشان‌دهنده تعداد ترکیبات مورد بررسی می‌باشد.

FIT: به‌عنوان یک پارامتر آماری برای مقایسه مدل‌های مختلف با تعداد متفاوت متغیرهای مستقل به کار می‌رود که هر چه مقدار این پارامتر برای یک مدل بیش‌تر باشد، آن مدل مناسب‌تر می‌باشد. مقدار آن از رابطه ۲۶-۱ به دست می‌آید [۱۸، ۱۹].

$$FIT = \frac{R^2(n-p-1)}{(n+p^2)(1-R^2)} \quad \text{رابطه ۲۶-۱}$$

که R^2 مجذور ضریب همبستگی مدل، p تعداد توصیف کننده‌های مدل و n تعداد ترکیبات مربوط به مدل است.

مجموع توان دوم باقی‌مانده‌ها^۲ (PRESS): مجموع توان دوم تفاوت بین مقدار کمیت مشاهده

شده (y_i) و مقدار برآورد شده ($\hat{y}_i$) می‌باشد رابطه ۲۷-۱ :

$$PRESS = \sum_{i=1}^n (y_i - \hat{y}_i)^2 \quad \text{رابطه ۲۷-۱}$$

^۱ Adjusted determination coefficient

^۲ Predictive Residual Sum of Squares

خطای استاندارد پیش‌بینی^۱ (SEP): طبق رابطه ۲۸-۱ بیان می‌شود:

$$SEP = \sqrt{\frac{\sum_{i=1}^n (y_i - \hat{y})^2}{n}} \quad \text{رابطه ۲۸-۱}$$

خطای مطلق میانگین^۲ (MAE): طبق رابطه ۲۹-۱ بیان می‌شود:

$$MAE = \frac{\sum_{i=1}^n |y_i - \hat{y}|}{n} \quad \text{رابطه ۲۹-۱}$$

خطای نسبی پیش‌بینی^۳ (REP): طبق رابطه ۳۰-۱ بیان می‌شود:

$$REP(\%) = \frac{100}{\bar{y}} \times \sqrt{\frac{\sum_{i=1}^n (y_i - \hat{y})^2}{n}} \quad \text{رابطه ۳۰-۱}$$

میانگین توان‌های دوم خطا^۴ (MSE): طبق رابطه ۳۱-۱ بیان می‌شود:

$$MSE = \frac{\sum_{i=1}^n (y_i - \hat{y})^2}{n} \quad \text{رابطه ۳۱-۱}$$

میانگین خطای نسبی^۵ (MRE): طبق رابطه ۳۲-۱ بیان می‌شود:

$$MRE = \frac{\sum_{i=1}^n \left| \frac{y_i - \hat{y}}{y_i} \right|}{n} \times 100 \quad \text{رابطه ۳۲-۱}$$

تولرانس^۶: همان‌گونه که می‌دانیم وجود یا عدم وجود هم‌خطی در بین متغیرهای پیش‌بینی کننده

بسیار مهم است که باید بررسی شود. پدیده هم‌خطی می‌تواند به یکی از دو دلیل زیر باشد:

❖ حداقل دو تا از متغیرهای مستقل از لحاظ خطی خیلی به هم وابسته باشند. یعنی قدر

مطلق ضریب همبستگی بین آن دو خیلی به یک نزدیک باشد.

❖ حداقل یکی از متغیرهای مستقل وابستگی شدیدی به مجموعه سایر متغیرهای مستقل

داشته باشد که معمولاً این مورد مهم‌تر است.

لذا برای کشف پدیده هم‌خطی برای هر متغیر مستقل پارامتری محاسبه می‌شود که به آن تولرانس

گفته می‌شود. برای هر متغیر X_i تولرانس از رابطه ۳۳-۱ به دست می‌آید [۲۰، ۲۱]:

¹ Standard Error of Prediction

² Mean Absolute Error

³ Relative Error of Prediction

⁴ Mean Square Error

⁵ Mean Relative Error

⁶ Tolerance

$$T_i = 1 - R_i^2 \quad i = 1, 2, 3, \dots, p \quad \text{رابطه ۱-۳۳}$$

که در آن R_i^2 مجذور ضریب همبستگی چندگانه و p تعداد متغیرهای پیشگو است که از رگرسیون x_i بر تمام متغیرهای پیش‌بینی کننده دیگر به دست می‌آید. در نظر گرفته شده است که اگر برای هر یک از متغیرهای مستقل، T_i از ۰/۱ بیش‌تر باشد، پدیده هم‌خطی وجود نخواهد داشت.

عامل تورم واریانس^۱: در برخی مواقع به جای تلرانس از عامل افزایش واریانس استفاده می‌شود

که رابطه آن با تلرانس به صورت رابطه ۱-۳۴ است:

$$VIF = \frac{1}{T_i} = \frac{1}{1 - R_i^2} \quad i = 1, 2, 3, \dots, p \quad \text{رابطه ۱-۳۴}$$

بدیهی است که اگر x_i رابطه خطی با سایر متغیرها داشته باشد، آن‌گاه R_i^2 نزدیک به یک است و عامل افزایش واریانس بزرگ می‌شود. مقادیر VIF_i بین محدوده ۱ تا ۱۰ قابل قبول است و مقادیر بیش‌تر از ۱۰ اکثراً به عنوان علامتی از این‌که داده‌ها مشکل هم‌خطی دارند، تلقی می‌شود. اگر هیچ‌گونه رابطه خطی بین متغیرهای پیش‌بینی کننده وجود نداشته باشد، آنگاه R_i^2 صفر و VIF_i برابر یک می‌شود [۲۲-۲۵].

R_0^2 : پارامتر ضریب تعیین مدل با عرض از مبدأ صفر است، زمانی که نمودار فعالیت پیش‌بینی شده بر حسب فعالیت تجربی رسم شده است. در صورتی که مقدار این پارامتر نزدیک به مقدار R^2 باشد، قابل قبول است [۲۶].

$R_0'^2$: ضریب تعیین مدل با عرض از مبدأ صفر است، زمانی که نمودار فعالیت تجربی بر حسب فعالیت پیش‌بینی رسم شده است. در صورتی که مقدار این پارامتر نزدیک به مقدار R^2 باشد، قابل قبول است [۲۶].

^۱ Variance Inflation Factor (VIF)

R_m^2 : این پارامتر اختلاف بین R^2 و R_0^2 را در مدل محدود می‌کند، به‌طوری که مدل با R_m^2 بیش‌تر از ۰/۵ قابل قبول بوده و مدل توسعه یافته از قدرت پیش‌بینی بالاتری برخوردار است [۲۷, ۲۸] و از رابطه ۳۵-۱ محاسبه می‌شود:

$$R_m^2 = R^2 \times [1 - ((R^2 - R_0^2))^{1/2}] \quad \text{رابطه ۳۵-۱}$$

$R_m'^2$: این پارامتر اختلاف بین R^2 و $R_0'^2$ را در مدل محدود می‌کند، به‌طوری که مدل با $R_m'^2$ بیش‌تر از ۰/۵ قابل قبول بوده و مدل از قدرت پیش‌بینی بالاتری برخوردار است [۲۷, ۲۸] و از رابطه ۳۶-۱ محاسبه می‌شود:

$$R_m'^2 = R^2 \times [1 - ((R^2 - R_0'^2))^{1/2}] \quad \text{رابطه ۳۶-۱}$$

$R - R$: پارامتر حاصل از تفاضل دو پارامتر R_m^2 و $R_m'^2$ است. محدوده قابل قبول این پارامتر کم‌تر از ۰/۲ تعریف شده است و طبق رابطه ۳۷-۱ محاسبه می‌شود [۲۹]:

$$R - R = |R_m^2 - R_m'^2| \quad \text{رابطه ۳۷-۱}$$

$R_{0,R}^2$: این پارامتر نشان‌دهنده نزدیکی دو پارامتر R^2 و R_0^2 می‌باشد. هر چه این مقادیر به هم نزدیک‌تر باشند و نسبت تفاضل آن‌ها به R^2 کوچک‌تر از ۰/۱ باشد، مدل دارای قدرت پیش‌بینی بالاتری است [۲۷] از رابطه ۳۸-۱ محاسبه می‌شود:

$$R_{0,R}^2 = \frac{(R^2 - R_0^2)}{R^2} \quad \text{رابطه ۳۸-۱}$$

$R_{0,R}'^2$: این پارامتر نشان‌دهنده نزدیکی دو پارامتر R^2 ، $R_0'^2$ می‌باشد. هر چه این مقادیر به هم نزدیک‌تر باشند و نسبت تفاضل آن‌ها به R^2 کوچک‌تر از ۰/۱ باشد، مدل دارای قدرت پیش‌بینی بالاتری است [۲۷] از رابطه ۳۹-۱ محاسبه می‌شود:

$$R_{0,R}'^2 = \frac{(R^2 - R_0'^2)}{R^2} \quad \text{رابطه ۳۹-۱}$$

K: شیب معادله رگرسیون برای نمودار فعالیت پیش‌بینی شده برحسب فعالیت تجربی با عرض از مبدأ صفر می‌باشد. پارامتر K باید در محدوده ۰/۸۵ تا ۱/۱۵ باشد تا مدل توسعه یافته از نظر آماری مورد تأیید واقع گردد [۲۷].

K': شیب معادله رگرسیون برای نمودار مقادیر تجربی برحسب مقادیر پیش‌بینی شده با عرض از مبدأ صفر می‌باشد و K' باید در محدوده ۰/۸۵ تا ۱/۱۵ باشد تا مدل قدرت پیش‌بینی بالایی داشته باشد [۲۷].

۱-۴-۲ ارزیابی مدل با استفاده از نتایج پیش‌بینی شده برای مجموعه آزمون

برای ارزیابی مدل توسعه یافته، از پیش‌بینی داده‌های مجموعه آزمون که در هیچ یک از مراحل مدل‌سازی حضور نداشته‌اند، استفاده می‌شود. به‌طوری‌که ترکیبات مجموعه آزمون با استفاده از مدل توسعه یافته، پیش‌بینی می‌شود. مقادیر خطای حاصل از پیش‌بینی و نتایج پارامترهای آماری قابل قبول، صحت و دقت مدل ساخته شده را اثبات می‌کند.

۱-۴-۳ ارزیابی مدل با استفاده از تکنیک رد مرحله ای تک تک (LOO)

تکنیک LOO نیز برای بررسی قدرت پیش‌بینی مدل برای کل مجموعه داده‌ها به کار گرفته می‌شود. به‌طوری‌که از مدل توسعه یافته برای پیش‌بینی همه داده‌ها، مورد استفاده قرار می‌گیرد. با این تفاوت که در این تکنیک هر داده مورد مطالعه یک‌بار به‌عنوان داده آزمون خارج شده و مدل با مابقی داده‌ها آموزش داده می‌شود و پیش‌بینی پاسخ مربوط به داده خارج شده با استفاده از مدل به‌دست می‌آید. این عملیات به تعداد کل داده‌ها تکرار شده تا پاسخ همه داده‌ها یک‌بار پیش‌بینی شود. با محاسبه پارامترهای آماری برای نتایج حاصل از تکنیک رد مرحله‌ای تک تک قدرت پیش‌بینی و تعمیم‌پذیری مدل توسعه یافته تأیید می‌شود.

۱-۴-۴ استفاده از نمودار خطای باقی مانده ها

عبارت خطای باقی مانده، اختلاف بین مقادیر پیش بینی شده و مقادیر تجربی می باشد، که اگر پراکندگی مقادیر در دو طرف نمودار صفر باشد، این مسئله نشان دهنده تصادفی بودن خطا است. اما اگر عمده نقاط در این نمودار، در یک طرف خط صفر باشند، به این معنی است که خطای جهت داری اتفاق افتاده است.

۱-۴-۵ استفاده از دامنه کاربرد

دامنه کاربرد فضای ویژگی های مولکولی را که مدل بر روی آن ها آموزش دیده است و باید در مورد آنها استفاده شود، تعیین می کند [۳۰]. یک منطقه تئوری در فضای شیمیایی است که به وسیله توصیف کننده های نهایی و پاسخ مدل تعریف می شود. توسط این روش بررسی این که آیا یک ماده شیمیایی درون دامنه ساختاری مدل قرار می گیرد یا خارج آن امکان می یابد. اگر ماده شیمیایی خارج دامنه ساختاری مدل قرار گیرد عدم اطمینان افزایش یافته و مقدار پیش بینی شده غیر واقعی است. مقیاس دوری یک ترکیب از دامنه کاربرد مدل با درجه نفوذ^۱ (H) آن سنجیده می شود که به صورت رابطه ۱-۴۰ محاسبه می گردد:

$$H_i = x_i (X^T X)^{-1} x_i \quad (i = 1, \dots, n) \quad \text{رابطه ۱-۴۰}$$

در این رابطه H_i درجه نفوذ مربوط به ترکیب i در فضای توصیف کننده ها، x_i بردار ردیفی توصیف کننده ترکیب i و X ماتریس توصیف کننده ها می باشد. بالاوند T نیز به ترانهاده^۲ بردار و ماتریس اشاره می کند. مقدار بحرانی $H > h^*$ است که $h^* = 3p/n$ و p تعداد توصیف کننده های مدل به علاوه یک و n تعداد ترکیبات در مجموعه آموزش است. وقتی H یک ترکیب بیش تر از مقدار بحرانی است یعنی آن ترکیب از نظر ساختاری دور از داده های مجموعه آموزشی است و خارج دامنه کاربرد مدل در نظر گرفته می شود.

^۱ Leverage

^۲ transpose

نمودار ویلیام^۱ برای تشخیص ساده گرافیکی مقادیر پیش‌بینی شده پرت (باقی‌مانده‌های استاندارد شده آن‌ها بیش‌تر از ± 3 می‌باشد) و همین‌طور ترکیباتی که از لحاظ ساختاری بر مدل تأثیر گذارند ($H > h^*$) به کار می‌رود. این نمودار از رسم مقادیر باقی‌مانده‌های استاندارد شده در برابر مقادیر H تمامی ترکیبات به دست می‌آید. مقادیر باقی‌مانده‌های استاندارد شده نیز با استفاده از رابطه ۴۱-۱ محاسبه می‌شود:

$$r_i = \frac{e_i}{s_{e_i}} = \frac{(y_i - \hat{y}_i)}{s_{e_i}} \quad \text{رابطه ۴۱-۱}$$

که در آن y_i و $\hat{y}_i$ به ترتیب مقادیر واقعی و پیش‌بینی پاسخ هر مشاهده است ($i = 1, \dots, n$)، e_i تفاوت بین پاسخ‌های واقعی و پیش‌بینی شده برای هر مشاهده است و s_{e_i} انحراف استاندارد مقادیر باقی‌مانده است.

۱-۴-۶ استفاده از آزمون Y-تصادفی

برای اثبات قدرت مدل ایجاد شده و این که نتایج به دست آمده تصادفی نبوده است از آزمون Y-تصادفی استفاده می‌شود. بدین منظور pIC_{50} (فعالیت دارویی به عنوان متغیر پاسخ) به صورت تصادفی در محدوده داده‌های متغیر پاسخ ایجاد شد.

۱-۴-۷ بررسی نحوه تأثیر متغیرهای مستقل بر متغیر وابسته

از آن جایی که واحدهای اندازه‌گیری متغیرهای مستقل (ضرایب رگرسیون) یکسان نمی‌باشد، میزان تأثیر و اهمیت یک متغیر مستقل بر روی متغیر وابسته، با استفاده از ضرایب رگرسیون غیر استاندارد مشخص نمی‌شود. برای حل این مسئله، از ضریب استاندارد شده که طبق رابطه ۴۲-۱ محاسبه می‌شود، استفاده گردید [۳۱].

$$\beta'_k = \left(\frac{s_K}{s_Y} \right) \beta_K \quad \text{رابطه ۴۲-۱}$$

^۱ William's Plot

در این رابطه β'_k ضریب رگرسیون استاندارد شده متغیر مستقل مورد نظر (توصیف کننده k ام)، S_K انحراف استاندارد توصیف کننده k ام، S_Y انحراف استاندارد متغیر وابسته (فعالیت دارویی) و β_K ضریب رگرسیون غیراستاندارد همان توصیف کننده است، که به صورت رابطه ۴۳-۱ و رابطه ۴۴-۱ زیر بیان می شود:

$$S_K = \sqrt{\frac{\sum_{i=1}^n (x_{ik} - \bar{x}_k)^2}{n - 1}} \quad \text{رابطه ۴۳-۱}$$

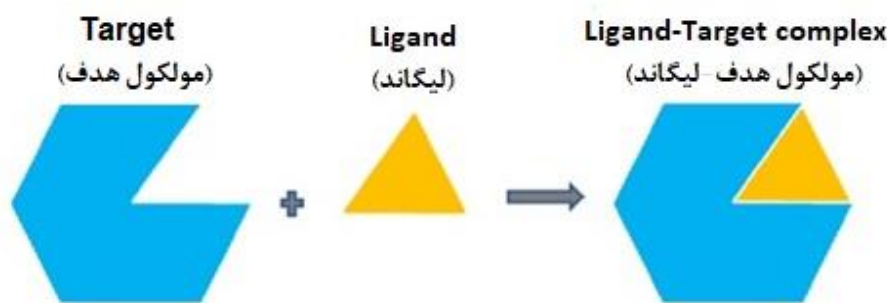
$$S_Y = \sqrt{\frac{\sum_{i=1}^n (Y_i - \bar{Y})^2}{n - 1}} \quad \text{رابطه ۴۴-۱}$$

X_{ijk} مقدار متغیر مستقل (توصیف کننده) k ام برای ترکیب i ام و $\bar{X}_k$ مقدار میانگین توصیف کننده، n تعداد ترکیبات، Y_i مقدار متغیر وابسته برای ترکیب i ام و $\bar{Y}$ مقدار میانگین متغیر وابسته می باشد.

۵-۱ شبیه سازی داکینگ مولکولی

داکینگ مولکولی یک روش شبیه سازی رایانه ای است که به طور گسترده ای برای پیش بینی ترکیب یک مجموعه گیرنده-لیگاند استفاده می شود، جایی که گیرنده معمولاً یک پروتئین است و لیگاند یک مولکول کوچک است [۳۲]. داکینگ مولکولی یکی از ابزارهای کلیدی در زیست شناسی مولکولی ساختاری و طراحی دارو است. هدف از داکینگ پروتئین-لیگاند، پیش بینی جایگاه اتصال مناسب لیگاند به پروتئینی با ساختار سه بعدی شناخته شده است [۳۳]. در شکل ۳-۱ ایده روش داکینگ مولکولی به صورت اجمالی نشان داده شده است. در این فرآیند لیگاند در حفرات مولکول هدف (پروتئین هدف) قرار می گیرد. بنابراین، محاسبه اتصال به عنوان یک رویکرد مهم برای درک برهم کنش پروتئین-لیگاند در نظر گرفته می شود [۳۴].

[۳۵]



شکل ۱-۳ اتصال لیگاند به مولکول هدف

داکینگ مولکولی به دو روش انجام می‌شود:

۱. داکینگ لیگاند درون پروتئین با پروتئین برای اعتبار سنجی فرآیند داکینگ مولکولی
۲. داکینگ لیگاندهای همانند با لیگاند درون پروتئین در شرایط مشخص شده توسط فرآیند اعتبارسنجی برای ایجاد مدل‌سازی

۱-۵-۱ فرآیند اعتبار سنجی

با توجه به این که ابزار و روش‌های متفاوتی برای داکینگ مولکولی وجود دارد و همچنین در نظر گرفتن این نکته که برخی از این روش‌ها نتایج بهتری برای بعضی دسته‌بندی مولکول‌های هدف ارائه می‌دهند باید در انتخاب روش مناسب داکینگ مولکولی دقت کرد. به بیان دیگر بایستی روش‌های متفاوت را به کار گرفت و در نهایت روشی که بهترین نتیجه را برای مولکول هدف مورد نظر ارائه می‌دهد انتخاب نمود. یکی از معیارهای کمک کننده به انتخاب روش مناسب، اندازه گیری انحراف ریشه میانگین مربع^۱ (RMSD) اتم‌های لیگاند موجود در ساختار کریستالوگرافی و داک شده است [۳۲].

موفقیت الگوریتم‌های داکینگ در پیش‌بینی حالت اتصال لیگاند به‌طور معمول بر اساس انحراف ریشه میانگین مربع بین موقعیت‌های اتم سنگین لیگاندها و حالت پیش‌بینی شده توسط الگوریتم، اندازه‌گیری

^۱ Root Mean Square Deviation

می‌شود. انعطاف پذیری سیستم یک چالش عمده در جستجوی حالت مناسب است [۳۶]. معمولاً وقتی RMSD کمتر از ۲ آنگستروم باشد عملکرد خوب در نظر گرفته می‌شود.

۱-۵-۲ آماده‌سازی پروتئین

ساختار کریستالوگرافی پروتئین‌ها از بانک اطلاعات پروتئین^۱ استخراج می‌شود. به ساختار همراه با لیگاند کمپلکس شده با پروتئین، ساختار "هولو"^۲ و ساختار بدون لیگاند، کوفاکتور و فقط دارای بخش پروتئینی در جایگاه فعال، ساختار "آپو"^۳ نامیده می‌شود [۳۷]. ساختارهای هولو بهترین روش شناسایی جایگاه فعال پروتئین می‌باشند. مرکز ثقل^۴ لیگاندی که در پروتئین وجود دارد همان مرکز تقریبی جایگاه فعال پروتئین است. درستی نتایج داکینگ مستقیماً به کیفیت ساختار کریستالوگرافی جایگاه فعال پروتئین بستگی دارد. از آنجایی که حتی بهترین ساختارهای کریستالوگرافی اغلب قدرت تفکیک ۱ آنگستروم یا بیش‌تر دارند، مولکول‌های آب در ساختار را حذف کرده و سپس بعد از ورود به برنامه، هیدروژن‌هایی را که به‌وسیله کریستالوگرافی پروتئین دیده نمی‌شوند، اضافه نموده و با فرمت pdb ذخیره می‌شود [۳۸].

۱-۵-۳ ساختن لیگاند

بعد از این‌که، لیگاند در برنامه داکینگ فراخوانی می‌شود، در جایگاه فعال پروتئین هدف قرار می‌گیرد و بعد برای همه لیگاندها، شبیه‌سازی داکینگ به‌طور جداگانه توسط کاربر انجام می‌شود.

۱-۵-۴ تنظیم کردن جعبه شبکه‌ای^۵

در داکینگ مولکولی نمی‌توان کل فضا را به‌عنوان جایگاه فعال گیرنده در نظر گرفت. باید بخشی از فضا که لیگاند درون جایگاه فعال قرار می‌گیرد را برای فرآیند داکینگ تعریف کرد تا هم دقت بالا باشد و

¹ Protein Data Bank

² Soaked structure

³ Apo structure

⁴ Gravity center

⁵ Grid box

هم‌زمان انجام فرآیند داکینگ به‌صرفه باشد. به‌جهت سرعت بخشیدن به محاسبات، یک فاصله محدود کننده^۱ تنظیم می‌شود، این فاصله محدود کننده معمولاً یک جعبه مستطیل شکل به‌نام جعبه شبکه‌ای است [۳۸]. بخش‌هایی از پروتئین که دور از جایگاه فعال هستند، به‌طور معمول هیچ اثر قابل اندازه‌گیری بر نتایج امتیازدهی ندارند و از این‌رو ایجاد شبکه با ابعاد یکسان برای پوشش‌دهی اسیدهای آمینه جایگاه فعال پروتئین توصیه می‌شود. اگر فایل مربوط به ساختار دارای لیگاند کمپلکس شده باشد (هولو)، جایگاه فعال پروتئین مشخص خواهد شد و بنابراین مختصات جعبه شبکه‌ای همان مختصات جایگاه فعال است که به شکل مکعبی با ابعاد یکسان تعریف می‌شود. ابعاد جعبه باید به‌گونه‌ای انتخاب شود که اسیدهای آمینه جایگاه فعال را در بر بگیرد [۳۹، ۴۰].

۱-۵-۵ گزینه داکینگ

به‌منظور محاسبه داکینگ، ورودی‌ها نظیر روش‌های نمره‌دهی، جایگاه فعال انعطاف پذیر، حلال پوشی، برخورد با جایگاه فعال احاطه شده، روش‌های جستجو و ... تنظیم می‌شوند.

۱-۵-۶ انجام محاسبه داکینگ

زمانی که ورودی‌ها تنظیم شدند محاسبات داکینگ انجام می‌شود. اغلب برنامه‌ها امکان اجراهای داکینگ به‌صورت تکی و تعداد اندکی قادر به اجرای محاسبات داکینگ در یک حالت مجموعه‌ای می‌باشند. این محاسبات در مواقعی توسط همان کامپیوتری که صفحه رابط گرافیکی از آن استفاده شده، انجام می‌گیرد و در مواقعی هم می‌تواند به یک سرور دیگر فرستاده شود.

¹ Cut off

۱-۵-۷ آنالیز و تحلیل نتایج

از مهم‌ترین نتایج محاسبات داکینگ، انرژی اتصال لیگاند به جایگاه فعال می‌باشد. این انرژی در بین ترکیبات مختلف برای تعیین بهترین مهار کننده، استفاده می‌شود. چند حالت و وضعیت از لیگاند در جایگاه فعال که با بهترین انرژی اتصال همراه است، به‌طور دقیق بررسی شده تا از مناسب بودن آن اطمینان حاصل شود. در بعضی موارد، وضعیتی که توسط یک محاسبه داکینگ تولید شده، به کاربر طرحی برای چگونگی تغییر ترکیبات در مراحل بعدی محاسبات می‌دهد.

در برنامه داکینگ، دو مؤلفه کلیدی الگوریتم رتبه‌بندی یا امتیازدهی^۱ و الگوریتم جستجو^۲ وجود دارد. الگوریتم جستجو، مولکول را در موقعیت‌ها و صورت‌بندی‌های متفاوت در جایگاه فعال پروتئین قرار می‌دهد. انتخاب الگوریتم جستجو، مشخص می‌کند که برنامه با چه صحتی در طول فرآیند، موقعیت‌های ممکن مولکول را ارزیابی می‌کند و چه مدت زمانی برای آن نیاز است. الگوریتم رتبه‌بندی مسئول تعیین این است که آیا جهت‌گیری‌های انتخاب شده توسط الگوریتم جستجو، از لحاظ انرژی مناسب‌ترین هستند و مسئول محاسبه انرژی اتصال است. قابل ذکر است که الگوریتم جستجو صحت نتایج به‌دست آمده از برنامه داکینگ را مشخص نمی‌کند. بنابراین صحت نتایج در درجه اول به صحت الگوریتم رتبه‌بندی بستگی دارد. [۳۸].

۱-۵-۸ کاربردهای داکینگ مولکولی

داکینگ مولکولی کاربردهای زیادی در مطالعات مربوط به طراحی داروها دارد که شامل موارد زیر است.

۱. غربالگری مجازی^۳

^۱ The scoring algorithm

^۲ The search algorithm

^۳ Virtual screening

۲. بهبود طراحی و کشف دارو

۳. پیش‌بینی انرژی آزاد اتصال

۴. بررسی برهم‌کنش لیگاند-گیرنده

متداول‌ترین کاربرد داکینگ مولکولی در اتصال پروتئین-لیگاند است. هدف نهایی اتصال پروتئین-لیگاند پیش‌بینی فعالیت دارویی لیگاند و برهم‌کنش لیگاند با پروتئین می‌باشد و در این تحقیق از این هدف (اتصال لیگاند-پروتئین) جهت بررسی برهم‌کنش‌های ترکیبات پیشنهادی بالقوه در جایگاه فعال گیرنده استفاده می‌شود.

۱- ۶ سرطان

سرطان یک بیماری ژنتیکی است که در آن کنترل رشد سلول‌ها از بین رفته است. بر اساس آمار سازمان جهانی بهداشت^۱ (WHO)، سرطان دومین علت مرگ‌ومیر در جهان پس از بیماری‌های قلبی و عروقی به‌شمار می‌آید و تقریباً ۱۰ میلیون مرگ در سال ۲۰۲۰ در سراسر جهان به خود اختصاص داده است [۴۱]. در حالت نرمال بین میزان تقسیم سلول، مرگ سلولی و تمایز سلولی یک تعادل وجود دارد. سلول‌ها به‌طور دقیقی به‌وسیله فاکتورهای رشد، پیام‌رسان‌ها، پیام‌های محیطی و برخی پروتئین‌ها کنترل می‌شوند. یک مجموعه جهش‌های متوالی در ژن‌ها اتفاق می‌افتد که منجر به از دست رفتن مسیر رشد طبیعی سلول‌ها و بروز سلول‌های سرطانی می‌شود. سلول‌های سرطانی کنترل رشد خود بر چرخه سلولی را از دست داده و به‌طور مداوم و بدون توجه به پیام‌های محیطی و فاکتورهای رشد، به تکثیر ادامه می‌دهند. یکی از فاکتورهای رشد، فاکتور رشد هپاتوسیت^۲ می‌باشد که به‌همراه گیرنده خود (c-Met)^۳ نقش‌های مهمی در اعمال بیولوژیکی بدن و سرطان انجام می‌دهد. گیرنده فاکتور رشد هپاتوسیت یکی از اعضای

¹ World health organization

² Hepatocyte Growth Factor

³ Mesenchymal-epithelial transition factor

خانواده گیرنده تیروزین کیناز (RTK) محسوب می‌شود. تیروزین کینازها گیرنده‌های سطح سلول هستند و اتصال لیگاندهای فاکتور رشد باعث فعال شدن بیش‌تر مسیرهای انتقال سیگنال می‌شود. با این حال، تغییرات ساختاری بیش‌ازحد فاکتورهای رشد و تیروزین کینازهای آن‌ها از طریق چندین مکانیسم با پیشرفت تومور در ارتباط است. مهار این تیروزین کینازهای فاکتور رشد به‌عنوان یک استراتژی درمان سیستماتیک برای سرطان‌های مختلف ایجاد شده است [۴۲، ۴۳]. از دسته ترکیبات بازدارنده c-Met، ترکیبات سولفون دار است. سولفون‌ها (R-SO₂-R)، مولکول‌های دارای واحد سولفون، که واسطه‌های چند منظوره در شیمی آلی هستند. کاربردهای مختلفی در زمینه‌هایی مانند شیمی دارویی و پلیمرها دارند. مشتقات سولفون دسته مهمی از ترکیبات فعال زیستی هستند و شامل طیف وسیعی از فعالیت‌ها می‌باشند. این ترکیبات دارای فعالیت‌های متنوعی به‌عنوان عوامل ضد باکتری، حشره‌کش، ضد قارچ، علف‌کش، ضد هپاتیت، ضد تومور، ضد التهاب، ضد سرطان، ضد HIV-1 و ضد سل می‌باشند [۴۴]. در واقع، اولین داروی مبتنی بر سولفون، سولفونال^۱، به سال ۱۸۸۸ باز می‌گردد [۴۵].

۷-۱ مروری بر کارهای انجام شده و ضرورت تحقیق

با توجه به این‌که ترکیبات مورد مطالعه در این تحقیق از دسته ترکیبات c-Met سولفون دار هستند، مروری بر مطالعات منتشر شده در راستای ارائه مدل‌های QSAR برای پیش‌بینی فعالیت دارویی این دسته از ترکیبات انجام شد و اهمیت ارائه مدل QSAR برای این نوع بازدارنده‌ها مورد بررسی قرار گرفت [۴۶-۴۹]. با توجه به جستجوی انجام شده مدل‌های QSAR توسعه یافته در این مقالات از قدرت پیش‌بینی مناسبی برخوردار هستند.

در این تحقیق، از روش مدل‌سازی انقباضی انحراف قدر مطلق به‌طور هموار کوتاه شده (SCAD) به‌عنوان روشی برای انتخاب متغیر و مدل‌سازی استفاده شد. روش انقباضی SCAD در مطالعات QSAR به

¹ sulfonal

دو شکل طبقه‌بندی^۱ و رگرسیونی استفاده شده است. موارد مربوط به کاربرد روش‌های انقباضی در مطالعات طبقه‌بندی مد نظر اهداف این پروژه نبوده است و جزئیات آن در بخش مرور بر کارها ذکر نشده است [۱۲]. ۵۴-۵۰]. در ادامه به مرور مطالعات منتشر شده از به‌کارگیری روش‌های انقباضی SCAD در مدل رگرسیونی در طی سال‌های ۲۰۱۰-۲۰۲۱ پرداخته شد.

الجمال و همکاران^۲ در سال ۲۰۱۶، مدل QSAR را برای ۱۱۱ بازدارنده‌های استیل کولین استراز ارائه دادند و از روش غربالگری مستقل مطمئن^۳ (SIS) به‌عنوان پیش‌پردازش برای کاهش بعد داده‌ها استفاده شد. در نهایت مدل‌های QSAR با به‌کارگیری روش‌های انقباضی LASSO، ALASSO^۴، SCAD و BRIDGE توسعه یافت. نتایج ضرایب تعیین مجموعه آزمون به‌ترتیب ۰/۸۰، ۰/۸۱، ۰/۸۷ و ۰/۹۱ به‌دست آمد. نتایج گواهی بر قدرت پیش‌بینی مناسب مدل‌های توسعه یافته است [۱۱].

یونگ زیبا^۵ و همکاران در سال ۲۰۱۸، مدل‌های QSAR را با استفاده از ۴ مجموعه از ترکیبات مختلف و به‌کارگیری روش‌های انقباضی SCAD، LASSO، Elastic Net و BRIDGE ارائه دادند. پارامترهای آماری متفاوتی برای ارزیابی مدل به‌کار گرفته شد که نتایج مدل‌های توسعه یافته از اعتبار قابل قبولی برخوردار بودند [۵۵].

ماجومدار و همکارانش^۶ در سال ۲۰۱۸ مدل‌های QSAR را برای مجموعه داده همگن متشکل از ۹۵ آمین، ارائه نمودند. با استفاده از توصیف‌کننده‌های محاسبه شده، فعالیت جهش‌زایی این ترکیبات را با استفاده از روش‌های مدل‌سازی متفاوتی همچون PCR^۷، PLS^۸، RF^۹، LASSO و SCAD پیش‌بینی

^۱ Classification

^۲ Algamal and et al.

^۳ Sure Independence Screening

^۴ Adaptive Least Absolute Shrinkage and Selection Operator

^۵ Yong Xia and et al.

^۶ Majumdar and et al.

^۷ Principal Component Regression

^۸ Partial Least Squares

^۹ Random Forest

کردند. میانگین مربعات خطای پیش‌بینی برای مدل‌های QSAR مربوطه به‌دست آمد و به‌ترتیب برابر با ۲۹/۱۱، ۱۸/۸۶، ۱۷/۲۵، ۲۶/۸۵ و ۲۵/۸۱ است. نتایج نشان می‌دهد که مدل‌های ساخته شده از اعتبار مناسبی برخوردار هستند [۵۶].

مظفری و همکاران^۱ در سال ۲۰۲۱ برای ۵۷ ترکیب از مشتقات آریلازولیل تیواستامید/ استانیلید به‌عنوان مهارکننده HIV یک مدل QSAR با روش شبکه عصبی مصنوعی^۲ (ANN) گزارش دادند که از روش انقباضی SCAD به‌عنوان انتخاب متغیر استفاده شد و نتیجه پارامتر ضریب تعیین برای مجموعه آزمون ۰/۹۲ به‌دست آورده شد [۵۷].

علاوه بر این طبق جستجوی کتابخانه‌ای انجام شده، در سال‌های اخیر گزارشی مبنی بر استفاده از روش Ridge به‌عنوان غربالگری در مطالعات QSAR وجود نداشته است و تنها یک گزارش از به‌کارگیری روش غربالگری مستقل مطمئن (SIS) قبل از مدل‌سازی SCAD توسط الجمال و همکارش در سال ۲۰۱۶ منتشر شده است [۱۱]. بنابراین با توجه به این‌که این روش در مطالعات QSAR استفاده نشده است، در این پروژه برای اولین بار از روش رگرسیون Ridge به‌عنوان روش غربالگری جفت شده با روش انقباضی SCAD برای پیش‌بینی فعالیت دارویی بازدارنده‌های سرطان استفاده شده است. لازم به ذکر است تکنیک غربالگری رد مرحله‌ای تک تک Ridge مورد استفاده در این مطالعه به‌طور ابتکاری و برای اولین بار در مطالعات QSAR پیشنهاد شده است.

۸-۱ نوآوری تحقیق

جستجوی کتابخانه‌ای نشان می‌دهد که مدل‌های QSAR توسعه یافته، از روش SCAD برای ساخت مدل خطی QSAR استفاده شده است. در داده‌هایی با ابعاد بالا به‌کارگیری روش‌های غربالگری قبل

^۱ Mozafari and et al.

^۲ Artificial Neural Network

از انتخاب متغیر از اهمیت ویژه‌ای برخوردار است. ابعاد داده‌ها باید کاهش یابد تا علاوه بر این که مدل قدرت پیش‌بینی بالایی داشته باشد، قدرت تفسیرپذیری روشی هم داشته باشد. در مواجهه با داده‌هایی با ابعاد بالا، به کارگیری روش غربالگری Ridge قبل از مدل‌سازی SCAD، مدل قدرتمندی با تفسیرپذیری مناسب-تری به دست می‌آید. استفاده از آنالیز دو مرحله‌ای Ridge-SCAD می‌تواند از نظر محاسباتی کارآمدتر باشد و برآورد مدل بهتری را نسبت به روش‌های جریمه‌ای یک مرحله‌ای ارائه دهد. بنابراین ابتدا از روش رگرسیون Ridge برای کاهش بعد داده‌ها و سپس از روش انقباضی SCAD برای انتخاب متغیر و مدل‌سازی به‌منظور پیش‌بینی فعالیت دارویی بازدارنده‌های c-Met استفاده شد و گزارشی مبنی بر ترکیب کردن روش رگرسیون Ridge با روش انقباضی SCAD وجود ندارد. در نهایت کارایی مدل توسعه یافته برای پیشنهاد برخی از ترکیبات جدید c-Met با فعالیت دارویی ضد سرطان استفاده شد.

فصل دوم

نتایج تجربی

۲- ۱ نرم افزارهای مورد استفاده

برای انجام کلیه مراحل مدل سازی از بسته های نرم افزاری مختلفی استفاده گردید که در ادامه به معرفی مختصر نرم افزارهای مورد استفاده در این تحقیق پرداخته می شود.

۲- ۱- ۱ نرم افزار هایپرکم^۱

برای رسم ساختار مولکول ها و بهینه سازی آن ها با استفاده از روش های کوانتومی و مکانیکی، از نرم افزار هایپرکم استفاده می شود [۵۸]. طول پیوند، زاویه پیوندی و زوایای پیچشی در مولکول با این برنامه بهینه می شود. داده های حاصل از این نرم افزار به عنوان ورودی سایر نرم افزارها از جمله دراگون^۲ استفاده می شود.

۲- ۱- ۲ نرم افزار دراگون

نرم افزار دراگون به گونه ای طراحی شده است که انواع توصیف کننده های مولکولی مشتق شده از ساختار ترکیبات مختلف را در اختیار کاربر قرار دهد و به کاربر این امکان را می دهد تا توصیف کننده های مولکولی را انتخاب کند که برای تحقیقات مناسب تر هستند. اولین نسخه دراگون در سال ۱۹۹۴ توسعه یافت [۵۹]. نرم افزار دراگون نسخه ۵/۵ [۶۰] امکان محاسبه ۳۲۲۴ توصیف کننده مختلف را برای شیمیدان های محاسباتی فراهم می کند. جهت محاسبه توصیف کننده ها لازم است از ساختارهای هندسی بهینه ترکیبات استفاده شود.

^۱ Hyperchem

^۲ Dragon

۲-۱-۳ نرم افزار SPSS^۱

نرم افزار SPSS [۶۱] امکان تجزیه و تحلیل آماری داده‌ها را فراهم می‌کند. برخی از قابلیت‌های این بسته نرم افزاری عبارتند از: محاسبه میانگین ساده برای داده‌ها، نمایش اطلاعات به صورت متنوع در قالب نمودار و انجام رگرسیون تک متغیره و چند متغیره. در این تحقیق از نرم افزار SPSS نسخه ۲۱ برای محاسبه هم‌خطی موجود بین متغیرها استفاده شد.

۲-۱-۴ نرم افزار R و R-Studio

نرم افزار R [۶۲] یک زبان برنامه نویسی برای محاسبات آماری و گرافیکی است. طیف گسترده‌ای از مدل‌های آماری (مدل‌سازی خطی و غیرخطی، آزمون‌های آماری کلاسیک، طبقه‌بندی و غیره) و تکنیک‌های گرافیکی را فراهم می‌کند. یکی از نقاط قوت نرم افزار R آسان بودن طراحی نمودارهای با کیفیت است که شامل نمادها و فرمول‌های ریاضی مورد نیاز می‌باشد. نرم افزار R-Studio یک محیط گرافیکی مناسب‌تر برای اجرای برنامه‌های تحت R است. استفاده از R-Studio برای اجرای بسته‌های نرم افزاری آماری موجود در نرم افزار R ضروری می‌باشد. در این تحقیق از نرم افزارهای R و R-Studio، برای پیش پردازش توصیف کننده با استفاده از بسته نرم افزاری Caret، برای غربالگری توصیف کننده‌ها به کمک روش انقباضی Ridge با استفاده از بسته نرم افزاری glmnet و انتخاب مؤثرترین توصیف کننده‌ها با استفاده از بسته نرم افزاری ncvreg استفاده شد.

۲-۱-۵ نرم افزار متلب^۲

نرم افزار متلب [۶۳] یکی از جامع‌ترین و کارآمدترین نرم افزارهای علمی و محاسباتی است که طی چند سال گذشته تهیه و به بازار عرضه شد و در طی سال‌های اخیر با تدوین نسخه‌های جدیدتر و کامل‌تر

^۱ Statistical Package for the Social Science

^۲ MATLAB

روزبه‌روز مورد توجه محققین قرار گرفته است. MATLAB^۱ به معنای آزمایشگاه ماتریس است. متلب یک محیط محاسباتی عددی و زبان برنامه‌نویسی است و امکان به‌کارگیری ماتریس، ترسیم توابع و داده‌ها و اجرای الگوریتم‌ها را فراهم می‌کند. در این تحقیق برای محاسبه همبستگی بین توصیف‌کننده‌ها با استفاده از دستور corecoeff و برخی از ارزیابی‌های مدل برتر از نرم افزار متلب ۲۰۱۷a استفاده شد.

۲-۱-۶ نرم افزار اتوداک^۲

نرم افزار اتوداک نسخه ۴/۲ یکی از نرم افزارهای مدل‌سازی مولکولی می‌باشد که قابل اجرا بر روی سیستم عامل ویندوز است [۶۴]. این نرم افزار برای پیش‌بینی چگونگی اتصال مولکول‌ها به‌عنوان سوبسترا یا دارو به یک مولکول گیرنده با ساختار سه‌بعدی طراحی شده است.

۲-۱-۷ نرم افزار ویورلایت^۳

نرم افزار ویورلایت نسخه ۵/۰ امکان بررسی جزئی ماکرومولکول و لیگاند را به کاربر می‌دهد و می‌توان با استفاده از این نرم افزار ماکرومولکول را مشاهده و ویرایش کرد. همچنین جایگاه فعال پروتئین را می‌توان مشاهده و تعیین مختصات نمود. همین‌طور می‌توان طول پیوندهای هیدروژنی و هیدروفوبی مربوط به لیگاند و اسیدهای آمینه موجود در جایگاه فعال را نیز بررسی کرد [۶۵]. در این مطالعه جهت آماده‌سازی لیگاند و پروتئین برای انجام داکینگ مولکولی و بررسی برهم‌کنش‌های لیگاند-پروتئین از نرم افزار ویورلایت نسخه ۵/۰ استفاده گردید.

۲-۱-۸ نرم افزار Discovery Studio

این نرم افزار امکان بررسی کمپلکس لیگاند و گیرنده را پس از اجرای داکینگ مولکولی فراهم می‌کند. به‌گونه‌ای که تمامی پیوندهای آب‌گریز و هیدروژنی و همین‌طور طول پیوندهای هیدروژنی و آب‌گریز

^۱ Matrix Laboratory

^۲ Autodock

^۳ ViewerLite

را نشان می‌دهد. بررسی بر هم‌کنش‌های اسیدهای آمینه موجود در جایگاه فعال با اجزای اتمی لیگاند با استفاده از این نرم افزار امکان پذیر است.

۲-۲ مدل سازی فعالیت دارویی مشتقات سولفون دار به عنوان بازدارنده

های c-Met

کمی کردن ویژگی‌های ساختاری و فعالیت دارویی ترکیبات دارویی و پیدا کردن رابطه حاکم بر کمیت‌های اندازه‌گیری شده، همواره موضوع در حال توسعه و مورد توجه بوده و از طریق ساخت مدل‌های QSAR صورت می‌گیرد. در سال‌های اخیر، به دلیل توسعه و پیشرفت تکنولوژی، استفاده از شیمی محاسباتی، شبیه‌سازی مولکولی و طراحی و ساخت مدل‌های QSAR متفاوت به کمک رایانه به‌طور گسترده تری مورد توجه شیمیدان‌های محاسباتی قرار گرفته است.

روش‌های محاسباتی دارای مزایایی همچون کاهش سنتزهای آزمایشگاهی و صرفه‌جویی در زمان، هزینه و نیروی انسانی می‌شوند. بنابراین محققین شیمی محاسباتی همواره در تلاش تا روش‌های آماری و محاسباتی پیشرفته‌تری را برای بررسی ارتباط بین متغیرهای وابسته و متغیر مستقل به کار ببرند که بتوان متغیرهایی که ارتباط بیش‌تری با پاسخ دارند را شناسایی کرد، بنابراین ارائه یک مدل QSAR با قدرت تفسیر پذیری و پیش‌بینی مناسب مورد توجه است. به‌منظور بهبود مدل‌های QSAR، به کارگیری روش‌های انتخاب متغیر و مدل‌سازی جدید و کارآمد نظیر روش‌های انتخاب متغیر جریمه‌ای به دلیل داشتن مزایایی چون تنگی، پایداری و بایاس کم، دارای اهمیت است. در داده‌های با ابعاد بالا بین توصیف کننده‌ها مشکل همبستگی بالا وجود دارد و همبستگی بالا منجر به ایجاد پدیده هم‌خطی می‌شود. لذا انتظار می‌رود با به کارگیری یک روش غربالگری متغیرها به همراه روش و انتخاب متغیر جریمه‌ای بتوان توصیف کننده‌هایی با تعداد محدودتر ولی با اطلاعات مفید را از تعداد بسیار زیادی توصیف کننده‌ها استخراج نمود. به‌منظور ارائه

یک مدل تنک و پایدار از روش رگرسیون Ridge برای غربالگری و کاهش تعداد توصیف کننده‌های محاسبه-ای استفاده شد و روش رگرسیون جریمه‌ای SCAD برای انتخاب مؤثرترین توصیف کننده‌ها بر فعالیت دارویی و ایجاد یک مدل خطی بین آن‌ها و فعالیت دارویی دسته‌ای از مشتقات سولفون دار به‌عنوان بازدارنده‌های سرطان به‌کار گرفته شد. لازم به‌ذکر است که روش رگرسیون Ridge که اساساً یک روش رگرسیون با تمام متغیرها و بدون قابلیت انتخاب متغیر می‌باشد به کمک یک روش ابتکاری به‌عنوان روش غربالگری قبل از توسعه مدل‌های QSAR به‌کار گرفته شد. به‌منظور بررسی عملکرد روش Ridge برای غربالگری توصیف کننده‌ها قبل از مدل‌سازی با SCAD (Ridge-SCAD)، مدل QSAR با استفاده از تمام متغیرها بدون هیچ نوع غربالگری با به‌کارگیری روش SCAD توسعه یافت و نتایج پیش‌بینی مدل، با هم مقایسه شد.

از آنجایی که سنتز و ارزیابی تجربی فعالیت دارویی ترکیبات فعال زمان‌بر و پرهزینه است، توسعه روش‌های تئوری از قبیل ساخت مدل QSAR برای ارزیابی و پیش‌بینی فعالیت‌های دارویی ترکیبات مشابه بسیار مهم است و باعث صرفه‌جویی در زمان و هزینه و به‌کارگیری نیروی انسانی می‌شود. با توجه به اهمیت موضوع، در این بخش مدل QSAR برای پیش‌بینی فعالیت ترکیبات مورد مطالعه و پیشنهاد ترکیبات مشابه جدید توسعه یافت. از این‌رو مدل Ridge-SCAD برای پیش‌بینی فعالیت دارویی (pIC_{50}) برخی از مشتقات سولفون دار به‌عنوان بازدارنده‌های بیماری سرطان ساخته شد. فعالیت دارویی این ترکیبات و پتانسیل بالقوه آن‌ها به‌عنوان کاندیدهای پیشنهادی با فعالیت ضد سرطان با استفاده از داکینگ مولکولی و محاسبه پارامترهای قاعده پنج لیپینسکی^۱ مورد بررسی قرار گرفت و نتایج حاصل امکان استفاده از آن‌ها به‌عنوان کاندیدهای جدید فعالیت ضد سرطان را اثبات نمود. در ادامه این بخش با جزئیات بیش‌تری به مراحل

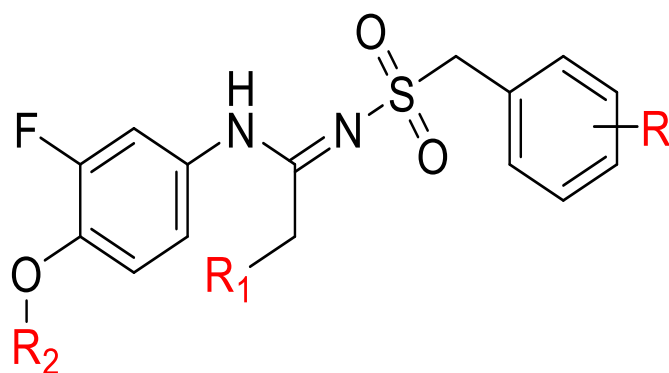
¹ Lipinski's rule of five

ساخت مدل‌های Ridge-SCAD و SCAD (برای بررسی کارایی مدل Ridge-SCAD) این دسته از بازدارنده‌ها پرداخته شده است.

۲-۱ معرفی مجموعه داده‌های مورد استفاده در این مطالعه

مجموعه داده‌ها شامل فعالیت دارویی ۶۳ ترکیب از مشتقات سولفون دار می‌باشد که توسط نان و همکارانش^۱ در سال‌های ۲۰۱۹ و ۲۰۲۰ ارائه شده است [۶۶، ۶۷]. اسکلت اصلی لیگاندها (از این پس ترکیبات مورد مطالعه که همان مشتقات سولفون دار می‌باشند به نام لیگاندها بیان می‌شوند) در شکل ۱-۲ و شکل ۲-۲ نشان داده شده است. جزئیات استخلاف‌ها و مقادیر منهای لگاریتم IC_{50} (pIC_{50}) هر لیگاند در جدول ۱-۲ و جدول ۲-۲ آورده شده است. IC_{50} غلظتی از دارو برحسب مولار است که تا ۵۰٪ تکثیر پروتئین c-Met را مهار می‌کند. لازم به ذکر است که هر چقدر دارو قوی‌تر باشد، غلظت کم‌تری از آن، برای داشتن ۵۰٪ اثر مهارکنندگی، لازم بوده و در نتیجه مقدار pIC_{50} بزرگ‌تری دارد. در کلیه مراحل ساخت مدل QSAR از pIC_{50} به عنوان متغیر وابسته و توصیف کننده‌های مولکولی به عنوان متغیر مستقل در نظر گرفته شدند. تقسیم‌بندی مجموعه داده‌ها با استفاده از الگوریتم KS انجام شد و در کلیه مراحل ساخت مدل QSAR مجموعه آموزش (متشکل از ۵۳ ترکیب) برای ساخت مدل‌های خطی Ridge-SCAD و SCAD و مجموعه آزمون (متشکل از ۱۰ ترکیب) برای ارزیابی نهایی مدل‌ها به کار گرفته شده‌اند.

¹ Nan et al.



شکل ۱-۲ اسکلت اصلی ساختار ترکیبات مورد مطالعه مربوط به ترکیبات ۱ تا ۳۹

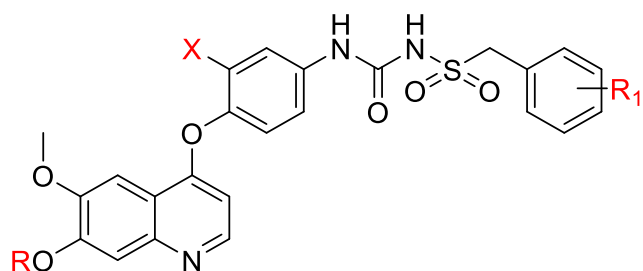
جدول ۱-۲ جزئیات ساختار ۱ تا ۳۹ به همراه مقادیر pIC_{50} ترکیبات مورد مطالعه

ردیف	R ₂	R	R ₁	(pIC_{50}) فعالیت دارویی
۱	thieno[3,2-b]pyridine	H	n-butyl	۵/۰۲
۲ ^t	thieno[3,2-b]pyridine	4-F	n-butyl	۵/۱۵
۳	thieno[3,2-d]pyrimidine	H	n-butyl	۴/۷۳
۴ ^t	thieno[3,2-d]pyrimidine	4-F	n-butyl	۴/۹۵
۵	6,7-dimethoxyquinazoline	H	n-butyl	۴/۵۲
۶	6,7-dimethoxyquinazoline	4-F	n-butyl	۴/۶۶
۷	6,7-dimethoxyquinoline	H	n-butyl	۵/۰۸
۸	6,7-dimethoxyquinoline	4-F	n-butyl	۵/۲۹
۹	6,7-dimethoxyquinoline	H	Phenyl	۴/۴۹
۱۰ ^t	6,7-dimethoxyquinoline	4-F	Phenyl	۴/۶۳
۱۱	6,7-dimethoxyquinoline	H	Methoxymethyl	۵/۴۵
۱۲	6,7-dimethoxyquinoline	4-F	Methoxymethyl	۵/۷۸
۱۳	6,7-dimethoxyquinoline	H	H	۵/۲۳
۱۴	6,7-dimethoxyquinoline	4-F	H	۵/۳۹
۱۵	6,7-dimethoxyquinoline	H	1-Cyclohexenyl	۴/۸۴
۱۶	6,7-dimethoxyquinoline	4-F	1-Cyclohexenyl	۴/۹۹
۱۷	6,7-dimethoxyquinoline	H	3-Pyridyl	۴/۶۹
۱۸	6,7-dimethoxyquinoline	4-F	3-Pyridyl	۴/۸۱
۱۹ ^t	6,7-dimethoxyquinoline	H	3-Thienyl	۴/۷۸
۲۰	6,7-dimethoxyquinoline	4-F	3-Thienyl	۴/۹۵
۲۱	6,7-dimethoxyquinoline	2-Methyl	Methoxymethyl	۵/۳۱
۲۲	6,7-dimethoxyquinoline	3-Methyl	Methoxymethyl	۵/۱۰
۲۳	6,7-dimethoxyquinoline	4-Methyl	Methoxymethyl	۵/۰۶

ادامه جدول ۱-۲

ردیف	R ₂	R	R ₁	(pIC ₅₀) فعالیت دارویی
۲۴	6,7-dimethoxyquinoline	2-F	Methoxymethyl	۵/۵۰
۲۵	6,7-dimethoxyquinoline	3-F	Methoxymethyl	۵/۶۱
۲۶ [†]	6,7-dimethoxyquinoline	4-Cl	Methoxymethyl	۶/۰۴
۲۷	6,7-dimethoxyquinoline	4-Br	Methoxymethyl	۵/۱۰
۲۸ [†]	6,7-dimethoxyquinoline	3,4-di-Cl	Methoxymethyl	۶/۰۶
۲۹	6,7-dimethoxyquinoline	4-CF ₃	Methoxymethyl	۴/۸۸
۳۰	7-butoxy-6-methoxyquinoline	4-Cl	Methoxymethyl	۵/۵۵
۳۱	7-butoxy-6-methoxyquinoline	3,4-di-Cl	Methoxymethyl	۵/۳۹
۳۲ [†]	4-(3-((6-methoxyquinolin-7-yl)oxy)propyl)morpholine	4-Cl	Methoxymethyl	۶/۵۰
۳۳	4-(3-((6-methoxyquinolin-7-yl)oxy)propyl)morpholine	3,4-di-Cl	Methoxymethyl	۶/۲۳
۳۴	6-methoxy-7-(3-(piperidin-1-yl)propoxy)quinoline	4-Cl	Methoxymethyl	۶/۳۹
۳۵	6-methoxy-7-(3-(piperidin-1-yl)propoxy)quinoline	3,4-di-Cl	Methoxymethyl	۶/۱۱
۳۶	6-methoxy-7-(3-(4-methylpiperazin-1-yl)propoxy)quinoline	4-Cl	Methoxymethyl	۶/۴۴
۳۷	6-methoxy-7-(3-(4-methylpiperazin-1-yl)propoxy)quinoline	3,4-di-Cl	Methoxymethyl	۶/۱۳
۳۸	6-methoxy-7-(3-(4-methylpiperidin-1-yl)propoxy)quinoline	4-Cl	Methoxymethyl	۶/۲۶
۳۹	6-methoxy-7-(3-(4-methylpiperidin-1-yl)propoxy)quinoline	3,4-di-Cl	Methoxymethyl	۶/۰۶

[†] نمایانگر داده‌های مجموعه آزمون می‌باشد. سایر داده‌ها مربوط به مجموعه آموزش می‌باشند.



شکل ۲-۲ اسکلت اصلی ساختار ترکیبات مورد مطالعه مربوط به ترکیبات ۴۰ تا ۶۳

جدول ۲-۲ جزئیات ساختار ۴۰ تا ۶۳ به همراه مقادیر pIC_{50} ترکیبات مورد مطالعه

ردیف	X	R	R ₁	فعالیت دارویی (pIC_{50})
۴۰	H	Methyl	H	۵/۳۶
۴۱	H	Methyl	4-F	۵/۴۹
۴۲	F	Methyl	H	۵/۵۲
۴۳	F	Methyl	4-F	۵/۵۷
۴۴	F	Methyl	3-F	۵/۵۱
۴۵	F	Methyl	2-F	۵/۵۱
۴۶ ^t	F	Methyl	4-Me	۵/۱۵
۴۷	F	Methyl	3-Me	۵/۲۲
۴۸	F	Methyl	2-Me	۵/۲۱
۴۹ ^t	F	Methyl	4-Cl	۵/۵۴
۵۰	F	Methyl	4-Br	۵/۵۱
۵۱	F	Methyl	4-CF ₃	۵/۰۹
۵۲	F	Methyl	2,4-Cl	۵/۵۳
۵۳	F	n-butyl	H	۵/۵۱
۵۴	F	n-butyl	4-F	۵/۵۴
۵۵ ^t	F	n-butyl	4-Me	۵/۰۳
۵۶	F	3-morpholinopropyl	H	۵/۸۹
۵۷	F	3-morpholinopropyl	4-F	۶/۴۷
۵۸	F	3-morpholinopropyl	4-Me	۵/۴۳
۵۹	F	3-(piperidin-1-yl)propyl	H	۵/۹۱
۶۰	F	3-(piperidin-1-yl)propyl	4-F	۶/۵۰
۶۱	F	3-(piperidin-1-yl)propyl	4-Me	۵/۶۶
۶۲	F	3-(pyrrolidin-1-yl) propyl	H	۶/۰۳
۶۳	F	3-(pyrrolidin-1-yl) propyl	4-F	۶/۸۰

^t نمایانگر داده‌های مجموعه آزمون می‌باشد. سایر داده‌ها مربوط به مجموعه آموزش می‌باشند.

۲-۲-۲ رسم و بهینه سازی ساختار هندسی ترکیبات

ساختار شیمیایی ترکیبات مورد مطالعه با استفاده از نرم افزار هایپرکم رسم شد و بهینه سازی ساختارها مطابق با روش کار (بخش ۱-۳-۲)، با استفاده از روش نیمه تجربی کوانتومی AM1، به عنوان یکی از کارآمدترین روش های نیمه تجربی، بهینه شد و بهینه سازی تا زمانی ادامه یافت که جذر میانگین مربعات گرادیان^۱ انرژی به ۰/۰۰۱ کیلو کالری بر مول برسد.

۲-۲-۳ محاسبه توصیف کننده ها

به منظور کمی کردن ساختارهای شیمیایی بهینه شده و استخراج توصیف کننده ها، از نرم افزار دراگون استفاده شد. به این منظور ترکیبات مورد مطالعه در نرم افزار دراگون فراخوانی شدند و تعداد ۳۲۲۴ توصیف کننده برای ساختار هر ترکیب در ۲۲ دسته از جمله نوع و تعداد اتم ها، گروه های عاملی، توصیف کننده های مکانی و غیره محاسبه شد.

۲-۲-۴ پیش پردازش اولیه

با توجه به تعداد زیاد توصیف کننده های محاسبه شده و احتمال وجود توصیف کننده های اضافی و فاقد اطلاعات مفید، فرآیند پیش پردازش توصیف کننده ها به منظور کاهش تعداد توصیف کننده ها و در نتیجه بهبود صحت پیش بینی، کاهش زمان مدل سازی و افزایش تفسیرپذیری مدل بر روی کل توصیف کننده های محاسبه شده اعمال شد. به این منظور توصیف کننده هایی که مقادیر ثابت برای تمام ترکیبات داشتند حذف گردید و ۱۵۹۱ توصیف کننده به دست آمد. سپس، ضرایب همبستگی بین تمام توصیف کننده ها محاسبه شد. از بین دو توصیف کننده همبسته ($R > 0/90$) که دارای اطلاعات تقریباً یکسانی هستند، توصیف کننده با بیشترین همبستگی با پاسخ، نگه داشته شده و توصیف کننده بعدی حذف شد و در نتیجه تعداد ۴۰۱

^۱ Root Mean Square Gradient

توصیف کننده باقی ماند. در مرحله بعد توصیف کننده‌های نسبتاً ثابت (توصیف کننده‌هایی که انحراف استاندارد کمتر از ۰/۰۰۱ داشتند) حذف شدند و ۳۹۲ توصیف کننده باقی ماند. هر چند تعداد توصیف کننده‌ها به ۳۹۲ کاهش یافت، ولی اگر بتوان به گونه‌ای با حذف توصیف کننده‌های با تأثیر کم‌تر، تعداد توصیف کننده‌ها را کاهش بیش‌تری داد، در این صورت می‌توان به مدل‌های خطی SCAD تنک‌تر دست یافت. بنابراین، روش رگرسیون Ridge به همراه یک فرآیند چند مرحله‌ای جدید به عنوان روش غربالگری روی توصیف کننده‌های باقی مانده اعمال شد.

۲-۲-۵ استفاده از روش رگرسیون Ridge برای غربالگری توصیف کننده ها

روش رگرسیون Ridge یک روش مدل سازی است و به دلیل این که همه‌ی متغیرها را در مدل نگه می‌دارد به عنوان روش انتخاب متغیر استفاده نمی‌شود. در این پروژه با استفاده از تکنیکی که در ادامه به شرح آن پرداخته می‌شود، از روش Ridge به عنوان روش غربالگری استفاده شد. روش Ridge با استفاده از تکنیک رد مرحله‌ای تک تک روی ۳۹۲ توصیف کننده باقی مانده اجرا شد. با به کارگیری روش رد مرحله‌ای تک تک به کمک Ridge، هر بار یکی از داده‌های مورد مطالعه (۱ ترکیب از ۶۳ ترکیب) کنار گذاشته شد و مدل Ridge با استفاده از باقی مانده داده‌ها (۶۲ ترکیب) توسعه داده شد. این فرآیند برای همه داده‌ها تکرار شد، به طوری که در نهایت ۶۳ مدل Ridge با ۳۹۲ توصیف کننده با ضرایب رگرسیونی متفاوت به دست آمد. به منظور انجام فرآیند غربالگری، ابتدا مقادیر میانگین ضرایب توصیف کننده زام ($\bar{\beta}_j$) و انحراف استاندارد ضرایب توصیف کننده زام (σ_{β_j}) مطابق با رابطه ۲-۱ و رابطه ۲-۲ محاسبه شدند. در این روابط N و β_{ij} به ترتیب نمایانگر ضرایب توصیف کننده زام در مدل i ام و تعداد ضرایب توصیف کننده زام در مدل i ام است.

$$\bar{\beta}_j = \frac{\sum_{i=1}^N \beta_{ij}}{N} \quad \text{رابطه ۱-۲}$$

$$\sigma_{\beta_j} = \sqrt{\frac{\sum_{i=1}^N (\beta_{ij} - \bar{\beta}_j)^2}{N-1}} \quad \text{رابطه ۲-۲}$$

برای استفاده از تکنیک غربالگری Ridge، نیاز به یک معیار مناسب جهت تشخیص حد آستانه^۱ برای شناسایی توصیف کننده‌های زائد (R) و اصلی (X) از یکدیگر و ایجاد ماتریس‌های جدید X، R و XR می‌باشد (طبق تعریف بخش ۱-۳-۴-۲)، تا توصیف کننده‌های معنا دار از بی معنا متمایز گردد. معیار C (قدر مطلق نسبت میانگین ضرایب توصیف کننده Z ام به انحراف استاندارد همان توصیف کننده) به عنوان معیار شناسایی و ایجاد ماتریس‌های R، X و XR مطابق با رابطه ۲-۳ مورد استفاده قرار گرفت. مقدار ماکسیمم C برای ماتریس R به عنوان حد آستانه برای شناسایی و ایجاد ماتریس‌های جدید X، R و XR تعریف شد. به طوری که پس از محاسبه پارامتر C مربوط به توصیف کننده‌های X موجود در ماتریس XR، توصیف کننده‌هایی که مقدار C کمتری نسبت به ماکسیمم C ماتریس R دارند، به عنوان توصیف کننده‌های زائد (R) جدید شناخته می‌شوند و از ماتریس X حذف می‌گردند و به عنوان ماتریس R جدید شناخته می‌شوند و سایر توصیف کننده‌ها با C بزرگتر از C ماکسیمم ماتریس R به عنوان ماتریس X جدید شناخته می‌شوند. بدین ترتیب ماتریس X، R و XR جدید برای هر مرحله به این روش ایجاد شد.

$$C = \left| \frac{\bar{\beta}_j}{\sigma_{\beta_j}} \right| \quad \text{رابطه ۲-۳}$$

لازم به ذکر است در مرحله اول که روی داده‌های اصلی (۳۹۲) روش Ridge اجرا می‌شود، ماتریس زائدی (R) وجود ندارد که برای آن پارامتر C محاسبه گردد، از این رو توصیف کننده‌هایی که میانگین ضرایب رگرسیونی آن‌ها با انحراف استانداردشان برابر بود $(C_0 : \bar{\beta}_j = \sigma_{\beta_j})$ ، یا به عبارتی نسبت میانگین به انحراف استاندارد برابر یک یا کوچکتر از یک داشتند $(C_0 \leq 1)$ ، به عنوان توصیف کننده‌های بی معنای اولیه (R) شناخته شدند و به این صورت اولین ماتریس R ایجاد شد. توصیف کننده‌های باقی مانده به ماتریس X

¹ Threshold

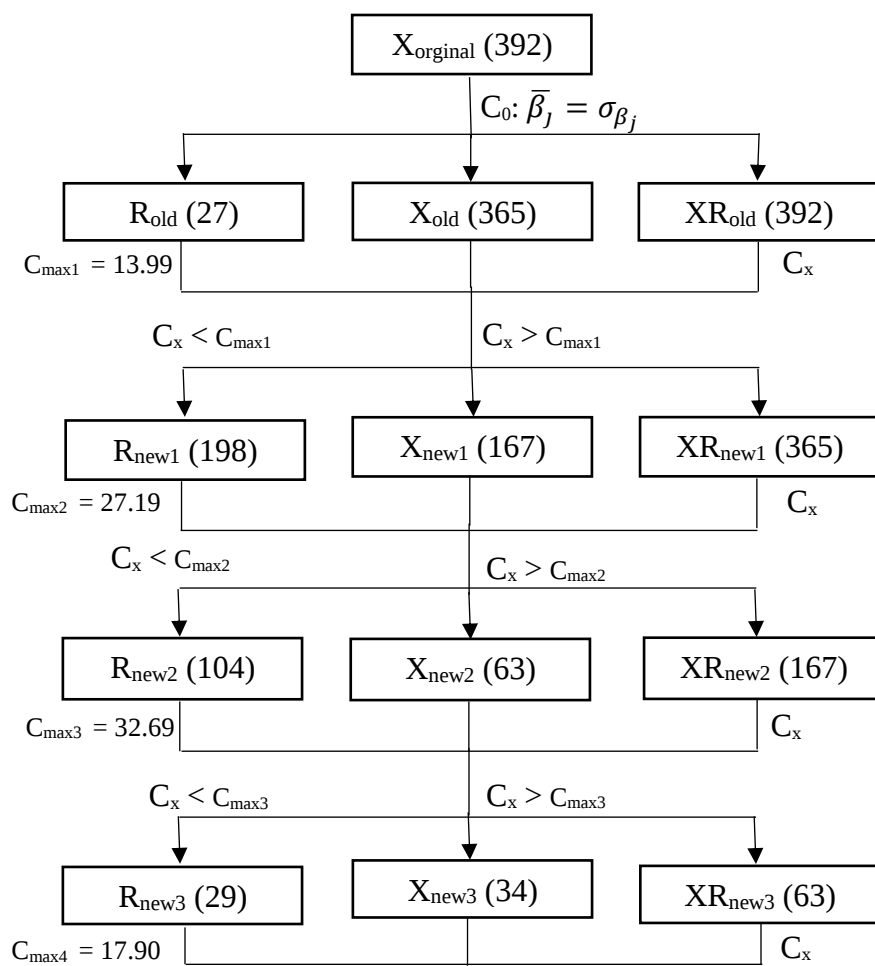
تعلق دارند و همچنین ماتریس XR به ترتیب از مجموع توصیف کننده های X و R تشکیل شد. طبق فلوچارت (شکل ۳-۲)، مدل Ridge با ماتریس های جدید X (ماتریس X اول در غیاب توصیف کننده های R جدید)، R (متشکل از توصیف کننده هایی با پارامتر C ماتریس X کوچکتر از C ماکسیمم ماتریس R) و XR (متشکل از توصیف کننده های به ترتیب ماتریس X جدید و توصیف کننده های R جدید) با روند اشاره شده در شکل ۳-۲ تشکیل شدند. ساخت مدل های Ridge و ایجاد ماتریس های R، X و XR جدید تا جایی ادامه می یابد که دیگر توصیف کننده ای C ماتریس X ای (C_x) کوچکتر از C ماکسیمم محاسبه شده از ماتریس R جدید (C_{maxR}) وجود نداشته باشد و در این مرحله غربالگری به اتمام می رسد، زیرا ساخت ماتریس R جدید دیگر امکان پذیر نیست. در ادامه بحث جزئیات مربوط به فلوچارت شکل ۳-۲ به طور مرحله ای توضیح داده شده است.

همان طور که در شکل ۳-۲ نشان داده شده است، در مرحله اول که ایجاد ماتریس R غیر ممکن است. روش Ridge با تکنیک رد مرحله ای تک تک با Z توصیف کننده شامل ۳۹۲ توصیف کننده ($X_{original}$) اجرا شد. پس از این که ضرایب رگرسیونی Ridge مربوط به ۳۹۲ توصیف کننده به دست آمد، از Z توصیف کننده، ۲۷ توصیف کننده (با شرط $C_0 \leq 1$)، جزء ماتریس R (R_{old})، ۳۶۵ توصیف کننده جزء ماتریس X (X_{old}) و ۳۹۲ توصیف کننده متعلق به ماتریس XR هستند (XR_{old}).

پس از ایجاد ماتریس های X، R و XR مرحله قبل، Ridge با تکنیک رد مرحله ای تک تک با Z توصیف کننده مرحله قبل (XR_{old}) اجرا شد، ضرایب رگرسیونی برای Z توصیف کننده محاسبه شد و پارامتر C (رابطه ۳-۲) برای ضرایب ۳۶۵ توصیف کننده X موجود در ماتریس XR و ۲۷ توصیف کننده ماتریس R به دست آمد. ماکسیمم پارامتر C ماتریس R به عنوان معیار غربالگری انتخاب شد. بنابراین از ۳۹۲ توصیف کننده مرحله قبل ۱۶۷ توصیف کننده متعلق به توصیف کننده های ماتریس X موجود در ماتریس XR، دارای پارامتر C بزرگتر از ماکسیمم پارامتر C ($C_x > C_{max1}$) ماتریس R است و به عنوان ماتریس X جدید

شناخته می‌شود (X_{new1}). در این مرحله، ماتریس R (R_{new1}) متشکل از ۱۹۸ توصیف کننده می‌باشد. در نهایت ماتریس XR جدید (XR_{new1}) این مرحله نیز با Z توصیف کننده شامل ۳۶۵ توصیف کننده ایجاد شد. در مرحله سوم، Ridge با تکنیک رد مرحله‌ای تک تک با استفاده از Z توصیف کننده (۳۶۵) توصیف کننده ماتریس (XR_{new1}) اجرا شد و ضرایب رگرسیونی برای Z توصیف کننده محاسبه شد و پارامتر C (رابطه ۳-۲) برای ضرایب ۱۶۷ توصیف کننده X موجود در ماتریس XR (شامل ۳۶۵ توصیف کننده ماتریس XR_{new1}) و ۱۹۸ توصیف کننده ماتریس R (R_{new1}) به دست آمد. از ۳۶۵ توصیف کننده مرحله قبل ۶۳ توصیف کننده دارای پارامتر C بزرگتر از ماکسیمم پارامتر C ($C_x > C_{max2}$) ماتریس R (R_{new2}) است و به عنوان ماتریس X جدید شناخته می‌شود (X_{new2}). در این مرحله، ماتریس R (R_{new2}) متشکل از ۱۰۴ توصیف کننده می‌باشد. ماتریس XR جدید این مرحله (XR_{new2}) نیز با Z توصیف کننده شامل ۱۶۷ توصیف کننده به دست آمد.

در مرحله بعد مجدداً Ridge با تکنیک رد مرحله‌ای تک تک با Z توصیف کننده مرحله چهارم (شامل ۶۳ توصیف کننده (XR_{new3})) اجرا شد و ضرایب رگرسیونی برای Z توصیف کننده محاسبه شد و پارامتر C (رابطه ۳-۲) برای ضرایب ۳۴ توصیف کننده X موجود در ماتریس XR و ۲۹ توصیف کننده ماتریس R (R_{new3}) به دست آمد. در این مرحله هیچ توصیف کننده‌ای دارای پارامتر C کوچک‌تر از ماکسیمم پارامتر C ($C_x < C_{max3}$) ماتریس R (R_{new3}) نبوده و تشکیل ماتریس R جدید غیر ممکن است. بنابراین اجرای Ridge با تکنیک رد مرحله‌ای تک تک پس از اجرای ۴ مرحله متوالی متوقف شد و در نهایت ۳۴ توصیف کننده (توصیف کننده‌های موجود در ماتریس X_{new3}) به عنوان متغیرهای مستقل باقی مانده با بیش‌ترین تأثیر با پاسخ (فعالیت دارویی) شناسایی شد. این توصیف کننده‌ها به عنوان ورودی مدل سازی انقباضی SCAD به کار گرفته شدند.



در این مرحله دیگر متغیری دارای شرط $C_x < C_{max4}$ برای ساخت ماتریس R_{new4} وجود ندارد و غربالگری در این مرحله متوقف می شود.

شکل ۳-۲ فلوچارت نمایش روند غربالگری Ridge (زیروند Original مربوط به توصیف کننده های مرحله پیش پردازش،

زیروند old مربوط به توصیف کننده های اولین مرحله غربالگری Ridge و زیروند new مربوط به توصیف کننده های

حاصل از هر مرحله اجرای روند غربالگری Ridge است).

۲-۲-۶ انتخاب متغیر و مدل سازی SCAD

در این بخش به منظور یافتن مؤثرترین توصیف کننده‌ها و ساختار مدل QSAR از روش انقباضی SCAD استفاده شد. ۳۴ توصیف کننده‌ی باقی‌مانده از مرحله غربالگری Ridge به‌عنوان ورودی مدل در نظر گرفته شد. لازم به ذکر است که در مراحل پیش‌پردازش اولیه و غربالگری از ۶۳ داده برای انتخاب بهترین توصیف کننده‌ها استفاده شد و ضرایب استخراج شد. در این مرحله ۶۳ داده به دو مجموعه تقسیم می‌شود. پس از تقسیم بندی داده‌ها به دو مجموعه آموزش و آزمون (مطابق با روش بخش ۱-۳-۱) از مجموعه آموزش برای انتخاب مؤثرترین توصیف کننده‌ها و مدل‌سازی استفاده شد. بنابراین روش SCAD با به-کارگیری بسته نرم افزاری ncvreg در نرم افزار R-Studio با تکنیک ۱۰ فولد-ارزیابی تقاطعی^۱، به‌طوری که برای یک λ داده شده، مجموعه داده‌های آموزش (۵۳) به‌طور تصادفی به ۱۰ فولد اختصاص می‌یابد. سپس یک فولد (فولد اول) به‌عنوان مجموعه آزمون و سایر فولدها (۹ فولد دیگر) به‌عنوان مجموعه آموزش در نظر گرفته می‌شود. پس از برآزش مدل با استفاده از مجموعه آموزش (۹ فولد)، مقدار متغیر پاسخ (فعالیت دارویی) به‌ازای داده‌های آزمون (فولد اول) پیش‌بینی می‌شود و خطای پیش‌بینی (PE) مجموعه آزمون (فولد اول) طبق رابطه ۲-۴ به‌دست می‌آید.

$$PE_1 = \sum_{i \in I_1} (y_i - \hat{y}_i)^2 \quad \text{رابطه ۲-۴}$$

در رابطه ۲-۴، $i \in I_1$ است، که I_1 اندیس ترکیبات گروه اول است. در مرحله‌ی بعد، یک فولد دیگر به‌عنوان مجموعه آزمون کنار گذاشته می‌شود و مابقی فولدها به‌عنوان مجموعه آموزش استفاده می‌شوند و خطای پیش‌بینی (PE) محاسبه می‌شود. این روند k بار تکرار می‌شود و در هر تکرار یکی از فولدها به‌عنوان داده‌های مجموعه آزمون و بقیه فولدها به‌عنوان داده‌های مجموعه آموزش در نظر گرفته می‌شوند. در نتیجه

^۱10- fold Cross-validation

k برآورد خطای $PE_1, \dots, PE_k$ آزمون به دست آمده و در نهایت، امتیاز ارزیابی تقاطعی (CV) به صورت رابطه ۵-۲ محاسبه می شود:

$$CV(\lambda) = \frac{1}{k} \sum_{i=1}^k PE_i \quad \text{رابطه ۵-۲}$$

این کار برای λ های مختلف تکرار شده و سپس λ ای انتخاب می شود که مقدار خطای پیش بینی را مینیمم کند. با توجه به $\lambda_{\min}$ مدل با حداقل خطای پیش بینی ارزیابی تقاطعی انتخاب شد [۶۸]. در نهایت با اجرای روش SCAD، تعداد ۹ توصیف کننده انتخاب شد. مدل SCAD (رابطه ۶-۲) با استفاده از توصیف کننده های منتخب توسعه یافت. نماد، طبقه بندی و معنای هر یک از این توصیف کننده ها در جدول ۳-۲ آورده شده است.

$$\begin{aligned} \text{رابطه ۶-۲} \quad pIC_{\Delta} = & 7/4 + 0/57 DECC + 0/25 MATS4e - 0/04 C025 - 0/96 Mor26p - 0/03 H046 - 0/02 \\ & N075 - 0/001 SEigZ + 0/52 Mor23m - 2/4 GATS4e \end{aligned}$$

به منظور بررسی دقیق تر توصیف کننده های منتخب مدل Ridge-SCAD، احتمال وجود همبستگی با محاسبه مربع ضریب همبستگی بین دو توصیف کننده در جدول ۴-۲ مورد بررسی قرار گرفت. نتایج نشان می دهد که بین توصیف کننده های منتخب همبستگی معنا داری وجود ندارد. علاوه بر این احتمال وجود هم خطی بین توصیف کننده های منتخب با محاسبه پارامتر مقادیر افزایش تورم واریانس (VIF) (مطابق رابطه ۱-۳۴) مورد بررسی قرار گرفت و نتایج آن در ستون آخر جدول ۴-۲ گزارش شده است. مقادیر VIF کمتر از ۱۰ توصیف کننده های مدل Ridge-SCAD بیانگر عدم وجود هم خطی بین توصیف کننده ها می باشد [۲۲-۲۵].

جدول ۳-۲ توصیف کننده های انتخاب شده توسط روش Ridge-SCAD

ردیف	نماد	طبقه بندی	معنا
۱	DECC	topological	Eccentric
۲	C-025	Atom-centred fragment	R-CR-R
۳	GATS4e	2D autocorrelation	Geary autocorrelation – lag 4
۴	Mor26p	3D-MoRSE	3D-MoRSE – signal 26
۵	Mor23m	3D-MoRSE	3D-MoRSE – signal 23
۶	H-046	Atom-centred fragment	H attached to C (sp ³) no X attached to next C
۷	MATS4e	2D autocorrelation	Moran autocorrelation – lag 4
۸	N-075	Atom-centred fragment	R-N-R / R-N-X
۹	SEigZ	Eigenvalue-based indices	Eigenvalue sum from Z

جدول ۴-۲ ماتریس همبستگی و مقادیر عامل افزایش واریانس توصیف کننده های انتخاب شده مدل Ridge-SCAD

توصیف کننده ها	DECC	C-025	GATS4e	Mor26p	Mor23m	H-046	MATS4e	N-075	SEigZ	VIF
DECC	۱									۳/۷۷۹
C-025	-۰/۰۰۶	۱								۱/۲۵۳
GATS4e	-۰/۳۰۶	-۰/۱۳۶	۱							۱/۵۷۵
Mor26p	-۰/۵۳۵	-۰/۱۲۱	۰/۳۱۵	۱						۱/۷۱۸
Mor23m	۰/۵۳۲	-۰/۱۵۸	-۰/۰۵	-۰/۳۵۶	۱					۱/۹۴۷
H-046	-۰/۲۱۷	-۰/۳۱	۰/۴۴۳	۰/۲۹۳	-۰/۰۵۷	۱				۱/۵۴۹
MATS4e	۰/۶۹۱	۰/۰۷۲	-۰/۱۴۴	-۰/۵۶۳	۰/۶۴۹	-۰/۲۲۵	۱			۴/۶۱۳
N-075	-۰/۳	-۰/۱۲۴	۰/۳۸۳	۰/۳۶۷	-۰/۲۶۶	۰/۳۴۵	-۰/۴۸۴	۱		۱/۶۳
SEigZ	۰/۵۳	-۰/۱۲۲	-۰/۴۲۷	-۰/۳۰۹	۰/۱۸۵	-۰/۳۱۹	۰/۰۷۶	-۰/۱۹۴	۱	۲/۴۲۳

به منظور بررسی عملکرد روش Ridge برای غربالگری توصیف کننده ها، مدل SCAD بدون استفاده

از روش غربالگری Ridge ساخته شد. پس از پیش پردازش های اولیه توصیف کننده ها، روش SCAD با

به کارگیری بسته نرم افزاری ncvreg در نرم افزار R-Studio با تکنیک ۱۰ فولد- ارزیابی تقاطعی اجرا شد.

به طوری که برای یک λ داده شده، مجموعه داده های آموزش (۵۳) به طور تصادفی به ۱۰ فولد اختصاص

می یابد. سپس فولد اول به عنوان مجموعه آزمون و سایر فولدها (۹) به عنوان مجموعه آموزش در نظر گرفته

می‌شود. پس از برازش مدل با استفاده از مجموعه آموزش، مقدار متغیر پاسخ (فعالیت دارویی) به‌ازای داده‌های آزمون پیش‌بینی می‌شود و خطای پیش‌بینی طبق رابطه ۲-۴ به‌دست می‌آید. این روند k بار تکرار می‌شود و در هر تکرار یکی از فولدها به‌عنوان داده‌های آزمون و بقیه فولدها به‌عنوان داده‌های آموزش در نظر گرفته می‌شوند. در نتیجه k برآورد خطای $PE_1, \dots, PE_k$ آزمون به‌دست آمده و در نهایت، امتیاز ارزیابی تقاطعی (CV) به‌صورت رابطه ۲-۵ محاسبه می‌شود. این کار برای λ های مختلف تکرار شده و سپس λ ای انتخاب می‌شود که مقدار خطای پیش‌بینی را مینیمم کند. با توجه به $\lambda_{\min}$ مدل با حداقل خطای پیش‌بینی ارزیابی تقاطعی انتخاب شد. با اجرای روش SCAD، تعداد ۱۲ توصیف کننده انتخاب شدند. نام و نوع توصیف کننده‌های منتخب روش SCAD در جدول ۲-۵ خلاصه شده است. مدل SCAD در رابطه ۲-۷ آورده شده است.

$$pIC_{50} = 4/2 + 1/1 DECC - 2/3 MATS2p - 0/55 Mor22u + 0/27 Mor09m + 0/05 Mor23m - 0/02 Mor30m - 0/09 Mor32v + 0/43 G3m + 0/04 E1m - 0/27 G3p - 0/43 C025 - 0/24 LAI$$

جدول ۲-۵ توصیف کننده های انتخاب شده توسط روش SCAD

ردیف	نماد	طبقه‌بندی	معنا
۱	DECC	topological	eccentric
۲	C-025	Atom-centred fragment	R-CR-R
۳	Mor22u	3D-MoRSE	3D-MoRSE – signal 22
۴	Mor09m	3D-MoRSE	3D-MoRSE – signal 09
۵	MATS2p	2D autocorrelation	Moran autocorrelation – lag 2
۶	LAI	molecular properties	Lipiniski Alert index.
۷	Mor23m	3D-MoRSE	3D-MoRSE – signal 23
۸	Mor32v	3D-MoRSE	3D-MoRSE – signal 32
۹	G3m	WHIM	3st component symmetry directional WHIM index
۱۰	Mor30m	3D-MoRSE	3D-MoRSE – signal 30
۱۱	E1m	WHIM	1st component accessibility directional WHIM index
۱۲	G3p	WHIM	3st component symmetry directional WHIM index

۲-۲-۷ ارزیابی مدل Ridge-SCAD

از مرحله‌های اساسی ساخت مدل QSAR، ارزیابی مدل است. در این پروژه، قدرت پیش‌بینی و تعمیم پذیری مدل توسعه یافته Ridge-SCAD با استفاده از پیش‌بینی فعالیت دارویی داده‌های مجموعه آزمون، فعالیت دارویی کل داده‌های پیش‌بینی شده به وسیله تکنیک رد مرحله‌ای تک تک (LOO)، محاسبه پارامترهای آماری، آزمون‌های دامنه کاربرد و Y-تصادفی مورد ارزیابی قرار گرفت.

۲-۲-۷-۱ ارزیابی مدل Ridge-SCAD و مدل SCAD با استفاده از پیش‌بینی داده‌های مجموعه

آزمون

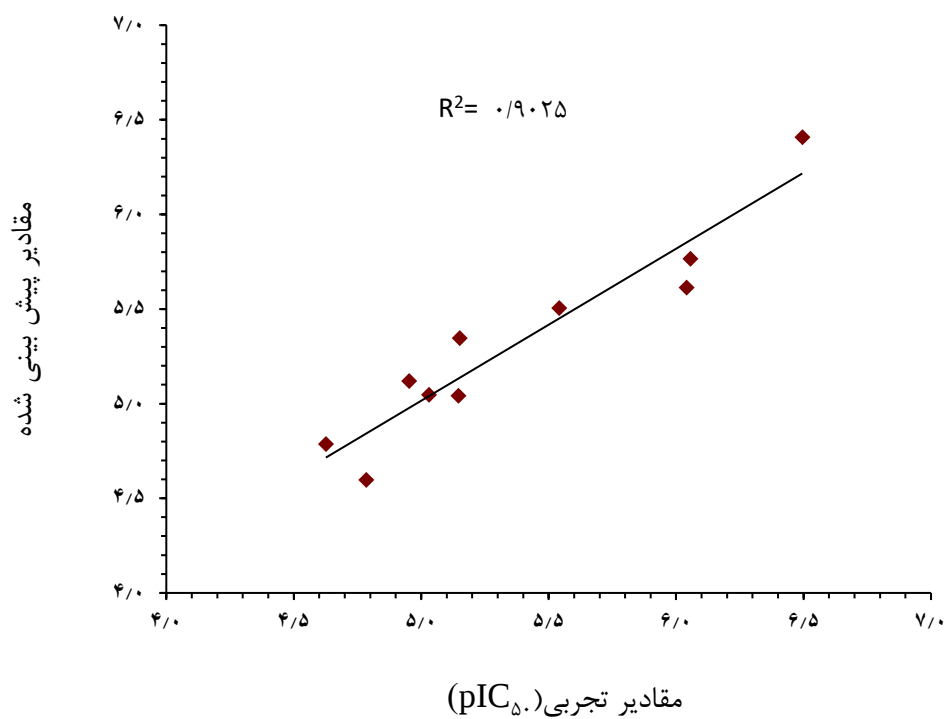
اعتبار و اعتمادپذیری^۱ مدل توسعه یافته Ridge-SCAD با استفاده از داده‌های مجموعه آزمون که در ساخت مدل شرکت نداشته‌اند، سنجیده شد. بنابراین فعالیت دارویی (pIC_{50}) ترکیبات مجموعه آزمون (۱۰ ترکیب) مشخص شده در جدول ۶-۲ با استفاده از مدل Ridge-SCAD (رابطه ۶-۲) پیش‌بینی شدند. مقادیر پیش‌بینی شده در جدول ۶-۲ خلاصه شده است. خطای کم پیش‌بینی شده، دلالت بر قدرت پیش‌بینی مدل Ridge-SCAD دارد. علاوه بر این مقادیر پیش‌بینی شده فعالیت دارویی برحسب مقادیر واقعی متناظر آن‌ها رسم شد (شکل ۴-۲). R^2 مربوط به داده‌های مجموعه آزمون برابر با ۰/۹۰۲۵ به دست آمد. مقدار R^2 بزرگ‌تر از ۰/۶ نشان‌دهنده قدرت پیش‌بینی مدل Ridge-SCAD است. نزدیکی R^2 مربوط به مجموعه آزمون ($R^2_{test} = ۰/۹۰۲۵$) به R^2 مربوط به مجموعه آموزش ($R^2_{train} = ۰/۹۳۱۷$) تعمیم‌پذیری مدل توسعه یافته Ridge-SCAD را اثبات می‌کند. علاوه بر این مقادیر باقی مانده‌های پیش‌بینی شده از محاسبه تفاضل مقادیر پیش‌بینی شده و تجربی به دست آمد. نمودار باقی مانده‌ها بر حسب مقادیر تجربی رسم شد (شکل ۵-۲). توزیع مقادیر باقی مانده‌ها حول محور صفر، نشان‌دهنده عدم وجود خطای سیستماتیک در مدل توسعه یافته Ridge-SCAD است.

^۱ Reliability

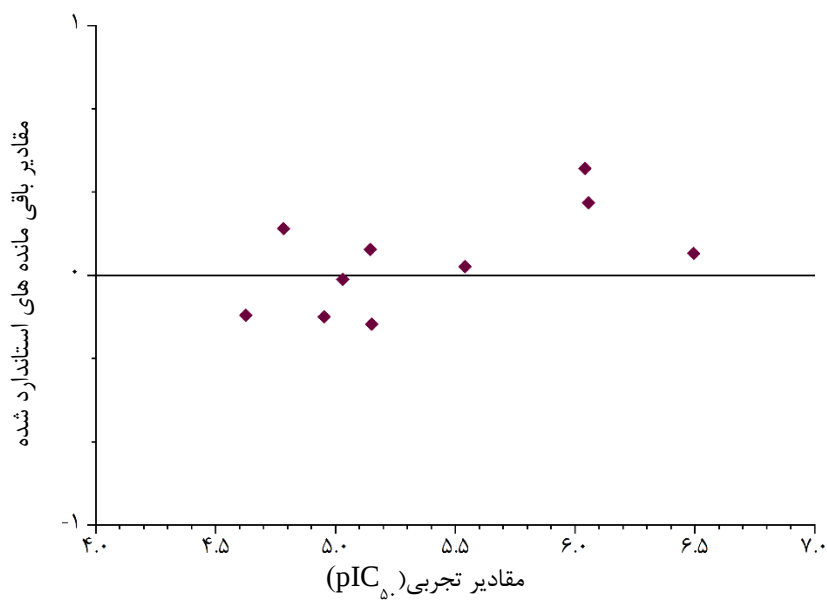
همچنین، به منظور بررسی عملکرد روش Ridge برای غربالگری توصیف کننده‌ها، ارزیابی مدل SCAD بدون استفاده از روش غربالگری Ridge با استفاده از ترکیبات مجموعه آزمون (۱۰ ترکیب) مشخص شده در جدول ۶-۲ و با مدل SCAD (رابطه ۲-۷) پیش‌بینی شدند. مقادیر خطای پیش‌بینی مدل SCAD نسبت به مدل Ridge-SCAD در برخی از ترکیبات مجموعه آزمون افزایش یافته است که این امر دلالت بر قدرت پیش‌بینی کم‌تر مدل SCAD دارد. علاوه بر این مقادیر پیش‌بینی شده فعالیت دارویی برحسب مقادیر واقعی متناظر آن‌ها رسم شد (شکل ۲-۶). R^2 مربوط به داده‌های مجموعه آزمون ($R^2_{test}=0/8042$) و مجموعه آموزش ($R^2_{train}=0/9491$) به دست آمد که هر دو دارای مقادیر قابل قبول هستند، اما نزدیکی مقادیر R^2 دو مجموعه کم‌تر از روش Ridge-SCAD است. با مقایسه نتایج قابلیت بالای Ridge در کاهش ابعاد داده که منجر به انتخاب متغیرهای مؤثرتر و بهتر می‌شود، مشهود است.

جدول ۶-۲ نتایج حاصل از پیش‌بینی فعالیت دارویی داده‌های مجموعه آزمون توسط مدل Ridge-SCAD و مدل SCAD در مجموعه داده‌ها

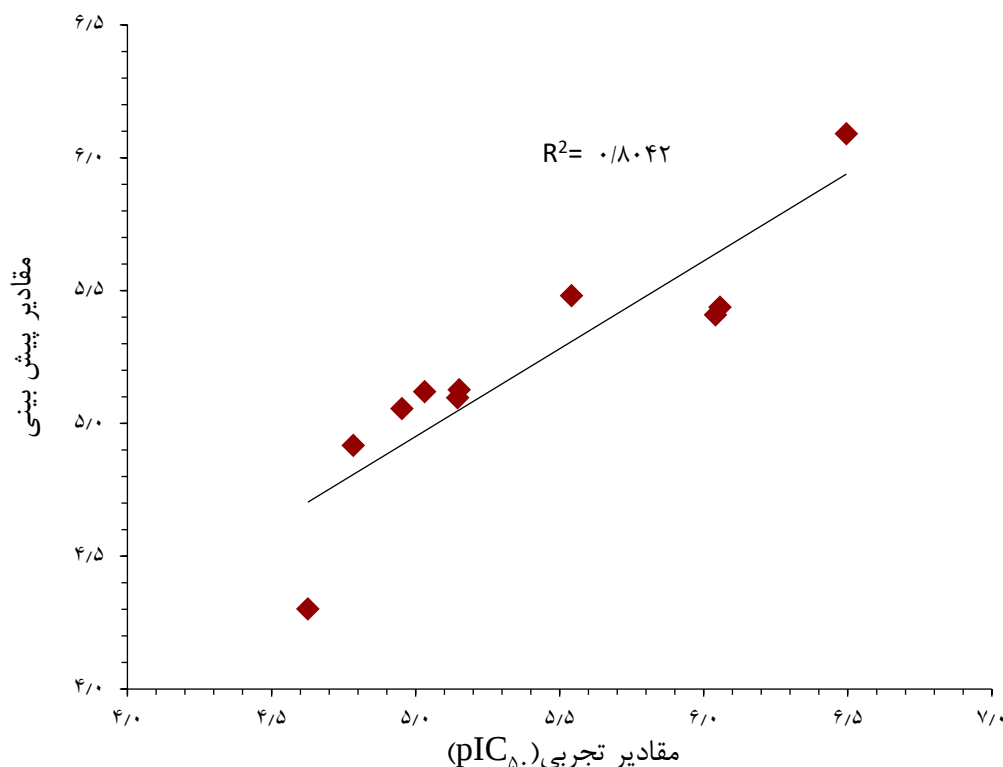
شماره ترکیب در مجموعه داده‌ها	مقدار تجربی (pIC_{50})	مقدار پیش‌بینی شده (pIC_{50})				درصد خطا
		Ridge-SCAD	SCAD	Ridge-SCAD	SCAD	
۲	۵/۱۵	۵/۳۵	۵/۱۳	-۳/۸	۰/۴۸	
۴	۴/۹۵	۵/۱۲	۵/۰۶	-۳/۳	-۲/۱	
۱۰	۴/۶۳	۴/۷۹	۴/۳۰	-۳/۵	۷/۰	
۱۹	۴/۷۸	۴/۶۰	۴/۹۲	۳/۹	-۲/۸	
۲۶	۶/۰۴	۵/۶۱	۵/۴۱	۷/۱	۱۰	
۲۸	۶/۰۶	۵/۷۶	۵/۴۴	۴/۸	۱۰	
۳۲	۶/۴۹	۶/۴۱	۶/۰۹	۱/۴	۶/۲	
۴۶	۵/۱۵	۵/۰۴	۵/۱۰	۲/۰	۰/۹۷	
۴۹	۵/۵۴	۵/۵۱	۵/۴۸	۰/۶۳	۱/۱	
۵۵	۵/۰۳	۵/۰۵	۵/۱۲	-۰/۳۲	-۱/۸	



شکل ۴-۲ نمودار تغییرات مقادیر پیش بینی شده (pIC_{50}) مدل Ridge-SCAD در مقابل مقادیر تجربی برای مجموعه آزمون



شکل ۵-۲ نمودار مقادیر باقی مانده های پیش بینی فعالیت دارویی با استفاده از مدل Ridge-SCAD و مقادیر تجربی بر حسب مقادیر تجربی



شکل ۶-۲ نمودار تغییرات مقادیر پیش بینی شده (pIC₅₀) مدل SCAD در مقابل مقادیر تجربی برای مجموعه آزمون

۲-۲-۷ ارزیابی مدل Ridge-SCAD با استفاده از تکنیک رد مرحله ای تک تک

در این بخش، یکی از تکنیک‌های قدرتمند برای ارزیابی مدل بهینه و بررسی قدرت پیش‌بینی مدل به نام تکنیک رد مرحله‌ای تک تک (LOO) برای پیش‌بینی همه داده‌های مورد مطالعه به عنوان داده آزمون مورد استفاده قرار گرفت. در این مطالعه، هر بار یکی از داده‌های مورد مطالعه کنار گذاشته شد و مدل Ridge-SCAD با سایر داده‌ها آموزش داده شد و فعالیت دارویی ترکیب کنار گذاشته با مدل آموزش داده شده، پیش‌بینی گردید. این فرآیند برای همه ترکیبات انجام شد و همه داده‌ها یکبار به عنوان داده آزمون در نظر گرفته شدند. از این رو فعالیت دارویی همه ترکیبات تکرار شد، به گونه‌ای که تکنیک LOO توسط مدل Ridge-SCAD پیش‌بینی شدند و نتایج حاصل در جدول ۷-۲ آورده شده‌اند. برای مطالعه بیشتر پایداری قدرت پیش‌بینی مدل، نمودار مقادیر پیش‌بینی شده pIC₅₀ تمام ترکیبات بر حسب مقادیر تجربی

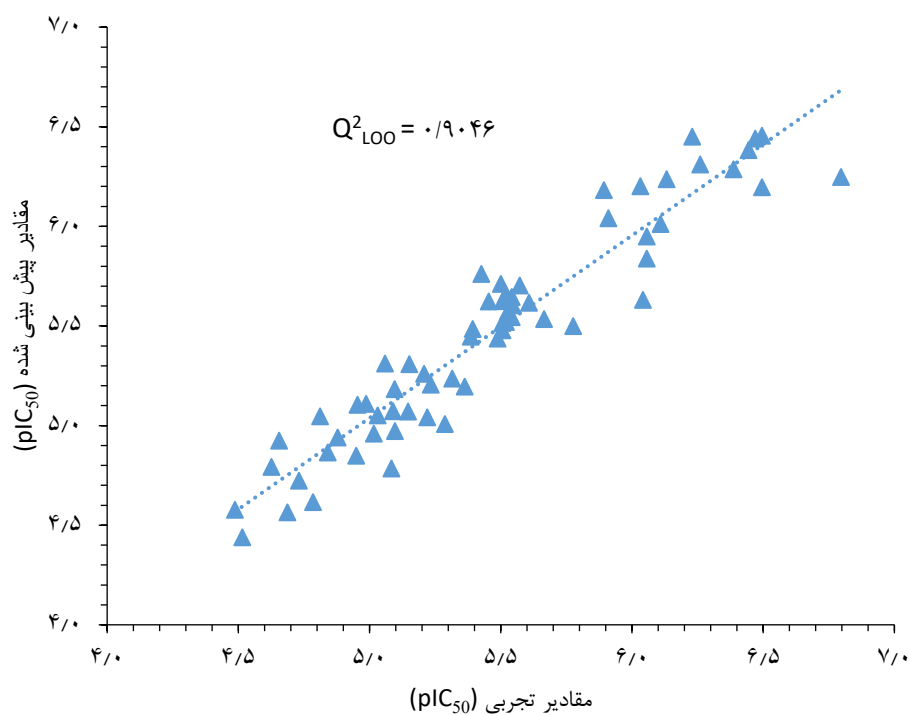
آن‌ها رسم شد. مقدار Q^2_{LOO} که در شکل ۷-۲ نشان داده شده است، بیش‌تر از مقدار قابل قبول ($Q^2 > 0.5$) بوده که نشان‌دهنده این است که مدل توسعه یافته از پایداری قدرت پیش‌بینی و تعمیم پذیری مناسبی برخوردار است. علاوه بر این، به‌منظور بررسی عدم وجود خطای سیستماتیک در مدل توسعه یافته، مقادیر باقی‌مانده بر حسب مقادیر تجربی رسم شد (شکل ۸-۲). نتایج حاصل نشان‌دهنده توزیع یکنواخت و تصادفی داده‌ها حول محور صفر می‌باشد که بیانگر عدم وجود خطای سیستماتیک در مدل Ridge-SCAD است.

جدول ۷-۲ نتایج حاصل از ارزیابی مدل Ridge-SCAD به روش رد مرحله ای تک تک برای کل داده‌ها

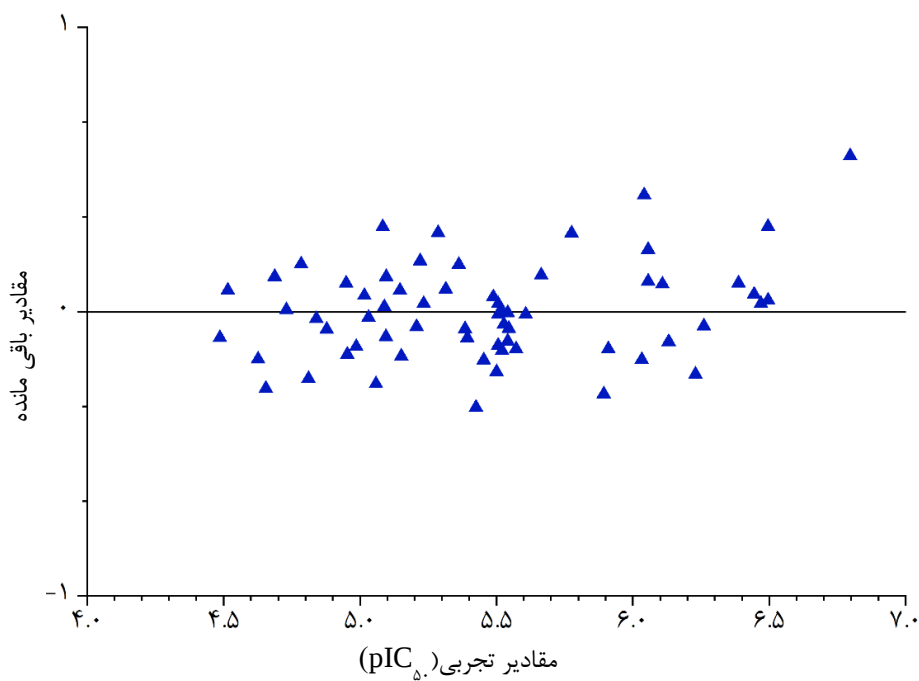
شماره ترکیب	pIC _{۵۰}			شماره ترکیب	pIC _{۵۰}		
	مقدار واقعی	مقدار پیش‌بینی شده	درصد خطا		مقدار واقعی	مقدار پیش‌بینی شده	درصد خطا
۱	۵/۰۲	۴/۹۶	۱/۱۵	۱۹	۴/۷۸	۴/۶۲	۳/۵۱
۲	۵/۱۵	۵/۳۱	-۳/۰۳	۲۰	۴/۹۵	۴/۸۵	۲/۰۳
۳	۴/۷۳	۴/۷۲	۰/۱۵	۲۱	۵/۳۱	۵/۲۳	۱/۴۹
۴	۴/۹۵	۵/۱۰	-۳/۰۳	۲۲	۵/۱۰	۴/۹۷	۲/۴۲
۵	۴/۵۲	۴/۴۴	۱/۶۸	۲۳	۵/۰۶	۵/۳۱	-۵/۰۲
۶	۴/۶۶	۴/۹۲	-۵/۷۹	۲۴	۵/۵۰	۵/۷۱	-۳/۸۴
۷	۵/۰۸	۴/۷۸	۵/۸۸	۲۵	۵/۶۱	۵/۶۲	-۰/۱۶
۸	۵/۲۹	۵/۰۱	۵/۲۷	۲۶	۶/۰۴	۵/۶۳	۶/۸۰
۹	۴/۴۹	۴/۵۸	-۲/۰۴	۲۷	۵/۱۰	۵/۱۸	-۱/۷۲
۱۰	۴/۶۳	۴/۷۹	-۳/۶۱	۲۸	۶/۰۶	۵/۸۴	۳/۶۱
۱۱	۵/۴۵	۵/۶۲	-۳/۱۱	۲۹	۴/۸۸	۴/۹۴	-۱/۲۶
۱۲	۵/۷۸	۵/۵۰	۴/۷۹	۳۰	۵/۵۵	۵/۶۰	-۱/۰۶
۱۳	۵/۲۳	۵/۲۰	۰/۵۵	۳۱	۵/۳۹	۵/۴۸	-۱/۷۰
۱۴	۵/۳۹	۵/۴۵	-۱/۱۲	۳۲	۶/۵۰	۶/۴۵	۰/۶۳
۱۵	۴/۸۴	۴/۸۶	-۰/۵۰	۳۳	۶/۲۳	۶/۴۵	-۳/۵۳
۱۶	۴/۹۹	۵/۱۱	-۲/۴۶	۳۴	۶/۳۹	۶/۲۹	۱/۵۸
۱۷	۴/۶۹	۴/۵۶	۲/۶۳	۳۵	۶/۱۱	۶/۰۱	۱/۵۹
۱۸	۴/۸۱	۵/۰۵	-۴/۸۸	۳۶	۶/۴۴	۶/۳۸	۰/۹۵

ادامه جدول ۷-۲

شماره ترکیب	pIC _{۵۰}		درصد خطا	شماره ترکیب	pIC _{۵۰}		درصد خطا
	مقدار واقعی	مقدار پیش‌بینی شده			مقدار واقعی	مقدار پیش‌بینی شده	
۳۷	۶/۱۳	۶/۲۴	-۱/۷۲	۴۹	۵/۵۴	۵/۵۴	-۰/۰۶
۳۸	۶/۲۶	۶/۳۱	-۰/۸۰	۵۰	۵/۵۱	۵/۵۲	-۰/۱۷
۳۹	۶/۰۶	۵/۹۵	۱/۷۷	۵۱	۵/۰۹	۵/۰۷	۰/۳۴
۴۰	۵/۳۶	۵/۲۰	۳/۱۰	۵۲	۵/۵۳	۵/۵۷	-۰/۸۰
۴۱	۵/۴۹	۵/۴۴	۰/۹۵	۵۳	۵/۵۱	۵/۴۸	۰/۵۱
۴۲	۵/۵۲	۵/۵۲	۰/۰۰	۵۴	۵/۵۴	۵/۶۴	-۱/۸۷
۴۳	۵/۵۷	۵/۷۰	-۲/۳۵	۵۵	۵/۰۳	۵/۰۵	-۰/۳۹
۴۴	۵/۵۲	۵/۶۶	-۲/۴۷	۵۶	۵/۸۹	۶/۱۸	-۴/۹۱
۴۵	۵/۵۱	۵/۶۲	-۲/۱۴	۶۱	۵/۶۶	۵/۲۳	۲/۲۸
۴۶	۵/۱۵	۵/۰۷	۱/۴۶	۶۲	۶/۰۳	۶/۲۰	-۲/۸۰
۴۷	۵/۲۲	۵/۰۴	۳/۴۳	۶۳	۶/۸۰	۶/۲۵	۸/۰۷
۴۸	۵/۲۱	۵/۲۶	-۱/۰۱				



شکل ۲-۷ نمودار تغییرات مقادیر پیش بینی شده همه داده ها بر اساس تکنیک LOO در مقابل مقادیر تجربی



شکل ۲-۸ نمودار باقی مانده های حاصل از پیش بینی فعالیت دارویی ترکیبات با استفاده از تکنیک LOO و مقادیر تجربی برحسب مقادیر تجربی

۲-۲-۷-۳ ارزیابی مدل Ridge-SCAD و مدل SCAD با استفاده از پارامترهای آماری

به منظور بررسی بیش تر کارایی مدل Ridge-SCAD پارامترهای آماری معرفی شده در بخش (۱-۴) محاسبه شد. کارایی مدل توسعه یافته Ridge-SCAD با محاسبه این پارامترهای آماری برای مجموعه آزمون و آموزش مورد بررسی قرار گرفت و نتایج حاصل از محاسبه پارامترهای آماری به همراه محدوده یا مقدار قابل قبول هر پارامتر در جدول ۸-۲ خلاصه شده اند. نتایج جدول ۸-۲ نشان می دهد که پارامترهای $(R_0^2, R_m^2, R_m'^2)$ از مقدار محدوده قابل قبول ۰/۵ بزرگ تر و به مقدار R^2 نزدیک هستند. شیب نمودار حاصل از مقادیر پیش بینی شده بر حسب مقادیر تجربی (و بالعکس) در عرض از مبدأ صفر نیز در محدوده ۰/۸۵ تا ۱/۱۵ قرار دارند که این نتیجه نیز حاکی از صحت مدل توسعه یافته Ridge-SCAD می باشد. از آن جا که R^2 (ضریب تعیین)، تعداد توصیف کننده های موجود در مدل را به حساب نمی آورد، در ارزیابی و مقایسه مدل های مختلف با تعداد متفاوت متغیرهای مستقل (توصیف کننده)، پارامترهای آماری ضریب تعیین تصحیح شده (R_{adj}^2) و FIT استفاده شد.

جدول ۲-۸ پارامترهای آماری برای مجموعه آموزش و آزمون داده‌های پیش بینی شده مدل

Ridge-SCAD				
ردیف	پارامتر	مجموعه آزمون	مجموعه آموزش	محدوده قابل قبول
۱	PRESS	۰/۱۵	۱/۱	-
۲	SEP	۰/۲۰	۰/۱۴	-
۳	MAE	۰/۱۷	۰/۱۱	-
۴	REP	۳/۸	۲/۶	-
۵	MSE	۰/۰۴	۰/۰۲	-
۶	MRE	۳/۱	۱/۱	-
۷	R^2	۰/۹۰	۰/۹۳	$> ۰/۶$
۸	R_{adj}^2	-	۰/۹۲	-
۹	FIT	-	۴/۳۸	-
۱۰	R_0^2	۰/۸۵	۰/۹۳	نزدیک به R^2
۱۱	R نسبی	۰/۰۶	۰/۰۰	$< ۰/۱$
۱۲	R_m^2	۰/۷۰	۰/۹۳	$> ۰/۵$
۱۳	$R_0'^2$	۰/۸۹	۰/۹۳	نزدیک به R^2
۱۴	R' نسبی	۰/۰۱	۰/۰۰	$< ۰/۱$
۱۵	$R_m'^2$	۰/۸۰	۰/۹۳	$> ۰/۵$
۱۶	R-R	۰/۰۴	۰/۰۰	$< ۰/۲$
۱۷	K	۰/۹۹	۱/۰۰	$۰/۸۵ < K < ۱/۱۵$
۱۸	K'	۱/۰۱	۱/۰۰	$۰/۸۵ < K < ۱/۱۵$

برای بررسی کارایی Ridge به عنوان روش غربالگری توصیف کننده قبل از مدل سازی با SCAD،

مدل SCAD ساخته شده با تمام توصیف کننده‌ها بدون غربالگری Ridge با محاسبه برخی از پارامترهای

آماري برای مجموعه آزمون و آموزش مورد ارزیابی و مقایسه با مدل Ridge-SCAD قرار گرفت و نتایج

حاصل در جدول ۲-۹ خلاصه شده‌اند.

جدول ۹-۲ مقایسه پارامترهای آماری محاسبه شده برای مجموعه آزمون داده های پیش بینی شده مدل Ridge-SCAD و SCAD

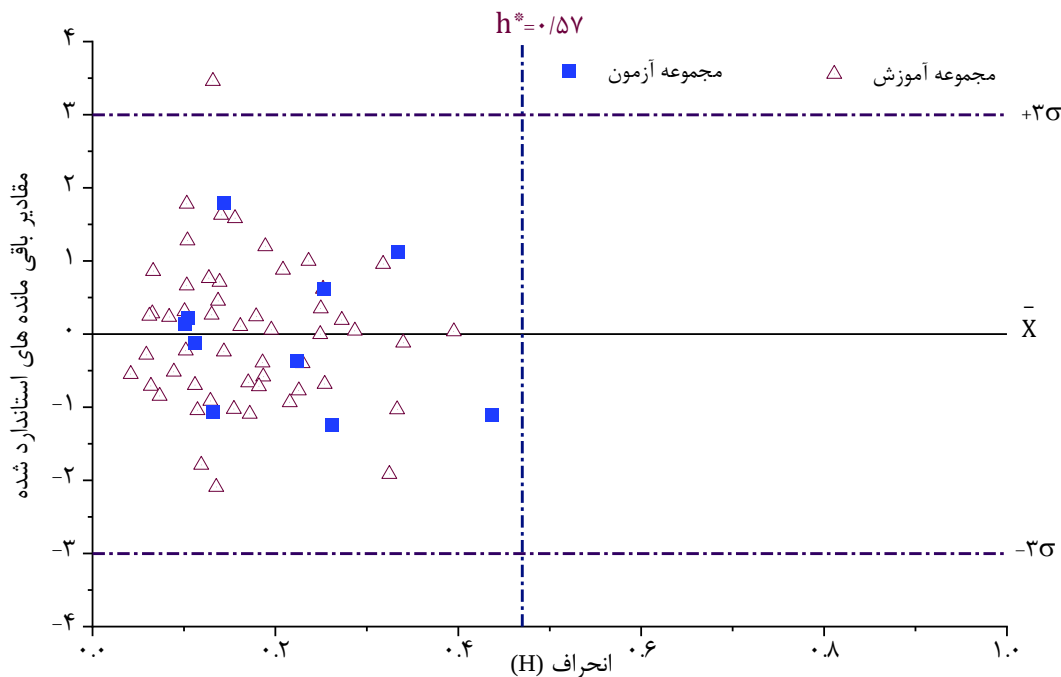
ردیف	پارامتر	Ridge-SCAD		SCAD		محدوده قابل قبول
		مجموعه آزمون	مجموعه آموزش	مجموعه آزمون	مجموعه آموزش	
۱	R^2	۰/۹۰	۰/۹۳	۰/۸	۰/۹۵	>۰/۶
۲	R_{adj}^2	-	۰/۹۲	-	۰/۹۱	-
۳	FIT	-	۴/۳۸	-	۳/۷۸	-

مطابق با نتایج جدول ۹-۲ مقدار R_{adj}^2 مدل Ridge-SCAD و مدل SCAD بزرگتر از مقدار قابل قبول ۰/۵ به دست آمد. مقدار R_{adj}^2 مدل Ridge-SCAD با تعداد توصیف کننده های کم تر نسبت به مدل SCAD با توصیف کننده های بیش تر، به مقدار R^2 خود نزدیک تر است. پارامتر FIT برای هر دو مدل Ridge-SCAD و SCAD محاسبه شد، نتایج نشان می دهد که FIT مدل Ridge-SCAD بیش تر از مدل SCAD می باشد، بنابراین مدل Ridge-SCAD با تعداد توصیف کننده کم تر (تنکی بیش تر) دارای قدرت پیش بینی بیش تری نسبت به مدل SCAD است. مقایسه نتایج پارامترهای آماری در مدل Ridge-SCAD و SCAD (جدول ۹-۲) نشان می دهد که مدل Ridge-SCAD ساده تر، تفسیرپذیرتر (توسعه یافته با توصیف کننده های کم تری) است، در حالی که نتایج آماری بهتری در پیش بینی فعالیت دارویی ترکیبات نسبت به مدل SCAD دارد و قدرت پیش بینی و تعمیم پذیری مدل Ridge-SCAD از مدل SCAD بیش تر است.

۲-۲-۴ ارزیابی مدل با استفاده از دامنه کاربرد

از آزمون دامنه کاربرد به عنوان یک معیار مناسب برای بررسی اعتمادپذیری و استحکام مدل توسعه یافته Ridge-SCAD استفاده شد. دامنه کاربرد یک مدل QSAR با استفاده از نمودار ویلیام نشان داده می شود. به منظور استخراج نمودار ویلیام، مقادیر انحراف (H) مجموعه داده های آموزش و آزمون (مطابق با رابطه ۴۰-۱) و توضیحات (بخش ۱-۴-۵) به دست آمد. مقادیر باقی مانده های استاندارد شده نیز با استفاده

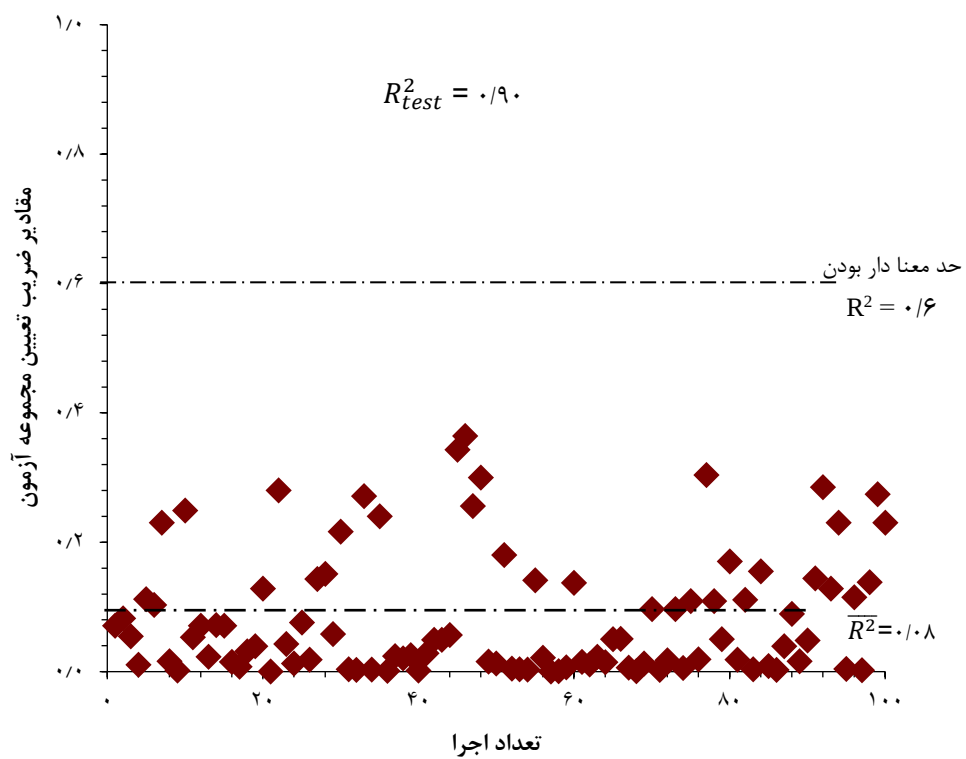
از رابطه ۱-۴۱ محاسبه شد. از رسم مقادیر باقی مانده های استاندارد شده بر حسب مقادیر انحراف (H) نمودار ویلیام به دست آمد. برای تفسیر و تحلیل نمودار دامنه کاربرد رعایت دو شرط الزامی است. ابتدا این که مقادیر باقی مانده های استاندارد شده (محور عمودی) باید در محدوده ۳ برابر بزرگ تر یا کوچک تر از انحراف استاندارد ($\pm 3\sigma$) قرار داشته باشند، تا نشان دهد باقی مانده به صورت کاملاً تصادفی توزیع شده اند. علاوه بر این باید مقادیر انحراف (H) در محدوده ای کوچک تر از h^* که برابر با $3p/n$ است، قرار گیرند. همان طور که شکل ۲-۹ نشان می دهد، همه داده ها دارای H کوچک تر از h^* هستند و تنها یک داده دارای باقی مانده ای است که از محدوده $\pm 3\sigma$ تجاوز کرده است و ۹۹٪ کل داده های مورد مطالعه در محدوده های فوق قرار دارند و داده پرت نمی باشند.



شکل ۲-۹ نمودار ویلیام (دامنه کاربرد) مدل Ridge-SCAD، خطوط نقطه چین افقی و عمودی در دو انتها نمودار به ترتیب نمایانگر مقادیر $\pm 3\sigma$ و h^* می باشد.

۲-۷-۵ ارزیابی مدل Ridge-SCAD با استفاده از Y-تصادفی

برای حصول اطمینان از تصادفی نبودن ارتباط ایجاد شده بین توصیف کننده‌های نهایی و فعالیت دارویی به وسیله مدل Ridge-SCAD، از آزمون Y-تصادفی استفاده شد. برای انجام این آزمون، مقادیر فعالیت دارویی ترکیبات مورد مطالعه در محدوده ۴/۴۹ تا ۶/۸۰ به تعداد ۱۰۰ بار به طور کاملاً تصادفی تغییر داده شدند. مدل Ridge-SCAD با استفاده از پاسخ‌های تصادفی فعالیت دارویی ترکیبات آموزش ساخته شد و مدل توسعه یافته با استفاده از پاسخ‌های تصادفی شده برای پیش‌بینی فعالیت دارویی ترکیبات مجموعه آزمون مورد استفاده قرار گرفت و مقادیر R^2 برای نتایج حاصل محاسبه گردید. نتایج R^2 حاصل از پیش‌بینی pIC_{50} ترکیبات مجموعه آزمون با ۱۰۰ مدل توسعه یافته با متغیر وابسته تصادفی، در شکل ۲-۱۰ آورده شده است. همان‌طور که نتایج شکل ۲-۱۰ نشان می‌دهد، مقادیر R^2 حاصل از پیش‌بینی فعالیت دارویی ترکیبات توسط مدل توسعه یافته تصادفی مجموعه آزمون کوچک‌تر از مقدار قابل قبول ۰/۶ است و میانگین ۱۰۰ تکرار نیز برابر با ۰/۰۸ است که در نمودار مشخص شده است. بنابراین مقادیر بسیار کم R^2 ، عدم تصادفی بودن ارتباط ایجاد شده بین توصیف کننده‌های مدل و فعالیت دارویی را اثبات می‌کند و ارتباط ایجاد شده توسط مدل Ridge-SCAD کاملاً منطقی و معنا دار است.



شکل ۲-۱۰ نمودار مقادیر R^2 به دست آمده در آزمون Y-تصادفی بر حسب تعداد اجرا، برای ۱۰۰ بار اجرای آزمون Y-تصادفی و پیش بینی فعالیت ترکیبات مجموعه آزمون به وسیله مدل Ridge-SCAD توسعه یافته با استفاده از پاسخ تصادفی شده

فصل سوم

نتیجه‌گیری و آینده‌نگری

۳-۱ تجزیه و تحلیل توصیف کننده های مدل Ridge-SCAD و تأثیر آن‌ها

بر فعالیت دارویی با استفاده از ضرایب رگرسیون استاندارد شده

قابلیت تفسیر پذیری مدل QSAR توسعه یافته یکی از مهم ترین معیارها در مطالعات QSAR است. در این بخش از مطالعه، اثر توصیف کننده های منتخب مدل توسعه یافته Ridge-SCAD بر فعالیت دارویی، با توجه به تأثیر ضرایب رگرسیونی استاندارد شده توصیف کننده های مدل Ridge-SCAD و نمودار سهم مشارکت توصیف کننده ها مورد بررسی قرار گرفت. بنابراین ضرایب استاندارد شده ۹ توصیف کننده مدل Ridge-SCAD با استفاده از ضرایب غیر استاندارد و رابطه ۱-۴۲ محاسبه شد و در جدول ۳-۱ آورده شده است. علامت مثبت ضرایب توصیف کننده، نشان دهنده اثر افزایشی توصیف کننده بر فعالیت دارویی می باشد، در حالی که ضرایب با علامت منفی نشان دهنده اثر کاهشی توصیف کننده بر فعالیت دارویی می باشد. در ادامه رابطه بین این توصیف کننده ها و فعالیت دارویی ترکیبات بررسی خواهد شد.

جدول ۳-۱ اثر متوسط توصیف کننده های انتخاب شده توسط روش Ridge-SCAD

توصیف کننده ها	DECC	MATS4e	C-025	Mor26p	H-046	N-075	SEigZ	Mor23m	GATS4e
ضریب تأثیر	+۰/۴۵۵۵	+۰/۰۲۴۰	-۰/۳۹۸۰	-۰/۲۰۲۴	-۰/۱۷۱۰	-۰/۰۱۱۸	-۰/۰۰۱۳	+۰/۱۷۵۶	-۰/۲۶۴۳

۳-۱-۱ بررسی میزان و چگونگی تأثیر توصیف کننده ها بر فعالیت دارویی مجموعه

بازدارنده های سرطان و کاربرد مدل ارائه شده در پیشنهاد ترکیبات جدید

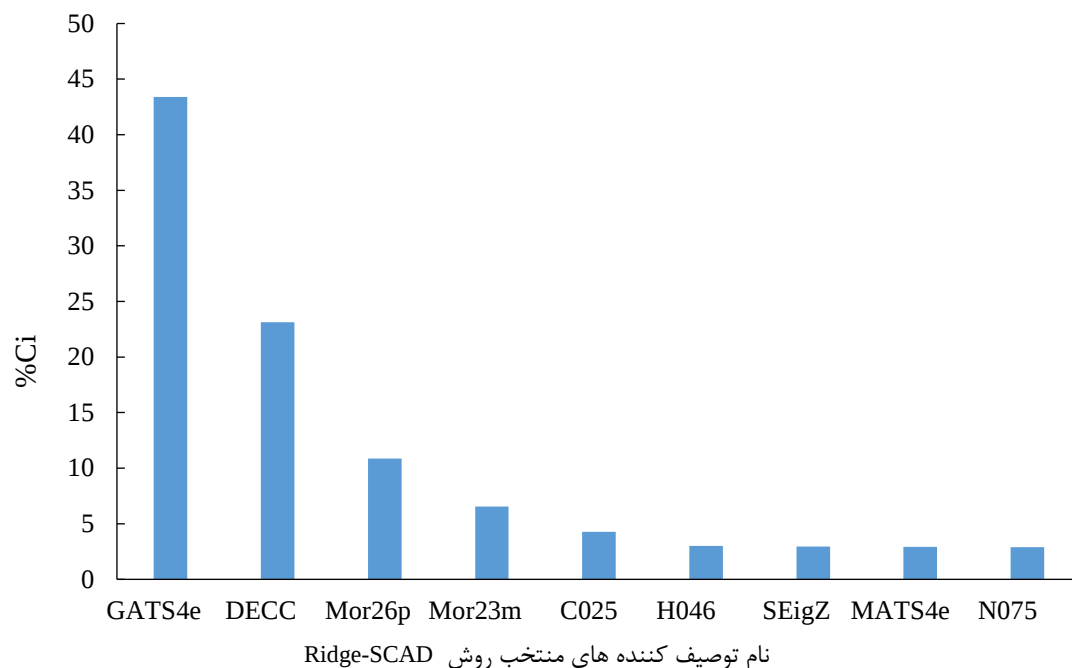
۳-۱-۱-۱ محاسبه سهم مشارکت هر توصیف کننده در مدل Ridge-SCAD

به منظور بررسی اهمیت و میزان مشارکت هر توصیف کننده در مدل Ridge-SCAD درصد سهم مشارکت هر توصیف کننده ($C_i/\%$) برآورد شد. روش کار به این صورت است که برای محاسبه سهم مشارکت

توصیف کننده i ، مقادیر آن توصیف کننده به صفر تغییر یافت. پس از پیش‌بینی فعالیت دارویی ترکیبات مجموعه آموزش، مقدار MSE مجموعه آموزش در حضور توصیف کننده i با مقادیر تصادفی شده، محاسبه شد (MSE_i). این فرآیند برای هر ۹ توصیف کننده تکرار شد و ۹ مقدار (MSE_i) برای مجموعه آموزش به‌دست آمد. سپس، مقادیر $C_i \%$ ($i=1,2,\dots,9$) برای همه توصیف کننده‌ها با استفاده از رابطه ۱-۳ به‌دست آمد.

$$C_i = \left(\frac{MSE_i}{\sum MSE_i} \right) \times 100 \quad \text{رابطه ۱-۳}$$

نتایج به‌دست‌آمده برای سهم مشارکت هر توصیف کننده در مدل Ridge-SCAD در شکل ۱-۳ آورده شده است. نتایج نشان می‌دهد که ۳ توصیف کننده GATS4e، DECC و Mor26p حدوداً ۸۰٪ در توسعه مدل Ridge-SCAD مشارکت دارند. بنابراین با توجه به نمودار سهم مشارکت (شکل ۱-۳) و تأثیر ضرایب این توصیف کننده‌ها در مدل Ridge-SCAD (جدول ۱-۳) چگونگی اثر توصیف کننده‌هایی با بیش‌ترین سهم مشارکت در مدل Ridge-SCAD بر فعالیت دارویی ترکیبات مورد بررسی بیش‌تر قرار گرفت، تا از نتایج آن برای پیشنهاد ترکیبات جدید با فعالیت دارویی مناسب استفاده گردد.



شکل ۱-۳ نمودار سهم مشارکت توصیف کننده ها در مدل Ridge-SCAD

۳-۱-۱-۲ توصیف کننده دو بعدی همبسته^۱

• توصیف کننده های گروه Geary Autocorrelation:

توصیف کننده های دو بعدی خود همبسته از گراف مولکول و از طریق محاسبه مجموع وزن های

اتم های انتهایی کل مسیرها با طول مسیر^۲ مورد نظر، به دست می آیند. نوعی از این توصیف کننده ها گروه

Geary Autocorrelation است که ضریب گری^۳ نام دارد و طبق (رابطه ۲-۳) محاسبه می شوند:

$$c(d) = \frac{\frac{1}{2\Delta} \sum_{i=1}^A \sum_{j=1}^A \delta_{ij} (W_i - W_j)^2}{\frac{1}{A-1} \sum_{i=1}^A (W_i - \bar{W})^2} \quad \text{رابطه ۲-۳}$$

^۱ D autocorrelation

^۲ Lag

^۳ Geary Autocorrelation

$c(d)$ ضریب گری می‌باشد، که W یک ویژگی‌اتم، $\bar{W}$ میانگین مقدار آن ویژگی روی مولکول، A تعداد اتم‌ها و d فاصله مکانی است، δ_{ij} نیز که به تابع کرونکر^۱ معروف است، در حالتی که $d = d_{ij}$ باشد، δ_{ij} برابر یک است، و در غیر این صورت صفر است و Δ هم مجموعه δ هاست. مقدار این توصیف کننده که از جنس فاصله است، از صفر تا بی‌نهایت متغیر است [۶۹]. از میان این توصیف کننده‌ها GATS4e در مرحله انتخاب توصیف کننده‌های مؤثر در مدل Ridge-SCAD پیشنهادی برگزیده شده است. این توصیف کننده براساس الکترونگاتیویته ساندرسون، وزن دار شده است. با توجه به منفی بودن ضریب رگرسیون برای این توصیف کننده (جدول ۳-۱) کاهش مقدار آن سبب افزایش مقدار فعالیت دارویی (pIC_{50}) می‌شود.

۳-۱-۱-۳ توصیف کننده مکانی^۲

توصیف کننده‌های مکانی که از گراف مولکولی به دست می‌آیند، به عنوان توصیف کننده‌های عمومی مطرح هستند، چون به طور مستقیم از ساختار مولکول محاسبه می‌شوند و مستقل از صورت بندی مولکول هستند. DECC جزء این دسته از توصیف کننده‌ها می‌باشند که در مدل Ridge-SCAD ظاهر شده‌اند. DECC به شاخه دار بودن ساختار ترکیب ارتباط دارد و با توجه به علامت ضریب رگرسیونی (جدول ۳-۱)، هرچه انشعابات و استخلاف‌های ترکیب شیمیایی بیش تر باشد، فعالیت دارویی نیز افزایش می‌باشد [۷۰].

۳-۱-۱-۴ توصیف کننده نمایش سه بعدی ساختار مولکول براساس تفرق الکترون^۳

توصیف کننده‌ها ۳D-MORSE (نمایش سه بعدی ساختار مولکول براساس تفرق الکترون) از طریق معادله تبدیلی که در پراش الکترون استفاده می‌شود، محاسبه می‌گردند:

$$I(s) = \sum_{i=2}^N \sum_{j=1}^{i-1} A_i A_j \frac{\sin(sr_{ij})}{sr_{ij}} \quad \text{رابطه ۳-۳}$$

^۱ Kroncker

^۲ Topological

^۳ D- Molecule Representation of Structure based on Electron diffraction

I شدت الکترون پراکنده شده، A_t و A_j خاصیت اتمی اتم‌های i و j، s زاویه پراکندگی، r_{ij} فاصله بین اتم‌های i و j و N تعداد کل اتم‌ها را نشان می‌دهد. این روش باعث می‌شود که ساختار سه بعدی مولکول به یک کد ثابت تبدیل شود [۶۹]. توصیف کننده‌های سه بعدی که در مدل Ridge-SCAD دیده می‌شوند، عبارتند از Mor26p که با استفاده از قطبش پذیری وزن دار شده است، و دارای ضریب تأثیر منفی می‌باشد (جدول ۳-۱). با توجه به این که مقادیر توصیف کننده Mor26p منفی است، بنابراین با منفی تر شدن مقدار این توصیف کننده مقدار فعالیت دارویی (pIC_{50}) ترکیب مورد نظر افزایش می‌یابد.

۳-۲ پیشنهاد ترکیبات جدید دارای فعالیت مناسب ضد سرطان با استفاده

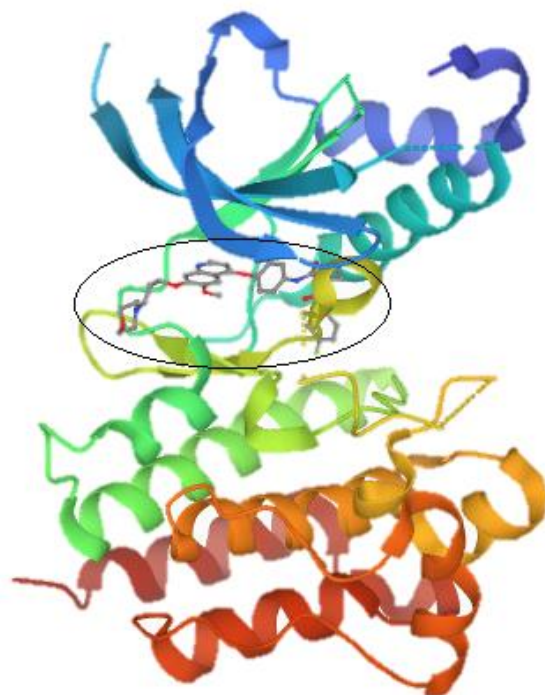
از مدل Ridge-SCAD و مطالعه داکینگ مولکولی ترکیبات پیشنهادی

یکی از اهداف مطالعات QSAR که اخیراً مورد توجه زیادی قرار گرفته است، پیشنهاد ترکیبات جدید و در نتیجه تسهیل فرآیند طراحی دارو است. مدل پیشنهادی نشان می‌دهد که چگونه تغییرات در توصیف کننده‌های ظاهر شده در مدل QSAR می‌تواند فعالیت دارویی را افزایش دهد. از این رو با توجه به اثر هر توصیف کننده در مدل Ridge-SCAD بر فعالیت دارویی مربوطه، می‌توان به چگونگی ایجاد یک اصلاح ساختاری مؤثر که منجر به افزایش فعالیت دارویی شود، پی برد. قابل توجه است که برای مثال، توصیف کننده‌های C-025 و H-046 با ضریب اثر منفی در مدل نشان‌دهنده این است که هرچه تعداد کربن‌های R-CR-R و تعداد هیدروژن‌های متصل به کربن (Sp^3)، که کربن بعدی به هتروآتمی متصل نباشد، کم‌تر باشد، فعالیت دارویی ترکیبات بیش‌تر می‌شود. بنابراین تلاش شد با در نظر گرفتن اثر توصیف کننده‌های مدل Ridge-SCAD بر فعالیت دارویی ترکیبات مورد مطالعه، ترکیبات جدید سنتز نشده با فعالیت دارویی مناسب پیشنهاد شود. به این منظور، مطابق با اثر توصیف کننده‌ها بر فعالیت دارویی، اصلاحات ساختاری روی ترکیب ضعیف (ترکیب شماره ۵) ایجاد شد. علاوه بر این تلاش شد تا با ایجاد تغییرات بر روی ترکیب فعال (ترکیب شماره ۶۳) نیز ترکیباتی با پتانسیل دارویی بیش‌تر پیشنهاد شود.

از این رو برای طراحی ترکیبات جدید با فعالیت دارویی مناسب، از معنای شیمیایی توصیف‌کننده‌های منتخب برای اصلاح ساختاری ترکیبات در راستای ایجاد ترکیبات با فعالیت دارویی مناسب، استفاده شد. ساختار ترکیبات پیشنهادی و مقادیر pIC_{50} پیش‌بینی‌شده مرتبط به آن‌ها در جدول ۳-۲ آورده شده است. ساختار این ترکیبات در هایپرکم رسم و بهینه‌سازی شدند. توصیف‌کننده‌های مدل برتر با استفاده از نرم افزار دراگون محاسبه شد و همچنین فعالیت این ترکیبات با استفاده از مدل برتر Ridge-SCAD پیش‌بینی شد، که فعالیت برخی از آن‌ها بیش‌تر از فعال‌ترین ترکیب مورد مطالعه (ترکیب ۶۳) بوده و فعالیت برخی ترکیبات در محدوده ترکیب فعال به دست آمد و نتایج در جدول ۳-۲ آورده شده است. در ادامه برای اثبات درستی فعالیت پیش‌بینی شده برای ترکیبات فعال پیشنهادی، برهم‌کنش ترکیبات در جایگاه فعال گیرنده آنزیمی با استفاده از نرم‌افزار داکینگ مولکولی مورد بررسی قرار گرفت که در ادامه چگونگی انجام کار آورده شده است.

۳-۲-۱ مطالعه داکینگ مولکولی برای بازدارنده‌های جدید پیشنهادی

به منظور انجام عملیات داکینگ مولکولی باید مراحل متفاوتی را انجام داد که در فصل دوم (بخش‌های ۱-۵ تا ۲-۵) به آن‌ها اشاره شده است. مطابق با مراحل گفته شده در قبل، ابتدا ساختار کریستالوگرافی گیرنده مورد استفاده برای ترکیبات مورد مطالعه به پیشنهاد مقالات شناسایی شد [۷۱]، [۷۲]. با مراجعه به مقالات، گیرنده با کد ۳LQ8 انتخاب و ساختار کریستالوگرافی آن از سایت بانک اطلاعاتی پروتئین با پسوند pdb دانلود شد. ارزش تفکیک مربوط به ساختار کریستالوگرافی گیرنده ۳LQ8 برابر ۲/۰۲ آنگستروم می‌باشد که این مقدار ارزش تفکیک برای انجام مطالعات داکینگ مولکولی مناسب است [۷۳]. ساختار کریستالوگرافی گیرنده ۳LQ8 و لیگاند کریستالوگرافی موجود در زنجیره A در شکل ۳-۲ نشان داده شده است.



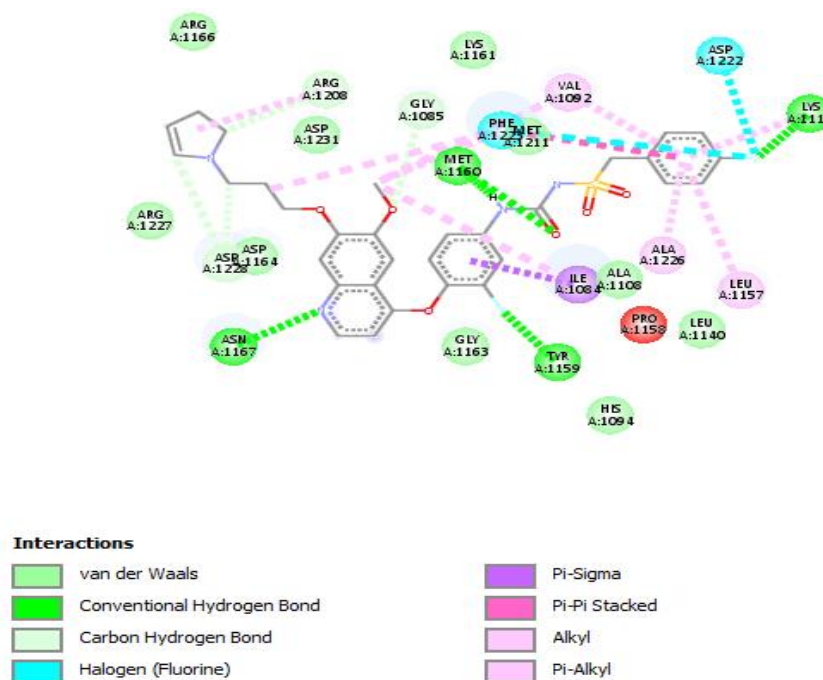
شکل ۳-۲ ساختار کریستالوگرافی ۳LQ8 (منطقه انتخاب شده نشان دهنده لیگاند کریستالوگرافی و مابقی زنجیره های اسید آمینه ای است)

در مرحله بعد ساختار کریستالوگرافی گیرنده ۳LQ8 در نرم افزار ویورلایت فراخوانی شد و عملیات آماده سازی متفاوتی از جمله حذف مولکول های آب، حذف کوفاکتورها و حذف زنجیره اسید آمینه ای فاقد لیگاند کریستالوگرافی انجام شد و باقی مانده ساختار که شامل زنجیره اسید آمینه ای A و لیگاند کریستالوگرافی است، با پسوند pdb ذخیره شد. سپس زنجیره اسید آمینه ای به عنوان فایل گیرنده و ترکیب موجود در ساختار کریستالوگرافی به عنوان لیگاند به طور مجزا با پسوند pdb ذخیره شدند و به عنوان ورودی های نرم افزار داکینگ برای انجام مرحله اعتبار سنجی داکینگ مورد استفاده قرار گرفتند. در مرحله بعد، مختصات جایگاه فعال گیرنده با توجه به مختصات مرکز ثقل لیگاند کریستالوگرافی و با استفاده از نرم افزار ویورلایت با مختصات X، Y و Z برابر با ۰/۸۲، ۴/۸۶ و ۳۰/۶۱۲ به دست آمد.

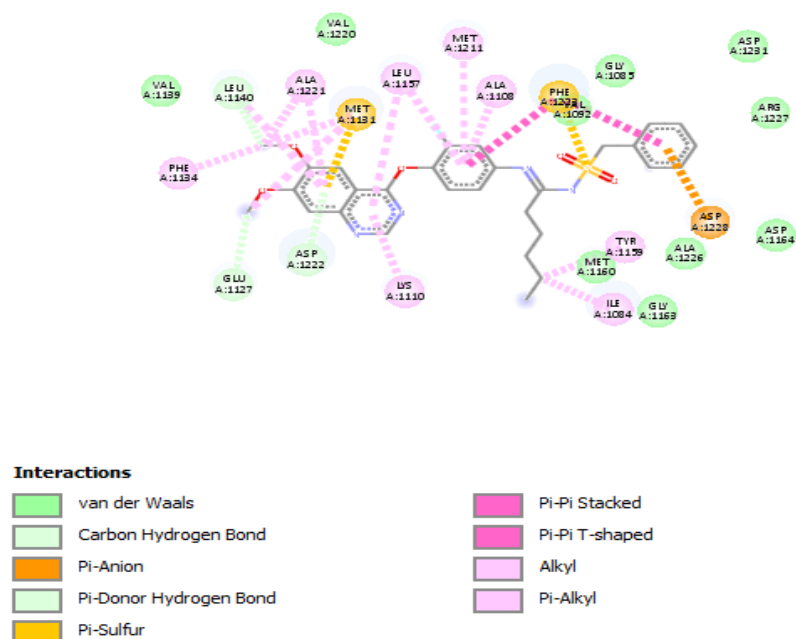
پس از آماده‌سازی فایل‌های ورودی (زنجیره اسید آمینه حاوی ساختار لیگاند کریستالوگرافی به‌عنوان گیرنده و لیگاند کریستالوگرافی به‌عنوان لیگاند)، فرآیند اعتبار سنجی داکینگ مولکولی با استفاده از نرم‌افزار اتوداک ۴/۲ انجام شد. به این منظور ابتدا فایل گیرنده در محیط نرم‌افزار وارد شد. آماده‌سازی‌های بیش‌تر روی ساختار گیرنده از جمله افزودن هیدروژن برای جبران کمبود هیدروژن ساختار کریستالوگرافی، ادغام هیدروژن‌های غیر قطبی متصل به هر اتم کربن و افزایش بار الکتریکی کلمن^۱ انجام شد [۳۹]. در مرحله بعد به‌منظور انجام فرآیند اعتبار سنجی، ساختار لیگاند کریستالوگرافی به‌عنوان لیگاند در نرم‌افزار اتوداک ۴/۲ فراخوانی شد و با پسوند pdbqt ذخیره شد. در مرحله بعد خروجی pdbqt ساختار گیرنده نیز ذخیره شد و مختصات یک شبکه مشتمل بر جایگاه فعال گیرنده، با پیشنهاد مقالات منتشر شده در حوزه مطالعات داکینگ مولکولی، با ابعاد ۶۰×۶۰×۶۰ آنگستروم و فاصله بین نقاط ۰/۳۷۵ آنگستروم (در حدود یک چهارم طول پیوند یگانه کربن-کربن) ایجاد شد و بر هم‌کنش‌های گیرنده-لیگاند در این شبکه مورد بررسی قرار گرفت [۴۰، ۷۴، ۷۵]. داکینگ مولکولی با استفاده از اجرای الگوریتم ژنتیک لامارکین (LGA)، در تعداد اجراهای متفاوت ۱۰۰، ۱۵۰ و ۲۰۰ انجام شد. لازم به ذکر است که سایر پارامترهای الگوریتم ژنتیک در حالات پیش‌فرض پیشنهادی نرم‌افزار در نظر گرفته شدند. عملیات اعتبارسنجی در شرایط ذکر شده، انجام شد و فایل‌های خروجی با پسوند dlga استخراج و ذخیره شدند. با توجه به نتایج داکینگ مولکولی در مرحله اعتبار سنجی در ۳ حالت با تعداد اجراهای متفاوت الگوریتم LG، تعداد اجرای الگوریتم LG برابر با ۱۰۰، کم‌ترین مقدار ریشه میانگین مربعات انحراف (RMSD) و بیش‌ترین تعداد کنفورماسیون در خوشه اول را دارا بوده و در نتیجه، برای انجام فرآیند داکینگ مولکولی ترکیبات مورد نظر از تعداد اجرای الگوریتم LG برابر با ۱۰۰ استفاده شد. در فرآیند انجام داکینگ مولکولی ترکیبات مورد مطالعه، همانند فرآیند اعتبار سنجی، تمامی شرایط (مختصات جایگاه فعال، ابعاد شبکه و پارامترهای پیش‌فرض الگوریتم ژنتیک) در

¹ Kollman

شرایط داکینگ مولکولی ترکیبات جدید نیز تعریف شد. بنابراین داکینگ مولکولی در شرایط بهینه برای ترکیب ۶۳ با بیشترین فعالیت ($pIC_{50} = 6/80$)، ترکیب ۵ با کمترین فعالیت ($pIC_{50} = 4/52$) و تمامی ترکیبات پیشنهادی در جایگاه فعال گیرنده انجام شدند. بهترین پیکربندی‌ها از ترکیبات مورد مطالعه و پیشنهادی با توجه به حداقل انرژی آزاد اتصال و بیشترین تعداد پیکربندی در خوشه اول، شناسایی شدند و برهم‌کنش آن‌ها با اسید آمینه‌های کلیدی با استفاده از نرم‌افزار Discovery Studio مورد بررسی قرار گرفت. شکل‌های دوبعدی برهم‌کنش‌های ترکیب فعال ۶۳ (شکل ۳-۳) و ترکیب ضعیف ۵ (شکل ۴-۳) آورده شد.



شکل ۳-۳ برهم‌کنش ترکیب ۶۳ با اسید آمینه‌های کلیدی



شکل ۳-۴ بر همکنش ترکیب ۵ با اسید آمینه های کلیدی

همان طور که شکل ها و نتایج جدول ۳-۲ نشان می دهد، ترکیب فعال ۶۳ برهم کنش های هیدروژنی

O-H Met1160 (2.49Å) H-O Met1160 (2.05Å), H-F Tyr1159 (2.58Å), H-N Asn1167 (2.29Å)

و H-F Lys1110 (1.96Å) (موارد داخل پرانتز طول پیوند اسید آمینه کلیدی با یک اتم در ساختار مورد

مطالعه را نشان می دهد) را با اسید آمینه داشته و ترکیب ضعیف هیچ پیوند هیدروژنی با اسید آمینه برقرار

نکرده و همه پیوندها از نوع هیدروفوبی بودند. بنابراین ترکیب فعال (ترکیب شماره ۶۳) پیوندهای هیدروژنی

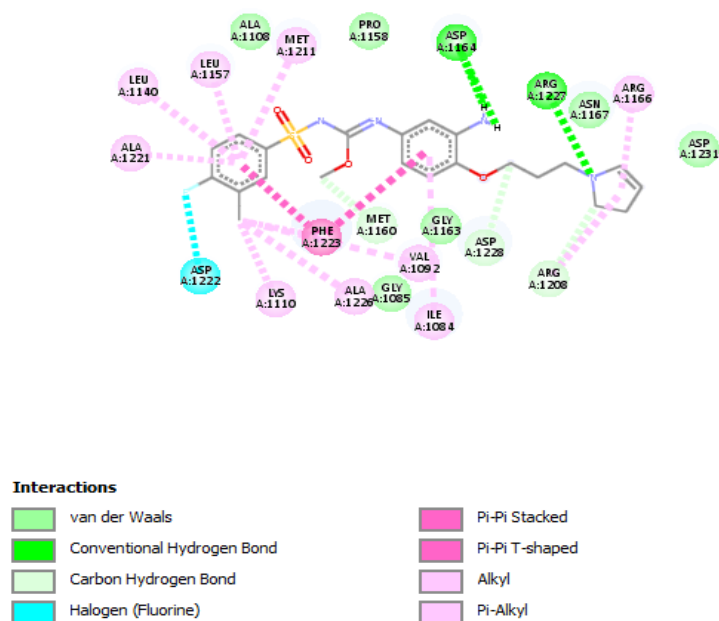
مناسبی را با اسید آمینه ها برقرار کرده است، که به طور معنا داری از نظر بر هم کنش متفاوت از ترکیب

ضعیف تر (ترکیبات شماره ۵) است. این به آن معنا است که بر هم کنش های حاصل از داکینگ مولکولی

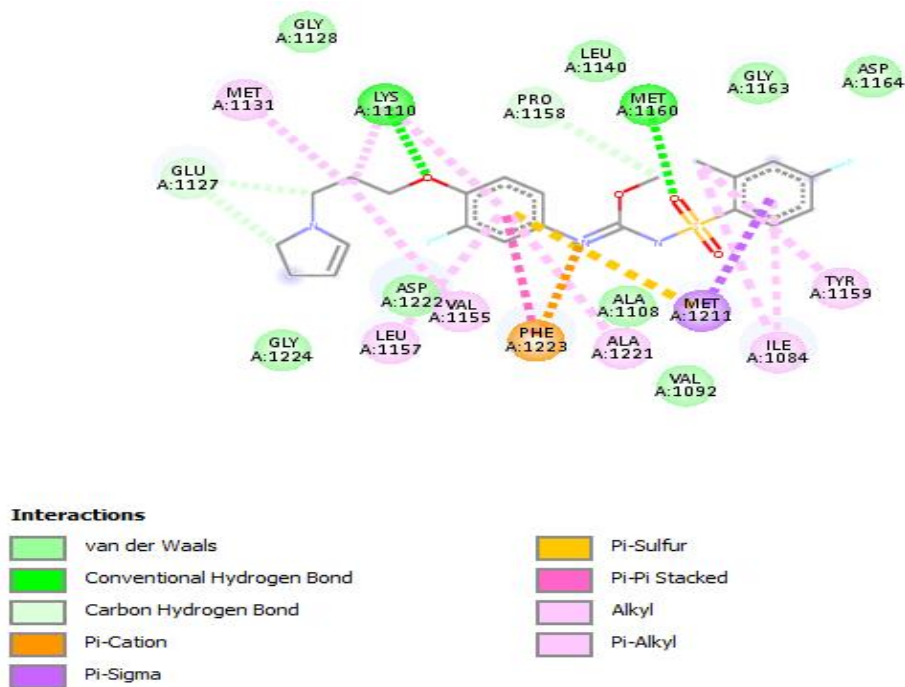
معیار مناسبی برای ارزیابی میزان فعالیت دارویی ترکیبات مورد مطالعه هستند. بنابراین برهم کنش های

لیگاند-گیرنده برای تمام ترکیبات پیشنهادی استخراج و مورد تجزیه و تحلیل قرار گرفت و نتایج برخی از

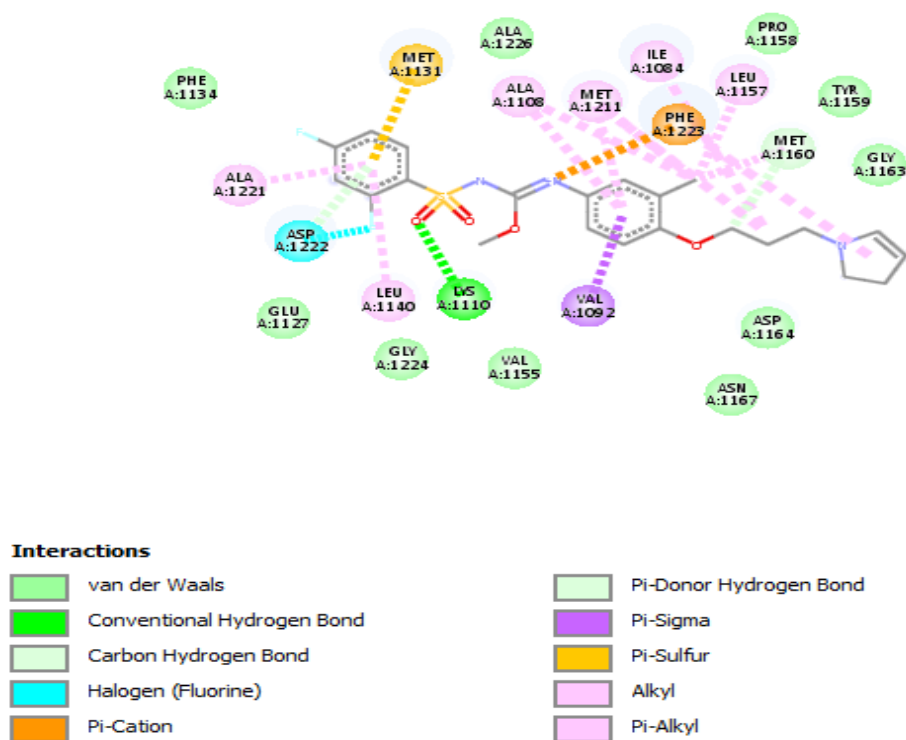
ترکیبات پیشنهادی (۱، ۲ و ۳) در شکل ۳-۵ تا شکل ۳-۷ آورده شده است.



شکل ۳-۵ بر هم کنش ترکیب جدید ۱ با اسید آمینه های کلیدی



شکل ۳-۶ بر هم کنش ترکیب جدید ۲ با اسید آمینه های کلیدی



شکل ۷-۳ بر همکنش ترکیب جدید ۳ با اسید آمینه های کلیدی

جدول ۲-۳ پارامترهای محاسبه شده برای ترکیبات مورد مطالعه و ترکیبات پیشنهادی

شماره ترکیب	ساختار شیمیایی	پیوند هیدروژنی	طول پیوند هیدروژنی (Å)	پیوند هیدروفوبی	pIC ₅₀
۶۳		H Asn1167- N H Tyr1159- F O Met1160- H H Met1160- O H Lys1110- F	۲/۲۹ ۲/۵۸ ۲/۴۹ ۲/۰۵ ۱/۹۶	Ala1108 Asp1164 Asp1231 Arg1166 Arg1227 Gly1163 Lys1161 Leu1140 Met1211 Phe1223	۶/۷۹۶
۵		-	-	Ala1226 Asp1231 Asp1164 Arg1227 Gly1163 Gly1085 Met1160 Val1092 Val1139 Val1220 Phe1223	۴/۵۱۵
NC۱		O Asp1164- H O Asp1164- H H Arg1227- N	۲/۱۴ ۲/۵۱ ۳/۰۳	Ala1108 Asn1167 Asp1231 Gly1163 Gly1085 Pro1158 Phe1223	۶/۹۴۹
NC۲		H Lys1110- O H Met1160- O	۲/۳۵ ۲/۰۴	Ala1108 Asp1164 Asp1222 Gly1163 Gly1224 Gly1128 Leu1140 Val1092 Phe1223	۶/۹۳۶
NC۳		H Lys1110- O	۱/۶۳	Ala1226 Asn1167 Asp1164 Gly1163 Gly1224 Glu1127 Pro1158 Phe1134 Tyr1159 Val1155	۶/۹۱۶
NC۴		H Lys1110- O	۱/۷۸	Ala1226 Gly1163 His1094 Pro1158 Phe1134 Val1155 Val1220	۶/۷۵۴

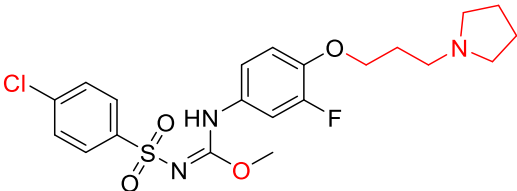
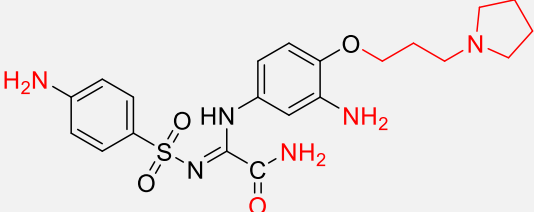
ادامه جدول ۲-۳

شماره ترکیب	ساختار شیمیایی	پیوند هیدروژنی	طول پیوند هیدروژنی (Å)	پیوند هیدروفوبی	pIC ₅₀
NC۵		O Asp1222- H H Lys1110- O H Lys1110- O O Phe1223- H	۲/۱۰ ۱/۹۶ ۲/۸۵ ۲/۲۳	Ala1226 Gly1163 Gly1224 His1202 Ile130 Leu1195 Leu1140 Met1211 Met1160 Phe1200 Tyr1159 Val1139 Val1155 Phe1134	۶/۷۱۵
NC۶		H Lys1110- O	۱/۸۹	Ala1226 Gly1163 Gly1224 Glu1127 Lys1161 Leu1140 Met1160 Pro1158 Phe1134 Val1155 Phe1223	۶/۷۰۴
NC۷		O Asp1224- H H Lys1110- O S Met1131- H	۲/۳۵ ۲/۹۹ ۲/۸۳	Ala1226 Asn1209 Gly1163 Gly1224 Gly1128 Pro1158 Tyr1159 Val1139 Val1092 Val1155 Phe1223	۶/۵۹۸
NC۸		H Asp1164- O H Met1160- O	۲/۳۳ ۲/۳۵	Ala1226 Asn1167 Asp1222 Asp1228 Arg1227 Gly1163 Gly1085 Lys1110 Leu1140 Pro1158 Phe1223	۶/۵۹۵
NC۹		O Lys1161- H O Met1160- H O Pro1158- H	۲/۳۴ ۲/۱۵ ۲/۱۴	Ala1108 Asp1222 Gly1163 His1162 Lys1110 Phe1200 Val1092 Val1220 Tyr1159	۶/۵۸۹

ادامه جدول ۲-۳

شماره ترکیب	ساختار شیمیایی	پیوند هیدروژنی	طول پیوند هیدروژنی (Å)	پیوند هیدروفوبی	pIC ₅₀
NC۱۰		H Lys1110- O	۱/۶۵	Ala1226 Gly1163 Gly1224 Glu1127 Leu1140 Met1160 Val1155	۶/۵۵۱
NC۱۱		H Asp1222- O	۱/۹۹	Asp1164 Gly1163 Met1160 Pro1158 Tyr1159 Val1092 Val1109	۶/۵۴۳
NC۱۲		H Lys1110- O H Met1160- O	۱/۷۹ ۲/۲۶	Gly1163 Lys1161 Pro1158 Val1092 Val1155	۶/۴۷۹
NC۱۳		O Asp1228- H H Asp1164- O H Tyr1159- O O Met1160- H O Pro1158- H O Pro1158- H	۲/۰۸ ۱/۸۱ ۲/۴۲ ۲/۲۵ ۲/۲۲ ۲/۲۲	Ala1221 Asp1222 Asn1167 Arg1086 Gly1163 Gly1085 His1094 Val1092 Tyr1159	۶/۴۷۷
NC۱۴		O Phe1223- H H Lys1110- O S Met1131- H	۲/۰۲ ۲/۶۲ ۲/۷۲	Ala1221 Ala1226 Asp1222 Gly1163 Gly1224 Val1155 Val1092 Phe1223	۶/۴۴۶
NC۱۵		O Asp1222- N H Asp1222- O O Glu1127- H	۲/۷۱ ۲/۰۰ ۲/۰۸	ALA1221 ASN1209 GLY1163 GLY1228 LEU1140 LYS1161 PHE1124	۶/۳۸۳

ادامه جدول ۲-۳

شماره ترکیب	ساختار شیمیایی	پیوند هیدروژنی	طول پیوند هیدروژنی (Å)	پیوند هیدروفوبی	pIC ₅₀
NC۱۶		H Lys1110- O H Met1160- O	۱/۹۷ ۲/۴۰	Gly1163 Glu1127 His1094 Leu1140 Pro1158 Val1155	۶/۳۶۲
NC۱۷		H Asp1222-N O Pro1158- H H Met1160- O	۲/۲۲ ۲/۱۴ ۱/۸۴	Ala1108 Asn1167 Asn1209 Asp1164 Gly1163 Lys1110 Tyr1159 Val1220 Val1139 Val1092 Phe1223	۶/۳۳۵

NC نماد ترکیبات جدید پیشنهادی است.

نتایج به وضوح نشان می‌دهد که ترکیبات پیشنهادی پیوندهای هیدروژنی مناسبی را با اسید آمینه‌های کلیدی برقرار نموده‌اند و درستی فعالیت دارویی پیش‌بینی شده توسط مدل Ridge-SCAD با مطالعات داکینگ و نتایج برهم‌کنش لیگاند-گیرنده همخوانی دارد. علاوه بر مطالعه داکینگ مولکولی، پارامترهای قاعده لیپینسکی نیز برای ترکیبات پیشنهادی محاسبه شدند و نتایج در جدول ۳-۳ خلاصه شده‌اند. نتایج نشان می‌دهد که تمامی پارامترهای مورد نظر از جمله وزن مولکولی ($MW < 500$), چربی‌دوستی ($MLOGP < 4/15$), تعداد گیرنده‌های پیوند هیدروژنی (کمتر از ۱۰), تعداد دهنده‌های پیوند هیدروژنی (کمتر از ۵) و تعداد پیوندهای قابل چرخش (کمتر از ۱۰) برای ترکیبات پیشنهادی دارای مقادیر قابل قبولی هستند [۷۶]. بنابراین ترکیبات پیشنهادی معیارهای شباهت به ترکیبات دارویی را داشته‌اند و

به عنوان ترکیبات بالقوه فعالی که شباهت به ترکیب رهبر^۱ دارند، می توانند به عنوان کاندیدهای فعال مناسبی برای سنتز و آزمایش های دارویی مورد توجه قرار گیرند.

علاوه بر این فاکتورها، پارامتر سهولت سنتز^۲ نیز با استفاده از سایت Swiss-ADME [۷۶] استخراج شد و در جدول ۳-۳ آورده شد. فاکتور سهولت سنتز می تواند دارای مقداری در محدوده ۱ (سنتز بسیار آسان) تا ۱۰ (ترکیب پیچیده و چالش برانگیز) باشد [۷۷]. فاکتور سهولت سنتز برای ترکیبات پیشنهادی دارای مقدار قابل قبولی می باشد و در نتیجه سنتز آنها در محیط آزمایشگاهی امکان پذیر خواهد بود.

^۱ Lead compound

^۲ Synthetic Accessibility

جدول ۳-۳ پارامترهای محاسبه شده قاعده ۵ لیپینسکی

سهولت سنتز	تعداد دهنده‌های پیوند هیدروژنی	تعداد گیرنده‌های پیوند هیدروژنی	MlogP	وزن مولکولی g/mol	پیوندهای قابل چرخش	شماره ترکیبات
۴/۲۸	۲	۱۰	۲/۵۵	۶۲۶/۷	۱۴	۶۳
۴/۲۹	۱	۹	۴/۰۶	۵۸۵/۷	۱۳	۵
۳/۹۵	۲	۷	۲/۸۷	۴۸۴/۱	۱۰	NC۱
۳/۷۲	۱	۸	۳/۷۶	۴۸۷/۱	۱۰	NC۲
۳/۷۹	۱	۸	۳/۷۶	۴۸۷/۱	۱۰	NC۳
۳/۷۶	۱	۸	۳/۷۶	۴۸۷/۱	۱۰	NC۴
۳/۸۰	۳	۸	-۰/۶۹	۴۶۴/۵	۱۰	NC۵
۳/۷۰	۱	۸	۳/۲۸	۴۵۳/۵	۱۰	NC۶
۳/۸۳	۳	۸	-۰/۲	۴۹۸/۱	۱۰	NC۷
۳/۷۵	۱	۸	۳/۷۶	۴۸۷/۱	۱۰	NC۸
۳/۷۱	۳	۷	-۰/۱۷	۴۸۰/۱	۱۰	NC۹
۳/۵۲	۲	۶	۲/۹۸	۴۶۸/۱	۱۰	NC۱۰
۳/۸۰	۲	۶	۲/۹	۴۶۶/۱	۱۰	NC۱۱
۳/۳۹	۲	۶	۲/۷۶	۴۵۴/۱	۹	NC۱۲
۳/۸۷	۴	۷	-۱/۱۶	۴۶۱/۵	۱۰	NC۱۳
۳/۵۸	۲	۸	۰/۷۲	۴۸۳/۱	۱۰	NC۱۴
۳/۷۶	۲	۷	۲/۷۹	۴۵۰/۵	۱۰	NC۱۵
۳/۶۸	۱	۷	۳/۳۹	۴۶۹/۱	۱۰	NC۱۶
۳/۸۸	۴	۶	۰/۶۵	۴۶۰/۶	۱۰	NC۱۷

NC نماد ترکیبات جدید پیشنهادی است.

۳-۳ نتیجه گیری نهایی

در این تحقیق یک مدل QSAR معتبر جهت پیش‌بینی فعالیت دارویی ترکیبات شیمیایی ضد سرطان ارائه شد. ارزیابی‌های متفاوت مدل Ridge-SCAD توسعه یافته نشان داد که مدل از اعتبار و تعمیم پذیری مناسبی برخوردار است. در این تحقیق از روش Ridge به‌عنوان یک روش غربالگری نوین برای کاهش ابعاد داده‌ها استفاده شد. سپس با اعمال روش انقباضی SCAD روی مجموعه آموزش مؤثرترین توصیف کننده‌ها انتخاب شدند و مدل SCAD ساخته شد. مدل خطی توسعه یافته Ridge-SCAD برای پیش‌بینی فعالیت دارویی ترکیبات مجموعه آزمون مورد استفاده قرار گرفت.

علاوه بر این تأثیر توصیف کننده‌های مدل Ridge-SCAD بر فعالیت دارویی با تجزیه و تحلیل ضرایب رگرسیونی استاندارد شده و نمودار سهم مشارکت توصیف کننده‌ها استخراج شد. در نهایت با توجه به میزان تأثیر توصیف کننده‌ها بر فعالیت دارویی، ترکیبات جدیدی با فعالیت دارویی مناسب ارائه شد. به‌منظور بررسی بیش‌تر ترکیبات جدید ضد سرطان و اثبات صحت پیش‌بینی فعالیت دارویی آن‌ها توسط مدل Ridge-SCAD از مطالعه داکینگ مولکولی برای استخراج نوع برهم‌کنش آن‌ها با گیرنده پروتئینی مربوطه استفاده شد. نتایج مربوط به برهم‌کنش ترکیبات پیشنهادی با گیرنده، صحت فعالیت پیش‌بینی شده با استفاده از مدل Ridge-SCAD را اثبات نمود. از طرفی ویژگی‌های فارماکوکینتیک این ترکیبات با قاعده ۵ لیپینسکی نیز بررسی شد. نتایج نشان‌دهنده پتانسیل بالای این ترکیبات به‌عنوان بازدارنده‌های سرطان است. بنابراین مدل توسعه یافته Ridge-SCAD به‌عنوان یک مدل معتبر می‌تواند برای پیش‌بینی فعالیت دارویی ترکیبات ضد سرطان c-Met کاربرد داشته باشد. همچنین، به‌منظور بررسی عملکرد روش Ridge برای غربالگری توصیف کننده‌ها، ارزیابی مدل SCAD بدون استفاده از روش غربالگری انجام شد. پارامترهای محاسباتی دلالت بر قدرت پیش‌بینی بیش‌تر مدل Ridge-SCAD دارد.

۳-۴ آینده نگری

- ✓ به کارگیری روش‌های غربالگری دیگری برای کاهش ابعاد داده‌ها
- ✓ اجرای روش‌های مدل‌سازی انقباضی جدید برای پیش‌بینی فعالیت دارویی متفاوت
- ✓ استفاده از مجموعه داده‌های متنوع دیگر برای ساخت مدل QSAR توسعه یافته با روش‌های انقباضی

٣- ٥ فهرست منابع

- [١] Massart D.L., Vandeginste B.G., Buydens L., De Jong S., et al. (1998) "Handbook of chemometrics and qualimetrics: Part A", pp
- [٢] Gemperline P. (2006), "Practical guide to chemometrics", CRC press ,
- [٣] Kowalski B.R. (1980) "Chemometrics", Analytical Chemistry. **52**, 5, pp 112-122
- [٤] Tropsha A. (2010) "Best practices for QSAR model development, validation, and exploitation", Molecular informatics. **29**, 6-7, pp 476-488
- [٥] Yasri A., and Hartsough D. (2001) "Toward an optimal procedure for variable selection and QSAR model building", Journal of chemical information and computer sciences. **41**, 5, pp 1218-1227
- [٦] Leonard J.T., and Roy K. (2006) "On selection of training and test sets for the development of predictive QSAR models", QSAR & Combinatorial Science. **25**, 3, pp 235-251
- [٧] Puzyn T., Mostrag-Szlichtyng A., Gajewicz A., Skrzyszowski M., et al. (2011) "Investigating the influence of data splitting on the predictive ability of QSAR/QSPR models", Structural Chemistry. **22**, pp 795-804
- [٨] Hdoufane I., Ounaissi D., Dermoune A., and Cherqaoui D. (2021) "Development of QSAR Models Using Singular Value Decomposition Method: A Case Study for Predicting Anti-HIV-1 and Anti-HCV Biological Activities", Biointerface Research in Applied Chemistry. pp
- [٩] Golbraikh A., and Tropsha A. (2000) "Predictive QSAR modeling based on diversity sampling of experimental datasets for the training and test set selection", Molecular diversity. **5**, 4, pp 231-243
- [١٠] Roy K., Kar S., and Das R.N. (2015), "A primer on QSAR/QSPR modeling: fundamental concepts", Springer ,
- [١١] Algamal Z., Lee M., Al-Fakih A., and Aziz M. (2016) "High-dimensional QSAR modelling using penalized linear regression model with L 1/2-norm", SAR and QSAR in Environmental Research. **27**, 9, pp 703-719
- [١٢] Algamal Z., and Lee M. (2017) "A new adaptive L1-norm for optimal descriptor selection of high-dimensional QSAR classification model for anti-hepatitis C virus activity of thiourea derivatives", SAR and QSAR in Environmental Research. **28**, 1, pp 75-90
- [١٣] Saber S., Mohamad H., and Aziz M., Development a QSAR Model of 1, 3, 4-Triazole Derivatives for Antioxidant Activity Prediction, in: 2018 International Conference on Advanced Science and Engineering (ICOASE), IEEE, 2018, pp. 380-383.
- [١٤] Hoerl A.E., and Kennard R.W. (1970) "Ridge regression: Biased estimation for nonorthogonal problems", Technometrics. **12**, 1, pp 55-67
- [١٥] Xu L., and Zhang W.-J. (2001) "Comparison of different methods for variable selection", Analytica Chimica Acta. **446**, 1-2, pp 475-481
- [١٦] Tibshirani R. (1996) "Regression shrinkage and selection via the lasso", Journal of the Royal Statistical Society: Series B (Methodological). **58**, 1, pp 267-288

- [١٧] Fan J., and Li R. (2001) "Variable selection via nonconcave penalized likelihood and its oracle properties", *Journal of the American statistical Association*. **96**, 456, pp 1348-1360
- [١٨] Mercader A.G., and Pomilio A.B. (2010) "QSAR study of flavonoids and biflavonoids as influenza H1N1 virus neuraminidase inhibitors", *European journal of medicinal chemistry*. **45**, 5, pp 1724-1730
- [١٩] Todeschini R., and Consonni V. (2008), "Handbook of molecular descriptors", John Wiley & Sons ,
- [٢٠] Asadi L., Gholivand K., and Zare K. (2016) "Phosphorhydrazides as urease and acetylcholinesterase inhibitors: biological evaluation and QSAR study", *Journal of the Iranian Chemical Society*. **13**, 7, pp 1213-1223
- [٢١] Sadeghi F., Afkhami A., Madrakian T., and Ghavami R. (2021) "A new approach for simultaneous calculation of pIC 50 and logP through QSAR/QSPR modeling on anthracycline derivatives: a comparable study", *Journal of the Iranian Chemical Society*. pp 1-16
- [٢٢] Goodarzi M., Deshpande S., Murugesan V., Katti S.B., et al. (2009) "Is feature selection essential for ANN modeling?", *QSAR & Combinatorial Science*. **28**, 11-12, pp 1487-1499
- [٢٣] Gholivand K., Valmoozi A.A.E., Salahi M., Taghipour F., et al. (2017) "Bisphosphoramidate derivatives: synthesis, crystal structure, anti-cholinesterase activity, insecticide potency and QSAR analysis", *Journal of the Iranian Chemical Society*. **14**, 2, pp 427-442
- [٢٤] Saleh A.M.E., Arashi M., and Kibria B.G. (2019), "Theory of ridge regression estimation with applications", John Wiley & Sons ,
- [٢٥] Montgomery D.C., Peck E.A., and Vining G.G. (2021), "Introduction to linear regression analysis", John Wiley & Sons ,
- [٢٦] Golbraikh A., and Tropsha A. (2002) "Beware of q²!", *Journal of molecular graphics and modelling*. **20**, 4, pp 269-276
- [٢٧] Veerasamy R., Rajak H., Jain A., Sivadasan S., et al. (2011) "Validation of QSAR models-strategies and importance", *Int. J. Drug Des. Discov*. **3**, pp 511-519
- [٢٨] Roy P.P., and Roy K. (2008) "On some aspects of variable selection for partial least squares regression models", *QSAR & Combinatorial Science*. **27**, 3, pp 302-313
- [٢٩] Gramatica P., and Sangion A. (2016) "A historical excursus on the statistical validation parameters for QSAR models: a clarification concerning metrics and terminology", *Journal of chemical information and modeling*. **56**, 6, pp 1127-1131
- [٣٠] Fatemi M.H., and Izadiyan P. (2011) "Cytotoxicity estimation of ionic liquids based on their effective structural features", *Chemosphere*. **84**, 5, pp 553-563
- [٣١] Bring J. (1994) "How to Standardize Regression Coefficients", *The American Statistician*. **48**, pp 209-213
- [٣٢] Morris G.M., and Lim-Wilby M., Molecular docking, in: *Molecular modeling of proteins*, Springer, 2008
- [٣٣] Huang S.-Y., and Zou X. (2010) "Advances and challenges in protein-ligand docking", *International journal of molecular sciences*. **11**, 8, pp 3034-3016

- [୩୧] Halperin I., Ma B., Wolfson H., and Nussinov R. (2002) "Principles of docking: An overview of search algorithms and a guide to scoring functions", *Proteins: Structure, Function, and Bioinformatics*. **47**, 4, pp 409-443
- [୩୨] Hernández-Santoyo A., Tenorio-Barajas A.Y., Altuzar V., Vivanco-Cid H., et al. (2013) "Protein-protein and protein-ligand docking", *Protein engineering-technology and application*. pp 63-81
- [୩୩] Sousa S.F., Fernandes P.A., and Ramos M.J. (2006) "Protein–ligand docking: current status and future challenges", *Proteins: Structure, Function, and Bioinformatics*. **65**, 1, pp 15-26
- [୩୪] Seeliger D., and De Groot B.L. (2010) "Conformational transitions upon ligand binding: holo-structure prediction from apo conformations", *PLoS computational biology*. **6**, 1, pp e1000634
- [୩୫] McGovern S.L., and Shoichet B.K. (2003) "Information decay in molecular docking screens against holo, apo, and modeled conformations of enzymes", *Journal of medicinal chemistry*. **46**, 14, pp 2895-2907
- [୩୬] Goodsell D.S., Morris G.M., and Olson A.J. (1996) "Automated docking of flexible ligands: applications of AutoDock", *Journal of molecular recognition*. **9**, 1, pp 1-5
- [୩୭] Feinstein W.P., and Brylinski M. (2015) "Calculating an optimal box size for ligand docking and virtual screening against experimental and predicted binding pockets", *Journal of cheminformatics*. **7**, 1, pp 1-10
- [୩୮] <https://www.who.int/news-room/fact-sheets/detail/cancer>
- [୩୯] Traxler P. (2003) "Tyrosine kinases as targets in cancer therapy—successes and failures", *Expert opinion on therapeutic targets*. **7**, 2, pp 215-234
- [୪୦] Witsch E., Sela M., and Yarden Y. (2010) "Roles for growth factors in cancer progression", *Physiology*.
- [୪୧] Xu W.-M., Han F.-F., He M., Hu D.-Y., et al. (2012) "Inhibition of tobacco bacterial wilt with sulfone derivatives containing an 1, 3, 4-oxadiazole moiety", *Journal of agricultural and food chemistry*. **60**, 4, pp 1036-1041
- [୪୨] Liu N.-W., Liang S., and Manolikakes G. (2016) "Recent advances in the synthesis of sulfones", *Synthesis*. **48**, 13, pp 1939-1973
- [୪୩] Caballero J., Quiliano M., Alzate-Morales J.H., Zimic M., et al. (2011) "Docking and quantitative structure–activity relationship studies for 3-fluoro-4-(pyrrolo [2, 1-f][1, 2, 4] triazin-4-yloxy) aniline-**୩**, fluoro-4-(1H-pyrrolo [2, 3-b] pyridin-4-yloxy) aniline, and 4-(4-amino-2-fluorophenoxy)-2-pyridinylamine derivatives as c-Met kinase inhibitors", *Journal of Computer-Aided Molecular Design*. **25**, 4, pp 349-369
- [୪୪] Huang D., Zhu X., Tang C., Mei Y., et al. (2012) "3D QSAR pharmacophore modeling for c-Met kinase inhibitors", *Medicinal chemistry*. **8**, 6, pp 1117-1125
- [୪୫] Lee J.Y., Lee K., Kim H.R., and Chae C.H. (2013) "3D-QSAR Studies on Chemical Features of 3-(benzo [d] oxazol-2-yl) pyridine-2-amines in the External Region of c-Met Active Site", *Bulletin of the Korean Chemical Society*. **34**, 12, pp 3553-3558
- [୪୬] Balasubramanian P.K., Balupuri A., Bhujbal S.P., and Cho S.J. (2019) "3D-QSAR-aided design of potent c-Met inhibitors using molecular dynamics simulation and binding free energy calculation", *Journal of Biomolecular Structure and Dynamics*. **37**, 8, pp 2165-

[٥٠] Algamal Z.Y., and Lee M.H. (2017) "A novel molecular descriptor selection method in QSAR classification model based on weighted penalized logistic regression", *Journal of Chemometrics*. **31**, 10, pp e2915

[٥١] Algamal Z., Qasim M., and Ali H. (2017) "A QSAR classification model for neuraminidase inhibitors of influenza A viruses (H1N1) based on weighted penalized support vector machine", *SAR and QSAR in Environmental Research*. **28**, 5, pp 415-426

[٥٢] Algamal Z.Y., Lee M.H., Al-Fakih A.M., and Aziz M. (2017) "High-dimensional QSAR classification model for anti-hepatitis C virus activity of thiourea derivatives based on the sparse logistic regression model with a bridge penalty", *Journal of Chemometrics*. **31**, 6, pp e2889

[٥٣] Qasim M., Algamal Z., and Ali H.M. (2018) "A binary QSAR model for classifying neuraminidase inhibitors of influenza A viruses (H1N1) using the combined minimum redundancy maximum relevancy criterion with the sparse support vector machine", *SAR and QSAR in Environmental Research*. **29**, 7, pp 517-527

[٥٤] Al-Thanoon N.A., Qasim O.S., and Algamal Z.Y. (2019) "A new hybrid firefly algorithm and particle swarm optimization for tuning parameter estimation in penalized support vector machine with application in chemometrics", *Chemometrics and Intelligent Laboratory Systems*. **184**, pp 142-152

[٥٥] Xia L.-Y., Wang Y.-W., Meng D.-Y., Yao X.-J., et al. (2018) "Descriptor selection via log-sum regularization for the biological activities of chemical structure", *International journal of molecular sciences*. **19**, 1, pp 30

[٥٦] Majumdar S., Basak S.C., Lungu C., Diudea M., et al. (2018) "Mathematical structural descriptors and mutagenicity assessment: a study with congeneric and diverse datasets", *SAR and QSAR in Environmental Research*. **29**, 8, pp 579-590

[٥٧] Mozafari Z., Arab Chamjangali M., Arashi M., and Goudarzi N. (2021) "Performance of smoothly clipped absolute deviation as a variable selection method in the artificial neural network-based QSAR studies", *Journal of Chemometrics*. **35**, 5, pp e3338

[٥٨] HyperChem(TM) Professional 8.0.5 H.I., 1115 NW 4th Street, Gainesville, Florida 32601, USA., 2008 pp

[٥٩] Todeschini R., Lasagni M., and Marengo E. (1994) "New molecular descriptors for 2D and 3D structures. Theory", *Journal of chemometrics*. **8**, 4, pp 263-272

[٦٠] Todeschini R., Consonni V., Mauri A., and Pavan M., DRAGONs software for the calculation of molecular descriptors, version 5.5 for Windows. Milan, Italy, in, 2007.

[٦١] IBM SPSS statistics 25 for Windows student G.D., and Mallery P., 2017 pp

[٦٢] Dessau R.B., and Pipper C.B. (2008) ""R"--project for statistical computing", *Ugeskrift for laeger*. **170**, 5, pp 328-330

[٦٣] Matlab Version R2017a T.M.I., Massachusetts, 2017 pp

[٦٤] Joy S., Nair P.S., Hariharan R., and Pillai R.M. (2006) "Detailed Comparison of the Protein-Ligand Docking Efficiencies of GOLD, a Commercial Package and ArgusLab, a Licensable Freeware", *In silico biology*. **6**, 6, pp 601-605

[٦٥] ViewerLite v 5.0 A.I., San Diego, CA, 2001 pp

[۶۶] Nan X., Jiang Y.-F., Li H.-J., Wang J.-H., et al. (2019) "Design, synthesis and evaluation of sulfonylurea-containing 4-phenoxyquinolines as highly selective c-Met kinase inhibitors", *Bioorganic & medicinal chemistry*. **27**, 13, pp 2801-2812

[۶۷] Nan X., Zhang J., Li H.-J., Wu R., et al. (2020) "Design, synthesis and biological evaluation of novel N-sulfonylamidine-based derivatives as c-Met inhibitors via Cu-catalyzed three-component reaction", *European Journal of Medicinal Chemistry*. **200**, pp 112470

[۶۸] کاظمی م، (۱۳۹۷)، "انتخاب متغیر و تشخیص ساختار در مدل های نیمه پارامتری"، دانشگاه صنعتی

شاهرود

[۶۹] Todeschini R., and Consonni V., *Handbook of Molecular Descriptors*, in, 2002.

[۷۰] Sun M., Chen J., Wei H., Yin S., et al. (2009) "Quantitative Structure–Activity Relationship and Classification Analysis of Diaryl Ureas Against Vascular Endothelial Growth Factor Receptor-2 Kinase Using Linear and Non-Linear Models", *Chemical biology & drug design*. **73**, 6, pp 644-654

[۷۱] Qian F., Engst S., Yamaguchi K., Yu P., et al. (2009) "Inhibition of tumor cell growth, invasion, and metastasis by EXEL-2880 (XL880, GSK1363089), a novel inhibitor of HGF and VEGF receptor tyrosine kinases", *Cancer research*. **69**, 20, pp 8009-8016

[۷۲] Rose P.W., Beran B., Bi C., Bluhm W.F., et al. (2010) "The RCSB Protein Data Bank: redesigned web site and web services", *Nucleic acids research*. **39**, suppl_1, pp D392-D401

[۷۳] Chiari L.P.A., da Silva A.P., de Oliveira A.A., Lipinski C.F., et al. (2021) "Drug design of new sigma-1 antagonists against neuropathic pain: A QSAR study using partial least squares and artificial neural networks", *Journal of Molecular Structure*. **1223**, pp 129156

[۷۴] Zhou Q., Zhang N., Zhang C., Huang L., et al. (2010) "Molecular mechanism of enantioselective inhibition of acetolactate synthase by imazethapyr enantiomers", *Journal of agricultural and food chemistry*. **58** 7, pp 4202-4206

[۷۵] Parida P.K., Deka P., Shankar B., and Yadav R.N.S. (2014) "In silico antigenic site evaluation and antiviral therapy against dengue serotypes", *Bangladesh Journal of Pharmacology*. **9**, pp 83-95

[۷۶] Daina A., Michielin O., and Zoete V. (2017) "SwissADME: a free web tool to evaluate pharmacokinetics, drug-likeness and medicinal chemistry friendliness of small molecules", *Scientific Reports*. **7**, pp

[۷۷] Ertl P., and Schuffenhauer A. (2009) "Estimation of synthetic accessibility score of drug-like molecules based on molecular complexity and fragment contributions", *Journal of cheminformatics*. **1**, 1, pp 1-11

Abstract

This study aims to use a screening technique to reduce the dimension of the data matrix. For this purpose, the new penalized Ridge preprocessing technique is used to reduce the data dimensions. Then, the Smoothly Clipped Absolute Deviation (SCAD) as a penalized variable selection method is used to select the most effective descriptors and modeling. Selected descriptors were used as input to construct the QSAR model using the SCAD method. For this purpose, the QSAR model was developed using multiple Ridge implementation as a preprocessing technique for data screening. The remained descriptors were entered into SCAD penalized model. After implementing SCAD to the train set data, some effective descriptors (9) were used in the model construction. The proposed optimum model (Ridge-SCAD) was evaluated using calculation of statistical parameter, applicability domain, and Y- Randomization test. The mentioned different tests show satisfactory results. The performance of the Ridge-SCAD model was applied for the suggestion of some potent c-Met inhibitors. The efficiency of the recommended compounds was investigated using molecular docking and calculation of Lipinski's rule of five. The obtained results prove the potency of the new compounds. For the comparison aim, the QSAR model was developed without using the Ridge method as a preprocessing technique, and all the calculated descriptors were defined as SCAD inputs. The SCAD model was constructed using the train set data, and 12 descriptors were obtained as significant descriptors. The biological activity of the test data was predicted using the optimum SCAD model. The statistical parameters of SCAD model were calculated and compared with the Ridge-SCAD model. The results proved that the high predictability of the Ridge-SCAD model against to the SCAD model.

Keywords: QSAR, c-Met, Preprocessing, Ridge, SCAD, Molecular Docking



Faculty of Chemistry

M.Sc. Thesis in Analytical Chemistry

Quantitative structure-activity relationship study of
some new sulfonated c-Met inhibitors with
anti-cancer activity

By:

Mohadeseh Lotfi

Supervisor:

Dr. Mansour Arab Chamjangali

March 2022