



دانشکده صنایع و سیستم ها
گروه مهندسی فناوری اطلاعات

رساله جهت دریافت مدرک دکتری

طراحی و ارائه مدلی هوشمند جهت پیش‌بینی خطر
پرفشاری خون مبتنی بر SNP در شرایط عدم قطعیت

نگارش

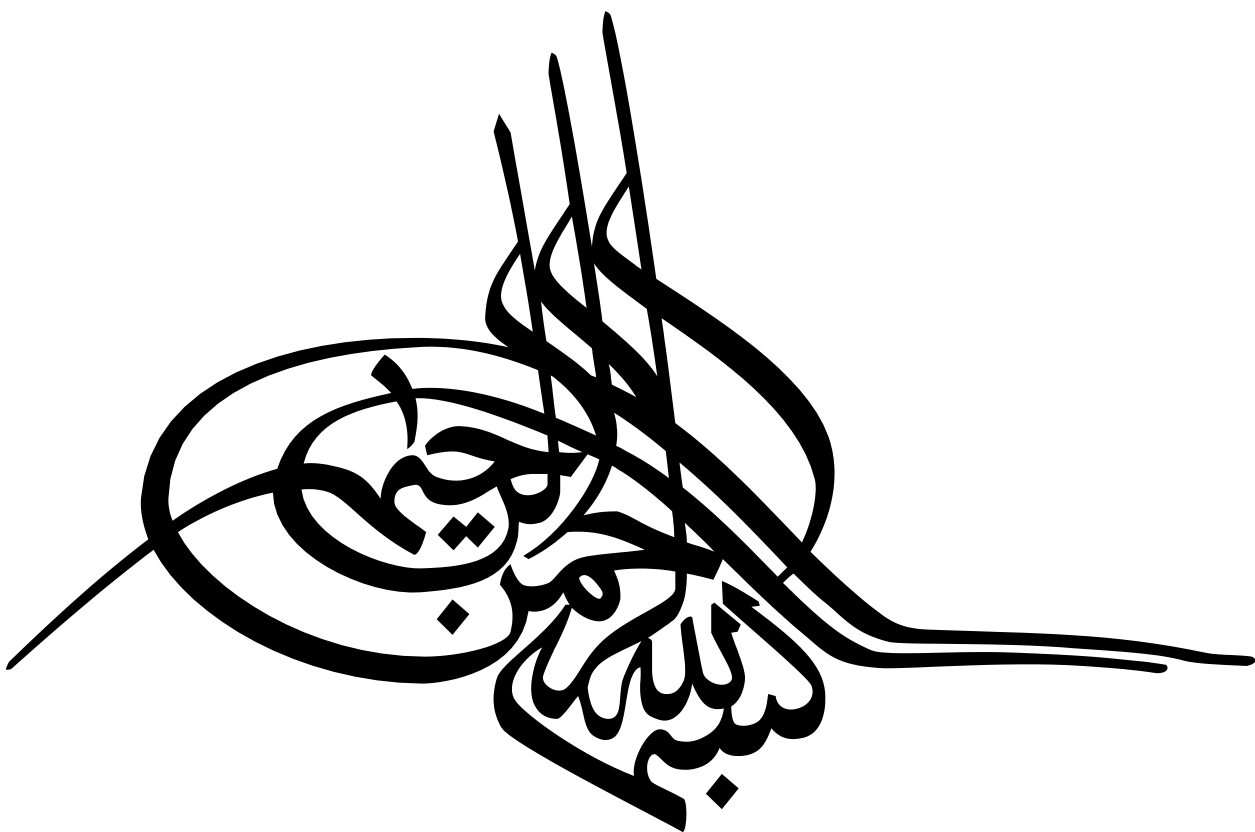
سید علی لاجوردی

استاد راهنما

دکتر مهرداد کارگری

استاد راهنما دوم

دکتر مریم السادات دانشپور



بسمه تعالی

تاییدیه اعضای هیأت داوران حاضر در جلسه دفاع از رساله دکتری

بدینوسیله گواهی می شود آقای سیدعلی لاجوردی به شماره دانشجویی ۹۶۶۶۲۳۲۰۰۳ رشته مهندسی فناوری اطلاعات - مدیریت سیستمهای اطلاعاتی در تاریخ ۱۴۰۲/۰۳/۲۷ از رساله دکتری ۱۹ واحدی خود با عنوان: طراحی و ارائه مدلی هوشمند جهت پیش بینی خطر پرفشاری خون مبتنی بر SNP در شرایط عدم قطعیت

Title: Design and presenting an intelligent model for predicting SNP-based hypertension in uncertainty

دفاع کرده است. اعضای هیأت داوران نسخه نهایی این رساله را از نظر فرم و محتوا بررسی نموده و پذیرش آنرا برای دریافت درجه دکتری تخصصی (Ph.D) تایید می نمایند.

اعضای هیأت داوران	نام و نام خانوادگی	رتبه علمی	امضاء
استاد راهنما	دکتر مهرداد کارگری	دانشیار	۱۶۰۹۲۷
استاد راهنمای دوم	دکتر مریم السادات دانشپور	دانشیار	د.ا
استاد مشاور	دکتر مهدی اکبرزاده	استادیار	
استاد داور داخلی	دکتر محمدمهدی سپهری	استاد	
استاد داور داخلی	دکتر توکتم خطیبی	دانشیار	
استاد داور خارجی	دکتر بهروز مینایی بیدگلی	دانشیار	
استاد داور خارجی	دکتر مهرداد روانشاد	دانشیار	
نماینده شورای تحصیلات تکمیلی	دکتر توکتم خطیبی	دانشیار	

آیین نامه حق مالکیت مادی و معنوی در مورد نتایج پژوهش های علمی دانشگاه تربیت مدرس

مقدمه: با عنایت به سیاست های پژوهشی و فناوری دانشگاه در راستای تحقق عدالت و کرامت انسان ها که لازمه شکوفایی علمی و فنی است و رعایت حقوق مادی و معنوی دانشگاه و پژوهشگران، لازم است اعضای هیئت علمی، دانشجویان، دانش آموختگان و دیگر همکاران طرح، در مورد نتایج پژوهش های علمی که تحت عناوین پایان نامه، رساله و طرح های تحقیقاتی با هماهنگی دانشگاه انجام شده است، موارد زیر را رعایت نمایند: ماده ۱ - حق نشر و تکثیر پایان نامه / رساله و درآمدهای حاصل از آن ها متعلق به دانشگاه هست ولی حقوق معنوی پدید آورندگان محفوظ خواهد بود.

ماده ۲ - انتشار مقاله یا مقالات مستخرج از پایان نامه / رساله به صورت چاپ در نشریات علمی و یا ارائه در مجامع علمی باید به نام دانشگاه بوده و با تأیید استاد راهنمای اصلی، یکی از اساتید راهنما، مشاور و یا دانشجو مسئول مکاتبات مقاله باشد. ولی مسئولیت علمی مقاله مستخرج از پایان نامه و رساله به عهده اساتید راهنما و دانشجو هست.

تبصره: در مقالاتی که پس از دانش آموختگی به صورت ترکیبی از اطلاعات جدید و نتایج حاصل از پایان نامه / رساله نیز منتشر می شود نیز باید نام دانشگاه درج شود.

ماده ۳ - انتشار کتاب، نرم افزار و یا آثار ویژه (اثری هنری مانند فیلم، عکس، نقاشی و نمایشنامه) حاصل از نتایج پایان نامه / رساله و تمامی طرح های تحقیقاتی کلیه واحدهای دانشگاه اعم از دانشکده ها، مراکز تحقیقاتی، پژوهشکده ها، پارک علم و فناوری و دیگر واحدها باید با مجوز کتبی صادره از معاونت پژوهشی دانشگاه و بر اساس آئین نامه های مصوب انجام شود.

ماده ۴ - ثبت اختراع و تدوین دانش فنی و یا ارائه یافته ها در جشنواره های ملی، منطقه ای و بین المللی که حاصل نتایج مستخرج از پایان نامه / رساله و تمامی طرح های تحقیقاتی دانشگاه باید با هماهنگی استاد راهنما یا مجری طرح از طریق معاونت پژوهشی دانشگاه انجام گیرد.

ماده ۵ - این آیین نامه در ۵ ماده و یک تبصره در تاریخ ۸۷/۴/۱ در شورای پژوهشی و در تاریخ ۸۷/۴/۲۳ در هیئت رئیسه دانشگاه به تأیید رسید و در جلسه مورخ ۸۷/۷/۱۵ شورای دانشگاه به تصویب رسیده و از تاریخ تصویب در شورای دانشگاه لازم الاجرا است.

این جانب سید علی لاجوردی دانشجوی رشته مهندسی فناوری اطلاعات - مدیریت دستگاه های اطلاعاتی ورودی سال تحصیلی ۱۳۹۶ مقطع دکتری دانشکده مهندسی صنایع و دستگاه ها متعهد می شوم کلیه نکات مندرج در آئین نامه حق مالکیت مادی و معنوی در مورد نتایج پژوهش های علمی دانشگاه تربیت مدرس را در انتشار یافته های علمی مستخرج از پایان نامه / رساله تحصیلی خود رعایت نمایم. در صورت تخلف از مفاد آئین نامه فوق به دانشگاه وکالت و نمایندگی می دهم که از طرف این جانب نسبت به لغو امتیاز اختراع بنام بنده و یا هر گونه امتیاز دیگر و تغییر آن به نام دانشگاه اقدام نماید. ضمناً نسبت به جبران فوری ضرر و زیان حاصله بر اساس برآورد دانشگاه اقدام خواهم نمود و بدین وسیله حق هر گونه اعتراض را از خود سلب نمودم.

امضا:



تاریخ: ۱۴۰۲/۴/۴

آیین‌نامه چاپ پایان‌نامه (رساله)‌های دانشجویان دانشگاه تربیت مدرس

نظر به اینکه چاپ و انتشار پایان‌نامه (رساله)‌های تحصیلی دانشجویان دانشگاه تربیت مدرس، مبین بخشی از فعالیت‌های علمی - پژوهشی دانشگاه است بنابراین به‌منظور آگاهی و رعایت حقوق دانشگاه، دانش‌آموختگان این دانشگاه نسبت به رعایت موارد ذیل متعهد می‌شوند:

ماده ۱: در صورت اقدام به چاپ پایان‌نامه (رساله)ی خود، مراتب را قبلاً به‌طور کتبی به دفتر نشر آثار علمی دانشگاه اطلاع دهد.

ماده ۲: در صفحه سوم کتاب (پس از برگ شناسنامه) عبارت ذیل را چاپ کند:

کتاب حاضر، حاصل رساله دکتری نگارنده سید علی لاجوردی در رشته مهندسی فناوری اطلاعات - مدیریت دستگاه‌های اطلاعاتی است که در سال ۱۴۰۱ در دانشکده مهندسی صنایع و سیستم‌ها دانشگاه تربیت مدرس به راهنمایی جناب آقای دکتر مهرداد کارگری و سرکار خانم دکتر مریم السادات دانشپور و مشاوره جناب آقای دکتر مهدی اکبرزاده از آن دفاع شده است.

ماده ۳: به‌منظور جبران بخشی از هزینه‌های انتشارات دانشگاه، تعداد یک درصد شمارگان کتاب (در هر نوبت چاپ) را به دفتر نشر آثار علمی دانشگاه اهدا کند. دانشگاه می‌تواند مازاد نیاز خود را به نفع مرکز نشر در معرض فروش قرار دهد.

ماده ۴: در صورت عدم رعایت ماده ۳، ۵۰٪ بهای شمارگان چاپ‌شده را به‌عنوان خسارت به دانشگاه تربیت مدرس، تأدیه کند.

ماده ۵: دانشجو تعهد و قبول می‌کند در صورت خودداری از پرداخت بهای خسارت، دانشگاه می‌تواند خسارت مذکور را از طریق مراجع قضایی مطالبه و وصول کند؛ به‌علاوه به دانشگاه حق می‌دهد به‌منظور استیفای حقوق خود، از طریق دادگاه، معادل وجه مذکور در ماده ۴ را از محل توقیف کتاب‌های عرضه‌شده نگارنده برای فروش، تأمین نماید.

ماده ۶: این‌جانب سید علی لاجوردی دانشجوی رشته مهندسی فناوری اطلاعات - مدیریت دستگاه‌های اطلاعاتی مقطع دکتری تعهد فوق و ضمانت اجرایی آن را قبول کرده، به آن ملتزم می‌شوم.

نام و نام خانوادگی: سید علی لاجوردی

تاریخ و امضا: ۱۴۰۲/۴/۴





دانشکده صنایع و سیستم ها
گروه مهندسی فناوری اطلاعات

رساله دکتری

طراحی و ارائه مدلی هوشمند جهت پیش‌بینی خطر
پرفشاری خون مبتنی بر SNP در شرایط عدم قطعیت

نگارش

سید علی لاجوردی

اساتید راهنما

دکتر مهرداد کارگری

دکتر مریم السادات دانشپور

استاد مشاور

دکتر مهدی اکبرزاده

خرداد ۱۴۰۲

مکاشفتہ از طحا شد نظر بر این در طحا صیقل شد

مکاشفتہ از یک بخش از طحا شد

مکاشفتہ از طحا شد در طحا شد

مکاشفتہ از یک بخش از طحا شد

من نشسته در کفاره در طحا شد

من صبر شد

من صبر شد

من صبر شد

من صبر شد

من صبر شد

من صبر شد

من صبر شد

چکیده

در دهه‌های اخیر، پزشکی شخصی بر مبنای ژنومیک برای ارائه خدمات مناسب و مؤثر برای بیماران بسته به ویژگی‌های ژنتیکی آن‌ها ارائه شده است. همچنین مطالعات گسترده ارتباط ژنومی انواع عوامل خطر ژنتیکی مبتنی بر جمعیت را برای بیماری‌های رایج و پیچیده شناسایی کرده است. جهت ارتقای سطح پزشکی شخصی، پژوهش‌های صورت گرفته در تلاش‌اند تا علاوه بر افزایش درک عمیق ژنومیک، به بهبود مراقبت‌های پزشکی شخصی با استفاده از مدل‌های پیش‌آگهی دقیق‌تر بیماری‌ها کمک کنند. امتیاز خطر چندژنی و یادگیری ماشین، دو روش اولیه برای پیش‌آگهی خطر بیماری هستند. علیرغم پیشرفت‌های اخیر، نتایج به‌دست‌آمده از خطرهای چندژنی به دلیل روش‌هایی که در حال حاضر مورد استفاده قرار می‌گیرند، محدود است. این در حالی است که در طول زمان، الگوریتم‌های یادگیری ماشین توانایی‌شان برای پیش‌آگهی خطر بیماری‌های پیچیده بیشتر شده است. این افزایش توان قدرت پیش‌آگهی از توانایی الگوریتم‌های یادگیری ماشین برای بررسی داده‌های چندبعدی است.

در این تحقیق، روش‌های مختلف یادگیری ماشین و یادگیری عمیق جهت مدل‌سازی خطر بیماری پرفشاری خون بر اساس SNP بر روی داده‌های ۲۰ ساله طرح ژنتیک کاردیومتابولیک تهران، پژوهشکده علوم غدد درون‌ریز و متابولیسم دانشگاه علوم پزشکی شهید بهشتی که شامل داده‌های ژنومی به‌صورت تراشه با ۶۵۲۹۱۹ SNP برای ۱۱۰۸۸ نفر در دوره‌های مختلف هست، پیاده‌سازی و در ادامه نتایج و دستاوردهای آن بیان می‌گردد.

روش پردازش اولیه و مدل‌سازی داده به خصوص در شبکه یادگیری عمیق از ادبیات موضوع متفاوت و نتایج تحقیق در دو بخش یادگیری ماشین و یادگیری عمیق دارای AUC به ترتیب ۰.۸۷ و ۰.۸۹ هست و ضمناً این مدل‌سازی امکان پیش‌آگهی در سنین جوانی را با دقت مناسبی فراهم کرده است.

کلمات کلیدی:

یادگیری عمیق، نشانگرهای ژنومی، پرفشاری خون، خطر بیماری

۱- مقدمه	۱
۱-۱ پزشکی شخصی	۱
۱-۱-۱ پزشکی دقیق	۳
۲-۱ ژنوم انسان	۳
۱-۲-۱ ژنومیک و ژنتیک	۳
۲-۲-۱ آلل	۴
۳-۲-۱ ژنوتیپ یا ژن نمود	۵
۴-۲-۱ فنوتیپ یا رخ نمود	۵
۵-۲-۱ تنوع رخ نمودی	۶
۶-۲-۱ تفاوت های ژنومی	۶
۷-۲-۱ چندریختی تک نوکلئوتید (SNP)	۷
۸-۲-۱ نقش عوامل ژنتیکی و غیر ژنتیکی در بروز صفات	۸
۳-۱ مطالعات گسترده ارتباط ژنومی	۱۰
۱-۳-۱ تحلیل GWAS و وراثت پذیری	۱۲
۴-۱ هم ترازای توالی ژنوم	۱۳
۱-۴-۱ روش های هم ترازای توالی چندگانه	۱۵
۵-۱ فشارخون	۱۵
۱-۵-۱ فشارخون سیستولیک	۱۵
۱-۵-۲ فشارخون دیاستولیک	۱۵
۶-۱ حوزه پژوهش	۱۶
۷-۱ ضرورت و اهداف پژوهش	۱۶
۱-۷-۱ سؤالات تحقیق	۱۷
۸-۱ نوآوری تحقیق	۱۷
۹-۱ جمع بندی	۱۷
۲- ادبیات نظری و پیشینه پژوهش	۱۸
۱-۲ استراتژی های افزایش بهره وری	۱۸
۱-۱-۲ انتخاب یک اندازه نمونه مناسب	۱۹
۲-۱-۲ رویکرد چندمرحله ای	۱۹
۳-۱-۲ جمعیت اولیه و نمونه های جمع آوری شده	۱۹
۴-۱-۲ اجتناب از گرفتن نتایج مثبت کاذب	۱۹
۲-۲ رویکرد آماری	۱۹
۱-۲-۲ مدل GRAMMAR	۲۰
۲-۲-۲ مدل GCTA	۲۰
۳-۲-۲ مدل ACTA	۲۰
۴-۲-۲ مدل REACTA	۲۱
۵-۲-۲ مدل PHEWAS	۲۲
۶-۲-۲ مدل OSCA	۲۲
۳-۲ رویکرد فرا تحلیل	۲۳
۱-۳-۲ مدل METAL	۲۴

۲۴	مقایسه مدل های فرا تحلیل.....	۲-۳-۲
۲۴	رویکرد یادگیری ماشین.....	۴-۲
۲۵	مدل GDL.....	۱-۴-۲
۲۶	روش های یادگیری ماشین سنتی.....	۵-۲
۲۷	الگوریتم درخت تصمیم.....	۱-۵-۲
۲۷	الگوریتم جنگل تصادفی.....	۲-۵-۲
۲۸	شبکه های بیزی.....	۳-۵-۲
۲۸	رده بند ماشین بردار پشتیبان.....	۴-۵-۲
۲۹	رده بند درخت تصمیم.....	۵-۵-۲
۲۹	شبکه های عصبی.....	۶-۵-۲
۳۰	یادگیری عمیق (DNN).....	۶-۲
۳۱	تحلیل داده های حجیم مبتنی بر یادگیری عمیق.....	۱-۶-۲
۳۱	الگوریتم های یادگیری عمیق.....	۲-۶-۲
۳۳	شبکه های عصبی پیچشی (CNN).....	۳-۶-۲
۳۴	پیش بینی امتیاز خطر.....	۷-۲
۳۵	امتیازدهی خطر چندرذنی.....	۱-۷-۲
۳۶	مدل های پیش بینی کننده بیماری بر پایه یادگیری ماشینی.....	۲-۷-۲
۳۸	بیماری پرفشاری خون (HTN).....	۲-۸
۴۵	مدل های پیش بینی کننده بیماری فشارخون.....	۲-۸-۲
۴۶	جمع بندی.....	۹-۲
۴۷	 ۳- روش شناسی تحقیق	
۴۸	داده های تحقیق.....	۱-۳
۵۰	پروژه ژنوم مرجع ایرانیان.....	۳-۱-۱
۵۰	ارزیابی مدل.....	۲-۳
۵۱	ماتریس درهم ریختگی.....	۳-۲-۱
۵۵	روش های بخش بندی داده جهت ارزیابی.....	۳-۲-۲
۵۷	کاهش خطا در یادگیری ماشین در بیوانفورماتیک.....	۳-۳
۵۸	محدوده های تحقیق و روش اجرا.....	۴-۳
۶۰	فرایندهای مدل سازی.....	۵-۳
۶۰	فرایند مدل سازی بر مبنای SNP کاندید.....	۳-۵-۱
۶۱	فرایند مدل سازی CNN.....	۳-۵-۲
۶۲	جمع بندی.....	۶-۳
۶۳	 ۴- نتایج و دستاوردها	
۶۳	بخش بندی و کنترل کیفی داده های ژنومی.....	۱-۴
۶۴	4-2- اجرای مدل بر مبنای SNP کاندیدا.....	
۶۶	محاسبه امتیاز خطر ژنتیکی مبتنی بر GRS.....	۳-۴
۶۷	مدل سازی فازی.....	۴-۴
۶۸	اجرای شبکه عمیق CNN.....	۵-۴
۶۸	تبدیل داده ژنومی به تصویر.....	۱-۵-۴
۶۹	اجرای شبکه CNN.....	۲-۵-۴

۶۹	آماده‌سازی داده‌های طولی رخ‌نمود.....	۳-۵-۴
۷۰	اجرای شبکه تمام متصل.....	۴-۵-۴
۷۳	جمع‌بندی نتایج.....	۴-۶
۷۳	شناسایی SNP های مهم.....	۴-۷
۷۵	ابزار سنجش خطر بیماری فشارخون بالا.....	۴-۸
۷۵	تحلیل ساختار علمی موضوع امتیاز خطر ژنتیکی.....	۴-۹
۷۶	هم رخدادی واژگان.....	۴-۹-۱
۷۷	نتایج و ارزیابی آن.....	۴-۹-۲
۷۹	جمع‌بندی.....	۴-۱۰
۸۰	۵- بحث و نتیجه گیری.....	
۸۰	بررسی مدل پیش‌بینی کننده خطر HTN مبتنی بر SNP.....	۱-۵
۸۰	مقایسه نتایج و معنی‌داری بهبود.....	۲-۵
۸۱	بررسی نتایج مدل‌سازی درخت تصمیم.....	۳-۵
۸۲	بررسی نتایج حاصل از رتبه‌بندی SNP.....	۴-۵
۸۲	خطاهای مؤثر در تحقیق.....	۵-۵
۸۲	پاسخ به سؤالات تحقیق.....	۵-۶
۸۳	نوآوری‌های تحقیق.....	۵-۷
۸۴	محدودیت‌ها و پیشنهادات آتی.....	۵-۸
۸۴	خلاصه فصل.....	۵-۹
۸۵	فهرست منابع.....	

فهرست جدول‌ها

صفحه	عنوان
۱۸	جدول ۱-۲ رویکردهای تشخیص تغییرات مبتنی بر صفات پیچیده در بیماری‌های رایج (Hirschhorn and Daly, 2005)
۲۴	جدول ۲-۲ مقایسه بسته‌های نرم‌افزاری مدل‌های فرا تحلیل (Evangelou and Ioannidis, 2013)
۲۶	جدول ۲-۳ خلاصه مدل‌های بررسی شده
۳۲	جدول ۲-۴ جدول مقایسه شبکه DNN (Shrestha and Mahmood, 2019)
۳۳	جدول ۲-۵ نقاط ضعف الگوریتم‌های یادگیری عمیق
۳۸	جدول ۲-۶ مقایسه روش GRS و یادگیری ماشین (Ho et al., 2019)
۴۳	جدول ۲-۷ اعداد مربوط به پرفشاری خون (Hengel, Benitah and Wenzel, 2022)
۵۲	جدول ۳-۱ ماتریس درهم‌ریختگی
۵۵	جدول ۳-۲ روش‌های قاعده‌مند سازی
۶۳	جدول ۴-۱ کنترل کیفی داده‌های SNP
۶۴	جدول ۴-۲ نتایج مدل‌سازی روش‌های یادگیری ماشین سنتی
۶۷	جدول ۴-۳ نتایج دسته‌بندی‌های مختلف برای امتیاز خطر ژنتیکی
۷۳	جدول ۴-۴ خلاصه نتایج مدل‌سازی‌های انجام شده در پیش‌بینی بیماری پرفشاری خون
۷۴	جدول ۴-۵ لیست مرتب شناسایی ۱۰ SNP مهم
۷۷	جدول ۴-۶ نام مجلات با بیشترین ارتباط با موضوع

فهرست شکل‌ها

عنوان	صفحه
شکل ۱-۱ اثر بخشی درمان تابع ژنوم و محیط در وضعیت حال و آینده (Strachan and Read, 2018).....	۲
شکل ۱-۲ چند ریختی تک نوکلئوتید برای دو آلل.....	۸
شکل ۱-۳ فواید کشف روابط ژنوم-صفت (McCarthy <i>et al.</i> , 2008).....	۹
شکل ۱-۴ حوزه فعالیت تحقیقات GWAS (McCarthy <i>et al.</i> , 2008).....	۱۲
شکل ۱-۵ نمودار حوزه پژوهش.....	۱۶
شکل ۲-۱ روند زمانی و تکامل تحقیقات آماری بررسی شده در GWAS.....	۲۲
شکل ۲-۲ روند تحقیقات PheWAS در مقایسه با تحقیقات GWAS (Hebbring, 2019).....	۲۲
شکل ۲-۳ مراحل مدل‌های فرا تحلیل (Evangelou and Ioannidis, 2013).....	۲۳
شکل ۲-۴ پارامترهای معرفی شده در مدل METAL (Willer, Li and Abecasis, 2010).....	۲۴
شکل ۲-۵ دیاگرام عملکرد مدل GDL (Sun <i>et al.</i> , 2019).....	۲۵
شکل ۲-۶ چارچوب یادگیری ماشین در داده‌های ژنومی (Zou <i>et al.</i> , 2019).....	۲۷
شکل ۲-۷ استفاده از ماشین بردار پشتیبان در مسئله رده‌بندی.....	۲۹
شکل ۲-۸ معماری شبکه CNN.....	۳۳
شکل ۲-۹ ده دلیل اصلی عامل مرگ در جهان (Alizadehsani <i>et al.</i> , 2019).....	۳۹
شکل ۲-۱۰ عوامل مؤثر بر پرفشاری خون (Wang <i>et al.</i> , 2020).....	۴۴
شکل ۲-۱۱ نسخه اصلاح شده نظریه موزاییک فشارخون بالا. نمک و سلول‌های ایمنی در مرکز قرار می‌گیرند، زیرا بر سایر سلول‌ها تأثیر می‌گذارند (Hengel, Benitah and Wenzel, 2022).....	۴۴
شکل ۳-۱ شمای کلی تحقیق (آونگ تحقیق).....	۴۷
شکل ۳-۲ (الف) نمودار توزیع سن در فاز اول (بالا) تا فاز ششم (پایین) (ب) نمودار توزیع DBP (ج) نمودار توزیع SBP.....	۴۹
شکل ۳-۳ توزیع سن به تفکیک برچسب فشارخون.....	۵۰
شکل ۳-۴ منحنی ROC و نمایش AUC.....	۵۴
شکل ۳-۵ مقایسه AUROC و AURPC.....	۵۵
شکل ۳-۶ نحوه اجرای تحقیق.....	۵۹
شکل ۳-۷ مراحل اجرای مدل‌سازی یادگیری ماشین سنتی.....	۶۱
شکل ۳-۸ مسیر اجرای مدل یادگیری ماشین سنتی مبتنی بر SNP کاندید.....	۶۱
شکل ۳-۹ مسیر اجرایی شبکه عمیق CNN.....	۶۲
شکل ۴-۱ نمودار تغییرات تعداد SNP با مقادیر گمشده.....	۶۳

شکل ۴-۲ نمودار چگالی MAF.....	۶۴
شکل ۴-۳ نمودار ROC مدل سازی جنسیت، سن و داده ژنومی.....	۶۵
شکل ۴-۴ نمودار توزیع سنی بر اساس برچسب فشارخون.....	۶۵
شکل ۴-۵ تصویر جداول مقایسه ای روش های با نظارت مبتنی بر بخش بندی سن.....	۶۶
شکل ۴-۶ نمونه کد محاسبه امتیاز خطر ژنتیکی.....	۶۷
شکل ۴-۷ نمودار ROC مدل سازی امتیاز خطر ژنومی.....	۶۷
شکل ۴-۸ فازی سازی سن.....	۶۸
شکل ۴-۹ تصویر ساخته شده از داده ژنومی.....	۶۸
شکل ۴-۱۰ خروجی Normalize، MaxP، CNN.....	۶۹
شکل ۴-۱۱ جدول مقایسه ای روش های با نظارت در فشار دیاستولی بالا.....	۷۰
شکل ۴-۱۲ جدول مقایسه ای روش های با نظارت در فشار سیستولی بالا.....	۷۰
شکل ۴-۱۳ جدول مقایسه ای روش های با نظارت در فشار سیستولی یا دیاستولی بالا.....	۷۰
شکل ۴-۱۴ نمودار ROC مدل سازی CNN.....	۷۱
شکل ۴-۱۵ نمودار صحت و مقدار ازدست رفته.....	۷۱
شکل ۴-۱۶ درخت تصمیم ناشی از مدل سازی فشارخون.....	۷۲
شکل ۴-۱۷ تصویر ابزار ساخته شده جهت پیش بینی خطر HTN.....	۷۵
شکل ۴-۱۸ نمایی از شبکه ارتباط کلمات کلیدی در مقالات.....	۷۸
شکل ۵-۱ مقایسه نمودارهای ROC.....	۸۱

جدول اختصارات

معاذل فارسی	معاذل انگلیسی	اخصار
ناحیه زیر منحنی	Area Under Curve	AUC
شاخص توده بدن	Body Mass Index	BMI
فشارخون	Blood Pressure	BP
شبکه عصبی پیچشی	Convolutional Neural Network	CNN
تغییرات تعداد کپی	Copy number variation	CNV
فشارخون دیاستولیک	Diastolic Blood Pressure	DBP
یادگیری عمیق	Deep learning	DL
شبکه عصبی عمیق	Deep Neural Network	DNN
برنامه نویسی پویا	Dynamic Programming	DP
پرونده الکترونیک سلامت	Electronic Health Record	EHR
شبکه کاملاً متصل	Fully Connected Network	FCN
یادگیری عمیق ژنوم	Genome Deep Learning	GDL
ماتریس رابطه ژنتیکی	Genetic Relation Matrix	GRM
مطالعه ارتباط گسترده ژنوم	Genome Wide Association Study	GWAS
پروژه ژنوم انسان	Human Genome Project	HGP
پرفشاری خون	Hypertension	HTN
پروژه ژنوم مرجع ایران	Iranian Reference Genome Project	IRGP
یادگیری ماشین	Machine Learning	ML
مدل خطی مختلط	Mixed Linear Model	MLM
همترازی توالی چندگانه	Multiple Sequence Alignment	MSP
ارزش پیش‌بینی منفی	Negative Predictive Value	NPV
تحلیل مؤلفه‌های اصلی	Principal Component Analysis	PCA
واکنش‌های زنجیره‌ای پلیمرز	Polymerase chain reactions	PCR
مطالعه ارتباط گسترده پدیده	Phenome-wide association study	PheWAS
ارزش پیش‌بینی مثبت	Positive Prediction Value	PPV
امتیاز خطر چند ژنی	Polygenic Risk Score	PRS
همترازی توالی دوگانه	Pairwise Sequence Alignment	PSA
چند ریختی طول قطعه محدود	Restriction fragment length polymorphisms	RFLPs
شبکه عصبی گردشی	Recurrent neural network	RNN
ویژگی عملکرد گیرنده	Receiver Operating Characteristic	ROC
فشارخون سیستولیک	Systolic Blood Pressure	SBP
شبکه‌های اجتماعی	Social networks	SN
چند ریختی تک نوکلئوتیدی	Single Nucleotide Polymorphism	SNP
تغییرات تک نوکلئوتیدی	Single nucleotide variation	SNV
تکرارهای پشت سر هم کوتاه	short tandem repeats	STR
ماشین بردار پشتیبان	Support Vector Machine	SVM
نرخ منفی واقعی	True Negative Rate	TNR
نرخ مثبت واقعی	True Positive Rate	TPR
تعداد متغیر تکرارهای پشت سر هم	Variable number of tandem repeats	VNTRs
توالی یابی کل ژنوم	Whole genome sequencing	WGS

۱- مقدمه

با هر مرحله از پیشرفت علم پزشکی، دانش و ابزارهای مورد استفاده برای تشخیص بیماری‌ها تغییر نموده و از متافیزیک به فیزیک، از سلول به مولکول و اخیراً هم به سطحی رسیده است که می‌توان واکنش بین مولکول‌ها را مطالعه و بررسی نمود. در زمان سقراط علم پزشکی در ارتباط نزدیکی با فلسفه و علوم طبیعی بود اما بقراط این شیوه را تغییر داد و پزشکی را به صورت حرفه‌ای و به‌دوراز فلسفه طبیعی بررسی کرد. بقراط تشخیص بیماری را بر پایه تحقیق در مورد احوالات بیمار قرارداد و اولین فردی بود که بسیاری از بیماری‌ها را توضیح و تشخیص داد، همچنین استعداد ابتلا به بیماری‌های مختلف در جمعیت‌های مختلف آسیایی و اروپایی را که از طریق وراثت باقی می‌مانند، را مورد توجه قرارداد (Adams, 1922). بنابراین احتمالاً بقراط اولین فردی بود که زمینه‌هایی را ایجاد کرد که در نهایت نظریه اصلی پزشکی منحصربه‌فرد بر پایه تشخیص بیماری برای هر فرد، شکل گیرد. اما علاوه بر تشخیص بیماری، تجویز داروی مناسب هم در درمان بیماری تأثیر به سزایی دارد. به‌طور کلی داروهایی که برای بیماران تجویز می‌شوند اصطلاحاً به صورت "یکی برای همه" است و همواره این سؤال مطرح می‌شد که چرا افراد مختلف در پاسخ به داروهایی که برایشان تجویز شده است، اثرات مطلوب و نامطلوب متفاوتی نشان می‌دهند؟ دانشمندان متوجه شدند که راز این امر احتمالاً در ژن‌ها نهفته است.

۱-۱- پزشکی شخصی

برای بالا بردن کیفیت درمانی باید خصوصیات ژنتیکی افراد را هم در نظر گرفت (Rafiei and Amjadi, 2013). با این دیدگاه واژه پزشکی شخصی^۱ وارد علم شد. در سال ۱۹۹۸، زمانی که برای اولین بار رساله‌ای با عنوان شخصی کردن پزشکی انتشار یافت، توجه کمی به این موضوع شد (Jain, 2006). پیشرفت‌های روزافزون علم پزشکی سبب شد که تشخیص و درمان بیماری‌ها کارآمدتر شود. یکی از این پیشرفت‌ها پروژه ژنوم انسان بود که با تکمیل شدن آن دریچه تازه‌ای از دیدگاه علم ژنتیک به سمت دنیای پزشکی باز شد. تکمیل این پروژه سرآغاز یک دوره جدیدی به نام پسا - ژنومی بود که سبب شد مفاهیم علمی جدید هم چون ژنومیک، پروتئومیک، فارماکوژنومیک و موارد مشابه وارد دنیای پزشکی شود. توالی یابی ژنوم انسان به همراه پیشرفت و توسعه فناوری‌های پربازده که عمدتاً به واژه Omics ختم می‌شود، دانش بشری را در مورد بیماری و سلامتی بهبود بخشیده و پایه‌ای را برای پزشکی انفرادی فراهم کردند.

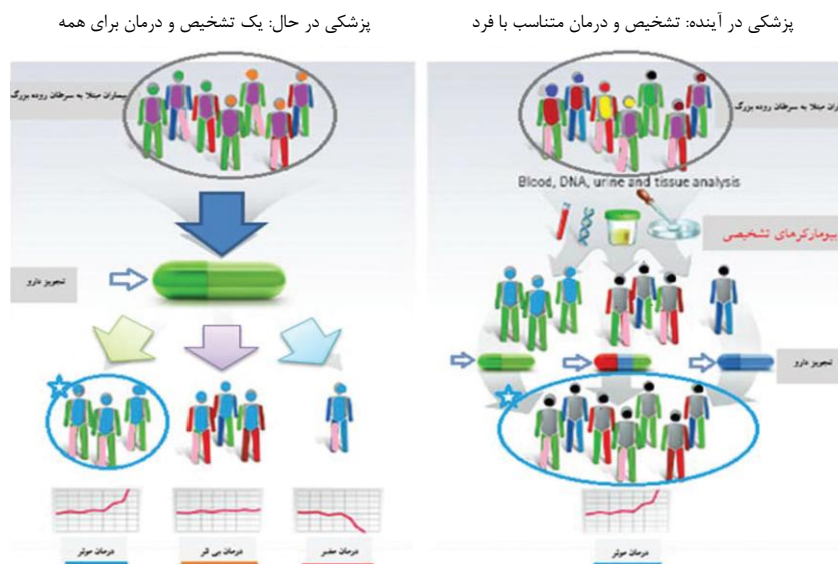
¹ Personalized Medicine

دامنه پزشکی فردی بسیار گسترده‌تر از آن است که با واژه طب ژنومی تعریف شود. علاوه بر ژنتیک، باید اپی ژنتیک و عوامل زیست‌محیطی نیز در نظر گرفته شود. بهترین راه برای تعبیر پزشکی فردی ادغام فن‌آوری‌های جدید برای عملکرد بالینی و تأثیر مهم آن در تشخیص مولکولی است. سامانه‌های زیست‌شناسی با رویکرد پزشکی برای توسعه اختصاصی نمودن درمان مهم است (Bauer et al., 2015).

در حال حاضر، دارویی که برای درمان یک فرد بیمار تجویز می‌شود تقریباً برای تمام مبتلایان به آن بیماری یکسان است. این در حالی است که واکنش بیمار به دارو اثرات مطلوب و نامطلوب و یا اصطلاحاً ناسازگاری دارویی وجود دارد. از طرفی بروز و شدت بیماری‌ها در افراد مختلف با توجه به دخالت عوامل ژنتیکی، محیطی و اپی ژنتیکی متفاوت است.

باین حال درمان‌های رایج برای همه این بیماران به صورت یکنواخت انجام می‌پذیرد. آمار نشان می‌دهد واکنش دارویی ناسازگار سالانه عامل مرگ بیش از صد هزار نفر در ایالات متحده هست و میلیون‌ها نفر به دلیل عوارض جانبی شدید بستری می‌شوند (Strachan and Read, 2018).

در پزشکی فردی الگوهای خاص سلولی و متابولیک در فاز تشخیص به عنوان زیست‌نشانگرهای^۱ تشخیصی استفاده می‌شود که بیماران را بر اساس شباهت‌های موجود معرفی می‌نماید و بهترین اطلاعات درمان‌های فردی را ارائه می‌دهد (شکل ۱-۱) (Strachan and Read, 2018). پیشرفت‌های چشمگیر در تشخیص و درمان بیماری بر اساس تفاوت‌های ژنوم افراد در حوزه پزشکی فردی ایجاد شده است (Ginsburg and Willard, 2009).



شکل ۱-۱ اثر بخشی درمان تابع ژنوم و محیط در وضعیت حال و آینده (Strachan and Read, 2018)

ژنوم انسان‌ها علیرغم وجود شباهت‌های زیاد دارای تفاوت‌های جزئی است که آن را منحصر به فرد می‌نماید. در نتیجه باعث بروز درمان فردی کاملاً با تحقیقاتی که برای تعیین تغییرات و تفاوت‌های ژنوم انسانی انجام می‌شود گره خورده است.

¹ Biomarker

۱-۱-۱- پزشکی دقیق

هدف پزشکی دقیق، ارائه درمان‌های پزشکی مناسب برای بیماران با توجه به ویژگی‌های ژنتیکی آن‌ها هست. این امر در درجه اول شامل شخصی‌سازی مراقبت‌های پیشگیرانه و فعال به‌منظور به حداکثر رساندن کارایی درمان و کاهش هزینه‌ها است. شخصی‌سازی از طریق استفاده از انواع مختلف اطلاعات اومیکس^۱ و ادغام آن‌ها باهدف درک خطرات بیماری حاصل می‌شود.

کاربرد پزشکی دقیق در فارماکوژنومیک باعث موفقیت‌های قابل‌توجهی در زمینه^۲ ارائه داروهای سفارشی و مصرف آن‌ها در اندازه‌های مناسب گردیده است. به‌عنوان مثال، اطلاعات ژنتیکی به‌طور مرتب در استراتژی‌های درمانی برای درمان سرطان‌های پستان HER2 مثبت با داروی trastuzumab، درمان سرطان‌های ریه دارای افزایش بیان EGFR با erlotinib، یا درمان لوسمی میلوژنی مزمن دارای کروموزوم فیلادلفیا با imatinib مورد استفاده قرار گرفته است.

بااین‌حال، در مورد این موضوع که آیا ژنومیک دقیق در نقطه‌ای قرار گرفته است که بتواند منافع بیشتری را در رویکردهای بهداشت عمومی استاندارد ارائه دهد یا خیر، در حوزه سلامت عمومی هنوز هم بحث وجود دارد.

۱-۲- ژنوم انسان

ژنوم انسان مجموعه‌ای کامل از DNA، شامل تمام اطلاعات موردنیاز برای حیات انسان است، که از ۳ میلیارد جفت باز تشکیل شده است. در هسته سلول‌های انسان ۲۳ جفت کروموزوم وجود دارد که از هر جفت کروموزوم یک نسخه از مادر (تخم) و دیگری از پدر (اسپرم) به فرزند می‌رسد. هر کروموزوم متشکل از یک مولکول درشت به نام داکسی ریبونوکلیک‌اسید یا همان DNA است. مجموع ۲۳ جفت مولکول DNA فرد حاوی تمام اطلاعات لازم برای تعیین صفاتی است که فرد را می‌سازد. DNA خود را به بخش‌های خاصی به نام ژن سازمان‌دهی می‌کند که پروتئین‌های ضروری برای رشد و عملکرد بدن ما را رمزگذاری می‌کند. کدگذاری اختصاصی درون ژن‌ها از طریق ترکیب چهار باز خاص به نام‌های آدنین، تیمین، گوانین و سیتوزین انجام شده است.

۱-۲-۱- ژنومیک و ژنتیک

توالی جفت بازها در رشته DNA، محتوی پیام ژنتیکی را تعیین می‌کند. به عبارتی ژن، توالی خاصی از این بازها است که حاوی اطلاعات لازم جهت ساخت یک محصول مانند یک هورمون پروتئینی یا آنزیم را فراهم می‌کند. انسان‌ها هزاران ژن دارند که در سراسر طول DNA و در قالب ۲۳ جفت کروموزوم بسته‌بندی شده است. ژنومیک، مطالعه تمام ژن‌های ژنوم انسان است.

ژنتیک معمولاً به بررسی تک‌تک ژن‌ها و نقش آن‌ها در بیماری یا وراثت اشاره دارد. ژنومیک به تمامیت اطلاعات ژنتیکی یک فرد اشاره دارد و به‌توالی ژنتیک ژن‌ها، ساختار و عملکرد آن‌ها و همچنین تعاملات بین ژن‌ها می‌پردازد.

¹ omics

امروزه تعیین توالی کل ژنوم انسانی ممکن شده است، تا بتواند اساساً "نقشه" ژنوم شخصی انسان را، تحت عنوان توالی کل ژنوم (WGS)، ارائه کند. دانستن بیشتر در مورد DNA، به انسان این امکان را می‌دهد که بالقوه به‌سوی زندگی سالم‌تری حرکت کند. بررسی کامل ژنوم انسان از لحاظ تاریخی مبتنی بر برنامه‌های تشخیصی بالینی بود و در آن توجه کمتری به پیشگیری از بیماری و مدیریت سلامت داده‌شده بود. با این حال، این نوع از اطلاعات در حال حاضر به‌عنوان بخشی از یک استراتژی مراقبت‌های بهداشتی پیشگیرانه فردی به‌سرعت در حال رشد است و با توجه به تشخیص زودهنگام بیماری، فرصتی برای درمان و نتایج بهبود بهتر به وجود می‌آورد.

۱-۲-۲-۱- آلل^۱

در هر رشته DNA مجموعه‌ای از اطلاعات به نام ژن‌ها قرار دارند. این ژن‌ها هستند که صفات مختلف جانوران یا انسان‌ها را تعیین می‌کنند. ژن‌ها برای همه انسان‌ها یکسان هستند بنابراین همه انسان‌ها ژن‌های یکسانی را به ارث می‌برند. اما این آلل‌ها هستند که تعیین می‌کنند ژن‌ها چگونه در یک فرد بروز پیدا کنند و آشکار شوند. به‌عنوان مثال وجود رنگ چشم به ژن‌ها مربوط است اما این که چه رنگ چشمی از والدین به ارث برده می‌شود به‌وسیله آلل‌ها تعیین می‌شود.

آلل‌ها به‌صورت جفت هستند و به عبارتی موجودات زنده برای هر صفت دو آلل دارد. (به‌عنوان مثال آلل قهوه‌ای و آبی برای چشم در یک نفر.) یک آلل روی یک کروموزوم و آلل دیگر روی کروموزوم دیگر قرار دارد. هر فرد یک آلل را از مادر و آلل دیگر را از پدر به ارث می‌برد.

تفاوت آلل‌ها و ژن‌ها در این است که ژن‌ها جفت نیستند چراکه چنان‌که گفته شد ژن‌ها در همه انسان‌ها یکسان هستند. به همین دلیل که گفته شد یک جفت آلل رخ‌نمودهای متفاوت به وجود می‌آورند اما این موضوع درباره ژن‌ها صادق نیست.

قانون وراثت مندل فرایندی که طی آن یک آلل انتقال می‌یابد را فرمول‌بندی کرده است.

چنان‌که گفته شد موجودات زنده دو آلل برای هر صفت دارند. به وضعیتی که دو آلل متفاوت برای یک صفت وجود داشته باشد هتروزیگوت^۲ گفته می‌شود. در نوعی از هتروزیگوت یکی از دو آلل غالب است و دیگری مغلوب و در نتیجه آلل غالب آشکار می‌شود و آلل مغلوب پنهان. به‌عنوان مثال، در انسان آلل برای رنگ چشم قهوه‌ای غالب است و آلل برای رنگ چشم آبی مغلوب است.

غالب بودن رنگ قهوه‌ای بر رنگ آبی در چشم یک نمونه از تسلط کامل است. در روابط هتروزیگوتی که در آن هیچ‌کدام از آلل‌ها غالب نیستند اما هر دو به‌طور کامل آشکارند، وضعیت هم بارز گفته می‌شود. نمونه وضعیت هم بارز، به ارث رسیدن گروه خونی AB است. وقتی یک آلل به‌طور کامل بر دیگری مسلط نیست وضعیت تسلط ناقص است.

^۱ Allele

^۲ Heterozygote

۱-۲-۳- ژنوتیپ^۱ یا ژن نمود

ژن نمود یکی از سه عاملی است که رخ نمود فرد را مشخص می نماید (دو عامل دیگر اپی ژنتیک ارثی و عوامل محیطی غیر ارثی می باشند). بنابراین شبیه بودن ژن نمود دو موجود زنده تضمین کننده شباهت رخ نمودی آن دو موجود نیست. ژن نمودها بیانگر نحوه چینش تفاوتها در افراد و این تفاوت می تواند باعث افتراق موجودات از یکدیگر شود.

هر انسان بیش از ۲۰ هزار ژن دارد. معمولاً، هر کدام از این ژنها چندین نوع آلل در میان افراد مختلف دارند. همین موجب متفاوت بودن ژن نمود آنها، و در نهایت گوناگونی می شود، چراکه احتمال این که دو فرد برای همه ژنهای خود یک آلل را به ارث برده باشند تقریباً صفر است! مگر دوقلوهای همسان که حاصل تقسیم میوز سلول تخم می باشند، و بنابراین اطلاعات ژنتیکی آنها کاملاً مشابه است.

۱-۲-۴- فنوتیپ^۲ یا رخ نمود

رخ نمود خصوصیات قابل مشاهده یا صفت یک ارگانیسم است. مانند خصوصیات بیوشیمیایی یا فیزیولوژیکی. رخ نمود از بیان ژنهای یک ارگانیسم و همچنین تأثیر عوامل زیست محیطی و اپی ژنتیکی و تعامل بین این سه نتیجه می شود. هیچ کدام از موجودات زنده با ژن نمود همانند به یک شیوه عمل نمی کنند به این دلیل که ظاهر و رفتار با شرایط زیست محیطی و توسعه تغییر می کند. به همین ترتیب، تمام موجودات زنده که به نظر یکسان می رسند لزوماً ژن نمودهای یکسانی ندارند. تمایز ژن نمود، رخ نمود توسط ویلهلم جانسون در سال ۱۹۱۱ پیشنهاد شد تا تفاوت بین وراثت موجود زنده و آنچه وراثت تولید می کند روشن شود که این تفاوت مشابه چیزی است که توسط اوت ویزمن پیشنهاد شده بود، که بین قسمت قابل توارث نطفه (وراثت) و سلولهای سوماتیک بدن تمایز قائل شده بود.

با وجود آن تعریف به ظاهر ساده، مفهوم رخ نمود پیچیدگیهای پنهانی دارد. بسیاری از مولکولها و ساختارهایی که با مواد ژنتیکی کد شده اند قابل مشاهده در ظاهر یک ارگانیسم نیستند، اما می توان آنها را قابل مشاهده کرد (برای مثال توسط وسترن بلات) و در نتیجه بخشی از رخ نمود هستند. گروههای خونی انسان جز این دسته هستند؛ بنابراین، رخ نمود باید شامل ویژگیهایی باشد که بتواند توسط برخی از روشهای فنی به صورت قابل مشاهده درآید. توسعه بعدی کردن رفتار به رخ نمود است، رفتارها نیز از ویژگیهای قابل مشاهده هستند. درواقع روی ارتباط بالینی رخ نمودهای رفتاری پژوهشهایی انجام شده است که به طیف وسیعی از سندرمها مربوط می شوند.

بعضی از رخ نمودهای در طول سن ثابت هستند مانند رنگ چشم. اما بعضی از رخ نمودها در طول عمر انسان در حال تغییر هستند. مثلاً فشارخون نرمال در سنین مختلف (بدون در نظر گرفتن تأثیر بیماریها) تغییر می نماید. به رخ نمودهایی که در طول سن انسان تغییر می نمایند رخ نمود طولی^۳ گفته می شود.

¹ Genotype

² Phenotype

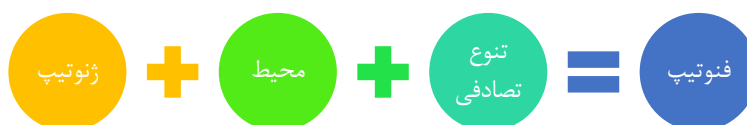
³ Longitudinal Phenotype

۱-۲-۵- تنوع رخ نمودی

تنوع رخ نمودی (ناشی از تنوع ژنتیکی ارثی) یک پیش نیاز اساسی برای تکامل توسط انتخاب طبیعی هست. این موجود زنده است که به عنوان یک کل در تولید نسل بعدی کمک می کند (یا نمی کند)، بنابراین انتخاب طبیعی به طور غیرمستقیم از طریق رخ نمودها بر ساختار ژنتیکی یک جمعیت تأثیر می گذارد. بدون تنوع رخ نمودی، هیچ تکاملی توسط انتخاب طبیعی وجود ندارد. تعامل بین ژن نمود و رخ نمود اغلب توسط رابطه زیر معرفی می شود:



به طور دقیق تر می توان گفت:



به رغم اینکه ژن نمود انعطاف پذیری بیشتری در بیان و اصلاح نسبت به رخ نمود دارد، اما بسیاری از موجودات زنده به دلیل قرار گرفتن در شرایط مختلف زیست محیطی دارای رخ نمودهای متنوعی شدند. گیاه آمبرلاتام^۱ می تواند در دو محیط مختلف در سوئد رشد کند. در زیستگاه های صخره ای، صخره های سمت دریا، گیاه با برگ انبوه و گل های بزرگ رشد پیدا می کند، ولی در میان تپه های شنی رشد گیاه همراه با برگ های باریک و گل های فشرده هست. این زیستگاه ها هستند که رخ نمود گیاه را مشخص می کنند. مفهوم رخ نمود را می توان به تغییراتی که در سطح ژن بر سازگاری موجود زنده تأثیر می گذارد گسترش داد.

۱-۲-۶- تفاوت های ژنومی

ژنوم هر فرد ویژه خود اوست و یک عامل مهم برای تعیین رخ نمودهای بیولوژیکی است. ۹۹.۹ درصد از توالی DNA ژنوم انسان ها یکسان است. ۰.۱ درصد باقی مانده باعث می شود یک شخص به طور مثال در ویژگی ها، رفتارها و بیماری ها منحصر به فرد شود (Ziegler et al., 2012).

در هنگام مقایسه ژنوم دو فرد، تفاوت های بسیاری (به لحاظ تعداد) دیده می شود که این تفاوت ها حاصل تغییرات ژنومی مانند چند مورفسم تک نوکلئوتیدی^۲ و تغییرات بزرگ تر مانند حذف^۳، جایگزینی^۴ و تغییر در تعداد کپی های ژنی^۵ هستند. هر کدام از این تغییرات می توانند منجر به بروز یک رخ نمود متفاوت در فرد شوند که این رخ نمود می تواند مربوط به صفات ظاهری فرد مانند قد و یا رنگ چشم و یا در ارتباط با خطر بروز یک بیماری باشد. SNP یک تغییر ژنتیکی کوچک در DNA است و رایج ترین نوع تنوع ژنتیکی هست که در میان افراد رخ می دهد.

^۱ Umbellatum Hieracium

^۲ Single nucleotide variation

^۳ Deletion

^۴ Insertion

^۵ Copy number variation

هر SNP نشان‌دهنده تفاوت در یک بلوک از ساختمان DNA تک‌رشته‌ای به نام نوکلئوتید است. DNA با چهار نوع نوکلئوتید که دارای باز آلی A، C، G و T مشخص می‌شود. SNP زمانی رخ می‌دهد که باز یک نوکلئوتید به‌عنوان مثال نوکلئوتید دارای باز A جایگزین یکی دیگر از سه باز نوکلئوتید (مثلاً G، C، یا T) می‌شود (Ganesalingam and Bowser, 2010).

SNPها به‌طور معمول در سراسر DNA افراد و به‌طور متوسط یک‌بار در هر ۳۰۰ نوکلئوتید رخ می‌دهد. بدین معنی که تقریباً ۱۰ میلیون SNP در ژنوم هر فرد وجود دارد (Diamandis, 2012) و یک تغییر باید حداقل در ۱ درصد از جمعیت اتفاق بیافتد تا جزء SNPهای رایج شود. حدود ۴۰ درصد SNP موجود در ژن‌ها باعث تغییر یک آمینواسید خواهد شد. SNP رایج‌ترین نوع یلی مورفسم است که باثبات تکاملی و فراوانی بالا در ژنوم توزیع شده‌اند.

البته این تغییرات می‌توانند: بی‌ضرر (تغییر در رخ‌نمود)، مضر (دیابت، سرطان، بیماری‌های قلبی، بیماری هانتینگتون و هموفیلی) و نهفته (تغییرات در ژنوم از جمله مناطق تنظیم‌کننده ژن به‌خودی‌خود مضر نیست، و تنها تحت شرایط خاصی مانند استعداد ابتلا به سرطان ریه و تغییر پاسخ فرد به یک درمان آشکار می‌شود) باشد.

بدین ترتیب کشف SNPها جهش عظیمی در زمینه ی ارتباط ژنتیک با اثربخشی دارویی و استعداد ابتلا به بیماری ایجاد کرده است. تعیین ژن‌نمود SNPها در درمان فردی اهمیت بسزایی دارد (Mahrokh et al., 2017).

۱-۲-۷- چندریختی تک نوکلئوتید^۱ (SNP)

یک نگاه کوتاه در زمان آغاز پلی مورفسم ژنتیکی انسان متعلق به ژن B-globin در سال ۱۹۷۸ بود که از آن برای تشخیص یک بیماری ارثی استفاده کرد. پس از ۲ سال، در سال ۱۹۸۰، تمایزهای کوتاهی در DNA کشف‌شده در کل ژنوم انسان گسترش یافت. این با استفاده از پلی مورفسم طول محدود^۲ استفاده شد. اطلاعات جالب پیچیده دیگری درباره پلی مورفسم DNA در سال ۱۹۸۵ گزارش شد. آن‌ها به‌عنوان ماهواره‌های کوچک^۳ شناخته شدند. به‌طور کلی، پلی مورفسم ژنتیکی می‌تواند در طرح‌های بی‌شماری موجود باشد، شامل: چندریختی تک نوکلئوتیدی، چندریختی تکرار پشت سرهم که شامل تعداد متغیر تکرارهای پشت سرهم^۴ و تکرارهای پشت سرهم کوتاه^۵. برای مطالعه انواع مختلف چندریختی‌های DNA، می‌توان از فن‌های مختلفی

¹ Single Nucleotide Polymorphism

² Restriction fragment length polymorphisms (RFLPs)

³ Mini satellites

⁴ Variable number of tandem repeats (VNTRs)

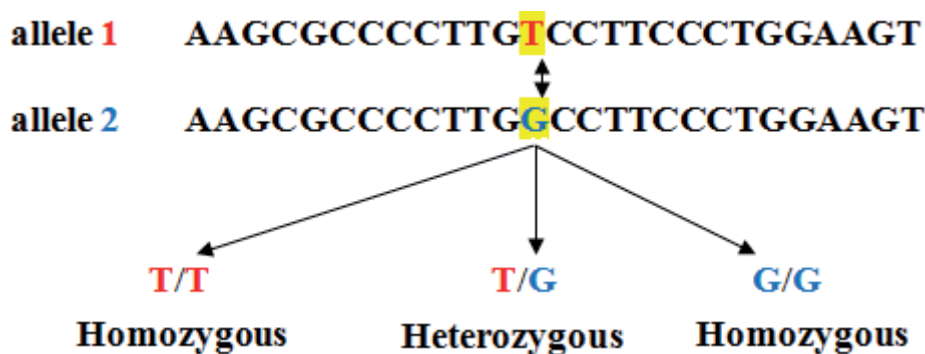
⁵ short tandem repeats (STR)

استفاده کرد، مانند پلی مورفیسم طول محدودکننده، واکنش‌های زنجیره‌ای پلیمرز^۱، روش‌های و ترتیب توالی ژنوم کل^۲ (Çalışkan, Erol and Öz, 2020).

SNP یک تغییر در واحد بلوکی ساختار دستورالعمل DNA است. این ساده‌ترین فرمول تفاوت ژنتیکی در بین افراد است. SNP شایع‌ترین اتفاق در تفاوت ژنتیکی در بین DNA افراد است.

به‌طورمعمول ممکن است در هر ۳۰۰ نوکلئوتید یک‌بار اتفاق بیفتد، یعنی به‌طور متوسط حدود ۱۰ میلیون SNP در ژنوم فرد وجود دارد. بیشترین تعداد، این SNP ها بین ژن‌ها یا درون ژن‌ها تنظیم می‌شوند. آن‌ها ممکن است به‌عنوان نشانه‌های زنده و یا شاخص‌های ارثی عمل کنند، متخصصان توالی‌هایی را پیدا می‌کنند که با بیماری در ارتباط هستند. به‌محض اینکه SNP ها در داخل یک ژن یا در یک ناحیه تنظیم‌کننده در نزدیکی یک ژن اتفاق می‌افتد، با تحریک کردن نقش ژن می‌توانند تأثیر بیشتری در بیماری نشان دهند. با این حال، SNP ها به‌طور کلی هیچ تأثیری در وضعیت عمومی سلامت ندارند. علاوه بر این، محققان اعلام کردند که SNP ها ممکن است به واکنش شخص در برابر داروهای قطعی کمک کند. هم‌چنین آن‌ها برای مسیری از عوامل ژنتیکی بیماری نادر در بستگان استفاده می‌شوند (Çalışkan, Erol and Öz, 2020).

شکل ۱-۲ نمونه‌ای از SNP برای دو آل را نشان می‌دهد.



شکل ۱-۲ چندریختی تک نوکلئوتید برای دو آل

۱-۲-۸- نقش عوامل ژنتیکی و غیر ژنتیکی در بروز صفات

تا مدت‌ها فقط در مورد صفات و شرایطی که توسط ژن‌ها در یک جایگاه منفرد در ژنوم انسان ایجاد می‌گردد، بحث می‌شد. اما امروزه طبق تحقیقات (Filshtein et al., 2019) و (Pattin and Moore, 2009) بسیاری از اختلالات شناخته‌شده با توارث مندلی مطابقت ندارند و هر صفت تنها توسط ژن‌هایی در یک لکوس تعیین نمی‌شود. درواقع، بسیاری از صفات توسط ترکیبی از ژن‌ها در بیش از یک محل در ژنوم انسان و همچنین عوامل غیر ژنتیکی تعیین می‌شوند. صفات چندعاملی صفاتی هستند که توسط عوامل متعدد ایجاد می‌شوند. در اغلب موارد، هر دو عوامل ژنتیکی و غیر ژنتیکی در ایجاد یک رخ‌نمود کمک می‌کند. به‌طور مثال قد انسان یک نمونه از یک صفت چندعاملی است. با این حال، بسیاری از متخصصان ژنتیک انسانی در مورد قد انسان نگران نیستند، در عوض، در عرصه‌ی ژنتیک انسانی به‌طور عمده در مورد بیماری‌های انسان نگران هستند که

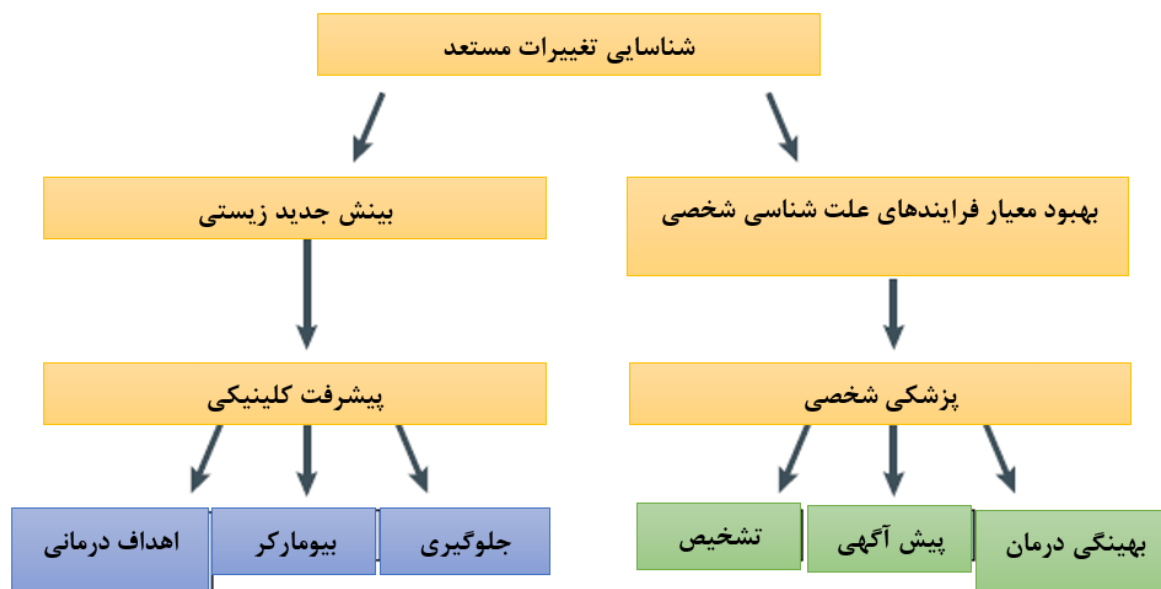
¹ Polymerase chain reactions (PCR)

² Whole genome sequencing (WGS)

در حال حاضر رایج‌ترین آن‌ها دیابت، بیماری قلبی، آلزایمر، آرتریت روماتوئید و سرطان هستند که تصور می‌شود در همه این بیماری‌ها چند جایگاه ژن و همچنین عوامل غیر ژنتیکی متعدد درگیر باشند و یک متخصص ژنتیک ممکن است آن‌ها را تحت رده‌بندی بیماری‌هایی که توسط عوامل متعدد ایجاد می‌شود به‌عنوان بیماری‌های چندعاملی مطرح نماید. یکی از ویژگی‌های مشترک بیماری‌های چندعاملی، طیف وسیعی از شدت و سن بیماری است. اضافه بر این پیچیدگی این واقعیت که برخی از این بیماری‌ها درواقع به وقوع یکدیگر کمک می‌کنند را نیز باید عنوان کرد، به‌عنوان مثال، دیابت خود یک عامل خطر برای بیماری قلبی است.

یک مشکل اساسی در ژنتیک کشف ارتباط بین ژن‌نمود و رخ‌نمود است. حصول اطمینان از وراثت یک صفت یک گام مهم به سمت امکان احتمال پیش‌بینی و ارزیابی خطر بیماری‌های انسان و درمان آن‌ها با استفاده از اطلاعات ژنتیکی هست. دانستن تمام عناصری که در وراثت دخالت دارند امکان وجود یک رابطه معلولی بین اطلاعاتی که نسل به نسل منتقل و منجر به رخ‌نمود موردنظر می‌شود را تسهیل می‌کند. مطالعات مرتبط با گستره ژنوم موفق به شناسایی بسیاری از چند مورفیسم‌های انسان وابسته به صفات قد، رنگ چشم، و یا استعداد ابتلا به بیماری‌های شایع را نشان داده‌اند اما این واریخت‌ها تنها بخش کوچکی از وراثت وابسته به یک صفت را شرح می‌دهند (Ligon et al., 2000; Trerotola et al., 2015).

کشف ارتباط بین عامل ژنتیکی و رخ‌نمود مزایایی مختلفی دارد که در تصویر ۱-۳ آمده است (McCarthy et al., 2008).



شکل ۱-۳ فواید کشف روابط ژنوم-صفت (McCarthy et al., 2008)

۱-۳- مطالعات گسترده ارتباط ژنومی^۱

در حدود سال ۲۰۰۰ میلادی، قبل از آنکه مطالعات گسترده ارتباط ژنومی معرفی شوند، جهت بررسی ارتباط بین ژن نمود و رخ نمود از بررسی توارث^۲ در خانواده و یا پیوستگی ژنتیکی^۳ استفاده می شد. این روش ها در بررسی بیماری هایی که یک ژن عامل بروز آنها بود، بسیار موفق بودند. اما این یافته ها در بیماری های چندعاملی پیچیده تعمیم پذیر نبودند.

بررسی ارتباطی گسترده ژنومی اولین بار در سال ۱۹۹۶ توسط Merikangas جهت بررسی ژنتیکی بیماری های چندعاملی پیشنهاد شد (Risch and Merikangas, 1996) اما چون ایده آن اجرایی نبود، ۱۰ سال طول کشید تا اولین مطالعه از این نوع انجام شود. پروژه تعیین توالی ژنوم انسان^۴ نیز در سال ۱۹۹۰ آغاز شد اما ۱۳ سال به طول انجامید تا نهایتاً در سال ۲۰۰۳، دقیقاً ۵۰ سال بعد از شناسایی اولیه DNA توسط واتسون و کریک (WATSON and CRICK, 1953) گزارش کامل توالی ژنوم انسان منتشر شد. این گزارش نشان داد که ژنوم انسان حاوی حدود سه بیلیون جفت باز هست که این توالی می تواند حاوی ۲۰ تا ۲۵ هزار ژن کد کننده پروتئین باشند.

پیشرفت های اخیر در پایگاه داده های اختصاصی جایگاه ژنی، در این زمینه بسیار اثرگذار بوده است. در مطالعات ارتباط ژنوم گسترده کشف نشانگرهای زیستی، شناسایی افراد دارای ژن بیماری های خاص و شناسایی نوع واکنش به درمان به طور وسیع تری صورت گرفته است (Ginsburg, Donahue and Newby, 2005). نشانگرهای زیستی را می توان به نشانگرهای زیستی DNA، تومور و عمومی تقسیم کرد. البته این رده بندی دارای اشتراک نیز هست. از بین این سه گروه، نشانگرهای زیستی DNA مزایای زیادی نسبت به سایر نشانگرهای زیستی دارد.

به طوری که: در طول زندگی فرد ثابت است، نگهداری DNA آسان است، در هر زمانی قابل اندازه گیری هست، قابل تکثیر است، می تواند برای مطالعات گذشته و آینده نگر به کار رود و در بانک های زیستی هم معتبر است. در بین نشانگرهای زیستی، رایج ترین نوع تغییرات DNA پلی مورفیسم های تک نوکلئوتیدی هستند که عمدتاً دو آللی اند و سه نوع ژنوتیپ (هوموزیگوت مرجع، هوموزیگوت تغییر یافته و هتروزیگوت) را ایجاد می کنند. البته شناسایی و استفاده از یک نشانگر زیستی بستگی به نوع مطالعه دارد (Ziegler et al., 2012).

این مطالعات اغلب از نوع مورد شاهده^۵ی انجام می گردد؛ بدین معنا که گروهی از بیماران در مقابل افراد سالم انتخاب می شوند و تمامی این افراد برای قسمت عمده چند مورفیسم های تک نوکلئوتیدی شناخته شده (تعداد آن بستگی به فناوری مورد استفاده دارد اما به طور کلی یک میلیون و یا بیشتر است) تعیین ژن نمود می شوند، سپس میزان فراوانی الی هر کدام از این نشانگرها در گروه های مورد و شاهد مقایسه می شود. اساس گزارش

^۱ Genome Wide Association Study (GWAS)

^۲ Inheritance

^۳ Genetic linkage

^۴ Human Genome Project (HGP)

^۵ Case - Control

میزان اثر^۱ بر مبنای محاسبه نسبت شانس^۲ است بدین معنا که حضور یا عدم حضور یک آلل چقدر به حضور یا عدم حضور رخ‌نمود موردبررسی در یک جامعه‌ی مشخص و معلوم، مربوط است. هنگامی که فراوانی اللی در گروه مورد بیشتر از گروه شاهد است، نسبت شانس بیشتر از یک می‌شود. برای گزارش این نسبت آزمون کای‌دو^۳ جهت محاسبه p-value استفاده می‌شود. این نوع مطالعه به روش‌های دیگری نیز قابل انجام است. یکی از این روش‌ها بررسی رخ‌نمودهای کمی مانند قد، غلظت، نشانگرهای زیستی و یا بیان ژن است. در مسائل طراحی مورد شاهدی موضوعات مهمی وجود دارند (McCarthy et al., 2008):

۱. اندازه نمونه و در این موضوع دیدگاه روشن است: نمونه‌های بیشتر بهتر است که البته افزایش زیاد آن‌هم در مواردی راهگشا نیست.

۲. گرایش^۴ به زیر ساختار پنهان جمعی (رده‌بندی جمعیت و رابطه پنهان) که نرخ خطا نوع ۱ را افزایش می‌دهد و ادعاهای غلطی را درباره انواع مختلفی روابطی که در این زیر ساختار وجود دارد، بیان می‌کند. انتخاب جمعیت بهینه باید زیر ساختار جمعیت را محدود کند. در صورت امکان، چندین جمعیت کنترل باید مورد استفاده قرار گیرد و به‌صورت جداگانه انتخاب شوند تا بیشترین میزان شباهت به جمعیت بیماری را انتخاب کنند. به‌طور کلی، مطالعات آینده‌نگر، کنترل بیشتری را نسبت به مطالعات موردی، کنترل می‌کنند (Cardon and Bell, 2001).

۳. شایستگی نسبی روش‌های مرتبط با خانواده و ارتباطات مورد شاهدی. استراتژی‌هایی که از کنترل‌های خانوادگی استفاده می‌کنند و ارتباط را شناسایی می‌کنند می‌توانند ارزشمند باشند اما در استفاده از اطلاعات موجود ناکارآمد هستند. نتایج مثبت داده‌های معنی‌دار را فراهم می‌کنند، اما کمبود اهمیت آماری را ایجاد می‌کنند (Cardon and Bell, 2001).

۴. پتانسیل استفاده از ژن‌نمودهای کنترل تاریخ^۵ برای جایگزینی یا تکمیل کنترل‌های جدید تایپ‌شده در مطالعات GWA آینده

روش‌های آماری مورد استفاده بسته به میزان نفوذ بیماری و الگوی غالب و یا مغلوبی آن متفاوت است که معمولاً با ابزارهای بیوانفورماتیک مانند PLINK و SNPTEST انجام می‌شود. بسیاری از برنامه‌های استاندارد محاسبات آماری (R و SPSS) برای محاسبات اولیه قابل بهره‌برداری است اما در مواجهه با داده‌های GWAS با توجه به حجم بالای این نوع داده‌ها، بسیاری از برنامه‌های رایج غیر قابل استفاده خواهند بود. در این شرایط بسیاری از محققین از برنامه‌های متن‌باز و قابل تغییر استفاده می‌کنند که از جمله می‌توان به برنامه‌های PLINK و GenABEL اشاره کرد (Aulchenko, De Koning and Haley, 2007).

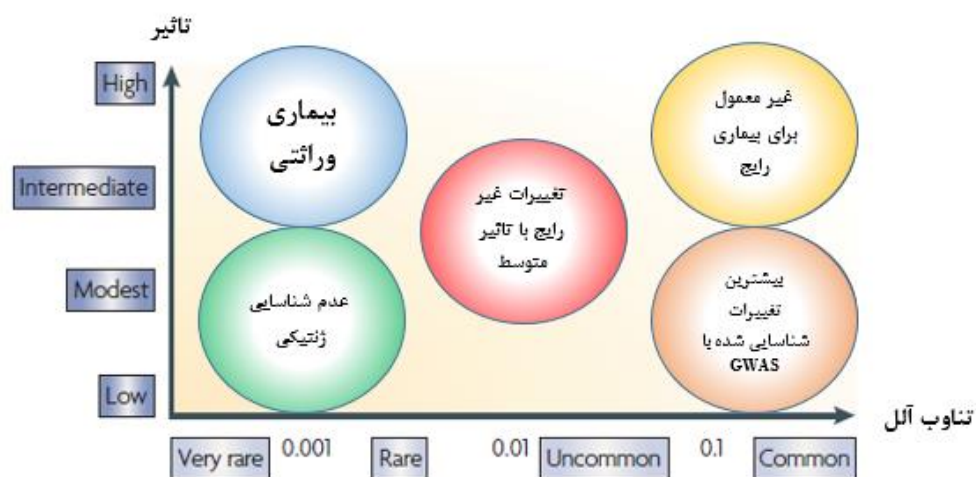
¹ Effect size

² Odd ratio

³ Chi squared

⁴ propensity

⁵ historical control genotypes



شکل ۱-۴ حوزه فعالیت تحقیقات GWAS (McCarthy et al., 2008)

حوزه فعالیت تحقیقات GWAS در شکل ۱-۴ آمده است (McCarthy et al., 2008). همان طور که در تصویر بیان شده است حوزه تحقیقات GWAS بر روی آللهایی با احتمال رخداد بالا (آلل رایج) و تأثیر متوسط به پایین هست. می توان گفت GWAS معمولاً برای جفت پایه DNA با فراوانی آلل جزئی (MAF) بیشتر از ۵٪ برای جمعیت مورد مطالعه تعریف می شود (Bush, 2019).

در نهایت می توان گفت صدها جهش در یک ژن منفرد منجر به بیماری های نادر می شود، اما اینکه چرا جهش ها حالت های شدید یا ضعیفی ایجاد می کنند، هنوز درک نشده است. ساخت مدل های بیماری ژنتیکی انسان، بخشی اساسی در درک بیماری، طراحی درمان ها و آزمایش روش های درمانی است. مطالعات ارتباط گسترده ژنوم در شناسایی صدها ژن زمینه ای و جهش های مرتبط قدرتمند بوده است؛ اما تجزیه و تحلیل عملکردی اکثریت قریب به اتفاق ژن ها و جهش های آن ها (حتی در ژن های کاملاً مشخص) وجود ندارد (Farris et al., 2021).

۱-۳-۱- تحلیل GWAS و وراثت پذیری

وراثت پذیری به نسبت وراثتی بودن یک رخ نمود به یک ژن نمود خاص است که می توان از طریق تقسیم واریانس ژن نمود به واریانس رخ نمود آن را محاسبه کرد. این عددی که با نماد h^2 شناخته می شود عددی بین صفر تا یک است که هر چه عدد به یک نزدیک باشد قابلیت توارث پذیری بیشتری را برای رخ نمود مورد نظر نشان می دهد.

اندازه گیری معیار h^2 یکی از مشکلاتی است که در روش های GWAS نیز بررسی شده اند. همان طور که اشاره شد، GWAS برای کشف ژن ها و مسیرهای مهم پزشکی به منظور شناسایی علل مولکولی و ژنتیکی بیماری ها طراحی شده است، اما به دلیل ناتوانی در توضیح وراثت پذیری بیشتر صفات های پیچیده مورد انتقاد قرار گرفته است (Maher, 2008). تجزیه و تحلیل داده های GWAS نشان داد که نسبت زیادی از وراثت پذیری برای صفات کمی مانند قد، شاخص توده بدنی^۱ و توانایی شناختی توسط SNP های رایج توضیح داده می شوند (Yang et

¹ Body Mass Index

(al., 2010, 2011; Davies et al., 2011) این نتایج نشان می‌دهد که بیشتر وراثت‌پذیری به‌جای گم‌شدن، پنهان‌شده است (Gibson, 2010) و GWAS، SNP‌هایی را شناسایی نکرده که این نسبت وراثت پنهان را توضیح می‌دهند، زیرا اندازه اثرات SNP‌های فردی برای رسیدن به سطح معنی‌داری دقیق ژنوم بسیار کم است (Yang et al., 2010, 2017). به‌جز GWAS، روش دیگری برای پراکندگی واریانس ژنتیکی بر روی کروموزوم‌ها و بخش‌های ژنومی وجود دارد. در این روش، واریانس مربوط به یک کروموزوم یا بخش DNA متناسب با طول آن، مخصوصاً برای قد و اسکیزوفرنی محاسبه‌شده است (Yang et al., 2013) و SNP‌هایی که در مناطق ژنی یافت می‌شوند، تغییرات بیشتری را نسبت به مناطقی که در ناحیه‌های بین ژنی مشاهده می‌شوند، نشان می‌دهند. تمام نتایج مطابق با الگوی ارثی چندژنی برای بیشتر ویژگی‌های پیچیده است.

علاوه بر این، طبق فرض بسیاری از ژن‌ها که دارای اثر کوچک هستند، انتظار می‌رود که نشانگر ژنتیکی بیشتری با افزایش اندازه نمونه مشخص شوند (Visscher et al., 2012). طیف وسیعی از همکاری‌های بین‌المللی وجود دارد که نمونه‌های بسیاری از گروه‌های GWAS را در یک فرا تحلیل ترکیب می‌کند تا قدرت آماری کارآمد را برای تشخیص اثرات کوچک فراهم کند. با این حال، برای یک فرا تحلیل، ژن‌نمود و اطلاعات رخ‌نمود سطح فردی به‌طور معمول در دسترس نیستند. برای یک فرا تحلیل بزرگ در مقیاس بزرگ، سازمان‌دهی و انجام یک تجزیه و تحلیل شرطی و غیرعملی برای انجام تجزیه و تحلیل شرطی به‌طور تکراری^۱ بسیار دشوار است. لازم به ذکر است در روش‌های آماری GWAS به‌جای محاسبه h^2 مقدار h^2_{SNP} را اندازه‌گیری می‌نمایند (Yang et al., 2013).

۱-۴- هم‌ترازی توالی ژنوم

در چند دهه اخیر، پیشرفت در زیست‌شناسی مولکولی و تجهیزات مورد نیاز آن باعث افزایش سریع تعیین توالی ژنوم بسیاری از گونه‌های موجودات شده است تا جایی که پروژه‌های تعیین توالی ژنوم‌ها از پروژه‌های بسیار رایج به حساب می‌آیند و هم‌ترازی توالی^۲ (SA) معمولاً اولین گامی برای درک فیلوژنی مولکولی یک توالی ناشناخته در بیوانفورماتیک است (Chowdhury and Garai, 2017).

این کار با هم‌ترازی توالی ناشناخته با یک یا چند توالی شناخته‌شده در پایگاه داده برای پیش‌بینی بخش‌های رشته توالی برای انجام نقش‌های عملکردی و ساختاری بر مبنای تکامل انجام می‌شود. هم‌ترازی مطلوب، دو یا چند توالی را به‌گونه‌ای تنظیم می‌کند که حداکثر تعداد عناصر یکسان یا مشابه را با هم مقایسه کند. توالی‌ها می‌توانند توالی‌های نوکلئوتیدی (DNA یا RNA) یا توالی‌های اسید آمینه (پروتئین) باشند (Reddy and Fields, 2022).

فرایند باز چینی ممکن است یک یا چند فضا یا شکاف در هم‌ترازی ایجاد کند. یک فاصله نشان‌دهنده از دست دادن یا افزایش احتمالی یک عنصر است بنابراین، درجه تکامل یا حذف^۳، انتقال و معکوس شدن در SA

¹ Iteratively

² Sequence Alignment

³ indel

مشاهده می‌شود. دو نوع هم‌ترازی توالی ترتیب عبارت‌اند از: توازن جفتی^۱ (PSA) و هم‌ترازی توالی چندگانه^۲ (MSA). در PSA دو توالی را در نظر می‌گیرد درحالی‌که MSA چندین (بیش از دو) توالی مرتبط را تراز می‌کند. فواید MSA بیشتر از PSA است، زیرا چندین عضو یک خانواده توالی را در نظر می‌گیرد و در نتیجه اطلاعات بیولوژیکی بیشتری ارائه می‌دهد. همچنین MSA پیش‌نیاز تجزیه و تحلیل ژنوم مقایسه‌ای برای شناسایی و اندازه‌گیری مناطق محافظت‌شده و یا نقاط عملکردی در یک خانواده کامل توالی و برآورد کننده انحراف تکاملی بین توالی‌ها و حتی برای پروفایل توالی اجدادی هست. ترتیب توالی در سطح اسیدآمینه بیشتر مربوط به سطح نوکلئوتید است زیرا پروتئین کلید عملکردی مولکول‌های بیولوژیکی است و از این رو اطلاعات ساختاری و یا عملکردی را حمل می‌کند. به این ترتیب هم‌ترازی، دارای پیوند قوی با زیست‌شناسی ساختاری است. از این رو، SA، به‌ویژه MSA، نقطه شروع هر زمینه تحقیق بیولوژیکی است. که MSA به‌عنوان یک پنجره برای مشاهده دیدگاه تکاملی، کاربردی و ساختاری ماکرو مولکول‌های بیولوژیک در قالب مختصر هست (Reddy and Fields, 2022).

روش‌های مختلفی برای امتیازدهی در هم‌ترازی توالی به‌منظور شناخت یکسانی یا تشابه استفاده می‌شود. نمره نوکلئوتید یک طرح شناسایی ساده است که در آن موقعیت‌های یکسان در هر دو توالی، امتیازات مثبتی به دست می‌دهند. در مقابل، برای پروتئین، یک نمره مشابهت نیز شمارش می‌شود (همراه با نمره یکسانی) که نشان‌دهنده اسیدآمینه با خواص فیزیکی و شیمیایی مشابه است. روش SA می‌تواند دو نوع باشد: هم‌ترازی سرتاسری و هم‌ترازی محلی. هم‌ترازی سرتاسری زمانی انجام می‌شود که شباهت در تمام طول دنباله شمارش شود. چندین فن MSA هم‌ترازی سرتاسری را انجام می‌دهند اما مشکلات زمانی به وجود می‌آیند که توالی‌ها تنها در مناطق محلی همخوانی داشته باشند، جایی که بلوک واضحی از ناپیوستگی‌های مشترک در همه‌ی توالی‌ها وجود دارد یا دامنه‌های تغییر یافته در میان توالی‌های مرتبط وجود داشته باشد. در چنین مواردی، هم‌ترازی محلی برای شناخت مناطق مشابه در میان توالی‌ها انجام می‌شود. وقتی تفاوت زیادی در طول توالی‌هایی که باید مقایسه شود وجود دارد، هماهنگی محلی معمولاً انجام می‌شود.

برای PSA، برنامه‌ریزی پویا^۳ (DP) همواره جهت‌گیری مطلوب برای یک تابع هدف مشخص را با یافتن راه بهینه تراز کردن با استفاده از فرایند ردیابی فراهم می‌کند. تابع هدف برای ارزیابی کیفیت هم‌ترازی مجموعه‌ای از توالی‌های ورودی استفاده می‌شود.

مشکلاتی محاسباتی در این روش‌ها بالاست و نیاز به منابع رایانه‌ای بالا دارد. علاوه بر این، DP از مشکلات بزرگ در MSA رنج می‌برد، زیرا تعداد توالی‌ها برابر با تعداد ابعاد است. اگر دو یا چند مسیر بهینه در دسترس باشند و نیاز به فرایند ردیابی باشد، پیچیدگی افزایش می‌یابد و محاسبه یک MSA دقیق NP-Complete است و بنابراین روش دقیق در MSA برای مجموعه‌های داده‌های کوچک غیر واقعی استفاده شده است.

¹ Pairwise Sequence Alignment

² Multiple Sequence Alignment

³ Dynamic Programming

بنابراین روش‌های معمول MSA به صورت اکتشافی یا تقریبی است که راه حل ممکن را در یک دوره زمانی کوتاه و محدود فراهم می‌کند. روش‌های اکتشافی موجود، بهترین راه حل را ارائه نمی‌دهند و به دلیل سرعت بالای رشد اندازه پایگاه‌های داده، توسعه فن جدید برای MSA هنوز تحت تحقیق است.

۱-۴-۱- روش‌های هم‌ترازی توالی چندگانه

به دلیل ارتباط پیچیده میان توالی‌های مرتبط و گاهی اوقات به دلیل فقدان تاریخ تکاملی، اغلب MSAها به سختی می‌توانند به توالی نهایی دست یابند. اغلب سه دسته از رویکردها در MSA استفاده می‌شود: رویکرد دقیق، پیشروند و تکرارپذیر. الگوریتم‌های دقیق معمولاً کیفیت بالایی را که بسیار نزدیک به مطلوب نهایی است ارائه می‌دهند. این روش‌ها تلاش می‌کنند به طور هم‌زمان چندین توالی را ترتیب دهند و بنابراین نیاز به DP دارند. از روش‌های دیگر می‌توان روش‌های درختی، ستاره‌ای و پیش‌رونده را نام برد که البته زمانی که طول توالی‌ها و نیز تعداد آن‌ها افزایش پیدا می‌کند، این روش‌ها چندان موفق نیستند.

به جهت اهمیت مسئله مدل‌سازی MSA در مقاله (Lajevardy and Kargari, 2022) مدلی مبتنی بر مدل‌سازی ریاضی مسئله MSA به همراه الگوریتم ژنتیک آن منتشر شده است.^۱

۱-۵-۱- فشارخون^۲

به فشار مایع درون رگ‌های خون، فشارخون می‌گویند. هم‌چنین به فشارخون شریانی توزیع شده در کل محفظه شریانی، فشارخون سیستمیک می‌گویند که توسط عوامل مختلفی تعیین می‌شود: (۱) نیروی خونی که با انقباض بطن چپ قلب وارد محفظه شریانی می‌شود. (۲) میزان گردش خون از محفظه شریانی به مویرگ‌ها. (۳) حجم کل خون. (۴) تنش ایجاد شده توسط دیواره‌های بزرگ‌ترین رگ‌های خونی در مقاومت به پالس خون خارج شده از عروق توسط قلب.

۱-۵-۱- فشارخون سیستمیک^۳

به حداکثر فشارخون شریانی که در هنگام انقباض بطن چپ و خارج شدن خون به محفظه شریانی اندازه‌گیری می‌شود، فشارخون سیستمیک می‌گویند. این فشار در میلی‌متر جیوه سنجیده می‌شود.

۱-۵-۲- فشارخون دیاستولیک^۴

به حداقل فشارخون شریانی که در مرحله استراحت قلب در شریان‌های بزرگ اندازه‌گیری می‌شود، فشارخون دیاستولیک می‌گویند. این فشار در میلی‌متر جیوه اندازه‌گیری می‌شود و تا حد زیادی توسط میزان جریان خون از محفظه عروقی شریانی به مویرگ‌ها تعیین می‌شود.

^۱ <https://github.com/l8026070/MSA>

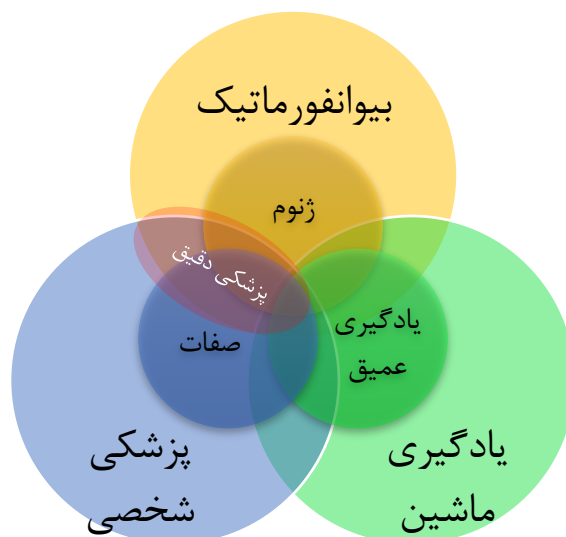
^۲ Blood Pressure

^۳ Systolic Blood Pressure (SBP)

^۴ Diastolic Blood Pressure (DBP)

۶-۱- حوزه پژوهش

این تحقیق بر آن است که بر اساس ژنوم انسانی که از فن‌های مختلف بیوانفورماتیک تهیه می‌گردد و داده‌های طولی صفت پرفشاری خون، مدلی مبتنی بر فن‌های یادگیری ماشین خصوصاً یادگیری عمیق طراحی و ارائه نماید. این مدل می‌تواند بر اساس اطلاعات جمعیت شناختی بهینه گردد. نمودار حوزه پژوهش در شکل ۵-۱ آمده است.



شکل ۵-۱ نمودار حوزه پژوهش

۷-۱- ضرورت و اهداف پژوهش

بر اساس تحقیقات اشاره شده در بخش‌های قبل، مطالعه گسترده ارتباط ژنومی اهمیت بالایی در پزشکی شخصی دارد که آینده پزشکی نوین و نشان از اهمیت بسزای آن دارد. بر همین اساس در این پژوهش اهداف کلی زیر پیگیری می‌شود:

✓ ارائه مدل مبتنی بر یادگیری ماشین در داده‌های حجیم جهت پیش‌بینی رخنمودهای پرفشاری خون بر اساس SNP در رخنمودهای:

○ فشارخون

▪ فشار سیستولی SBP و فشار دیاستولی DBP

▪ خطر پرفشاری خون

✓ بهبود مدل مبتنی بر یادگیری ماشین جهت پیش‌بینی داده‌های طولی

✓ بهبود مدل بر اساس اطلاعات جمعیت شناختی

✓ انتخاب SNP با اهمیت بر اساس دقت و تعداد

با توجه به داده‌های ژنومی که دارای حجم بالایی هست، هدف نهایی این پژوهش، پیش‌بینی خطر پرفشاری خون بر اساس SNP و با روش‌های یادگیری ماشین در داده‌های حجیم مانند داده‌های توالی ChIP (مجموعه‌ای مشخص از SNP‌های توالی یابی شده برای افراد نمونه) و یا توالی یابی کل ژنوم است.

۱-۷-۱- سؤالات تحقیق

بر اساس اهداف تعیین شده سؤال اصلی تحقیق به صورت زیر هست:

✓ آیا می توان از روش های یادگیری ماشین در پیش بینی خطر پرفشاری خون مبتنی بر SNP به شکل مؤثرتر از روش های موجود بهره برد؟

سؤال فرعی تحقیق نیز به صورت زیر هست:

✓ آیا می توان اطلاعات جمعیت شناختی و داده های طولی رخ نمود را در بهبود مدل پیش بینی کننده استفاده کرد؟

✓ با توجه به مدل ارائه شده، آیا می توان SNP های مهم ژنومی را استخراج کرد؟

۱-۸- نوآوری تحقیق

همان طور که در ادبیات موضوع اشاره گردید، تحقیقات موجود در حوزه GWAS در حوزه آماری متمرکز است که این امر به واسطه رده بندی جمعیت، وابستگی ژنتیکی، تشخیص^۱ و سایر منابع ناهمگونی منجر به سیگنال های جعلی و کاهش توان در مطالعات ارتباط ژنتیکی می شود (Runcie and Crawford, 2018). همچنین مطالعاتی که تعداد اندکی SNP مرتبط با بیماری را در نظر می گیرند نسبت به مطالعات سنتی، رده بندی خطر ناکافی و غیر معینی را به دست می آورند (Khan et al., 2020). به همین دلیل در این تحقیق تعداد زیادی SNP با روش های یادگیری ماشین در نظر گرفته می شوند.

رویکرد پردازش داده و یادگیری ماشین خصوصاً یادگیری عمیق (با توجه به داده های حجیم ژنومیک) باعث بهبود خطاهای ذکر شده است (Esteva et al., 2019). همچنین می توان با بهره گیری از لایه های مختلف در شبکه یادگیرنده داده های مختلف omic را به مدل افزود که باعث افزایش دقت مدل می گردد (Hebbring, 2019; Xu et al., 2019). دستاورد این تحقیق با توجه به راهکار یادگیری ماشین، امکان برطرف کردن مشکلات ذکر شده برای روش های آماری و افزایش دقت را خواهد داشت (Sun et al., 2019; Yin et al., 2019). لازم به ذکر است انتخاب فشارخون به دلیل اهمیت آن و همچنین چالش موجود در ادبیات موضوع در خصوص پیش بینی خطر بیماری در جوانان هست (Wu et al., 2020; Phulka et al., 2023).

۱-۹- جمع بندی

در این فصل به بیان کلیات و مقدمه موضوع پرداخته شد و در نهایت بایان اهداف و ضرورت تحقیق به پایان رسید. در فصل بعد با بررسی ادبیات موضوع روش های مختلف حل مسئله گفته شده بیان می گردد.

¹ ascertainment

۲- ادبیات نظری و پیشینه پژوهش

بسیاری از رویکردهای احتمالی برای ارزیابی ژن‌هایی که ریشه بیماری‌های شایع و صفات کمی هستند، به‌طور گسترده‌ای به دودسته تقسیم می‌شوند (Hirschhorn and Daly, 2005): مطالعات ژن‌های منتخب، که از روش‌های همبستگی یا باز توالی^۱ استفاده می‌کنند و مطالعات گسترده ژنومی شامل هر دو نوع نقشه‌سازی پیوند^۲ و مطالعات روابط گسترده ژنومی^۳ هست. رویکردهای اصلی و مزایا و معایب آن‌ها در جدول ۱-۲ خلاصه‌شده است (Hirschhorn and Daly, 2005).

جدول ۱-۲ رویکردهای تشخیص تغییرات مبتنی بر صفات پیچیده^۴ در بیماری‌های رایج (Hirschhorn and Daly, 2005)

مزایا	Association*	Resequencing*	Linkage†	Admixture†	Missense SNPs†	Association†	Resequencing†
هیچ اطلاعات پیشین در مورد عملکرد ژنی	-	-	+	+	+	+	+
مورد نیاز نیست	+	+	-	-	+	+	+
محلی سازی به منطقه ژنومی کوچک	+	-	+	+	+/-	-	Prohibitive
ارزان	+	+	-	+	+	+	+
خانواده‌ها مورد نیاز نیست	+	-	+	+	-	+	+
هیچ فرضیه‌ای در مورد نوع درگیر وجود ندارد	+/+	-/+	+	+	-/+	-/+	-/+
حساس به اثرات رده بندی نیست	+	+	+	-	+	+	+
نیازی به تغییر فرکانس آلل در میان جمعیت نیست	+	-	-/+	+	+	+	+
قدرت کافی برای تشخیص آلل‌های معمولی (%5<MAFs)	-	+	+	-	-	-	+
توانایی تشخیص آلل‌های نادر (%1>MAFs)	+	-/+	+/+	N/A	N/A	N/A	N/A
رکوردهای منطقی برای بیماری‌های شایع	+	+	+	+	+	+/+	-
ابزار برای تجزیه و تحلیل در دسترس است	+	+	+	+	+	+/+	-

باوجود روش‌های مختلف، اما بااین‌حال، آلل‌های نادر معمولاً بسیار سخت توسط روش‌های کشف روابط به دلایل مختلفی، شناسایی می‌شوند (Visscher *et al.*, 2014) که در ادامه به استراتژی‌های افزایش کارایی این روش‌ها پرداخته می‌شود.

۲-۱- استراتژی‌های افزایش بهره‌وری

در این بخش، در مورد استراتژی‌هایی که برای اجرای یک مطالعه گسترده ارتباط ژنومی پیشنهاد شده‌اند بحث می‌شود (Hirschhorn and Daly, 2005) که با به حداقل رساندن تعداد ژن‌نمود مورد نیاز، قدرت کافی درزمینه‌ی آزمایش آزمون‌های چندگانه را دارند. این استراتژی‌ها شامل طرح‌های چندمرحله‌ای برای به حداقل رساندن اندازه نمونه، استفاده از جمعیت خاص و جمع‌آوری نمونه‌ها هست. یکی دیگر از روش‌های بهبود، استفاده از اپی ژنتیک هست (Shah *et al.*, 2015) که در حوزه این تحقیق نیست.

¹ Resequencing

² Linkage Mapping

³ GWAS

⁴ Complex Trait

۲-۱-۱- انتخاب یک اندازه نمونه مناسب

واضح‌ترین محدودیت در مطالعات گسترده ژنومی هزینه‌های بالا و تلاش‌های قابل‌توجهی است که نیاز به ژن‌نمود صدها هزار SNPs در هر فرد است. به دلیل این هزینه بالا، فشار برای محدود کردن حجم نمونه وجود دارد که همراه با کاهش در قدرت نتایج است. برای درک پیامدهای این آستانه، یک آلل با بسامد ۱۵٪ و یک نسبت شانس از ۱.۲۵ را در نظر بگیرید. برای چنین نوعی، حتی با فرض اینکه SNP علی (یا SNP دیگری که به‌عنوان یک پروکسی کامل عمل می‌کند) تایپ‌شده است، تقریباً ۶۰۰۰ مورد و ۶۰۰۰ کنترل برای ۸۰ درصد قدرت آماری تشخیص ارتباط با مقدار ارزش P برابر 5×10^{-8} برای ۵۰۰,۰۰۰ SNP مستقل موردنیاز است که این اندازه نمونه ۶ میلیارد ژن‌نمود نیاز دارد که هزینه‌های بالایی دارد.

۲-۱-۲- رویکرد چندمرحله‌ای

یک روش ساده برای بهبود، استفاده از یک فرآیند غربالگری دو یا سه مرحله‌ای است که در آن، آستانه معکوس بیشتری برای مثبت شدن نتیجه، در هنگام ارزیابی اسکن اولیه ژنوم ارتباط استفاده می‌شود.

۲-۱-۳- جمعیت اولیه^۱ و نمونه‌های جمع‌آوری‌شده

روش‌های دیگر پیشنهادشده که باعث افزایش کارایی مطالعات ارتباطی می‌شود استفاده از جمعیت اولیه هست که باعث کاهش تعداد نشانگرهایی که نیاز به ژن‌نمود دارند، و نمونه‌های جمع‌آوری‌شده که تعداد نمونه‌های ژن‌نمود را کاهش می‌دهند، می‌شود. جمعیت اولیه به‌صورت جمعیت‌هایی که از یک مجموعه محدود از افراد در ۱۰۰ نسل گذشته یا کمتر ایجادشده‌اند، تعریف می‌شود. جمعیت اولیه مزایای خوبی برای ژن‌هایی که منجر به اختلالات مندلی می‌شوند دارند، اما ممکن است برای بیماری‌های رایج کمتر مؤثر باشند.

۲-۱-۴- اجتناب از گرفتن نتایج مثبت کاذب

در یک مطالعه که صدها هزار نشانگر را شامل می‌شود، به حداقل رساندن مثبت کاذب ضروری است. منابع ایجاد ارتباطات مثبت کاذب را می‌توان به سه دسته اصلی تقسیم کرد: نوسانات آماری که به‌وسیله شانس به وجود می‌آیند و در نتیجه باعث کم شدن p-value می‌شوند، انحرافات سامانمند پایه با توجه به طراحی مطالعه؛ و تجهیزات فنی.

روش‌های موجود در مطالعات گسترده ژنومی را می‌توان در سه رویکرد: آماری، فرا تحلیل و یادگیری ماشین رده‌بندی کرد که در ادامه به بررسی آن‌ها پرداخته می‌شود.

۲-۲- رویکرد آماری

رویکرد آماری، رویکرد غالب در تحلیل ارتباطات گسترده ژنومی هست. این تحلیل برآمده از رابطه زیر هست:

$$\text{رخ‌نمود} \rightarrow \text{تنوع تصادفی} + \text{ژن‌نمود} + \text{محیط}$$

در ادامه به بررسی معروف‌ترین مدل‌های ارائه‌شده در این حوزه پرداخته می‌شود.

¹ Founder populations

۲-۲-۱-۱ مدل GRAMMAR^۱

این مدل در سال ۲۰۰۷ ارائه گردید (Aulchenko, De Koning and Haley, 2007). در مرحله اول، داده‌ها تحت مدل ترکیبی^۲ رابطه ۱-۲ مورد تجزیه و تحلیل قرار می‌گیرند.

$$y_i = \mu + \sum_j \beta_j c_{ji} + G_i + e_i$$

رابطه ۲-۱ مدل GRAMMAR

که y_i رخ نمود فرد i ام است، μ میانگین رخ نمود، c_{ij} مقدار هم وردایی j ام یا اثر ثابت برای فرد i ام است (مثلاً جنسیت یا سن)، β_j برآورد از اثر ثابت یا هم وردایی است، G_i اثرات تصادفی چندژنی و e_i خطای باقی مانده است.

این روش مبتنی بر اخذ مجدد از یک مدل چندژنی است و بر اساس تجزیه و تحلیل ارتباط این باقیمانده‌ها با چند مورفیسیم‌های ژنتیکی با استفاده از روش‌های کلاسیک برای افراد مستقل است.

۲-۲-۲ مدل GCTA^۳

در تحلیل‌های آماری با رویکرد جزئی خطاهای متعددی ایجاد می‌گردد. یکی از این خطاها وراثت گم شده^۴ هست که به دلایل مختلفی من جمله تعداد زیاد از انواع تغییرات رایج با اثر کم و یا انواع نادر تغییرات با اثر بالا و یا تغییرات ساختاری DNA است. در روش GCTA (Yang et al., 2013) برخلاف تجزیه و تحلیل روابط تک SNP، به اثرات همه SNPs به عنوان اثرات تصادفی با یک مدل خطی آمیخته (MLM^۵): بر اساس فیشر (۱۹۱۸) نگاه می‌شود و رابطه ۲-۲ آن را بیان می‌کند.

$$y = X\beta + Wu + \varepsilon \text{ with } \text{var}(y) = V = WW'\sigma_u^2 + I\sigma_\varepsilon^2$$

رابطه ۲-۲ مدل GCTA

در این رابطه y یک بردار از رخ نمودها با n نمونه است، β یک بردار از اثرات ثابت مانند جنسیت، سن و (یا) یک یا چند بردار اطلاعاتی خاص از تجزیه و تحلیل مؤلفه‌های اصلی (PCA^۶)، u یک بردار از اثرات SNP، I یک ماتریس همانی است و e یک بردار از اثرات باقی مانده خطا است. W یک ماتریس ژن نمود استاندارد شده و برآورد شده با حداکثر احتمال (Gilmour, Thompson and Cullis, 1995) است.

۲-۲-۳ مدل ACTA^۷

صفات پیچیده با تعداد زیادی از عوامل ژنتیکی و محیطی و همچنین تعاملات آن‌ها تعیین می‌شود. شناسایی ژن‌های مشارکتی و اندازه‌گیری اثرات آن‌ها در زمینه‌ی یک یا چند محیط، اهمیت اساسی در اهداف درمانی برای بیماری حیوانات و انسان و درک نحوه انتخاب طبیعی و مصنوعی ژنوم حیوانات و انسان را می‌دهد.

^۱ Genome wide Rapid Association Using Mixed Model and Regression

^۲ Mixed Model

^۳ Genome-wide Complex Trait Analysis

^۴ Missing Heritability

^۵ Mixed Linear Model

^۶ Principal Component Analysis

^۷ Advanced Complex Trait Analysis

نرم افزار GCTA (Yang *et al.*, 2013) یک بسته نرم افزاری رایگان C++ برای اندازه گیری ارتباط تنوع ژنتیکی به تنوع رخ نمودی برای ویژگی های پیچیده است که شامل دو مرحله است: ارزیابی «ماتریس رابطه ژنتیکی»^۱ (GRM)، که حاوی اندازه گیری همبستگی ژنتیکی بین افراد است و با استفاده از داده های تک ژن نمود چند مورفیسیم (SNP) محاسبه و استفاده از حداکثر درست نمایی محدود شده^۲ (REML) برای به حداکثر رساندن احتمال رخ نمودها به GRM داده می شود. پیاده سازی الگوریتم GCTA، زمانی که در پردازنده های چند هسته ای مدرن اجرا می شود، کارآمد نیست.

در مدل ACTA (Gray, Stewart and Tenesa, 2012) بهبود عملکرد چشمگیری ارائه می شود. در این مدل ابتدا مدل خطی آمیخته^۳ (MLM) که توسط GCTA استفاده می شود، خلاصه می شود و از رابطه ۲-۳ استفاده می شود.

$$y = X\beta + Wu + \epsilon \text{ with } V = WW^T\sigma_u^2 + I\sigma_\epsilon^2$$

رابطه ۲-۳ مدل ACTA

در این رابطه، i و j همیشه نشانگرهایی برای مجموع افراد P هست، و k نشان دهنده یک SNP در مجموع SNP ها است. GCTA شامل تأثیرات مناسب SNP با استفاده از MLM است و y_i رخ نمود فرد i است و u یک بردار از اثرات SNP با توزیع نرمال است. U_k اثر رخ نمودی است که به k امین SNP مربوط می شود. W ماتریس رخداد است که اثرات SNP را به Y می دهد. این ماتریس شامل ژن نمود برای هر فرد است و توسط رابطه زیر محاسبه می شود:

$$W_{ik} = \frac{(r_{ik} - 2p_k)}{\sqrt{2p_k(1 - p_k)}}$$

رابطه ۲-۴

که R_{ik} تعداد کپی های آلل مرجع برای k امین SNP از فرد i ام، و P_k فراوانی آلل مرجع است. به طور مشابه، X یک ماتریس بروز برای بردار اثرات ثابت است. β یک بردار از اثرات باقی مانده با توزیع نرمال است.

۲-۲-۴ مدل REACTA^۴

این مدل بر پایه مدل ACTA بنا شده است و بر اساس استفاده از GPU سرعت پردازش اطلاعات را به طور متوسط از دو هفته در یک رخ نمود خاص، به ۱۰ دقیقه کاهش داده است (Cebamanos *et al.*, 2014). لازم به ذکر است حجم اطلاعات برای انجام یک GWAS در حدود ۴۰ گیگابایت هست که منجر به نیاز به حافظه بالا برای اجرای الگوریتم های آماری هست.

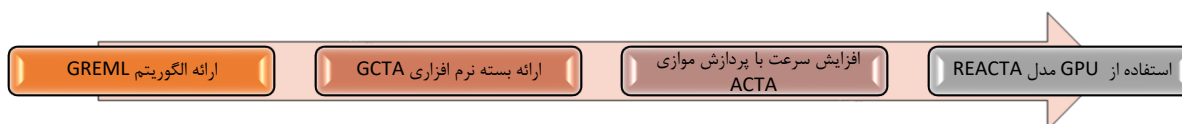
روند تحقیقات آماری در حوزه GWAS به صورت شکل ۲-۱ هست.

¹ Genetic Relation Matrix

² Restricted Maximum Likelihood

³ Mixed Liner Model

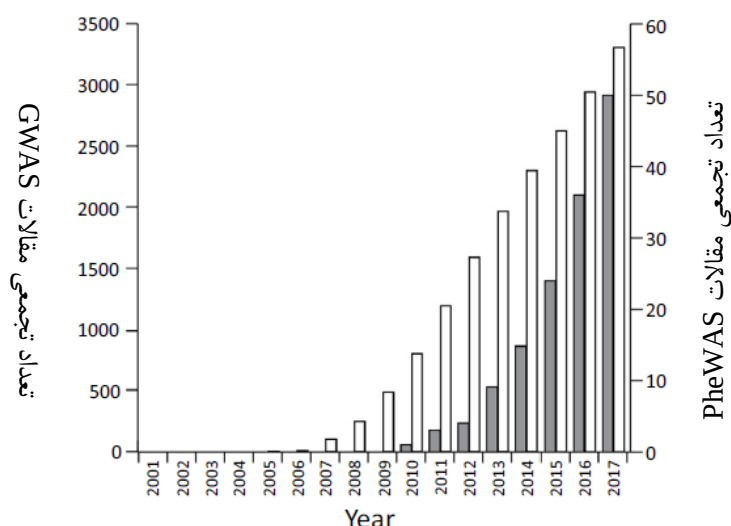
⁴ Regional heritability advanced complex trait analysis



شکل ۲-۱ روند زمانی و تکامل تحقیقات آماری بررسی شده در GWAS

۲-۲-۵- مدل PheWAS^۱

این مدل یک طرح کلی مطالعه یا روش آماری برای شناسایی رخنموده‌های مرتبط با یک ژن‌نمود هست. تحقیقات PheWAS اغلب به اطلاعات رخنمودی عمیق و کامل، از جمله داده‌های سامانه‌های سلامت الکترونیکی^۲ یا داده‌های شیوع شناسی، برای تعیین وضعیت مورد شاهی برای انواع رخنمودها تکیه می‌کنند (Hebbring, 2019). در این تحقیقات سؤال اصلی این است که چگونه می‌توان داده‌های رخنمودی را با ژن‌نمودی و دیگر لایه‌های داده‌های امیک ترکیب کرد تا بتوان به شناخت بهتر رسید و درمان بیماری‌های انسانی را فهمید و در این مسئله چالش‌های جدیدی در مدیریت داده‌ها، مسائل مربوط به حریم خصوصی و مراقبت‌های بالینی وجود دارد. روند این تحقیقات در شکل ۲-۲ آمده است.



شکل ۲-۲ روند تحقیقات PheWAS در مقایسه با تحقیقات GWAS (Hebbring, 2019)

۲-۲-۶- مدل OSCA^۳

افزایش سریع داده‌های اومیک در دهه‌های گذشته، بررسی ارتباطات بین پروفایل‌های اومیک مانند متیلاسیون در اپی ژنومیکس و ویژگی‌های پیچیده در گروه‌های بزرگ را تسهیل کرده است. این پیشرفت‌ها عبارت‌اند از شناسایی تعداد زیادی از انواع ژنتیکی مرتبط با بیان ژن، متیلاسیون DNA، اصلاح هیستون و فراوانی پروتئین؛ کشف اقدامات اومیک مربوط به صفات پیچیده؛ بهبود دقت در پیش‌بینی صفت با استفاده از داده‌های اومیک؛ و اولویت‌بندی ژن‌های هدف برای صفات پیچیده با تلفیق داده‌های ژنتیکی و

^۱ Phenome-wide association study

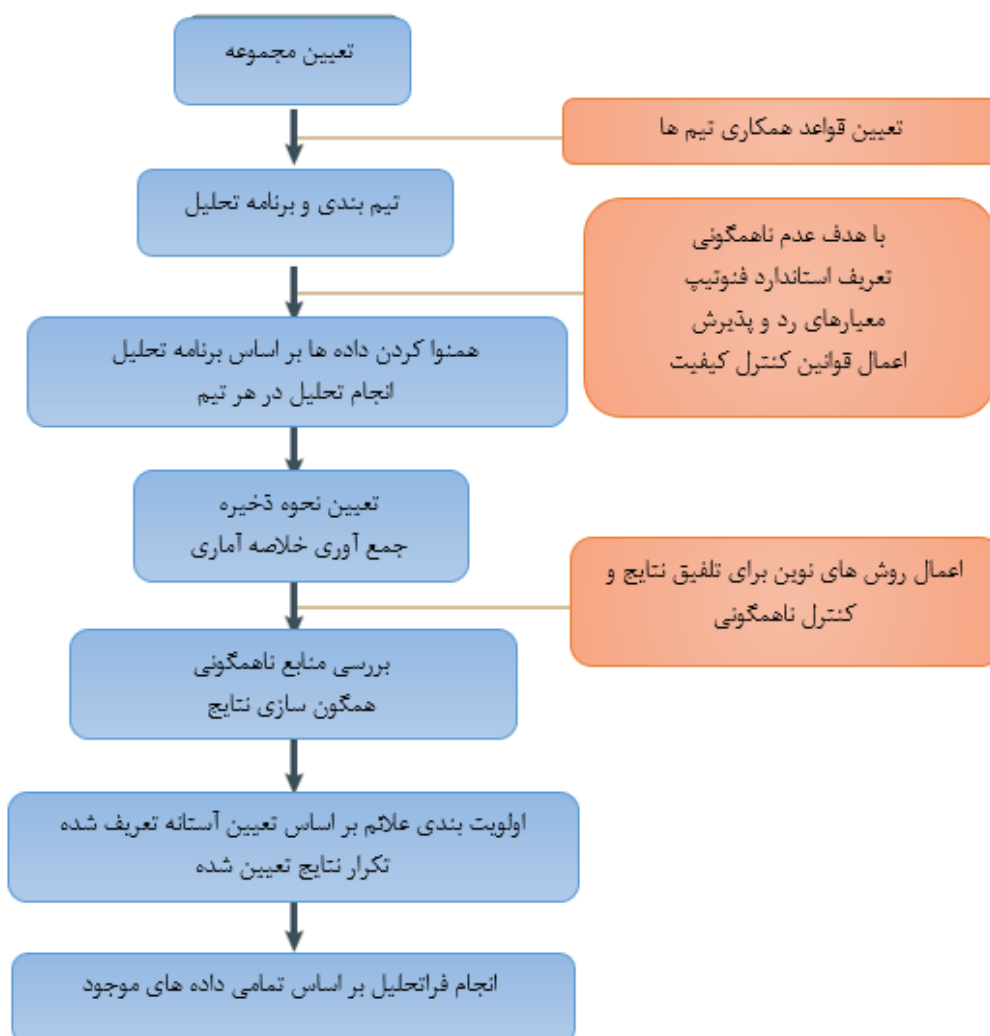
^۲ Electronic Health Record

^۳ omic-data-based complex trait analysis

امیک در نمونه‌های بزرگ هست (Zhang et al., 2018). در این مدل اطلاعات مختلف اومیک در مدل آماری خطی MLM پیاده‌سازی و تأثیر عوامل مختلف بر روی رخ‌نمود بررسی شده است.

۲-۳- رویکرد فرا تحلیل

با گسترش تحقیقات GWAS و گزارش روابط متعدد بین عناصر ژنوم و صفات، تحقیقاتی مبتنی بر GWAS های انجام‌شده شکل گرفت که به فرا تحلیل^۱ شناخته می‌شوند. مراحل فرا تحلیل در تصویر ۲-۳ آمده است. در این تصویر مراحل کار با کادر آبی و نکات مهم در هر مرحله با کادر قرمز نمایش داده شده است.



شکل ۲-۳ مراحل مدل های فرا تحلیل (Evangelou and Ioannidis, 2013)

در روش‌های فرا تحلیل از مقادیر مختلفی مانند p -value و اندازه نمونه استفاده می‌شود. لازم به ذکر است p -value تا دهه ۱۹۸۰ به‌طور گسترده‌ای در زمینه‌ی های مختلف علمی استفاده می‌شده است، اما پس از آن در علوم زیست پزشکی به شاخصی غیر رایج و تقریباً ممنوع شده، تبدیل شده است (Evangelou and Ioannidis, 2013). در ادامه مهم‌ترین انواع فرا تحلیل بررسی می‌شوند.

¹ Meta-Analysis

۲-۳-۱- مدل METAL

فرا تحلیل در حال تبدیل شدن به یک ابزار فزاینده‌ی مهم در مطالعات ارتباط ژنومی (GWAS) بیماری‌ها و صفات ژنتیکی پیچیده است (de Bakker *et al.*, 2008) که می‌تواند یک راهبرد کارآمد و عملی برای تشخیص انواع تغییر با ابعاد متوسط باشد (Skol *et al.*, 2007). مدل متال (Willer, Li and Abecasis, 2010) به معرفی متغیرهای مختلف در بررسی روابط ژن-صفت می‌پردازد که کاربرد فراوانی در مدل‌های فرا تحلیل دارد (شکل ۲-۴).

راهبرد تحلیل

	مبتنی بر معکوس واریانس	مبتنی بر حجم نمونه
ورودی	β_i - effect size estimate for study i	N_i - sample size for study i P_i - P -value for study i Δ_i - direction of effect for study i
آمار میانی	se_i - standard error for study i $w_i = 1/SE_i^2$ $se = \sqrt{1/\sum_i w_i}$ $\beta = \sum_i \beta_i w_i / \sum_i w_i$	$Z_i = \Phi^{-1}(P_i/2) * \text{sign}(\Delta_i)$ $w_i = \sqrt{N_i}$
Overall Z-Score	$Z = \beta/se$	$Z = \frac{\sum_i Z_i w_i}{\sqrt{\sum_i w_i^2}}$
Overall P-value	$P = 2\Phi(- Z)$	

شکل ۲-۴ پارامترهای معرفی شده در مدل METAL (Willer, Li and Abecasis, 2010)

۲-۳-۲- مقایسه مدل‌های فرا تحلیل

پس از انجام پروژه‌های GWAS در نقاط مختلف دنیا، مدل‌های متعددی به منظور فرا تحلیل منتشر شده است. در جدول ۲-۲ خلاصه‌ای از مهم‌ترین مدل‌هایی که تبدیل به بسته‌های نرم‌افزاری شده‌اند به همراه خواص آن‌ها بیان شده است (Evangelou and Ioannidis, 2013).

جدول ۲-۲ مقایسه بسته‌های نرم‌افزاری مدل‌های فرا تحلیل (Evangelou and Ioannidis, 2013)

	METAL	GWAMA	MetABEL	PLINK	R packages
قابلیت پردازش فایل‌های خروجی نرم افزارهای GWAS	No	Yes; SNPTTEST, PLINK	Yes; ABEL	Yes; PLINK	No
پیاده سازی تاثیر ثابت	Yes	Yes	Yes	Yes	Yes
پیاده سازی تاثیر تصادفی	No	Yes	No	No	Yes
تولید معیار ناهمگونی	Q, I^2	Q, I^2	Q, I^2	Q, I^2	Q, I^2
مصورسازی نتایج فراثتحلیل	No	Manhattan and QQ plots	Forest plots	No	Yes

۲-۴- رویکرد یادگیری ماشین

رویکردهای یادگیری ماشین برای تکمیل روش‌های تک SNP و چند SNP مؤثر در درک ساختار ژنتیک کلی بیماری‌های پیچیده انسان کاملاً امیدوارکننده است. بااین حال، به دلیل اینکه این روش‌ها برای داده‌های SNP

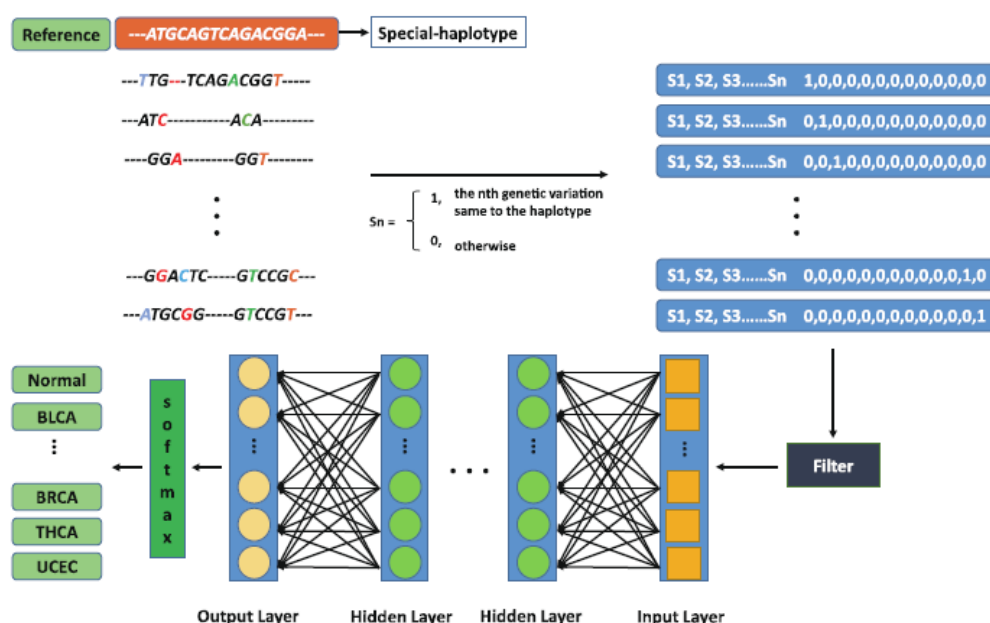
در ژنوم بهینه نشده‌اند، پیاده‌سازی‌های بهبود یافته و روش‌های انتخاب جدید متغیر مورد نیاز است (Szymczak et al., 2009).

در بررسی اثر مشترک چند متغیر ژنتیکی و محیطی، چالش‌های متعددی وجود دارد. اول، در GWAS، ژن‌نمودها تا یک میلیون SNP در چند هزار نفر تعیین می‌شود، که منجر به مشکل کوچک بودن نمونه و بزرگ بودن تعداد بعد می‌شود. دوم، هنگامی که تعداد زیادی از SNP‌های ژن‌نمود در یک مقیاس ژنومی هستند، باید عدم تعادل ارتباط (LD) بین SNP‌ها (که منجر به متغیرهای همبسته) شود، مورد توجه قرار گیرد. به همین علت، روش‌های آماری چند متغیره استاندارد مانند رگرسیون خطی چندگانه یا لجستیک برای داده‌های ژنومی مناسب نیستند (Szymczak et al., 2009).

در نتیجه، به رغم محدودیت‌های فعلی، روش‌های یادگیری ماشین برای تکمیل پایه ژنتیکی بیماری‌های پیچیده، تکمیلی معقول برای آزمون استاندارد SNP است.

۲-۴-۱- مدل GDL^۱

صفات بیولوژیک نتیجه تعاملات بین توالی ژن و تعامل ژن با شرایط خاص محیطی است. مدل یادگیری عمیق برای مطالعه رابطه بین این عوامل و رخ‌نمود مناسب است. مدل GDL مبتنی بر یادگیری عمیق بر اساس اطلاعات ژنومی در پیش‌بینی سرطان طراحی شده است (Sun et al., 2019). تصویر ۲-۵ دیاگرام عملکرد این مدل را نشان می‌دهد.



شکل ۲-۵ دیاگرام عملکرد مدل GDL (Sun et al., 2019)

در این مدل رشته ورودی اطلاعات ۱۰ هزار SNP مرتبط با ۱۲ نوع سرطان است که برای هر نمونه با ژن مرجع مقایسه و تبدیل به رشته ۱۰^۰ می‌شود. جهت برچسب نیز از روش یک داغ^۲ استفاده شده است. به این

^۱ Genome Deep Learning

^۲ One Hot

صورت که برچسب تولیدشده شامل ۱۲ صفر است که در صورتی که نمونه دارای سرطان خاصی باشد تنها یکی از صفرها به صورت یک ثبت می شود. دقت تحقیقاتی این مدل در تشخیص سرطان، ۷۰ درصد بوده است. خلاصه مدل های ارائه شده به تفکیک رویکرد و ارائه مدل در جدول ۲-۳ آمده است.

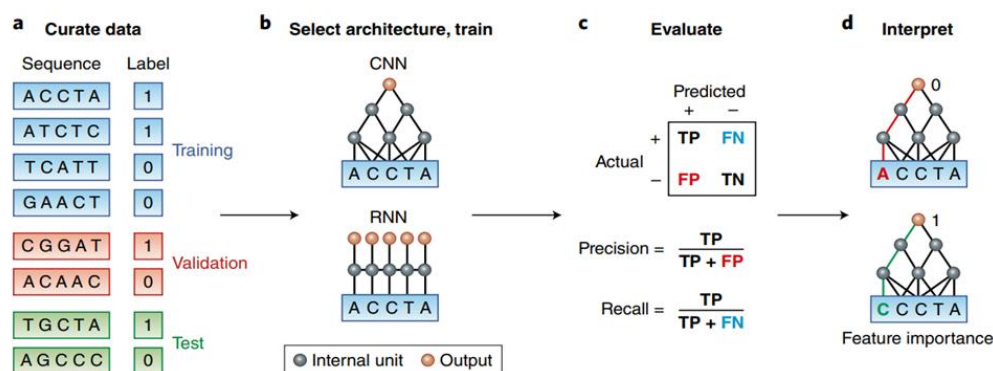
جدول ۲-۳ خلاصه مدل های بررسی شده

رویکرد	نام	سال	ارائه دهنده
آماري	GRAMMER	2007	et al.Aulchenko
	GCTA	2011	et al.Yang
	ACTA	2012	et al.Gray
	REACTA	2014	et al.Cebamanos
	PheWAS	2018	Hebbring
	OSCA	2018	et al.Zhang
فرا تحلیل	METAL	2010	et al.Willer
یادگیری	PENALIZED REGRESSION	2009	et al.Szymczak
ماشین	GDL	2019	et al.Sun

۲-۵- روش های یادگیری ماشین سنتی

یادگیری ماشین^۱ یکی از زمینه های مهم هوش مصنوعی است که دستگاه های رایانه ای قدرتمندی را برای توصیف یادگیری و اجرای الگوریتم ها و مدل ها با دقت بی نظیر فراهم می کند. یک مزیت عمده یادگیری ماشین توانایی یادگیری و ساخت الگوریتم های بدون دخالت انسان است. علاوه بر این، دقت تجزیه و تحلیل یادگیری ماشین با افزودن داده های آموزش بهبود می یابد (Khorraminezhad et al., 2020). فن های یادگیری ماشین به طور گسترده ای در تحقیقات ژنومیک استفاده شده است. وظایف یادگیری ماشین به دو دسته کلی با نظارت و بدون نظارت تقسیم می شود. در یادگیری با نظارت، هدف پیش بینی برچسب (رده بندی) یا پاسخ هر نقطه داده با استفاده از مجموعه ای ارائه شده از مثال های آموزشی دارای برچسب است. در یادگیری بدون نظارت هدف یادگیری الگوهای ذاتی درون داده ها است. هدف نهایی در بسیاری از کارهای یادگیری ماشین بهینه سازی عملکرد مدل نه بر روی داده های موجود (عملکرد آموزش) بلکه در عوض بر روی دادگان های مستقل (عملکرد تعمیم) است. با این هدف، داده ها به طور تصادفی به حداقل سه زیرمجموعه آموزش، اعتبار سنجی و مجموعه آزمون تقسیم می شوند. مجموعه آموزش برای یادگیری پارامترهای مدل، مجموعه اعتبار سنجی برای انتخاب بهترین مدل و مجموعه آزمون برای برآورد عملکرد تعمیم کنار گذاشته می شوند. یادگیری ماشینی باید به یک تعادل مناسب بین انعطاف پذیری مدل و میزان داده های آموزش برسد. چارچوب استفاده از یادگیری ماشین در داده های ژنومی به صورت شکل ۲-۶ است (Zou et al., 2019).

¹ Machine Learning



شکل ۶-۲ چارچوب یادگیری ماشین در داده‌های ژنومی (Zou et al., 2019)

الگوریتم‌های یادگیری نظارت‌شده را می‌توان به‌عنوان روش‌های مبتنی بر رگرسیون یا مبتنی بر درخت رده‌بندی کرد. رگرسیون لجستیک، رگرسیون خطی، شبکه‌های عصبی و SVM نمونه‌های رایج الگوریتم‌های یادگیری تحت نظارت رگرسیون هستند. در این بخش به بررسی روش‌های یادگیری با نظارت در مقالات متناسب با موضوع پرداخته‌شده است.

۲-۵-۱- الگوریتم درخت تصمیم^۱

درخت تصمیم‌گیری، روند نمایی^۲ است با ساختار درخت مانند به‌صورت سلسله‌مراتبی که از سه عنصر اصلی تشکیل‌شده است (Quinlan, 1986). بالاترین گره^۳ در درخت، گرهی ریشه است و گره داخلی (غیر برگ) آزمونی را بر روی یک ویژگی مشخص می‌کند و هر شاخه خارج‌شده از این گره، دستاورد این آزمون را نشان می‌دهد و گرهی پایانی را برگ می‌نامند که نشان‌دهنده کلاس‌ها هستند؛ ازجمله اشیائی که به یک طبقه تعلق دارند یا اینکه خیلی مشابه هستند و می‌توان آن‌ها را در یک طبقه قرارداد. به‌طور عادی، پیچیدگی یک درخت تصمیم با افزایش تعداد ویژگی^۴ افزایش می‌یابد. اگرچه در بعضی از شرایط دیده‌شده است که تنها تعداد کمی از ویژگی‌ها می‌توانند کلاسی را که هر شیء به آن تعلق دارد، تعیین کنند و بقیه ویژگی‌ها کم و یا بی‌تأثیرند. فرآیند دسته‌بندی هر نمونه با امتحان ویژگی بیان‌شده در گره ریشه شروع‌شده و سپس از شاخه‌های درخت با توجه به مقدار آن ویژگی پایین می‌آییم. سپس این فرآیند با امتحان گره بعدی که در انتهای شاخه انتخاب‌شده قرار دارد، ادامه می‌یابد تا نهایتاً به یک برگ برسیم. مهم‌ترین الگوریتم‌های درخت تصمیم شامل: ID3، C4.5، CART هست.

۲-۵-۲- الگوریتم جنگل تصادفی

روش جنگل تصادفی الگوریتمی است که توسط لئو بریمن در سال ۲۰۰۰ برای تولید مجموعه‌ای از پیش‌فرض‌ها و داده‌ها با استناد بر درختان جنگل که در فضای تصادفی رشد می‌یابند، مطرح شد (Breiman, 2000).

^۱ Decision Tree

^۲ Flowchart

^۳ Node

^۴ Feature

الگوریتم‌های جنگل تصادفی مدل‌های پیش‌بینی را با استفاده از یک روش تلفیقی و به کمک بسیاری از درختان تصمیم‌گیری می‌سازند. به‌ویژه، الگوریتم‌های جنگل تصادفی SNP‌هایی را انتخاب می‌کنند که در فرآیند تصمیم‌گیری ساخت درخت حاوی اطلاعات مفیدی هستند. قدرت مدل‌های جنگل تصادفی معادل با توانایی آن‌ها در به‌کارگیری مؤثر ساختار داده‌های به‌شدت چندبعدی و مفقودشده که حاوی تعاملات پیچیده هستند هست. به‌عنوان مثال، در یک مطالعه اخیر از یک الگوریتم جنگل تصادفی برای پیش‌بینی خطر T2D استفاده شد، که بهتر از هر دو مدل SVM و مدل‌های رگرسیون لجستیک عمل کرده است. در این مطالعه مجموعه‌ای از ۱۰۷۴ ژن‌نمود منفرد و ۱۰۱ SNP مربوط به T2D از پیش تعیین‌شده جمع‌آوری و تصحیح شدند. داده‌های تصحیح‌شده (۶۷۷ نمونه با ۹۶ SNP مرتبط) در یک الگوریتم یادگیری جنگل تصادفی قرار گرفتند و یک پیش‌بینی T2D با اعتبار $AUC = 0.85$ را با استفاده از متد اعتبارسنجی متقابل انجام دادند. با انجام این کار، همچنین مدل جنگل تصادفی SNP‌های از پیش تعیین‌شده را برای شناسایی زیرمجموعه‌هایی که به‌شدت با T2D همراهی دارند تصحیح کرده و می‌تواند به‌منظور تحقیق اتیولوژی بیماری مورد استفاده قرار گیرد. اجرای مدل جنگل تصادفی هنوز هم به‌عنوان یک روش یادگیری ماشینی برای مدل‌سازی خطر بیماری‌های پیچیده بسیار سودمند است.

۲-۵-۳- شبکه‌های بیزی^۱

فرض اصلی رده‌بندی BN، استقلال ویژگی‌ها است. یکی دیگر از فرض‌های دور از دسترس این است که همه ویژگی‌ها به‌طور کامل وابسته هستند. این فرض منجر به مدل شبکه بیزی می‌شود که یک گراف خطی جهت‌دار است که گره‌های آن متغیرهای تصادفی هستند و لبه‌ها وابستگی‌های شرطی را نشان می‌دهند. BN یک مدل کامل برای متغیرها و روابط آن‌ها در نظر گرفته است. بنابراین، یک توزیع احتمال احتمالی مشترک (JPD) بر روی تمام متغیرها برای یک مدل مشخص‌شده است. در استخراج متن، پیچیدگی محاسبات BN بسیار سنگین است و به همین دلیل اغلب استفاده نمی‌شود. BN توسط Hernandez و Rodríguez مورد استفاده قرار گرفت (Hernández-Rodríguez et al., 2016) با در نظر گرفتن یک مسئله در دنیای واقعی که در آن دیدگاه نویسنده با سه متغیر هدف اما مرتبط، موردنظر است. آن‌ها از رده‌بندی‌های چندبعدی بیزی استفاده کردند. به‌منظور استفاده از روابط بالقوه بین آن‌ها، متغیرهای هدف متفاوتی را در یک فعالیت رده‌بندی، به یکدیگر متصل نمودند. آن‌ها چارچوب رده‌بندی چندبعدی را به دامنه نیمه نظارتی گسترش دادند تا از مقدار زیادی از اطلاعات بدون برچسب در دسترس در این زمینه استفاده کنند.

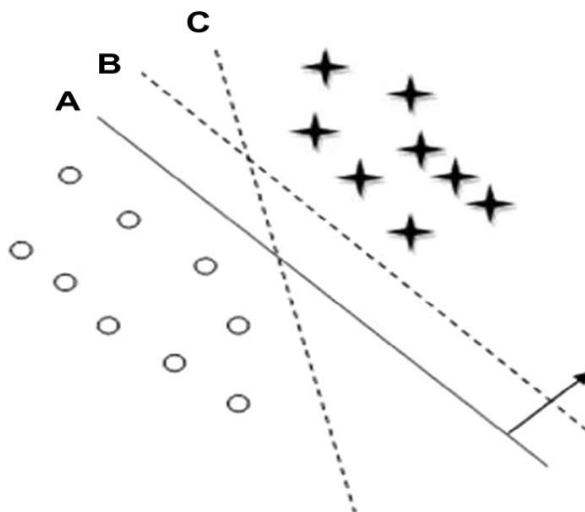
۲-۵-۴- رده‌بند ماشین بردار پشتیبان^۲

ماشین بردار پشتیبان (SVM) که نخستین بار توسط واپنیک در سال ۱۹۹۸ ارائه شد، از یک مدل خط برای اجرای حدود کلاس غیرخطی، به‌وسیله‌ی مسیره‌ی به بردارهای ورودی غیرخطی در داخل یک فضای

¹ Bayesian Network

² Support Vector Machine

چندبعدی استفاده می‌کند. در این فضای جدید برای ایجاد یک سطح افزار بهینه همواره (OSH) انجام می‌شوند. نمونه‌های آموزشی که به بیشینه حد سطح خیلی نزدیک هستند، بردارهای پشتیبان نامیده می‌شوند. سایر نمونه‌های آموزشی در تعریف حدود کلاس دودویی بی‌تأثیرند. روش SVM به قدری ساده است که می‌توان آن را به وسیله‌ی ریاضیات تجزیه و تحلیل کرد. از این جهت، SVM می‌تواند به عنوان یک جایگزین مطلوب برای ترکیب نقاط قوت روش‌های آماری مرسوم که به طور معمول دارای محیطی تئوری گرا هستند و به آسانی تجزیه و تحلیل می‌شوند و روش‌های فراگیری ماشین که محیطی بیشتر داده‌ای گرا هستند و توزیع آزاد و قوی دارند، مورد استفاده قرار گیرد. اصل مهم SVMها تعیین جدایی خطی در فضای جستجو است که می‌تواند بهترین کلاس‌های مختلف را جدا کند. در شکل ۲-۷، دو کلاس x و o ۳ صفحه A، B و C وجود دارد. صفحه A بهترین جدایی بین کلاس‌ها را فراهم می‌کند، زیرا فاصله نرمال هر یک از نقاط داده بزرگ‌ترین است، بنابراین نشان‌دهنده حداکثر حاشیه جدایی است.



شکل ۲-۷ استفاده از ماشین بردار پشتیبان در مسئله رده‌بندی

SVM دارای کاربردهای مختلفی در حوزه رده‌بندی هست. به عنوان مثال، یک مدل یادگیری ماشین غیر پارامتری مبتنی بر رگرسیون SVM در مورد ژنتیک دیابت نوع ۱ ساخته شده و بر اساس ۳۴۴۳ نمونه ژن‌نمود فردی آموزش داده شده است. این مدل به $AUC = 0.84$ دست یافته است، که به طور قابل توجهی بالاتر از مدل نمره دهی خطر چندژنی $AUC = 0.71$ هست.

۲-۵-۵- رده‌بند درخت تصمیم

رده‌بندی درخت تصمیم، تقسیم سلسله مراتبی فضای داده آموزش را فراهم می‌کند که در آن، یک شرط بر مقدار ویژگی، برای تقسیم داده‌ها استفاده می‌شود. تقسیم فضای داده به صورت بازگشتی انجام می‌شود تا گره‌های برگ حاوی حداقل تعداد رکوردها باشند که برای رده‌بندی استفاده می‌شوند.

۲-۵-۶- شبکه‌های عصبی

شبکه عصبی شامل بسیاری از نورون‌ها است که نورون واحد اصلی آن است. ورودی‌های نورون توسط بردار X_i نشان داده شده است. مجموعه‌ای از وزن A وجود دارد که با هر نورون مورد استفاده برای محاسبه یک تابع از

ورودی‌های آن f ، مورد استفاده قرار می‌گیرد. تابع خطی شبکه عصب به صورت مقابل است: $p_i = A \cdot X_i$. در یک مسئله رده‌بندی باینری، فرض بر این است که برچسب کلاس X_i با y_i مشخص می‌شود و علامت تابع پیش‌بینی شده p_i برچسب کلاس را نشان می‌دهد.

شبکه عصبی چندلایه برای مرزهای غیرخطی استفاده می‌شود. این لایه‌های چندگانه برای ایجاد چندین خط مرزی تقسیم‌بندی شده استفاده می‌شود. این لایه‌ها برای تقریب محدوده‌های محصور متعلق به یک کلاس خاص استفاده می‌شود. خروجی نورون‌ها در لایه‌های قبلی به نورون‌ها در لایه‌های بعدی وارد می‌شود. فرآیند آموزش پیچیده‌تر است، زیرا خطاها باید بر روی لایه‌های مختلف بازپخش شوند. پیاده‌سازی NN‌ها برای داده‌های متن وجود دارد که در این مقالات به آن اشاره شده است (Ng, Goh and Low, 1997; Ruiz and Srinivasan, 1999). مدل یادگیری عمیق ANN عملکرد پیش‌بینی کننده مهمی در مورد چاقی در مجموعه آزمایش‌ها با $AUC = 0.9908$ ارائه داد. مونتاز (Montañez, Fergus and Chalmers, 2018) توانایی الگوریتم یادگیری عمیق ANN را در ضبط اثرات ترکیبی SNP‌ها و همچنین پیش‌بینی بیماری‌های پیچیده پلی ژنی به وضوح نشان دادند.

۲-۶- یادگیری عمیق^۱ (DNN)

یادگیری عمیق که زیرمجموعه‌ای از یادگیری ماشین محسوب می‌گردد، رویکردی است که با استفاده از ساختار شبکه‌های عصبی عمیق، سعی بر یادگیری در چند سطح و یادگیری از آموخته‌های پیشین به صورت سلسله مراتبی دارد. این روش یک رشته نوظهور در حوزه یادگیری ماشین و نوعی رویکرد با سطوح مختلف یادگیری بازنمایی است که سعی دارد با تکیه بر روش‌های مبتنی بر یادگیری، همچون انسان اقدام به حل مسائل غیرخطی پیچیده در حوزه‌های مختلف مانند پردازش زبان‌های طبیعی، پردازش تصویر، پردازش صوت نماید (Lecun, Bengio and Hinton, 2015).

پیشرفت عمده اخیر در یادگیری ماشین، توسعه سریع الگوریتم‌های یادگیری عمیق^۲ است که می‌تواند از طریق یک فرایند یادگیری انباشته و سلسله مراتبی، ویژگی‌های معنی‌دار را از داده‌گان‌های چندبعدی و پیچیده استخراج کند. در حال حاضر، تصور می‌شود که یادگیری عمیق پیشرفته‌ترین فن یادگیری ماشین برای کاربردهای مختلف است (Marcos-Zambrano et al., 2021).

یادگیری عمیق برای پیش‌بینی دقیق داده‌های پیچیده مانند تصویر، متن یا فیلم به عنوان ابزاری قدرتمند ظهور کرده است. با این حال، توانایی آن در پیش‌بینی مقادیر رخ‌نمودی از داده‌های مولکولی کمتر مورد مطالعه قرار گرفته است. در بین معماری‌های یادگیری عمیق که در حال حاضر در پیش‌بینی ژنومی آزمایش شده‌اند، شبکه‌های عصبی پیچشی امیدوارکننده‌تر از گیرنده‌های چندلایه به نظر می‌رسند. محدودیت یادگیری عمیق

¹ Deep Neural Network

² Deep learning

در تفسیر نتایج است. این ممکن است برای پیش‌بینی ژنومی در اصلاح نژاد گیاهان و حیوانات مهم نباشد اما می‌تواند هنگام تعیین خطر ژنتیکی برای یک بیماری حیاتی باشد (Pérez-Enciso *et al.*, 2019). انواع مختلفی از شبکه‌های عصبی عمیق مانند شبکه‌های عصبی بازگشتی، شبکه‌های عصبی پیچشی، ماشین بولتزمن محدودشده، رمزگذار خودکار و موارد دیگری وجود دارد.

۲-۶-۱- تحلیل داده‌های حجیم مبتنی بر یادگیری عمیق

داده‌های حجیم^۱ به‌عنوان دادگان‌های بزرگ یا پیچیده شناخته‌شده‌اند که معمولاً سیستم پایگاه داده عادی برای انجام پردازش داده‌ها در یک‌زمان قابل‌قبول ناکافی است. این برنامه‌های پردازشی شامل تجزیه و تحلیل، ضبط، پردازش داده‌ها، جستجو، به اشتراک‌گذاری، ذخیره‌سازی، انتقال، مصورسازی، پرس‌وجو، به‌روزرسانی و مباحث حریم خصوصی اطلاعات می‌باشند.

داده‌هایی را می‌توان به‌عنوان داده‌های حجیم رده‌بندی کرد اگر داده‌ها شامل سه V است (Russom, 2018). اولین V حجم^۲ است که داده‌ها باید تعداد زیادی داده داشته باشند که شامل اطلاعات، پرونده‌ها یا تراکنش‌ها هست. V دوم تنوع^۳ داده است که داده‌ها در اشکال مختلف هستند و می‌توانند ساخت‌یافته یا غیر ساخت‌یافته باشند. آخرین V به‌سرعت^۴ اشاره می‌کند که به‌سرعت انتشار داده اشاره می‌کند.

داده‌های ژنوم انسان، دارای حجم بالایی است (داده ژنوم هر انسان حدود ۳ گیگابایت هست) و از این منظر می‌توان گفت داده‌های ژنوم از جنس داده‌های حجیم هست.

فن‌های مختلف یادگیری ماشین مانند باهم‌آیی، داده‌کاوی حجیم، تجزیه خوشه‌ای و تجزیه و تحلیل متن در تحلیل این نوع از داده‌ها استفاده می‌گردد (Maltby, 2011).

بر اساس (Hordri *et al.*, 2017) تحقیقاتی که در حوزه پردازش داده حجیم و یادگیری عمیق انجام‌شده است حول ۴ محور به ترتیب: انتزاع سطح بالا، لایه‌های سلسله‌مراتبی، پردازش حجم بالای داده و مدل عمومی هست که نشان‌دهنده مراحل اجرای یادگیری عمیق هست.

الگوریتم‌های یادگیری عمیق، مانند آلفاگو^۵ (Silver *et al.*, 2016) و تشخیص ابعاد (Russakovsky *et al.*, 2015)، از عملکرد انسان در وظایف بصری بهتر هستند و فن‌های تحلیل انعطاف‌پذیر و قدرتمند برای مقابله با مشکلات پیچیده دارند. یادگیری عمیق یک الگوریتم انتزاعی سطح بالا برای حل مسائل رده‌بندی و رگرسیون است. از طریق یادگیری عمیق و استخراج الگوهای داده‌ها، ساختارهای پیچیده را در مجموعه‌های داده‌های عظیم شناسایی می‌کند و پتانسیل زیادی برای برنامه‌های کاربردی در ژنتیک و ژنومیک دارد (Libbrecht and Noble, 2015).

۲-۶-۲- الگوریتم‌های یادگیری عمیق

¹ Big Data

² Volume

³ Variety

⁴ Velocity

⁵ Alpha Go

مروری بر الگوریتم‌های یادگیری عمیق در جدول ۲-۴ ادامه آمده‌است.

جدول ۲-۴ جدول مقایسه شبکه DNN (Shrestha and Mahmood, 2019)

نوع شبکه	معماری	مدل شبکه	نوع آموزش	الگوریتم آموزشی	نمونه پیاده‌سازی شده	کاربرد رایج
شبکه عصبی پیشرفته ^۱	CNN	متمايزکننده	با ناظر	شیب نزول بر اساس پردازش مجدد	شبکه سیامی، CNN عمیق	تشخیص عکس و رده‌بندی
	شبکه باقیمانده ^۲	متمايزکننده	با ناظر	شیب نزول بر اساس پردازش مجدد	DeepResNet HighwayNet DenseNet	تشخیص عکس
	خود رمزنگار ^۳	تولیدی	بدون ناظر	پردازش مجدد	خود رمزگذار حریص و خود رمزگذار متغیر	کاهش ابعاد، رمزگذاری
	شبکه‌های مخالف ^۴	متمايزکننده و تولیدی	بدون ناظر	پردازش مجدد	شبکه تولیدی مخالف	تولید داده جعلی واقعی، بازسازی مدل‌های 3D، بهبود تصویر
	RBM	تولیدی با متمايزکننده قوی	بدون ناظر	شیب نزول بر اساس واگرایی متضاد	شبکه عقیده عمیق و ماشین بلتزن عمیق	کاهش ابعاد، یادگیری ویژگی، مدل‌سازی موضوعی
شبکه عصبی گردشی ^۵	LSTM	متمايزکننده	با ناظر	شیب نزول و پردازش مجدد در طول زمان	Deep RNN GRU NMT	پردازش زبان طبیعی، ترجمه زبان
شبکه عصبی عملکرد پایه شعاعی ^۶	شبکه RBF	متمايزکننده	با ناظر و بدون ناظر	خوشه‌بندی k-means و تابع حداقل مربع	RNN	تقریب زدن تابع، پیش‌بینی سری زمانی
شبکه عصبی خودسازمان‌دهی کونون ^۷	گره‌های شکل یافته بر اساس شبکه‌های مستطیلی یا شش‌ضلعی	تولیدی	بدون ناظر	یادگیری رقابتی	شبکه عصبی خودسازمان‌دهی کونون	کاهش ابعاد، مسائل بهینه‌سازی، تحلیل خوشه‌بندی

نقاط ضعف الگوریتم‌های یادگیری عمیق و فن‌های بهبود در جدول ۲-۵ آمده‌است.

¹ Feed forward neural network

² Residual network

³ Auto encoder

⁴ Adversarial networks

⁵ Recurrent neural network

⁶ radial basis function neural network

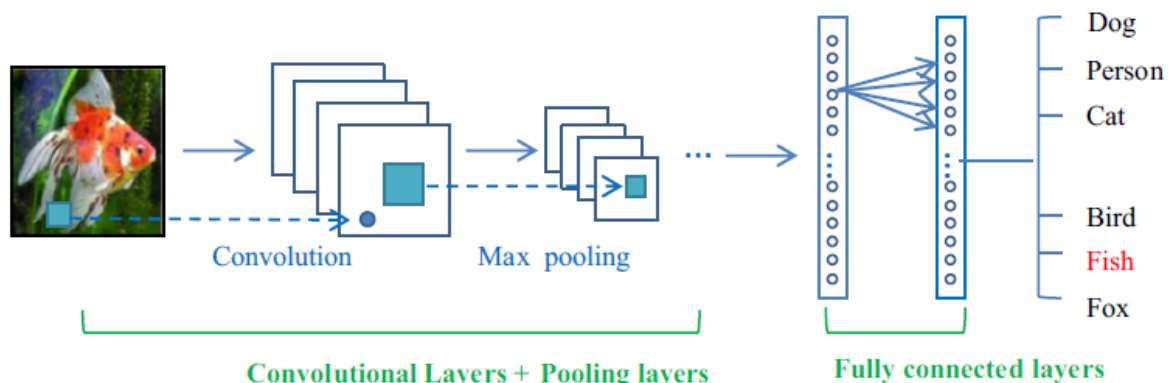
⁷ kohonen self organizing neural network

جدول ۵-۲ نقاط ضعف الگوریتم‌های یادگیری عمیق

نقاط ضعف	علت	اثر	بهینه‌سازی برای آدرس دادن نقاط ضعف
کاهش یا افزایش ناگهانی و شدید شیب	انتشار مشتق‌ها	زمان طولانی آموزش، پرت شدن کمینه عمومی، سخته آموزشی	تابع فعال‌سازی ReLU، وزن دهی ابتدایی، برقراری ارتباط بین فعال‌سازی درب فراموشی و محاسبه شیب در LSTM (RNN)، ارتباطات پرشی (ResNet)، سخت‌افزار سریع‌تر (GPUs)
کمینه محلی / نواحی مسطح	شیب نزول فضای مسئله غیر محدب	همگرایی کمینه محلی	نرخ یادگیری انطباقی، مقداردهی مجدد شروع به پارامترها، فن‌های مقداردهی،
لبه‌های تیز	شیب نزول فضای مسئله غیر محدب و فضای تیز	پرت شدن به کمینه عمومی اشتباه	نرخ یادگیری انطباقی
پوشش بیش‌ازحد	گره‌های با اندازه بزرگ‌تر (شبکه) متناسب با دادگان، دادگان ضعیف	دقت ضعیف مدل	مقررات، از قلم افتادن، انتخاب اندازه مناسب برای شبکه، داده‌های آموزشی بهتر
زمان آموزشی زیاد	شبکه بزرگ (وزن‌ها و پارامترهای زیاد)، داده با ابعاد بالا و ...	استفاده ضعیف از منابع محاسبه، زمان هدر شده	نرخ یادگیری انطباقی، استفاده از ابر و GPUs
انتخاب بیش‌ازحد پارامتر	فضای جواب بسیار بزرگ (مسئله NP-hard) برای مسائل با ابعاد بالا	همگرایی به کمینه محلی	فرا ابتکاری (مانند الگوریتم ژنتیک) برای بهینه‌سازی با پارامترهای بالا، جستجو تصادفی، الگوریتم مبتنی بر مدل ترتیبی

۲-۶-۳- شبکه‌های عصبی پیچشی^۱ (CNN)

شبکه‌های عصبی پیچشی یکی از مهم‌ترین روش‌های یادگیری عمیق هستند که در آن‌ها چندین لایه با روشی قدرتمند آموزش می‌بینند. این روش بسیار کارآمد بوده و یکی از رایج‌ترین روش‌ها در کاربردهای مختلف بینایی رایانه است. تصویر کلی یک معماری شبکه عصبی پیچشی در شکل ۲-۸ نمایش داده شده است.



شکل ۸-۲ معماری شبکه CNN

^۱ Convolutional Neural Network

به‌طور کلی، یک شبکه CNN از سه لایه اصلی تشکیل می‌شود که عبارت‌اند از: لایه پیچشی^۱، لایه ادغام^۲ و لایه تماماً متصل^۳ (FCN). لایه‌های مختلف وظایف مختلفی را انجام می‌دهد (Krizhevsky, Sutskever and Hinton, 2017). در هر شبکه عصبی پیچشی دو مرحله برای آموزش وجود دارد. مرحله جلوبرنده^۴ (FF) و مرحله پس انتشار^۵ (BP). در مرحله اول تصویر ورودی به شبکه تغذیه می‌شود و این عمل چیزی جز ضرب نقطه‌ای بین ورودی و پارامترهای هر نورون و نهایتاً اعمال عملیات پیچشی در هر لایه نیست. سپس خروجی شبکه محاسبه می‌شود. در اینجا به‌منظور تنظیم پارامترهای شبکه و یا به‌عبارت دیگر همان آموزش شبکه، از نتیجه خروجی جهت محاسبه میزان خطای شبکه استفاده می‌شود. برای این کار خروجی شبکه را با استفاده از یک تابع خطا^۶ (LF) با پاسخ صحیح مقایسه کرده و این‌طور میزان خطا محاسبه می‌شود. در مرحله بعدی بر اساس میزان خطای محاسبه‌شده مرحله BP آغاز می‌شود. در این مرحله گرایانت هر متغیر با توجه به‌قاعده زنجیره‌ای محاسبه می‌شود و تمامی پارامترها با توجه به تأثیری که بر خطای ایجادشده در شبکه دارند تغییر پیدا می‌کنند. بعد از بروز آوری شدن پارامترها مرحله بعدی FF شروع می‌شود. بعد از تکرار تعداد مناسبی از این مراحل آموزش شبکه پایان می‌یابد.

۲-۷- پیش‌بینی امتیاز خطر^۷

مدل‌های معمول خطر بیماری بر اساس همه‌گیرشناسی (با قدرت پیش‌بینی محدود) در درجه اول مبتنی بر خطر عامل‌های شیوه زندگی مانند وجود سابقه خانوادگی مثبت می‌باشند. اخیراً، در نظر گرفتن خطر عامل‌های ژنتیکی از جمله SNP‌های مرتبط با بیماری و یا مربوط به رخ‌نمود همراه با مدل‌سازی خطر، دقت پیش‌بینی فردی بیماری‌ها را بهبود بخشیده است.

شاید بهترین مدل‌های پیش‌بینی کننده خطر بر اساس پتانسیل آن‌ها به‌منظور راهنمایی در مورد پیشگیری از بیماری و درمان بدون نیاز به روش‌های غربالگری پزشکی پرهزینه و بالقوه نامطلوب (به‌عنوان مثال بیوپسی های تهاجمی) باشد.

در حال حاضر، توسعه مدل‌های خطر ژنتیکی بر روی دستیابی به قدرت پیش‌بینی دقیق برای شناخت افراد در خطر به شیوه‌ای قدرتمند تمرکز دارد. همان‌طور که قبلاً گفته شد، مطالعات GWA، SNP‌ها را با توجه به ارتباط آن‌ها با یک بیماری / رخ‌نمود در سطح جمعیت تعریف می‌کند. بنابراین، ترکیب SNP‌ها در یک مدل پیش‌بینی کننده‌ی خطر نیاز به ترکیب با مدل‌هایی دارد که ژن‌نمود یک فرد را به دستمی‌آورند تا بتوانند خطر را پیش‌بینی کنند. مدل‌های پیش‌بینی خطر ژنتیکی معمولاً توسط: (۱) امتیازدهی خطر چندژنی^۸ (PRS) یا (۲) یادگیری ماشینی ساخته می‌شوند.

¹ Convolution

² Pooling

³ Fully Connected Network

⁴ feed forward

⁵ back propagation

⁶ loss function

⁷ Risk Scroe

⁸ Polygenic Risk Score

عملکرد پیش‌بینی کننده هر دو نوع مدل با استفاده از منحنی‌های مشخصه عملکرد سیستم^۱ (ROC) ارزیابی می‌شود، که در آن حساسیت و اختصاصی پیش‌بینی‌ها در مقادیر مختلف رتبه‌بندی می‌شود. مساحت زیر منحنی^۲ ROC (AUC) برابر با این احتمال است که مدل موردبررسی، از میان یک جفت نمونه و شاهد (کنترل) تصادفی، یک مورد را به‌درستی شناسایی کرده است. نتایج AUC از ۰.۵ (به‌عنوان مثال تصادفی) تا ۱ (یعنی دقت ۱۰۰ درصد) متغیر است.

۲-۷-۱- امتیازدهی خطر چندژنی

مدل‌های معمول خطر بیماری بر اساس همه‌گیرشناسی (با قدرت پیش‌بینی محدود) در درجه اول مبتنی بر خطر عامل‌های شیوه زندگی مانند وجود سابقه خانوادگی مثبت می‌باشند. اخیراً، در نظر گرفتن خطر عامل‌های ژنتیکی از جمله SNP‌های مرتبط با بیماری و یا مربوط به رخ‌نمود همراه با مدل‌سازی خطر، دقت پیش‌بینی فردی بیماری‌ها را بهبود بخشیده است (Mosley et al., 2020). امتیازدهی خطر چندژنی از یک رویکرد مدل ثابت برای جمع‌آوری مجموعه‌ای از آلل‌های خطر در یک بیماری پیچیده خاص استفاده می‌کند. در این روش تعداد زیادی SNP برای اندازه‌گیری خطر در بدو تولد استفاده می‌شوند (Phulka et al., 2023). نمرات خطر چندژنی ممکن است وزنی یا غیر وزنی باشند. در نمرات خطر چندژنی وزنی، سهم آلل‌های خطر به‌طور معمول با نسبت شانس یا اندازه اثر آن‌ها اندازه‌گیری می‌شود. در مقابل، نمرات خطر چندژنی غیر وزنی برابر با مجموع تعداد گونه‌های اللی پیوسته در یک ژنوم است. مدل غیر وزنی فرض می‌کند که همه گونه‌ها دارای اندازه اثر متعادل هستند. این فرض ساده‌انگاران، استفاده از نمرات خطر چندژنی غیر وزنی را برای صفات پیچیده با ساختارهای ژنتیکی زیربنایی که شامل اثرات نامتعادل گونه‌ها هست را محدود می‌کند.

دو مرحله برای تدوین امتیاز خطر چندژنی وجود دارد: (۱) مرحله کشف و (۲) مرحله اعتبار سنجی. مرحله کشف در مورد رتبه‌بندی خطر چندژنی وزنی آزمون ارتباط آماری (برای مثل رگرسیون خطی یا لجستیک) استفاده می‌کند تا اندازه اثر یک مجموعه بزرگ از داده‌ها و موارد کنترل پروفایل‌های ژن‌نمود فردی را برآورد کند. مرحله کشف یک نمره خطر چندژنی غیر وزنی نیازمند معیارهای انتخاب دقیق SNP‌ها به‌منظور جلوگیری از اختلاط SNP‌ها با اندازه اثرهای جزئی هست. در هر دو مورد نمره خطر چندژنی وزنی و غیر وزنی، پس از توسعه، مدل اکتشافی به مرحله اعتبارسنجی منتقل می‌شود. اعتبارسنجی نمره خطر چندژنی مستلزم استخراج SNP‌هایی حاوی اطلاعات مفید و اندازه‌های اثر از مجموعه کشف، با استفاده از آستانه p-value پیوستگی (به‌عنوان مثال 5×10^{-8}) هست، که متعاقباً به مرحله اعتبار سنجی منتقل می‌شود. در طی این فرآیند، مدل نمره خطر چندژنی بر روی یک دادگان آزمایشی (یعنی مجموعه‌ای مستقل از داده‌های ژن‌نمودی مورد-شاهد) اعمال می‌شود. نمرات خطر چندژنی برای هر پروفایل ژن‌نمود فردی در داده‌های آزمایش محاسبه می‌شود.

¹ Receiver Operating Characteristic curves

²Area Under Curve

سپس قدرت پیش‌بینی نمرات خطر چندژنی فردی برای صفات پیچیده با استفاده از قدرت پیوستگی نمرات با نتایج کلینیکی (رخ‌نمودها) در دادگان‌های آزمایش محاسبه می‌شود.

تلاش‌های اولیه برای استفاده از نمرات خطر چندژنی وزنی مبتنی بر تعداد کمی از SNP های بسیار مهم بوده است که با استفاده از مطالعات GWA مشخص شده‌اند و در مورد بیماری‌های پیچیده فقط به ارزش پیش‌بینی کننده محدودی دست یافته‌اند. این امر یک محدودیت اصلی در مدل‌سازی نمره خطر چندژنی وزنی، به‌ویژه در مورد حد آستانه p-value برای انتخاب SNP در تأثیر دادگان‌های اکتشافی بر عملکرد مدل و قدرت پیش‌بینی را نشان می‌دهد. انتخاب تعداد محدودی از SNP ها با اندازه اثر زیاد، با نادیده گرفتن بخش عمده‌ای از گونه‌ها که تأثیرات بسیار کمتری بر روی رخ‌نمود دارند، زیربناهای بیولوژیکی بیماری‌های پیچیده را بیش‌ازحد ساده می‌کند.

به‌عنوان مثال، نسبت شانس متوسط در هر آلل خطر T2D از ۱.۰۲ تا ۱.۳۵ متغیر است. مدل‌های امتیازات چندژنی اخیر، به‌منظور دستیابی به نتایج پیش‌بینی کننده بهتر در مورد صفات پیچیده ژنتیکی، SNP های گسترده‌تری را انتخاب می‌کنند.

به‌عنوان مثال، استفاده از حد آستانه p-value (به‌اندازه ۰.۰۰۱، ۰.۰۱ و ۰.۲ و غیره) منجر به توسعه مدل‌های بهبودیافته خطر چندژنی را برای بیماری‌های روانی با حداقل افزایش در خطاهای مثبت کاذب (یعنی مدل‌ها از نسبت قدرت به نویز قابل قبولی برخوردار هستند) شده است.

رویکرد نمره خطر چندژنی وزنی، پیش‌بینی خطر اسکیزوفرنی را با دستیابی به اثربخشی معقول AUC 0.65 امکان‌پذیر کرده است. به‌طور مشابهی، نتایج قابل توجهی از پیش‌بینی نمره خطر چندژنی وزنی برای سایر صفات پیچیده از جمله دیابت نوع ۱ و بیماری سلیاک نیز به‌دست آمده است (Padmanabhan and Dominiczak, 2021). اخیراً، با توجه به اینکه SNP هایی که آستانه اهمیت بسیار سخت‌گیرانه برای ارتباط گسترده ژنومی را برآورده نمی‌کنند نیز می‌توانند پیش‌بینی کننده بیماری باشند، طیف وسیع‌تری از SNP ها، از هزاران تا میلیون‌ها SNP برای ایجاد یک GRS بهبودیافته به‌عنوان PRS استفاده شده‌اند (Torkamani, Wineinger and Topol, 2018). باید تأکید کرد اندازه خطر ارائه‌شده توسط GRS یک محدوده احتمالی است و این با اطلاعات خطر از نشانگرهای ژنتیکی اختلالات تک ژنی متفاوت است.

۲-۷-۲ - مدل‌های پیش‌بینی کننده بیماری بر پایه یادگیری ماشینی

روش‌های یادگیری ماشینی مجموعه‌ای از الگوریتم‌های آماری و محاسباتی پیچیده هست (به‌عنوان مثال ماشین بردار پشتیبان (SVM) یا جنگل تصادفی) که با استفاده از نقشه‌برداری ریاضی ارتباطات پیچیده بین یک مجموعه SNP های خطرناک و رخ‌نمودهای بیماری پیچیده، پیش‌بینی‌هایی انجام می‌دهد. این روش‌ها از رویکردهای نظارت‌شده یا نظارت‌نشده برای ترسیم ارتباط با بیماری‌های پیچیده استفاده می‌کنند. علیرغم سودمندی روش‌های یادگیری ماشین نظارت‌نشده و داده‌های غیر ژنتیکی در پیش‌بینی

بیماری‌ها، ما در این نوشتار بر روی مدل‌سازی نظارت‌شده که با داده‌های SNP اطلاع‌رسانی می‌شود، تمرکز خواهیم کرد.

مدل‌های یادگیری ماشینی نظارت‌شده به‌منظور پیش‌بینی بیماری‌ها با استفاده از آموزش الگوریتم‌های یادگیری از پیش تعیین‌شده برای ترسیم روابط بین داده‌های ژن‌نمودی فردی و بیماری مرتبط با آن‌ها ایجاد می‌شوند. قدرت پیش‌بینی بهینه برای بیماری موردنظر از طریق نقشه‌برداری از الگوی ویژگی‌های منتخب (متغیرها) در داده‌های ژن‌نمود آموزش حاصل می‌شود. برخی از مدل‌ها از روش‌های با شیب نزولی و دوره‌های تکراری برآورد پارامترها به‌منظور جستجو در میان داده‌های آموزشی برای قدرت پیش‌بینی کننده بهینه استفاده می‌کنند. این روند بازگشتی تا رسیدن به عملکرد بهینه پیش‌بینی کننده ادامه می‌یابد. در پایان مرحله آموزشی، مدل‌هایی با حداکثر قدرت پیش‌بینی کننده در دادگان‌های آموزش برای اعتبارسنجی انتخاب می‌شوند.

در مرحله اعتبار سنجی، عملکرد پیش‌بینی کننده‌ی مدل‌های یادگیری ماشین برای تعیین قدرت آن‌ها به جهت پیش‌بینی عمومی مورد ارزیابی قرار می‌گیرد. مشابه امتیازدهی خطر چندژنی، مرحله اعتبارسنجی با ارزیابی الگوریتم در یک دادگان‌های مستقل انجام می‌شود. مرحله اعتبار سنجی برای اطمینان از این‌که مدل‌های پیش‌بینی کننده داده‌های آموزش را بیش‌ازحد پوشش نمی‌دهند ضروری است. اعتبار سنجی متقابل روشی متداول برای اعتبارسنجی عملکرد مدل‌ها با استفاده از دادگان‌های اصلی است. با این وجود، اعتبارسنجی (آزمایش) خارجی با استفاده از یک دادگان مستقل لازم است تا در نهایت قدرت پیش‌بینی کننده یک مدل یادگیری ماشین را تأیید کند. سرانجام کارایی الگوریتم از طریق مقایسه‌های کنترل‌شده تصادفی با بهترین عملکرد بالینی فعلی مشخص می‌شود. تنها در صورتی که الگوریتم اطلاعاتی بیشتری را در جهت رده‌بندی دقیق‌تر جمعیت، پیش‌بینی خطر بیماری یا پاسخ‌های درمانی ارائه دهد، کارایی بالینی خود را در نهایت اثبات می‌کند (Ho et al., 2019).

پیشرفت‌های روزافزون فناوری به‌طور مداوم در حال بهبود کیفیت و کمیت دادگان‌های پیوسته پیچیده‌ای هست که در مورد رخ‌نمودها و بیماری‌های انسانی جمع‌آوری شده است. ادغام این داده‌های با ابعاد ژنومی بسیار بالا در مدل‌های یادگیری ماشین می‌تواند منجر به بهبودهایی در زمینه پیش‌بینی خطرات ژنتیکی فراتر از آنچه در مورد نمرات خطر چندژنی حاصل می‌شود گردد. پیش‌بینی نمره خطر چند ژن بر اساس یک مدل رگرسیون پارامتری خطی است که شامل فرضیات سخت‌گیرانه هست، که شامل اثرات پیش‌بینی کننده افزایشی و مستقل، یک توزیع نرمال برای داده‌های زمینه‌ای و همچنین عدم همبستگی داده‌ها است. این فرضیات لزوماً در مورد ساختارهای ژنتیکی پایه‌ای بیماری‌های پیچیده چندژنی صحیح نیستند و بنابراین منجر به کاهش زیادی در کارایی پیش‌بینی می‌شوند. نکته قابل توجه این است که مدل‌سازی رگرسیون افزایشی خطی قادر به محاسبه اثرات متقابل پیچیده بین آلل‌های مرتبط، که بنا بر گزارش‌ها به مقدار زیادی با رخ‌نمودها مرتبط است نیست. از این‌رو، مدل‌سازی مبتنی بر رگرسیون افزایشی خطی، نمرات خطر چندژنی را به سمت پیش‌بینی‌های جانب‌دارانه و با کارایی کمتر سوق می‌دهد. در مقابل، الگوریتم‌های یادگیری ماشین از روش‌های

چند متغیره غیر پارامتری استفاده می‌کنند که به طرز قدرتمندی الگوها را از داده‌هایی که توزیع غیر نرمال دارند و به شدت به یکدیگر پیوسته‌اند تشخیص می‌دهند. ظرفیت الگوریتم‌های بر پایه یادگیری ماشین در جهت مدل‌سازی داده‌هایی با ساختارهای بسیار پیچیده و مرتبط باعث شده است که این رویکردها با اهداف پیش‌بینی بیماری‌های پیچیده از محبوبیت بسیار زیادی برخوردار باشند.

نقاط قوت و ضعف هر دو روش در جدول ۲-۶ خلاصه شده است (Ho et al., 2019).

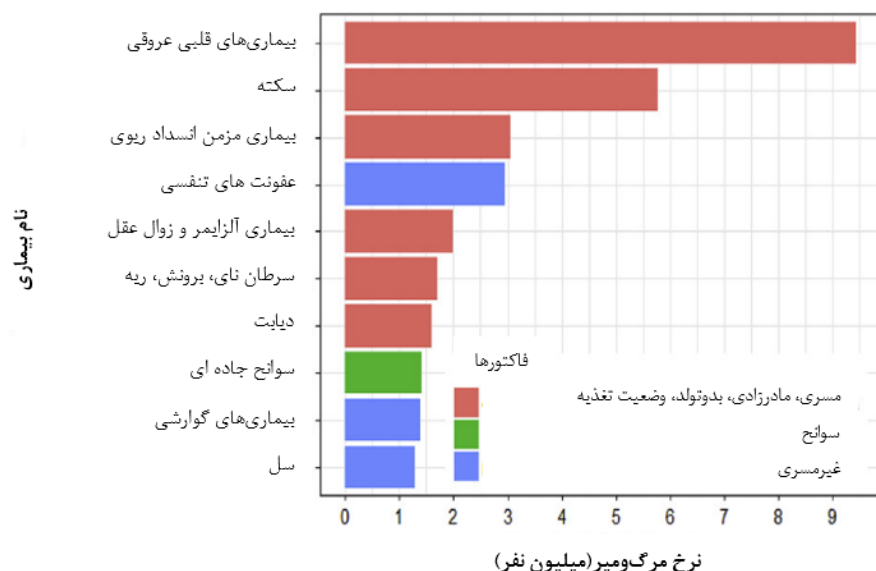
جدول ۲-۶ مقایسه روش GRS و یادگیری ماشین (Ho et al., 2019)

امتیاز خطر چندژنی	مدل یادگیری ماشین	
✓ سادگی اجرا	✓ مؤثر برای داده‌های چندبعدی	نقاط قوت
✓ سادگی تفسیر	✓ قابل استفاده در داده‌های پیچیده مرتبط	
	✓ عدم نیاز به فرض نرمال بودن	
✗ تأثیر افزایشی و مستقل پیش‌بینی کننده	✗ سختی در اجرا	نقاط ضعف
✗ فرض نرمال بودن	✗ سختی در تفسیر اجزای مؤثر	
✗ بی استفاده در داده پیچیده مرتبط	✗ نیاز به داده زیاد	

۸-۲- بیماری پرفشاری خون^۱ (HTN)

بیماری پرفشاری خون یکی از مهم‌ترین مشکلات سلامت عمومی و در حال افزایش در سطح دنیا به شمار می‌آید (تیموری، ابراهیمی و علوی نیا، ۱۳۹۴). فشارخون بالا یکی از عوامل مهم در بیماری قلبی عروقی است که توسط عوامل ژنتیکی و غیر ژنتیکی ایجاد می‌شود و بیش از ۱.۲۸ میلیارد نفر را در جهان مبتلا کرده است. (Loganathan et al., 2019) (Eshmurodova Dilrabo Khudayberdi kizi, Ergasheva Iroda Akbarali) (kizi and Siab, 2023). این بیماری با بیماری کلیوی، سکته مغزی نیز ارتباط دارد و شیوع بالایی در کشورهای در حال توسعه و کشورهای توسعه یافته دارد و سالانه جان بیش از ۱۰ میلیون نفر را می‌گیرد (Singh et al., 2016) (Garimella et al., 2023). به همین دلیل تشخیص بهنگام فشارخون برای درمان زودرس بسیار مهم است (Nour and Polat, 2020; Niu et al., 2021). شکل ۲-۹ سهم هر بیماری در مرگ و میر جهانی را نشان می‌دهد (Alizadehsani et al., 2019) که بیان کننده اهمیت بیماری‌های قلبی و عروقی است.

¹ Hypertension



شکل ۹-۲ ده دلیل اصلی عامل مرگ در جهان (Alizadehsani et al., 2019)

در مقاله "مدل پیش‌بینی فشارخون بالا مبتنی بر عوامل خطر ژنتیکی و محیطی در شمال هان چین" که در سال ۲۰۱۹ منتشر شد، پرفشاری خون به‌عنوان یک بیماری مزمن شناخته شد که شیوع بالایی در جهان دارد و از عوامل بیماری قلبی عروقی و مغزی عروقی است. در این مقاله علاوه بر عوامل و پارامترهای محیطی و کلینیکی، اطلاعات ژنتیکی مانند چندریختی نوکلئوتیدی^۱ (SNP) نیز مورد بررسی قرار گرفت. روش‌های آماری سنتی مانند آزمایش t و رگرسیون لجستیک چند متغیره برای تحلیل اطلاعات انجام شد. یک روش یادگیری ماشین مانند ماشین بردار پشتیبان نیز برای توسعه مدل پیش‌بینی بیماری فشارخون بالا استفاده شد. در نهایت دو مدل برای پیش‌بینی فشارخون سیستولیک و فشارخون دیاستولیک ساخته شد. مدل فشارخون سیستولیک شامل شش عامل محیطی مانند سن، شاخص توده بدن، دور مچ، فعالیت فیزیکی، سابقه خانوادگی فشارخون بالا و یک SNP و مدل فشارخون دیاستولیک شامل شش عامل محیطی مانند وزن، مصرف الکل، فعالیت فیزیکی، تری‌گلیسیرید، سابقه خانوادگی فشارخون بالا و سه SNP هست. ناحیه زیر نمودار برای مدل فشارخون سیستولیک ۰/۶۷۳ و مدل فشارخون دیاستولیک ۰/۸۱۷ به دست آمد (Li et al., 2019).

هم‌چنین می‌توان با استفاده از روش‌های مختلف یادگیری ماشین، یک مدل پیش‌بینی بر مبنای هزینه طراحی کرد تا بهترین عملکرد در پیش‌بینی افراد دیابتی در معرض خطر پرفشاری خون را داشته باشد. در مقاله‌ای تحت عنوان "مقایسه روش‌های مختلف یادگیری ماشین در تشخیص پرفشاری خون در بیماران دیابتی با و بدون در نظر گرفتن هزینه‌ها از الگوریتم‌های مختلف داده‌کاوی مانند الگوریتم درخت تصمیم، ماشین بردار پشتیبان، شبکه عصبی و رگرسیون لجستیک استفاده شد. یکی از یافته‌های تحقیق این بود که افزایش فشارخون سیستول به میزان ۱۳۰ میلی‌متر جیوه، فرد دیابتی را بیشتر در معرض پرفشاری خون قرار می‌دهد. با توجه به هدف حداقل سازی هزینه، درخت تصمیم و رگرسیون لجستیک بهترین عملکرد را دارند. در نتیجه

¹ Single Nucleotide Polymorphism

در مسائل پیش‌بینی بیماری‌ها حساس به هزینه کردن روش‌ها و در نظر گرفتن توزیع واقعی بیماری در جامعه اهمیت بیشتری دارد تا اینکه تنها هدف، کمینه کردن تعداد خطاهای دسته‌بندی روی دادگانی موجود باشد. سابقه خانوادگی، سن، جنس، مصرف نمک، مصرف الکل، مصرف دخانیات، کم‌تحركی و چاقی از عوامل خطر بیماری فشاری خون بالا به شمار می‌رود. این عوامل روی شاخص توده بدنی (BMI)، دیابت، چربی خون بالا (شامل کلسترول، لیپوپروتئین با تراکم بالا (HDL)، لیپوپروتئین با تراکم پایین (LDL) و تری گلیسرید) تأثیر دارند (تیموری، ابراهیمی و علوی نیا، ۱۳۹۴).

آقای دانگ‌دونگ و همکارانش مقاله‌ای را تحت عنوان توسعه مدل‌های پیش‌بینی ریسک برای پرفشاری خون در سال ۲۰۱۷ منتشر کردند. از آنجا که بیماری پرفشاری خون از تهدیدهای پیشرو سلامت جهانی و از عوامل مؤثر در بیماری قلبی عروقی هست و همچنین چون مداخلات کلینیکی در به تأخیر انداختن پیشرفت بیماری پرفشاری خون مؤثر است، مدل‌های پیش‌بینی تشخیصی برای شناسایی جمعیت بیمار پر ریسک فشارخون بالا مهم هست.

داده‌ها در ابتدا از پایگاه داده PubMed و Embase جمع‌آوری شدند، سپس عواملی ریسک، مدل‌های آماری، ویژگی‌های طراحی مطالعه و شرکت‌کنندگان و همچنین معیار عملکرد موردبررسی قرار گرفتند. از عواملی ریسک سنتی مانند شاخص توده بدن، سن، مصرف سیگار، سطح فشارخون، سابقه خانوادگی فشارخون و عوامل زیست‌شیمیایی می‌توان نام برد، درحالی‌که در شش گزارش از نمره ریسک ژنتیکی به‌عنوان عامل پیش‌بینی استفاده شده است. دامنه محدوده زیر نمودار از ۰/۶۴ تا ۰/۹۷، C-Statistic از ۰/۶۰ تا ۰/۹۰ به‌دست‌آمده آمده است.

از یافته‌ها می‌توان به این نتیجه رسید که مدل‌های سنتی هنوز جزء مدل‌های پیش‌بینی ریسک غالب برای پرفشاری خون هستند، اما اخیراً مدل‌های بیشتری شروع به ترکیب و وارد کردن عواملی ژنتیکی به‌عنوان قسمتی از پیش‌بینی کننده مدل استفاده می‌شود. هرچند که این پیش‌بینی کننده‌های ژنتیکی باید به‌خوبی انتخاب شوند. مدل‌های استفاده شده در این مقاله توانایی کالبراسیون و تفکیک قابل قبول تا خوبی را دارند، اما برای اینکه در فعالیت کلینیکی و عملی مورد استفاده قرار بگیرند، به اعتبارسنجی و تنظیم بیشتری نیاز دارند (Li et al., 2019).

در مقاله طبقه‌بندی خودکار انواع فشارخون بالا بر اساس ویژگی‌های شخصی توسط الگوریتم‌های یادگیری ماشین، به‌منظور شناسایی انواع فشارخون بالا بر اساس اطلاعات و ویژگی‌های شخصی، چهار روش یادگیری ماشین شامل طبقه‌بندی درخت تصمیم C4.5، جنگل تصادفی، تجزیه و تحلیل تفکیک خطی و پشتیبانی خطی ماشین بردار استفاده شده و سپس با یکدیگر مقایسه شده است. در ادبیات، ابتدا رده‌بندی انواع فشارخون بالا با استفاده از الگوریتم‌های طبقه‌بندی بر اساس داده‌های شخصی انجام شد. برای توضیح بیشتر تنوع نوع طبقه‌بندی، چهار الگوریتم طبقه‌بندی کننده مختلف برای حل این مسئله انتخاب شد. در دادگان‌های فشارخون، هشت ویژگی از جمله جنس، سن، قد (سانتی‌متر)، وزن (کیلوگرم)، فشارخون سیستولیک (میلی‌متر

جیوه)، فشارخون دیاستولیک (میلی متر جیوه)، ضربان قلب (دور در دقیقه) و BMI (کیلوگرم بر مترمربع) وجود دارد. برای توضیح وضعیت فشارخون بالا و سپس چهار کلاس وجود دارد که شامل فشارخون طبیعی (سالم)، پیش فشارخون، فشارخون مرحله ۱ و فشارخون مرحله ۲ است. نتایج به دست آمده نشان داده است که می توان با اطمینان از روش های یادگیری ماشین در تعیین خودکار انواع فشارخون بالا استفاده کرد (Nour and Polat, 2020).

همچنین در مقاله رویکرد طبقه بندی بر اساس داده های ژنتیکی برای پیش بینی فشارخون بالا شناخته شده است که فشارخون بالا دلیل اصلی مرگ در سراسر جهان است. میزان مرگ و میر ناشی از این بیماری از دهه گذشته تا ۹۴٪ افزایش یافته است. به دلیل این بیماری، شرایط خطرناک سلامتی مانند نارسایی قلبی، نارسایی کلیه و سکتة مغزی ایجاد می شود. برای جلوگیری از بین رفتن دائمی، نیازی به فن های خودکار برای کشف این بیماری وجود ندارد. اخیراً، اطلاعات ژنتیکی با روش های یادگیری ماشین برای تشخیص بیماری فشارخون بالا ترکیب شده است. تشخیص مبتنی بر ژنتیک نیز نقشی اساسی در تکامل این بیماری دارد. در این مقاله از داده های ژنتیکی و اطلاعات دموگرافیک برای تشخیص این بیماری استفاده شده است. برای شناسایی، دادگان های ژنتیکی ۲۵۰ بیمار از پایگاه های داده مختلف (مخازن آنلاین) جمع آوری شد. تحقیقات از یک مجموعه ویژگی شامل ژن ها، آلل های SNP، آلل های خطرناک، کروموزوم ها، منطقه و ملیت بیماران استفاده می کند. در اینجا، الگوریتمی برای ساختار داده ارائه شده است. نتایج به دست آمده نشان می دهد که بهترین دقت با استفاده از نایو بیزین یعنی ۹۸٪ حاصل می شود (Anis, Fahiem and Tauseef, 2016).

مقاله داده کاوی طبقه بندی-محور برای شناسایی الگوهای ریسک در ارتباط با پرفشاری خون در جمعیت خاورمیانه توسط آقای کبیر و همکارانش در سال ۲۰۱۶ منتشر شد. هدف، شناسایی الگوهای ریسک در ارتباط با بیماری پرفشاری خون با استفاده از داده کاوی بود و بررسی های مورد نیاز با استفاده از مطالعات هم گروهی انجام شده است.

داده از ۶۲۰۵ شرکت کننده که ۴۴٪ مردان بزرگتر از ۲۰ سال و بقیه زنان بودند، همچنین این افراد بدون فشارخون و سابقه بیماری قلبی عروقی بودند، به دست آمده آمد. برای توسعه یک مدل پیش بینی از ۳ نوع الگوریتم درخت تصمیم استفاده شد. اعتبارسنجی روی داده ها با سنجیدن و ارزیابی همه دسته بندی ها انجام شد. الگوریتم درخت آماری کارا سریع در میان مردان و زنان و همچنین درخت رگرسیون و طبقه بندی در میان کل جمعیت بهترین عملکرد را داشته و مقدار C-Static و حساسیت برای مدل های پیش بینی ۷۰٪ و ۷۱٪ برای مردان، ۷۹٪ و ۷۱٪ زنان و ۷۸٪ و ۷۲٪ در جمعیت کل بود.

در مدل های درخت تصمیم، فشارخون سیستولیک، فشارخون دیاستولیک، سن، دور کمر شرکت کنندگان در هر دو جمعیت زنان و مردان، دور مچ، گلوکز پلاسما در میان زنان و گلوکز پلاسما در میان مردان از عوامل عواملی استفاده شده در مدل ریسک است. در مردان، بیشترین ریسک فشارخون در آن هایی که فشارخون سیستولیک بیشتر از ۱۱۵ میلی متر جیوه و سنشان از ۳۰ سال به بالا بود، مشاهده شد. در زنان آن هایی که

فشارخون سیستولیک بیشتر از ۱۱۴ میلی‌متر جیوه و سن بالای ۲۳ سال داشتند جزء افراد با ریسک بالای پرفشاری خون بودند و در مجموع و کل جمعیت، ریسک بالا در فشارخون سیستولیک بالای ۱۱۴ میلی‌متر جیوه و سن بالای ۳۸ سال مشاهده شد (Ramezankhani et al., 2016).

آقای هیان‌کیم و همکارانش مقاله وضعیت تشخیصی و سن در تشخیص فشارخون بالا در پیروی از توصیه‌های شیوه زندگی را در سال ۲۰۱۹ منتشر کردند. در این مقاله فعالیت فیزیکی منظم، ترک سیگار و مصرف کم الکل از رفتارهای مهم در شیوه زندگی است که می‌تواند باعث تغییر در مدیریت فشارخون شود. این مقاله ارتباط وضعیت تشخیصی و سن در تشخیص فشارخون با رفتارهای سبک زندگی در میان افراد با پرفشاری خون بررسی می‌کند. داده‌ها از بررسی آزمایش‌ها ملی بهداشتی و تغذیه‌ای از سال ۲۰۰۷ تا ۲۰۱۲ ($n=5231$) به‌دست‌آمده آمده است. مدل‌های رگرسیون لجستیک چندجمله‌ای برای تخمین ریسک مرتبط^۱ با انطباق رفتارهای شیوه زندگی استفاده شده است.

تشخیص فشارخون بالا با سیگاری بودن فرد همراه بود. بین مدت‌زمان تشخیص از زمان تشخیص و سیگاری بودن ارتباط وجود دارد. نوشیدن بیش از حد ارتباط معکوسی با مدت‌زمان تشخیص دارد. سن بالا در تشخیص با ریسک سیگاری بودن در گذشته ارتباط مستقیم و با نوشیدن بیش از اندازه رابطه منفی دارد. افرادی که ورزش کرده‌اند، حتی اگر کمتر از زمان توصیه شده باشد، در مقایسه با افرادی که مشغول انجام فعالیت‌های بدنی نبودند، از نظر تشخیص سن کمتری دارند و مدت‌زمان تشخیص نیز کوتاه‌تر است.

افراد بیمار باید با ویژگی‌های منظم توسط پزشک بر کنترل فشارخون و اصلاح رفتارهای بهداشتی متمرکز باشند (Kim and Andrade, 2019).

مقاله پیش‌بینی فشارخون بالا بر اساس آنتروپومتری، پارامترهای خون و اسپیرومتری، نیز فشارخون بالا را از عوامل خطر در بیماری‌های قلبی عروقی بیان می‌کند. از آنجایی که ارتباط فشارخون بالا با آنتروپومتری، پارامترهای خون و اسپیرومتری بررسی نشده است، هدف از این مقاله مطالعه شناسایی عوامل ریسک فشارخون بالا در بزرگسالان کره‌ای و بررسی مدل‌های پیش‌بینی فشارخون بالا همراه با آنتروپومتری، پارامترهای خون و اسپیرومتری است.

تجزیه و تحلیل رگرسیون لجستیک دودویی برای ارزیابی اهمیت آماری فشارخون بالا انجام شد و مدل‌های پیش‌بینی با استفاده از رگرسیون لجستیک، بیز و درخت تصمیم‌گیری تهیه شد. از میان تمام عوامل خطر ساز فشارخون بالا، شاخص توده بدنی به‌عنوان بهترین شاخص در هر دو جنسیت شناخته شد. در مردان $CI = 1.204-1.453$ ، 95% ، $odd\ ratio = 1.428$ و زنان $CI = 1.304-1.462$ ، 95% ، $ratio = 1.429$ مدل پیش‌بینی فشارخون، مردان AUC برابر با ۰/۷۰۰ و زنان AUC برابر با ۰/۸۴۵ نشان دادند. این مطالعه عوامل خطر مختلفی را برای پیش فشارخون و فشارخون بالا پیشنهاد می‌کند و یافته‌ها می‌تواند به‌عنوان یک ابزار غربالگری در مقیاس بزرگ برای کنترل و مدیریت فشارخون بالا استفاده شود (Heo, 2018).

¹ Relative risk

مقاله مدل‌های ریاضی برای اثرات فشارخون بالا و استرس روی کلیه و عدم قطعیت در سال ۲۰۱۷ منتشر شد. در این مقاله با استفاده از رویکردهای قطعی و تصادفی مجموعه‌ای از مدل‌های ریاضی برای بیماری مزمن کلیه ارائه شده است. به‌طور خاص، فشارخون بالا و استرس به‌عنوان دلایل اصلی آسیب رساندن و از بین بردن نفرون‌ها در نتیجه کاهش میزان پالایه‌اسیون گلومرولی در نظر گرفته می‌شود که منجر به صدمات جبران‌ناپذیری به عملکرد کلیه می‌شود و ممکن است منجر به نارسایی کلیه یا حتی از دست دادن کلیه شود. در این مقاله تعادل برای هر دو مدل قطعی و تصادفی تجزیه و تحلیل می‌شود. در این راستا، قضیه‌های ریاضی مربوط به ثبات و ثبات تصادفی مدل‌ها ارائه شده است. مسئله مهم دیگر، که در کار ارائه شده مورد بحث قرار گرفته است، مشکل عدم اطمینان برای مدل‌های تصادفی است. برای محاسبه عدم قطعیت مدل‌های تصادفی از شبیه‌سازی‌ها در نرم‌افزار متلب استفاده شد. علاوه بر این، چارچوبی برای پیش‌بینی آینده مدل‌های پیشنهادی در موارد خاص ارائه گردید. سرانجام، مدل‌ها و نتایج نظری در برنامه با استفاده از آن‌ها در داده‌های بالینی واقعی مورد تأیید واقع می‌شوند (Waezizadeh et al., 2018). محدوده پرفشاری خون به‌صورت جدول ۲-۷ هست.

جدول ۲-۷ اعداد مربوط به پرفشاری خون (Hengel, Benitah and Wenzel, 2022)

سیستولی (mm Hg)	دیاستولی (mm Hg)	
کمتر از ۱۲۰	کمتر از ۸۰	عادی
۱۲۰-۱۲۹	کمتر از ۸۰	مرتفع (پرفشاری خون)
۱۳۰-۱۳۹	۸۰-۹۰	سطح ۱ پرفشاری خون
۱۴۰ و بالاتر	۹۰ یا بالاتر	سطح ۲ پرفشاری خون
بیشتر از ۱۸۰	بیشتر از ۱۲۰	نقطه بحرانی پرفشاری خون

۲-۸-۱- عوامل محیطی پرفشاری خون

یکی دیگر از عوامل مؤثر بر بیماری فشارخون بالا، عوامل محیطی مانند مصرف نمک، مصرف الکل، سیگار کشیدن، عدم تحرک هست؛ که می‌توان با تغییر در این موارد از ابتلا و بروز بیماری فشارخون بالا جلوگیری کرد (Wang et al., 2020). عوامل مؤثر بر بیماری فشارخون بالا را می‌توان به‌طور کلی به شکل ۲-۱۰ نمایش داد.

در این بررسی در خصوص سن به عنوان پارامتر مؤثر بیان شده است: افزایش نمک با افزایش سن، اختلال یا از بین رفتن توانایی تعادل کل سدیم بدن در طول زمان است. سدیم فراوان ترین کاتیون در مایع خارج سلولی است که از مایعات بینابینی و داخل عروقی تشکیل شده است و از بین رفتن این تعادل باعث افزایش نمک و پرفشاری خون هست.

۲-۸-۲- مدل های پیش بینی کننده بیماری فشارخون

شناخت بیماری فشارخون تاریخچه ای بیش از ۵۰ سال دارد. فریزر رابرتز با پیکرینگ و همکارانش (McKusick *et al.*, 1960) اشاره کرده بودند سطح فشارخون یک متغیر پیوسته است و همچنین با رسم نقشه لگاریتمی آن خاطرنشان کردند که چولگی مثبت آن از بین می رود. وی همچنین داده های بارنز و براون را در مورد فشارخون بستگان درجه یک بیماران مبتلا به مسمومیت خونی بارداری مجدداً تجزیه و تحلیل کرد و با استفاده از دستگاه امتیازدهی خود برای از بین بردن اثرات سن و جنس نتیجه گرفت که فشارخون در چنین اقوامی بیشتر از جمعیت در کل است.

مطالعات پیکرینگ و همکارانش (McKusick, 1960) آن ها را به این نتیجه رسانده است که توزیع فراوانی فشارخون پیوسته است و تعیین ژنتیکی فشارخون به صورت چندژنی است. در سال ۱۹۵۹ دو تحلیل گزارش داده شده است که یکی توسط پلات از منچستر و دیگری توسط موريسن و موريس از لندن هست که نشانگر دونمایی بودن^۱ فشارخون هست که با این دیدگاه یک ژن غالب عامل مهمی در فشارخون بالا است. کورتس و اسپنس (Kurtz and Spence, 1993) بیان می کنند که فشارخون یک ویژگی کمی پیچیده ای است که توسط عوامل محیطی چندگانه و عوامل ژنتیکی تعیین می شود. هرچند که اشکال ساده ای برای فشارخون بالا در نظر گرفته شده است، اما فشارخون اساسی بر اساس روش های پیچیده وراثت پذیری مشخص می گردد.

بر اساس این تحقیق امکان جستجوی مناطق کروموزومی که شامل ژن بیماری زا فشارخون بالا باشد، فراهم شده است. برای مثال، مطالعات پیوند و ارتباط اخیر، احتمال اینکه یک مکان تنظیم کننده فشارخون ممکن است در ژن آنژیوتانسینوزن^۲ روی کروموزوم یک یا در نزدیکی آن باشد را افزایش داده است.

واعظی زاده و دیگران (Waezizadeh *et al.*, 2018) به بررسی مدل های مختلف ریاضی برای بیماری مزمن کلیه که منجر به افزایش فشارخون می شود پرداخته اند و فشارخون بالا و استرس را به عنوان اصلی ترین علت صدمه و از بین بردن نفرون در نتیجه کاهش میزان پالایهاسیون گلومرولی بیان کرده اند که منجر به آسیب های جبران ناپذیری در عملکرد کلیه می شود و حتی ممکن است در نهایت باعث از دست دادن کلیه شود.

در سال ۲۰۲۰ وو و همکاران بر مبنای ملاک های سنتی پیش بینی فشارخون تحقیقی را مبتنی بر روش های یادگیری ماشین انجام دادند که نشان می دهد روش های یادگیری ماشین بهترین نتیجه را نسبت به روش های رگرسیون و FRS دارد (Wu *et al.*, 2020). همچنین در تحقیق دیگری با استفاده عوامل بالینی سن،

¹ bimodality

² angiotensinogen

مصرف سیگار، سابقه خانوادگی، فعالیت فیزیکی و شاخص جرمی بدن^۱ دقت ۰.۸۵۴ و با اضافه کردن پارامترهای ژنتیکی دقت ۰.۸۷۱ به دست آمده است (Niu et al., 2021; Chowdhury et al., 2022). در تحقیقات بر مبنای PRS دقت به دست آمده بر اساس AUC برابر ۰.۷۱ (با در نظر گرفتن سن، جنسیت و BMI و SNP ۸۵۸) و ۰.۷۶۴ (با در نظر گرفتن سن، نژاد، سابقه مصرف سیگار و دیابت و SNP ۳,۰۴۳,۶۳۷) و ۰.۷۵ (با در نظر گرفتن جنسیت، BMI و SNP ۲۱۶۶) هست (Giontella et al., 2021; Kurniansyah et al., 2022; Maj et al., 2022).

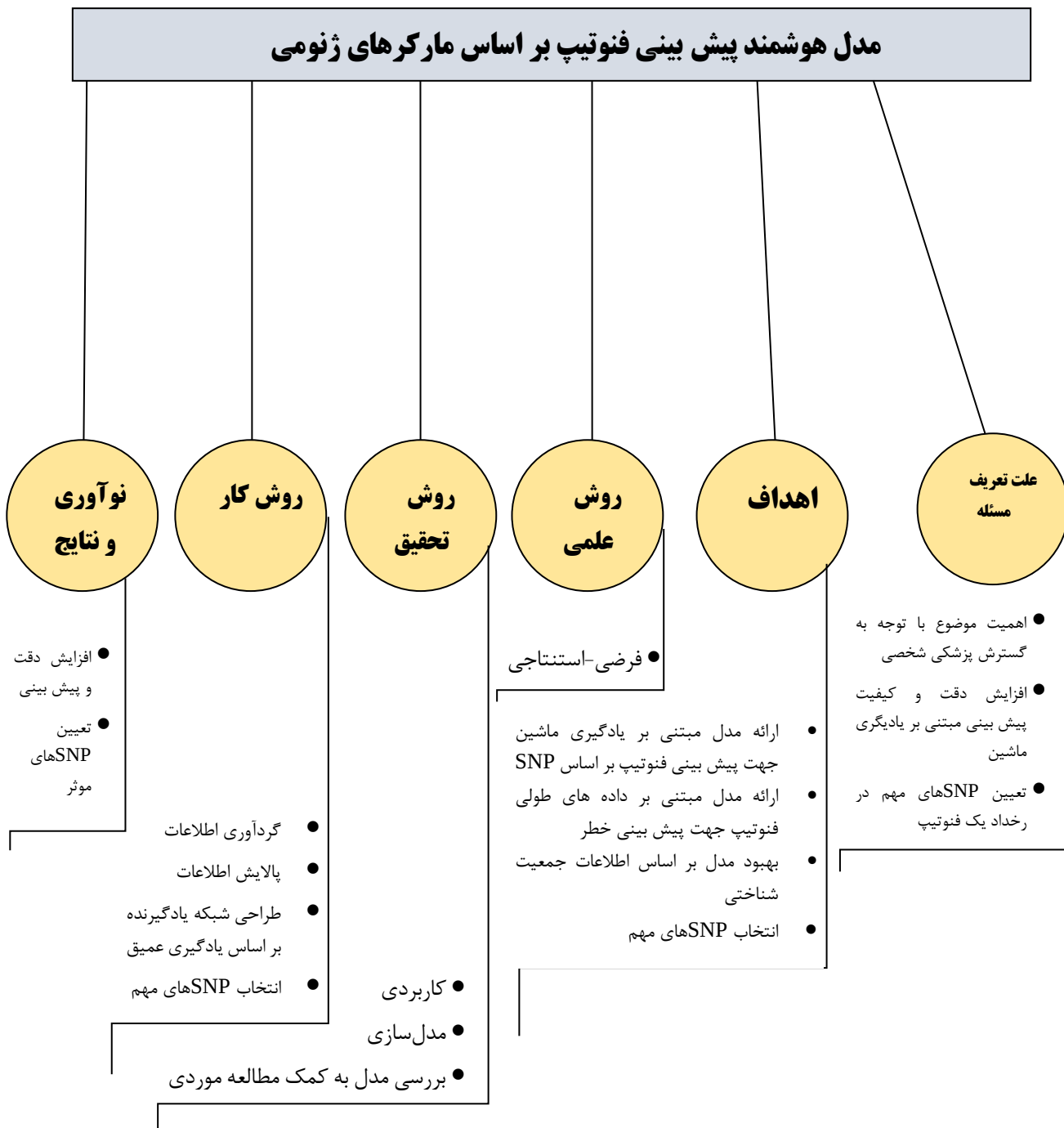
۲-۹- جمع بندی

در این فصل بایان تحقیقات مختلف حوزه ژنوم و فنوتیپ به بررسی مدل های متنوع پرداخته شد و بایان نقاط قوت و ضعف روش ها و اهداف آتی که در تحقیقات اشاره شده بود، ساختار ادبیات موضوع شکل گرفت. در ادامه بایان دقت های گزارش شده و نقاط ضعف مدل های اشاره شده، ساختار مدل جدید پیشنهاد گردید که در فصل بعدی به طور کامل توضیح داده می شود.

¹ Body Mass Index (BMI)

۳- روش‌شناسی تحقیق

در این فصل به روش‌شناسی تحقیق پیش رو می‌پردازیم. بر اساس نوع تحقیق، این تحقیق یک پژوهش کاربردی است، چراکه در این پژوهش بر آنیم تا از طریق مطالعه ادبیات موضوع و حوزه تحقیق، مدلی مبتنی بر یادگیری ماشین بر اساس داده‌های طولی رخ‌نمود و نشانگرهای ژنومی، طراحی و ارائه نماییم. نمونه‌گیری انتخابی در این تحقیق به صورت کاملاً انتخابی بر روی داده‌های ژنومی و رخ‌نمودهای ذخیره‌شده انجام می‌شود. شمای کلی تحقیق در شکل ۳-۱ آمده است.



شکل ۳-۱ شمای کلی تحقیق (آونگ تحقیق)

بر اساس ادبیات بررسی شده در حوزه مورد مطالعه، روش‌های آماری و تحلیل فراداده بیشترین کاربرد را نسبت به سایر روش‌ها دارند؛ اما همان‌طور که بررسی شد این روش‌ها دارای معایبی است که بر همین اساس روش‌های مبتنی بر یادگیری ماشین، پیشنهاد و دورنمای تحقیقات در این حوزه هست (Esteva et al., 2019). بر همین اساس در این فصل ابتدا داده‌ها و محدوده‌های تحقیق بیان شده و در نهایت بایان نوآوری‌ها فصل به اتمام می‌رسد.

۳-۱- داده‌های تحقیق

اگرچه بسیاری از ژن‌هایی که در ایجاد بیماری‌ها دخالت دارند، شناسایی شده‌اند، اما متخصصان بر این باور هستند که بسیاری از آن‌ها هنوز کشف نشده‌اند. داده‌های ژن‌نمود و رخ‌نمود مورد نیاز برای طراحی مدل، می‌تواند از پایگاه داده‌های موجود و یا به وسیله شبیه‌سازی به دست آید. با توجه به مسئله تحقیق، داده‌های تحقیق شامل سه بخش هست:

۱. داده‌های ژنومی

۲. داده‌های جمعیت شناختی و بالینی مانند سن، جنس، قومیت و موارد دیگر

۳. داده‌های طولی رخ‌نمود فشارخون در دوره‌های مختلف اندازه‌گیری شده

برای به دست آوردن داده‌های تحقیق از طرح ژنتیک کاردیومتابولیک تهران، پژوهشکده علوم غدد درون‌ریز و متابولیسم دانشگاه علوم پزشکی شهید بهشتی استفاده می‌شود و این تحقیق با کد IR.SBMU.ENDOCRINE.REC.1400.008 توسط کمیته ملی اخلاق در تحقیقات زیست پزشکی در اردیبهشت ۱۴۰۰ تأیید شده است.

مطالعه قند و لیپید تهران یک مطالعه گسترده در حوزه بهداشتی است که برای نخستین بار در کشور توسط مرکز تحقیقات غدد درون‌ریز و متابولیسم دانشگاه علوم پزشکی شهید بهشتی از سال ۱۳۷۶ به بررسی عوامل خطر ساز بیماری‌های غیر واگیر در شهروندان ساکن شرق کلان‌شهر تهران پرداخته است.

این مطالعه با رویکردی جامعه‌نگر، بر پایه اطلاعات اولیه مبنی بر شیوع پیشرونده بیماری‌های غیر واگیر، روی بیش از ۱۵۰۰۰ نفر از جمعیت ساکن در شرق تهران انجام شده است و تاکنون در ۶ مقطع زمانی با بازه‌های سه‌ساله از این افراد اطلاعات جمع‌آوری شده است. این مطالعه، قدیمی‌ترین، بزرگ‌ترین و معتبرترین مطالعه‌ی آینده‌نگر در ایران است (Kheradmand et al., 2015). جزئیات طراحی این مطالعه به تفصیل شرح داده شده است (Azizi et al., 2009). جزئیات مطالعه گسترده‌ی ژنومی در مطالعه ژنتیک کاردیومتابولیک تهران که روی بیش از ۱۴۰۰۰ نفر از افراد شرکت‌کننده در طرح قند و لیپید تهران انجام شده است (Daneshpour et al., 2017). در مطالعه مذکور، کلیه افرادی که تحت پوشش مطالعه قند و لیپید تهران قرار گرفته‌اند، به صورت خانوادگی بررسی شدند. در مجموع، ۳۹۵۹ خانواده شامل ۱۳۸۸۷ نفر در مطالعه گسترده ژنومی بررسی و هر نفر برای بیش از ۷۰۰ هزار نشانگر ژنتیکی تعیین ژن‌نمود شده است. از میان ۱۹۹۳ رخ‌نمود بررسی شده در کلیه GWAS‌هایی که تاکنون انجام شده‌اند و اطلاعات آن در کاتالوگ GWAS موجود است، ۲۰۹ رخ‌نمود آن

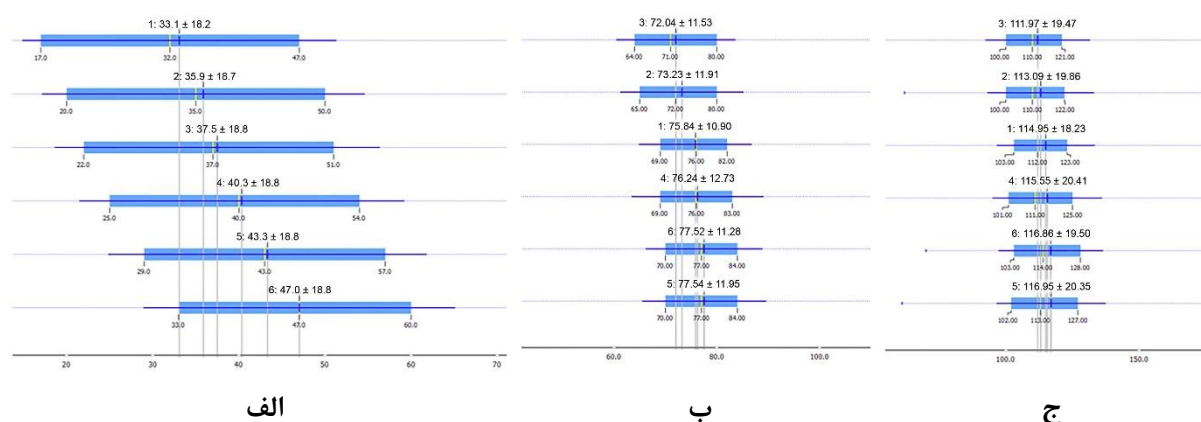
در این جمعیت بررسی شده و اطلاعات رخ‌نمودی و ژن‌نمودی آن موجود است. امکان تعریف و یا اندازه‌گیری ۴۰۵ رخ‌نمود دیگر نیز در نمونه‌های سرمی این افراد وجود دارد.

در طی سال‌های گذشته با استفاده از این داده‌ها، تحقیقات گسترده‌ای در زمینه‌ی تأثیر تغییرات رژیم غذایی، شیوع دیابت و تغییرات میزان لیپید، رخ‌نمودهای مرتبط با چاقی و عوامل ژنتیکی اثرگذار انجام شده است که باوجود این زیرساخت و افزوده شدن داده‌های حاصل از فن‌آوری امیکس (دانشپور ۱۳۹۵، *et al.*) می‌توان گام مؤثری در جهت پزشکی شخصی برداشت. تمامی شرکت‌کنندگان یک نمونه خون پایه کافی برای آنالیز پلاسما و DNA ارائه کرده‌اند و برای آنالیزهای مبتنی بر خون و پیگیری طولانی‌مدت رضایت کتبی دادند. ضمناً برای اندازه‌گیری فشارخون، ابتدا شرکت‌کنندگان به مدت ۱۵ دقیقه روی صندلی می‌مانند و پس‌از آن پزشک متخصص فشارخون را دو بار با استفاده از فشارسنج جیوه‌ای استاندارد اندازه‌گیری می‌کند. حداقل یک‌فاصله ۳۰ ثانیه‌ای بین این دو اندازه‌گیری جداگانه وجود دارد و میانگین دو اندازه‌گیری به‌عنوان فشارخون شرکت‌کننده ثبت می‌شود (Tohidi *et al.*, 2012).

با توجه به محدودیت‌های سخت‌افزاری در پردازش حجم بالای داده‌های کامل ژنوم انسان، در این تحقیق داده‌های ژنومی به‌صورت چپ^۱ شامل ۶۵۲۹۱۹ SNP که برای ۱۱۰۸۸ نفر در دوره‌های مختلف هست، استفاده شده است.

لازم به ذکر است معیار ورود داده ثبت اطلاعات ژن‌نمودی و رخ‌نمودی در این طرح هست و معیارهای خروج شامل موارد زیر است:

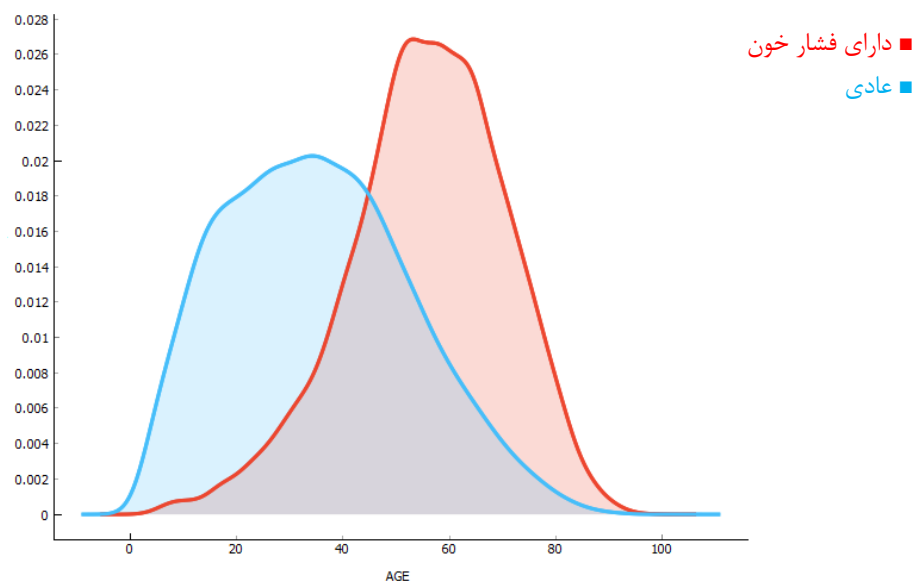
- ۱- SNP‌هایی که دارای بیش از ۵ درصد داده گمشده می‌باشند از ژن‌نمود خارج شده‌اند.
- ۲- افرادی که در ژن‌نمود خود دارای بیش از ۵ درصد داده گمشده هستند از نمونه‌ها خارج می‌شوند.
- ۳- افرادی که با مصرف داروی خاص تغییرات معنی‌داری در فشارخون داشته باشند (این تغییرات ممکن است باعث کاهش یا افزایش باشد) از نمونه‌ها خارج شده‌اند. شکل ۲-۳ توزیع سن و مقادیر DBP و SBP را در فازهای مختلف نشان می‌دهد.



شکل ۲-۳ (الف) نمودار توزیع سن در فاز اول (بالا) تا فاز ششم (پایین) (ب) نمودار توزیع DBP (ج) نمودار توزیع SBP

¹ Chip

در شکل ۳-۳ توزیع سن با توجه به برچسب فشارخون را نشان می‌دهد که نقطه آستانه برابر ۴۸ سال را نمایش می‌دهد.



شکل ۳-۳ توزیع سن به تفکیک برچسب فشارخون

۳-۱-۱- پروژه ژنوم مرجع ایرانیان^۱

پروژه ژنوم مرجع ایرانیان برای ۸ قومیت ایرانی با کمک مطالعه ژنتیک کاردیومتابولیک تهران و مرکز تحقیقات ژنتیک انسانی کوثر بر روی بیش از ۱۴۰۰۰ نفر انجام شده است. این مطالعه با همکاری یکی از معتبرترین مراکز تحقیقات ژنتیک جهان یعنی، کمپانی ايسلندی deCODE Genetics® در سال ۲۰۱۴ شروع و در سال ۲۰۱۷ با انتقال فناوری به ایران و در اختیار داشتن تمامی داده‌ها در پژوهشکده علوم غدد درون‌ریز و متابولیسم دانشگاه علوم پزشکی شهید بهشتی به انتها و بهره‌برداری رسیده است (عزیزی، هدایتی و دانشپور، ۱۳۹۶).

نمونه ژنوم از گلبول‌های سفید با روش استاندارد پروتئیناز K استخراج شده با تراشه HumanOmniExpress-24-v1-0 توسط شرکت deCODE مطابق با دستگاه Illumina ژن‌نمود شدند (Kolifarhood et al., 2021). داده‌های ژنومی برای هر فرد دارای ۶۰۰ هزار SNP می‌باشند که در پروتجای ped ذخیره‌سازی شده است.

۳-۲- ارزیابی مدل

دانشی که در مرحله یادگیری مدل تولید می‌شود، می‌بایست در مرحله ارزیابی مورد تحلیل قرار گیرد تا بتوان ارزش آن را تعیین نمود و در نتیجه، کارایی الگوریتم یادگیرنده مدل را نیز مشخص کرد. این معیارها را می‌توان هم برای دادگان‌های آموزشی در مرحله یادگیری و هم برای مجموعه رکوردهای آزمایشی در مرحله ارزیابی محاسبه نمود (Stehman, 1997).

¹ Iranian Reference Genome Project

در بحث دسته‌بندی^۱ یک دادگان^۲ با استفاده از روش‌های رده‌بندی، هدف دستیابی به بالاترین دقت ممکن در رده‌بندی و تشخیص دسته‌ها است. در برخی از مسائل، تشخیص صحیح نمونه‌های مربوط به یکی از دسته‌ها برای ما اهمیت بیشتری دارد. به عنوان مثال، تحقیقی را در نظر بگیرید که در آن، هدف شناسایی افراد مبتلابه یک نوع خاص از یک بیماری خطرناک است. فرض کنید برای افرادی که مبتلابه این بیماری هستند، خطر مرگ وجود دارد و جهت رفع این خطر، نیاز به دریافت نوعی داروی خاص دارند. در این شرایط، تشخیص درست بیماران دارای اهمیت بسیار زیادی است.

به این معنا که خطا در تشخیص افراد سالم قابل چشم‌پوشی است اما برای شناسایی افراد بیمار نمی‌توان این احتمال را به جان خرید. به عبارت دیگر، انتظار ما تشخیص تمام افراد بیمار است، بدون جا انداختن، حتی اگر فرد سالمی به اشتباه جز افراد بیمار رده‌بندی شود. در چنین مواقعی، که دقت تشخیص یک دسته در مقایسه با دقت تشخیص کلی، اهمیت بیشتری دارد، مفهوم ماتریس درهم‌ریختگی^۳ استفاده می‌شود (Powers and Ailab, 2011).

۱-۲-۳- ماتریس درهم‌ریختگی

بر اساس مثالی که پیش‌تر بیان شد، فرض کنید تعلق به دسته افراد بیمار را مثبت بودن^۴ و عدم تعلق به این دسته را منفی بودن^۵ در نظر بگیریم. هر نمونه یا فردی در واقعیت، متعلق به یکی از برچسب‌های مثبت یا منفی است و از سوی دیگر، از هر الگوریتمی که برای رده‌بندی داده‌ها استفاده شود، درنهایت هر نمونه عضو یکی از این دودسته^۶، رده‌بندی خواهد شد. بنابراین برای هر نمونه داده، یکی از چهارحالتی که در ادامه بیان شده، ممکن است اتفاق بیفتد.

نمونه عضو دسته مثبت باشد و عضو همین کلاس تشخیص داده شود: مثبت صحیح یا True Positive
نمونه عضو کلاس مثبت باشد و عضو کلاس منفی تشخیص داده شود: منفی کاذب یا False Negative
نمونه عضو کلاس منفی باشد و عضو همین کلاس تشخیص داده شود: منفی صحیح یا True Negative
و درنهایت، نمونه عضو کلاس منفی باشد و عضو کلاس مثبت تشخیص داده شود: مثبت کاذب یا False Positive
پس از اجرای الگوریتم رده‌بندی، با توجه به توضیحات و تعاریف ذکرشده، می‌توان عملکرد یک رده‌بند را به کمک جدولی به شکل زیر بررسی کرد.

¹ Classification

² Data Set

³ Confusion Matrix

⁴ Positive

⁵ Negative

⁶ Class

جدول ۳-۱ ماتریس درهم‌ریختگی

		برچسب پیش‌بینی شده	
		مثبت	منفی
برچسب شناخته شده	مثبت	TP	FN
	منفی	FP	TN

این جدول را اصطلاحاً ماتریس درهم‌ریختگی می‌گویند. جدول یا ماتریس درهم‌ریختگی، نتایج حاصل از رده‌بندی را بر اساس اطلاعات واقعی موجود، نمایش می‌دهد. حال بر اساس این مقادیر می‌توان معیارهای مختلف ارزیابی رده‌بند و اندازه‌گیری دقت را تعریف کرد. پارامتر دقت^۱، متداول‌ترین، اساسی‌ترین و ساده‌ترین معیار اندازه‌گیری کیفیت یک رده‌بند است و عبارت است از میزان تشخیص صحیح رده‌بند در مجموع دودسته. این پارامتر درواقع نشان‌گر میزان الگوهای است که درست تشخیص داده شده‌اند و بر اساس ماتریس ارائه شده در بالا، با رابطه زیر تعریف می‌شود.

$$\text{Accuracy} = (TP+TN) / (TP+FN+FP+TN) \quad \text{رابطه ۳-۱}$$

البته، پارامتر دقت معمولاً به صورت درصد بیان می‌شود. اما پارامترهای دیگری نیز علاوه بر معیار دقت وجود دارند که می‌توان به سادگی از این ماتریس استخراج کرد. یکی از متداول‌ترین آن‌ها، معیار حساسیت^۲ است که آن را نرخ پاسخ‌های مثبت درست^۳ نیز می‌گویند. حساسیت به معنی نسبتی از موارد مثبت است که آزمایش آن‌ها را به درستی به عنوان نمونه مثبت تشخیص داده است. این پارامتر به صورت زیر محاسبه می‌شود.

$$\text{Sensitivity (TPR)} = TP / (TP+FN) \quad \text{رابطه ۳-۲}$$

درواقع، «حساسیت» همان معیار بحث شده در مورد مثال بالا است. معیاری که مشخص می‌کند رده‌بند، به چه اندازه در تشخیص تمام افراد مبتلا به بیماری موفق بوده است. همان‌گونه که از رابطه فوق مشخص است، تعداد افراد سالمی که توسط رده‌بند به اشتباه به عنوان فرد بیمار تشخیص داده شده‌اند، هیچ تأثیری در محاسبه این پارامتر ندارد و درواقع زمانی که پژوهش‌گر از این پارامتر به عنوان پارامتر ارزیابی برای رده‌بند خود استفاده می‌کند، هدفش دستیابی به نهایت دقت در تشخیص نمونه‌های کلاس مثبت است.

در نقطه مقابل این پارامتر، ممکن است در مواقعی دقت تشخیص کلاس منفی حائز اهمیت باشد. از متداول‌ترین پارامترها که معمولاً در کنار حساسیت بررسی می‌شود، پارامتر خاصیت^۴ است که به آن نرخ پاسخ‌های منفی درست^۵ نیز می‌گویند. خاصیت به معنی نسبتی از موارد منفی است که آزمایش آن‌ها را به درستی به عنوان نمونه منفی تشخیص داده است. این پارامتر به صورت زیر محاسبه می‌شود.

$$\text{Specificity (TNR)} = TN / (TN+FP) \quad \text{رابطه ۳-۳}$$

^۱ Accuracy

^۲ Sensitivity

^۳ True Positive Rate

^۴ Specificity

^۵ True Negative Rate

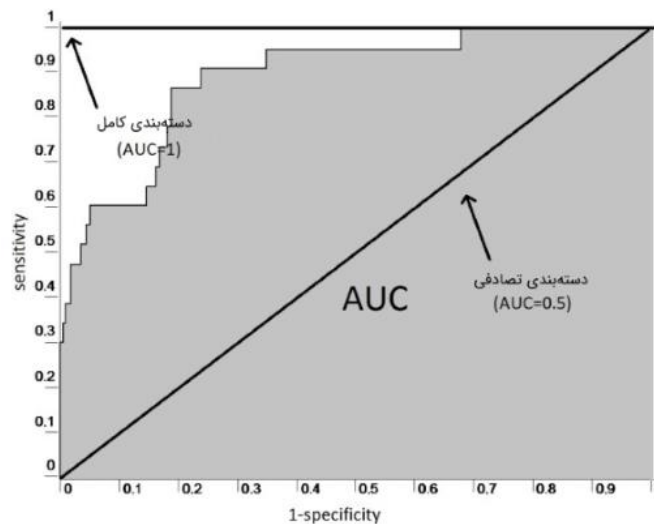
این دو پارامتر (حساسیت و خاصیت) نیز مشابه معیار دقت، معمولاً به صورت درصد بیان می‌شوند. واضح است که پیش‌بینی عالی، پیش‌بینی است که مقادیر حساسیت و ویژگی مربوط به آن، هر دو صد درصد باشند؛ اما احتمال وقوع این اتفاق در واقعیت بسیار کم است و همیشه یک حداقل خطایی وجود دارد. پارامترهای حساسیت و خاصیت، بنا بر ماهیتی که دارند همواره در رقابت با یکدیگر هستند. یعنی افزایش یکی با کاهش دیگری همراه است و برعکس. همین وضعیت منجر به تولید ابزاری دیگر برای ارزیابی کیفیت رده‌بندها شده است. منحنی مشخصه عملکرد سیستم^۱ (ROC)، عبارت است از منحنی که ارتباط بین دو پارامتر حساسیت و خاصیت را بیان می‌کند (Fawcett, 2006).

بسیاری از رده‌بندها همانند روش‌های مبتنی بر درخت تصمیم و یا روش‌های مبتنی بر قانون، به گونه‌ای طراحی شده‌اند که تنها یک خروجی دودویی (مبنی بر تعلق ورودی به یکی از دودسته ممکن) تولید می‌کنند. به این نوع رده‌بندها که تنها یک خروجی مشخص برای هر ورودی تولید می‌کنند، رده‌بندهای گسسته گفته می‌شود که این رده‌بندها تنها یک نقطه در فضای ROC تولید می‌کنند.

به طور مشابه رده‌بندهای دیگری نظیر رده‌بندهای مبتنی بر روش بیز و یا شبکه‌های عصبی نیز وجود دارند که یک احتمال و یا امتیاز برای هر ورودی تولید می‌کنند، که این عدد بیانگر درجه تعلق ورودی به یکی از دودسته موجود هست. این رده‌بندها پیوسته نامیده می‌شوند و به دلیل خروجی خاص این رده‌بندها یک آستانه جهت تعیین خروجی نهایی در نظر گرفته می‌شود.

چنانکه در شکل ۳-۴ مشاهده می‌کنید، محور عمودی این نمودار نشان‌دهنده نرخ مثبت صحیح، و محور افقی نشان‌دهنده مقدار نرخ مثبت غلط است. نتایج مختلف رده‌بندی نشانگر نقاط مختلف بر روی این نمودار هستند و در نهایت یک منحنی را تشکیل می‌دهند. با توجه به شکل زیر، در بهترین حالت و با فرض رده‌بندی صد درصد صحیح در هر دودسته، نقطه مربوطه عبارت است از نقطه گوشه بالای سمت چپ، یعنی نقطه (۰،۱) و نیز با فرض رده‌بندی به صورت تصادفی، نقطه متناظر در منحنی، یکی از نقاط موجود روی خط واصل نقطه (۰،۰) و نقطه (۱،۱) خواهد بود. در واقعیت، منحنی حاصل از یک رده‌بندی، منحنی بین این دو حالت است. یک طبقه بند کامل دارای AUC برابر با ۱ خواهد و یک رده‌بندی ضعیف دارای AUC ۰.۵ است. یک رده‌بندی کننده خوب دارای AUC بیشتر از ۰.۵ خواهد بود و با رسیدن به ۱ عالی در نظر گرفته می‌شود (Martinez-Ríos et al., 2021).

^۱ Receiver Operating Characteristic: ROC



شکل ۳-۴ منحنی ROC و نمایش AUC

مساحت زیر نمودار^۱ (AUC) که در شکل ۳-۸ آمده است، به عنوان یک معیار برای ارزیابی عملکرد رده‌بند مورد استفاده قرار می‌گیرد. با توجه به توضیحاتی که پیش‌تر ارائه شد، بدیهی است که در حالت ایدئال، مساحت زیر منحنی برابر با بیشترین مقدار خود، یعنی یک است. بنابراین، هر چه مساحت زیر نمودار به عدد یک نزدیک‌تر باشد، به معنای بهتر بودن عملکرد رده‌بند است. علاوه بر دو پارامتر حساسیت و خاصیت، پارامترهای دیگری هم از ماتریس درهم‌ریختگی استخراج می‌شوند که هر یک بیان‌کننده مفهومی هستند و کاربردهای متفاوتی دارند.

از جمله آن‌ها، دو پارامتر مقدار پیش‌بینی‌شده مثبت^۲ و مقدار پیش‌بینی‌شده منفی^۳ هستند، که برای بیان «نسبت پاسخ‌های درست در هر دسته» استفاده می‌شوند. ارزش اخباری مثبت، بیان‌کننده این است که چند درصد از الگوهایی که رده‌بند آن‌ها را مثبت تشخیص داده، در واقعیت هم مثبت هستند و به همین ترتیب، ارزش اخباری منفی نشان می‌دهد که چند درصد از نمونه‌هایی که عضو دسته منفی تشخیص داده شده‌اند، در واقعیت هم عضو همین دسته هستند. این دو پارامتر نیز به سادگی از روی ماتریس درهم‌ریختگی به صورت رابطه زیر است.

$$NPV = TN / (TN + FN), PPV = TP / (TP + FP) \quad \text{رابطه ۳-۴}$$

پارامتر مهم دیگری به نام معیار اف^۴ وجود دارد که برای ارزیابی عملکرد رده‌بندها بسیار مورد استفاده قرار می‌گیرد و از ترکیب دو پارامتر حساسیت و ارزش اخباری مثبت حاصل می‌شود. با این توضیح که پارامتر ارزش اخباری مثبت را اصطلاحاً دقت^۵ و حساسیت را اصطلاحاً بازخوانی^۶ می‌نامند، «معیار اف» به صورت رابطه زیر است.

$$F\text{-measure} = 2 * (\text{Recall} * \text{Precision}) / (\text{Recall} + \text{Precision}) \quad \text{رابطه ۳-۵}$$

^۱ Area Under Curve

^۲ Positive Prediction Value

^۳ Negative Predictive Values

^۴ F-Measure

^۵ Precision

^۶ Recall

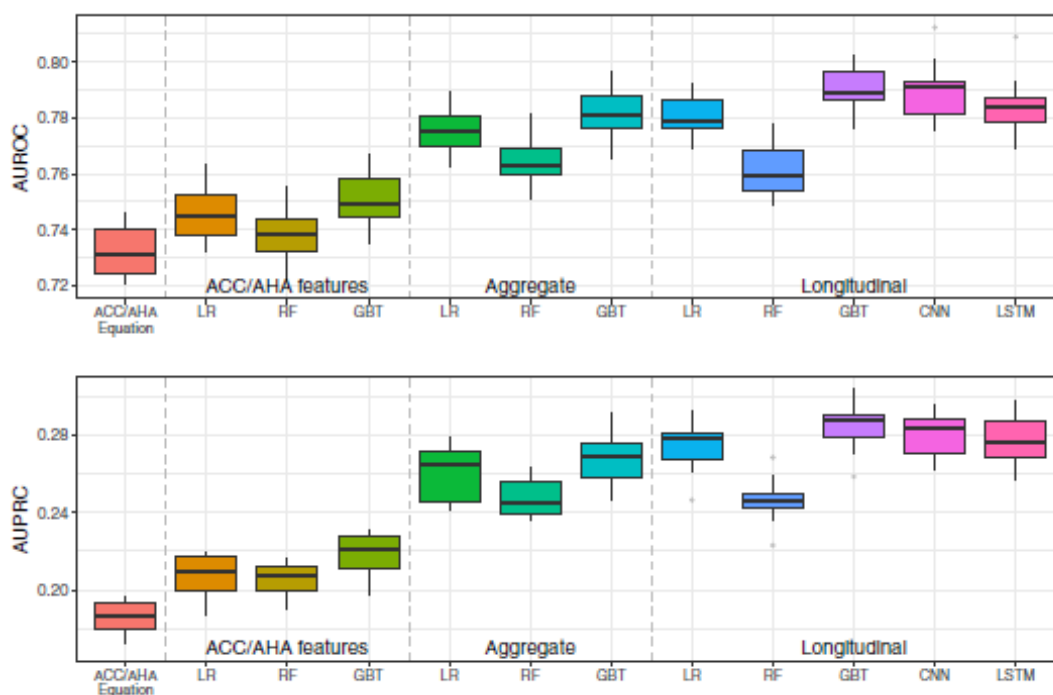
ماتریس درهم‌ریختگی، باوجود منطق و ساختار ساده‌ای که دارد، مفهومی قدرتمند است که در انواع تحقیقات، می‌تواند به‌تنهایی اطلاعاتی جامع از نحوه عملکرد رده‌بند ارائه کند.

یکی از دیگر از مشکلاتی که در ارزیابی مدل مؤثر است بیش‌برازش^۱ داده‌ها بر مدل هست. برای حل این مشکل از رویه قاعده‌مند سازی^۲ استفاده می‌کنند (Zhu et al., 2018). هدف از رویه قاعده‌مند سازی جلوگیری از بیش‌برازش به داده‌ها توسط مدل است و در نتیجه با مشکل واریانس بالا طرف است. جدول ۲-۳ خلاصه‌ای از انواع روش‌های متداول قاعده‌مند سازی را ارائه می‌دهد.

جدول ۲-۳ روش‌های قاعده‌مند سازی

LASSO	Ridge	Elastic Net
ضرایب را تا ۰ کاهش می‌دهد برای انتخاب متغیر مناسب است	ضرایب را کوچک‌تر می‌کند	بین انتخاب متغیر و ضرایب کوچک مصالحه می‌کند

یکی دیگر از معیارهایی که در ادبیات موضوع مطرح‌شده است معیار AUPRC هست که سطح زیر منحنی نمودار دقت و بازخوانی است که در داده‌هایی که از نظر برچسب نامتعادل هستند دارای کارایی بهتری هست. در شکل ۳-۵ نمودار مقایسه پیش‌بینی بیماری قلبی مبتنی بر داده‌های ژنومی در دو معیار AUROC و AUPRC آمده است (Zhao et al., 2018).



شکل ۳-۵ مقایسه AUPRC و AURPC

۲-۲-۳ روش‌های بخش‌بندی داده جهت ارزیابی

از روش‌های ارزیابی الگوریتم‌های رده‌بندی (که در این الگوریتم روال کاری بدین‌صورت است که مدل رده‌بندی توسط دادگان آموزشی ساخته‌شده و به‌وسیله دادگان آزمایشی مورد ارزیابی قرار می‌گیرد) می‌توان به روش

^۱ Overfit

^۲ Regularization

بسط یافته^۱ اشاره کرد که در این روش چگونگی نسبت تقسیم دادگان‌ها (به دو دادگان آموزشی و دادگان آزمایشی) بستگی به تشخیص تحلیلگر دارد که معمولاً دوسوم برای آموزش و یک‌سوم برای ارزیابی در نظر گرفته می‌شود. مهم‌ترین مزیت این روش سادگی و سرعت بالای عملیات ارزیابی است ولیکن روش بسط یافته معایب زیادی دارد از جمله اینکه دادگان‌های آموزشی و آزمایشی به یکدیگر وابسته خواهند شد، درواقع بخشی از دادگان اولیه که برای آزمایش جدا می‌شود، شانس برای حضور یافتن در مرحله آموزش ندارد و به‌طور مشابه در صورت انتخاب یک رکورد برای آموزش دیگر شانس برای استفاده از این رکورد برای ارزیابی مدل ساخته‌شده وجود نخواهد داشت. همچنین مدل ساخته‌شده بستگی فراوانی به چگونگی تقسیم دادگان اولیه به دادگان‌های آموزشی و آزمایشی دارد. چنانچه روش بسط یافته را چندین بار اجرا کنیم و از نتایج حاصل میانگین‌گیری کنیم از روشی موسوم به زیر نمونه‌گیری تصادفی^۲ استفاده نموده‌ایم. که مهم‌ترین عیب این روش نیز عدم کنترل بر روی تعداد دفعاتی که یک رکورد به‌عنوان نمونه آموزشی و یا نمونه آزمایشی مورد استفاده قرار می‌گیرد، است. به‌بیان دیگر در این روش ممکن است برخی رکوردها بیش از سایرین برای یادگیری و یا ارزیابی مورد استفاده قرار گیرند.

چنانچه در روش زیر نمونه‌گیری تصادفی به شکل هوشمندانه‌تری عمل کنیم به‌صورتی که هرکدام از رکوردها به تعداد مساوی برای یادگیری و تنها یک‌بار برای ارزیابی استفاده شوند، روش مزبور در متون علمی بانام اعتبارسنجی متقابل^۳ شناخته می‌شود.

همچنین در روش جامع اعتبارسنجی متقابل چند تا^۴ کل دادگان‌ها به k قسمت مساوی تقسیم می‌شوند. از k -قسمت به‌عنوان دادگان‌های آموزشی استفاده می‌شود و بر اساس آن مدل ساخته می‌شود و با یک قسمت باقی‌مانده عملیات ارزیابی انجام می‌شود. فرآیند مزبور به تعداد k مرتبه تکرار خواهد شد، به‌گونه‌ای که از هرکدام از k قسمت تنها یک‌بار برای ارزیابی استفاده‌شده و در هر مرتبه یک دقت برای مدل ساخته‌شده، محاسبه می‌شود. در این روش ارزیابی دقت نهایی رده‌بند برابر با میانگین k دقت محاسبه‌شده خواهد بود. معمول‌ترین مقداری که در متون علمی برای k در نظر گرفته می‌شود برابر با ۱۰ هست. بدیهی است هر چه مقدار k بزرگ‌تر شود، دقت محاسبه‌شده برای رده‌بند قابل‌اعتمادتر بوده و دانش حاصل‌شده جامع‌تر خواهد بود و البته افزایش زمان ارزیابی رده‌بند نیز مهم‌ترین مشکل آن هست. حداکثر مقدار k برابر با تعداد رکوردهای دادگان اولیه است که این روش ارزیابی بانام رها کردن یکی^۵ شناخته می‌شود.

در روش‌هایی که تاکنون به آن اشاره‌شده، فرض بر آن است که عملیات انتخاب نمونه‌های آموزشی بدون جایگذاری صورت می‌گیرد. به‌بیان دیگر یک رکورد تنها یک‌بار در یک فرآیند آموزشی مورد توجه واقع می‌شود.

¹ Holdout

² Random Sub-sampling

³ Cross Validation

⁴ k-Fold Cross Validation

⁵ Leaving One Out

چنانچه هر رکورد در صورت انتخاب شدن برای شرکت در عملیات یادگیری مدل بتواند مجدداً برای یادگیری مورداستفاده قرار گیرد روش مزبور بانام خود راهانداز^۱ شناخته می‌شود.

۳-۳- کاهش خطا در یادگیری ماشین در بیوانفورماتیک

در این بخش موضوعات مختلف و چالش‌های مربوط به روش‌های یادگیری ماشین در بیوانفورماتیک موردبحث قرار گرفته است. حضور داده‌های تکراری، مسئله مهمی در بیوانفورماتیک است. با ظهور اینترنت، داده‌ها به‌صورت آزاد در دسترس عموم است که تشخیص خطاها و اندازه‌گیری کیفیت داده‌ها را دشوار می‌کند. منبع داده‌ها، در زمان استفاده از آن‌ها، توسط محققان بیوانفورماتیک در نظر گرفته نشده است و اگر داده‌های موجود در منبع کثیف باشند، دقت رده‌بندی کاهش می‌یابد. داده‌های بیولوژیکی کثیف می‌تواند به دلیل عوامل زیادی از جمله موارد زیر باشد:

۱- خطاهای حین آزمایش

۲- تفسیر غلط داده‌ها توسط زیست‌شناسان

۳- تایپ اشتباه مربوط به خطای انسانی

۴- روش‌های غیراستاندارد در آزمایش‌ها استفاده شود.

الگوریتم‌های یادگیری، هنگام یادگیری در یک پایگاه داده‌های بیولوژیکی کثیف، باید با روندی پایدار داشته باشد. روش‌های ML (یادگیری ماشین) با تطابق‌پذیری خود باید بتوانند تصمیماتی بهینه بگیرند تا از بیش‌برازش داده‌ها جلوگیری کنند. ضمناً از آنجاکه پایگاه داده‌های بیولوژیکی روزانه آپدیت می‌شوند، الگوریتم‌های ML باید مدت‌زمان یادگیری و آموزش کوتاهی داشته باشند.

در حین تجزیه و تحلیل بیولوژیکی، همکاری با متخصصان این حوزه (در این مورد زیست‌شناسان) ضروری است که برای حفظ کیفیت داده‌ها، داده‌های موجود در پایگاه‌های داده باید دائماً مورد بازبینی قرار گیرند. انتخاب ویژگی داده‌ها مهم‌ترین عاملی است که بر عملکرد پیش‌بینی کننده مدل یادگیری ماشین تأثیر می‌گذارد. انتخاب ویژگی داده در مرحله آموزش یادگیری ماشین باهدف کاهش ابعاد داده‌ها، حذف نویز داده‌ها و داده‌های بی‌ربط و درنتیجه حفظ مفیدترین سیگنال‌ها از دادگان‌ها انجام می‌شود. فرآیند انتخاب ویژگی داده‌ها می‌تواند به‌طور گسترده با استفاده از روش‌های پالایش آن‌ها، ماژول‌های تعبیه‌شده یا روش‌های پوشانه^۲ انجام شود. انتخاب این روش‌ها به ویژگی‌های اصلی داده و معیارهای مدل پیش‌بینی بستگی دارد. در مورد بیماری‌های پیچیده چندژنی، در حال حاضر در میان داده‌های ژن‌نمودی، SNP‌ها دارای مفیدترین اطلاعات محسوب می‌شوند. فرض بر این است که SNP‌هایی که برای استفاده در مدل‌های پیش‌بینی کننده انتخاب می‌شوند با مکان‌هایی همراهی دارند که در اتیولوژی بیماری زمینه‌ای نقش مهمی ایفا می‌کنند.

¹ Bootstrap

² Wrapper

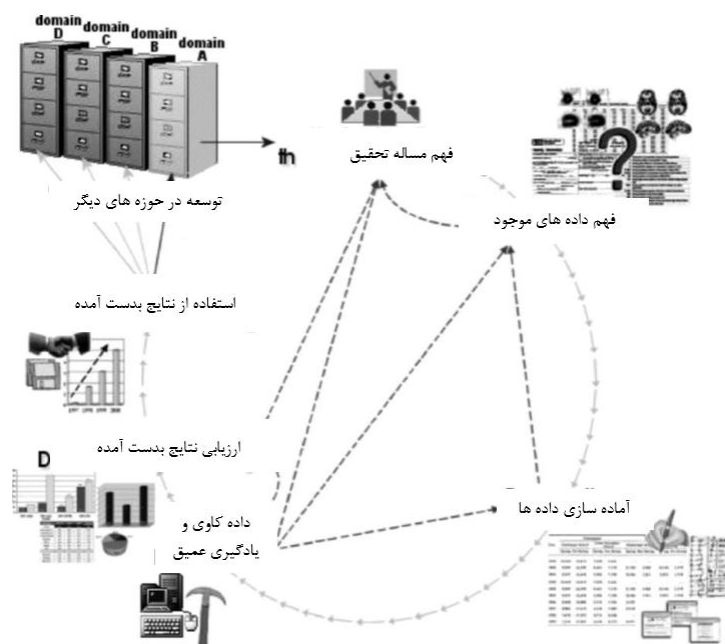
با این وجود، ممکن است چگونگی ارتباط SNP با بیماری قابل درک نباشد. معمولاً در مرحله اول ساخت مدل، گونه‌های موجود در داده‌های ژن‌نمودی بر اساس آستانه‌های p -value مطالعه GWA پالایش شده و دسته‌بندی می‌گردند. در الگوریتم ساخت و عملکرد مدل، برای انتخاب SNP ها و اثرات متقابل آن‌ها روش‌هایی دیده شده است که بدین ترتیب تنها به SNP های دارای اطلاعات مفید اجازه می‌دهد که در مدل پیش‌بینی کننده شرکت کنند. پوشش‌دهنده‌ها نیز همان هدفی را دارند که روش‌های تعبیه شده دنبال می‌کنند. با این حال، پوشش‌دهنده‌ها مازول‌های انتخاب SNP مستقل هستند که قبل از فرآیند ساخت مدل اجرا شده‌اند. بیش‌برازش^۱ پدیده‌ای است که در آن مدل‌ها به حدی یک دادگان را پوشش می‌دهند که نمی‌توان آن‌ها را به داده‌های دیگر تعمیم داد. احتمال به وجود آمدن مدل‌های با بیش‌برازش را می‌توان با تنظیم آن‌ها کاهش داد، این فرآیندی است که قدرت تعمیم پیش‌بینی مدل‌های یادگیری ماشین را به حداکثر می‌رساند.

۳-۴- محدوده‌های تحقیق و روش اجرا

اطلاعات ژنوم انسان شامل بخش‌های مختلفی است. پردازش این بخش‌ها بر اساس نوع داده متفاوت می‌باشند. در این تحقیق اطلاعات تغییرات یک SNP از ژنوم هر انسان بررسی می‌شود و متغیر مستقل اندازه‌گیری شده تغییرات SNP ژنوم انسان نسبت به ژن مرجع هست. همچنین برای بررسی بهتر نقش عوامل محیطی اطلاعات وضعیتی انسان مانند جنس، سن، قومیت و موارد دیگر به عنوان متغیرهای مستقل وارد مدل می‌شوند (Zhao et al., 2018).

متغیر وابسته در این تحقیق داده‌های طولی فشارخون (به صورت کمی) و خطر پرفشاری خون (به صورت برچسب ۱ و ۰) طی سال‌های اندازه‌گیری شده در تحقیق قند و لیپید تهران هست. منظور از داده‌های طولی، داده‌های ثبت شده برای هر فرد در ۶ دوره هست که ماشین یادگیرنده بر مبنای دوره‌های گذشته به پیش‌بینی دوره‌های بعدی می‌پردازد. با توجه به اینکه این تحقیق در حوزه داده‌کاوی در اطلاعات پزشکی هست، از روش گفته شده در (Cios and William Moore, 2002) استفاده شده که در تصویر ۳-۶ آمده است.

¹ Overfitting



شکل ۶-۳ نحوه اجرای تحقیق

در این روش ابتدا اطلاعات خام از منابع داده متناسب با متغیرهای مستقل و وابسته استخراج می شوند. برای این کار لازم است که اطلاعات و نوع مقادیر آن‌ها به خوبی شناخته شوند.

در مرحله بعدی آماده سازی داده ها انجام می گردد. در این مرحله داده هایی که حذف شده^۱ و یا ناشناخته^۲ می باشند باید معین شوند که روش های مختلفی مانند حذف کردن، جایگزینی با مقدار میانگین و یا پیش بینی مبتنی بر داده های هم گروه استفاده می گردد. ضمناً برای شروع مدل سازی بهتر است از روش های گسسته سازی^۳ که باعث کاهش احتمال بیش برآزش نیز هست استفاده کرد.

پس از آماده سازی داده ها، مدل سازی بر اساس ابزارهای مختلفی انجام می شود. با توجه به حجم داده و نوع الگوریتم استفاده شده از نرم افزار Orange جهت پیاده سازی مدل های اولیه استفاده می شود و برای پیاده سازی مدل نهایی از زبان برنامه نویسی php تحت Linux استفاده می شود. البته جهت نمایش اولیه داده ها و محاسباتی نظیر توزیع متغیرها، از نرم افزارهای Excel و SPSS استفاده می شود.

با توجه به مدل هایی که در حوزه ادبیات موضوع بررسی شده است، زیرساخت پیشنهادی جهت اجرای مدل نهایی، سروری شامل ۸ پردازشگر 3.2Ghz به همراه 64GB رم و 300GB هارد دیسک هست و زمان تخمینی جهت اجرای نهایی یک هفته هست.

با اجرای مدل و استفاده از روش های ارزیابی معرفی شده در ادبیات، خروجی های مدل نهایی حاصل می شود که این خروجی ها با توجه به اهداف تحقیق شامل موارد زیر است:

۱- مدل پیش بینی کننده پرفشاری خون بر اساس SNP های در نظر گرفته شده

۲- مدل پیش بینی کننده پرفشاری خون بر اساس SNP و داده های طولی

¹ Missing

² Unknown

³ Binding

۳- تعیین SNP های مهم و مؤثر بر پرفشاری خون و تعیین دقت پیش‌بینی مبتنی بر مجموعه پیشنهادی در خصوص مدل‌ها نتایج شامل جداول درهم‌ریختگی بر اساس برچسب واقعی نمونه هست.

۳-۵- فرایندهای مدل‌سازی

با توجه به اینکه در این تحقیق مدل‌های متنوعی ساخته و پیاده‌سازی شده است، فرایندهای مدل‌سازی در ادامه آمده است.

۱-۵-۳- فرایند مدل‌سازی بر مبنای SNP کاندید

مطالعات GWAS شامل تحقیقاتی است که حداقل ۱۰۰۰ نفر در آن مشارکت داشته باشند و با روش‌های علمی - آماری، SNP های مؤثر یک در رخ‌نمود شناسایی شده‌اند. به همین دلیل برای شناسایی بهتر SNP های مؤثر تمامی SNP های گزارش شده در تحقیقات فشارخون را که به‌عنوان SNP کاندیدا وارد مدل می‌کنیم. پروفایل رخ‌نمود نیز مبتنی بر سن، جنسیت، فشارخون که متناسب با هر فرد تعداد یک الی شش دوره هست، وارد مدل می‌گردد (داده جنسیت ثابت می‌ماند و داده سن و برچسب‌های فشارخون مبتنی بر دوره تعیین می‌شود). برای برچسب‌زنی داده‌ها، از محدوده‌های زیر برای فشارخون استفاده می‌کنیم (Basile et al., 2019).

Normal: sbp<120+10 and dbp<80+5

High: sbp<140 or dbp<90

G1 Hyper: sbp<160 or dbp<100

G2 Hyper: sbp<180 or dbp<110

G3 Hyper: sbp>180 or dbp>110

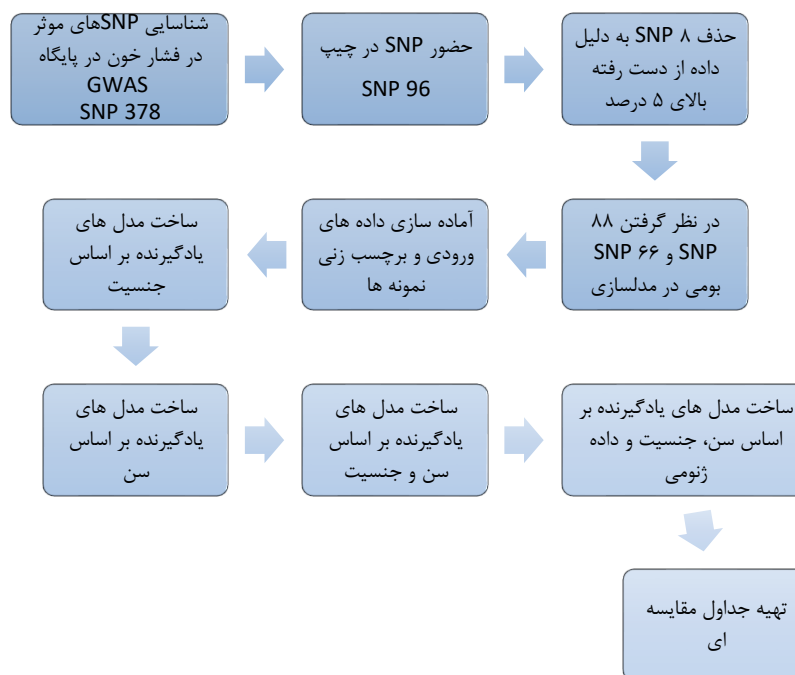
Isolated: sbp>140 and dbp<90

برای آنکه خطایی در برچسب‌ها ایجاد نشود داده‌هایی با برچسب Isolated از نمونه‌ها حذف شده است (Mahajan et al., 2019). ضمناً به جهت سهولت، در هر بار مدل‌سازی تنها یک برچسب وارد مدل می‌شود. تعداد کل نمونه‌ها برابر ۴۸۸۱۹ ردیف هست که نشان می‌دهد به‌طور متوسط افراد در ۳.۹۸ دوره شرکت کرده‌اند.

بر اساس فهرست GWAS نسخه GRCh38.p13 تمامی SNP های مرتبط با رخ‌نمود کلی فشارخون (صفت با کد EFO_1001132) برابر ۳۷۸، SNP هست که از این تعداد ۹۶ SNP در داده‌های ژنومی این تحقیق می‌باشند. ضمناً از این تعداد نیز ۸ SNP شرایط Miss<5% و MAF>5% را ندارند. ضمناً با توجه به تحقیق بومی که در انستیتو انجام شده است، علاوه بر ۸۸ SNP گفته شده، ۶۶ SNP را نیز از تحقیقات (Kolifarhood, M. Daneshpour, et al., 2019; Kolifarhood, M. S. Daneshpour, et al., 2019; Kolifarhood et al., 2021) استخراج شده است و درنهایت تعداد ردیف‌های نهایی ۴۸۸۱۹ و تعداد ستون‌های نهایی ۱۵۸ ستون ورودی (۱۵۶ SNP به همراه جنسیت و سن) و ۴ برچسب هست.

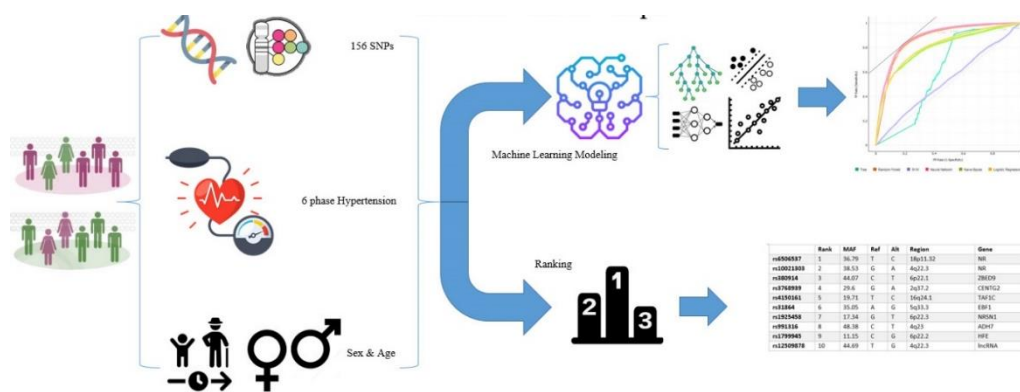
برای مقایسه روش‌های مختلف یادگیری ماشین در پیش‌بینی بیماری، از روش‌های شبکه عصبی، رگرسیون خطی، جنگل تصادفی، درخت تصمیم و ماشین بردار پشتیبان استفاده شده است و درنهایت ارزیابی نیز با روش

10-fold رده‌بندی شده انجام شده است. مراحل تعیین SNP‌هایی که در مدل‌سازی آمده‌اند و روند کار در شکل ۷-۳ آمده است.



شکل ۲-۳ مراحل اجرای مدل سازی یادگیری ماشین سنتی

مسیر اجرایی مدل‌سازی در شکل ۳-۸ آمده است.



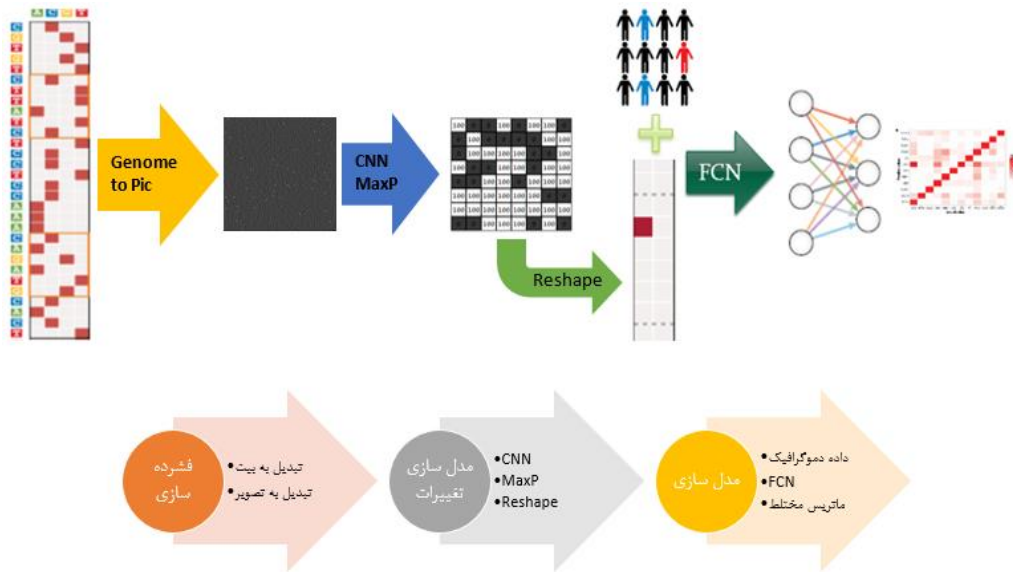
شکل ۸-۳ مسیر اجرای مدل یادگیری ماشین سنتی مبتنی بر SNP کاندید

پارامترهای اصلی برای مدل‌ها عبارت‌اند از: درخت تصمیم با حداقل اندازه زیرمجموعه ۵، حداقل نمونه در برگ دو، و آستانه توقف ۹۵٪ است. جنگل تصادفی دارای ۱۰۰ درخت با حداقل زیرمجموعه ۵ است. SVM با چندین تابع فعالیت استفاده‌شده و بهترین نتیجه تابع سیگموید^۱ با مقدار خطای ازدست‌رفته ۰.۱، هزینه ۱ و اغماض^۲ عددی ۰.۰۰۱ است. رگرسیون منطقی با تنظیم‌کننده مرز L2 و قدرت ۱ است. شبکه عصبی با سه لایه و تابع فعالیت ReLu توسط حل گر Adam، آلفا ۰.۰۰۰۱ و حداکثر تکرار ۲۰۰ است. لازم به ذکر است که مدل‌سازی با چهار و پنج لایه و سایر توابع فعال انجام‌شده و بهترین نتیجه با این پارامترها بوده است.

۲-۵-۳- فرایند مدل سازی CNN

¹ Sigmoid² tolerance

مسیر اجرایی شبکه عمیق CNN به صورت شکل ۳-۹ هست.



شکل ۳-۹ مسیر اجرایی شبکه عمیق CNN

۳-۶- جمع بندی

این فصل به بیان روش کامل تحقیق پرداخت و با ارائه نمودارهای تصویری، بیان چگونگی تحقیق، انجام شد. بر اساس نمودارهای گفته شده در فصل بعدی نتایج گزارش می شود.

۴- نتایج و دستاوردها

این فصل به بیان نتایج و دستاوردهای به دست آمده مبتنی بر روش‌شناسی و داده‌های گفته شده در فصل قبل می‌پردازد.

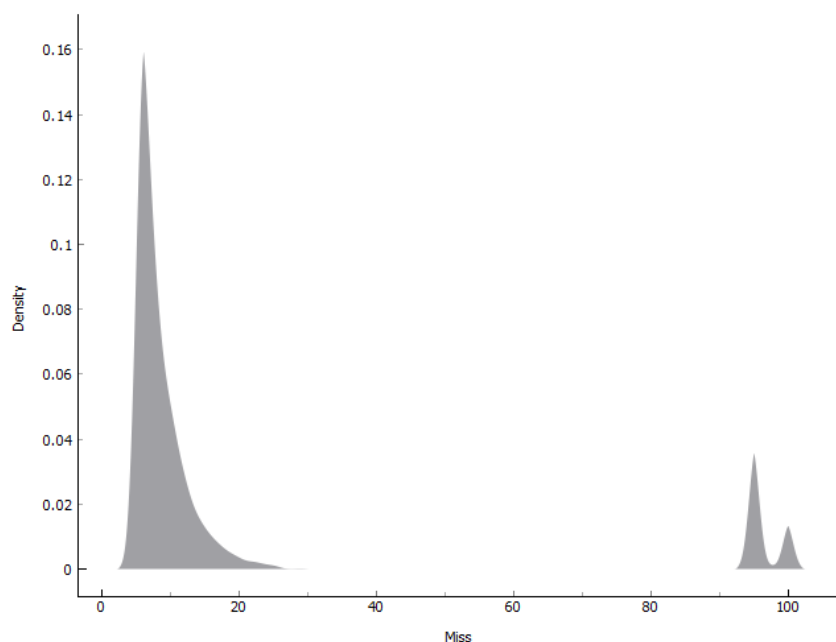
۴-۱- بخش‌بندی و کنترل کیفی داده‌های ژنومی

بر مبنای (Mattoo, 2019) جهت بهبود نتایج و کنترل کیفی، دسته‌بندی بر اساس سن (بالای ۱۸ سال و پایین‌تر از ۱۸ سال)، بخش‌بندی SNP از منظر نادر و رایج بودن انجام شد. برای کنترل کیفی داده‌های ژنومی، SNP‌هایی که دارای مقادیر گمشده^۱ کمتر از ۵ درصد هستند در نظر گرفته شد. در مرحله بعدی با توجه به اهمیت SNP های رایج، تنها SNP‌هایی با MAF بالاتر از ۵ درصد در نظر هست (Mattoo, 2019). جدول ۴-۱ تعداد SNP های نهایی را نمایش داده است.

جدول ۴-۱ کنترل کیفی داده‌های SNP

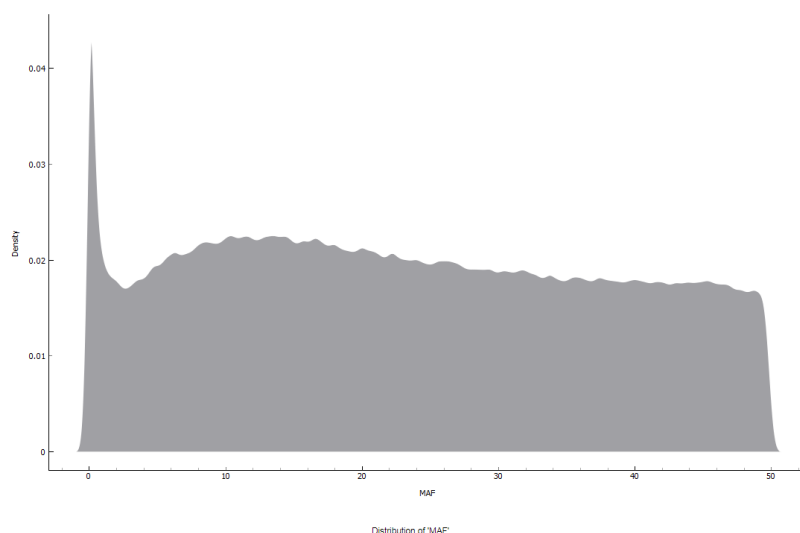
مجموع	652919	100.0%
Miss<5	647357	99.1%
MAF>5	571860	87.6%

تصویر ۴-۸ نمودار چگالی SNP های گمشده در شکل ۴-۱ و تصویر ۴-۲ نمودار چگالی MAF را بیان می‌کنند.



شکل ۴-۱ نمودار تغییرات تعداد SNP با مقادیر گمشده

¹ Missed Value



شکل ۲-۴ نمودار چگالی MAF

از روش‌هایی که باعث افزایش دقت می‌گردد، انتخاب SNP های مهم بر اساس عوامل مختلف هست. یکی از این عوامل انتخاب SNP های کاندیدا معرفی شده در GWAS های گذشته در خصوص فشارخون هست.

۲-۴- اجرای مدل بر مبنای SNP کاندیدا

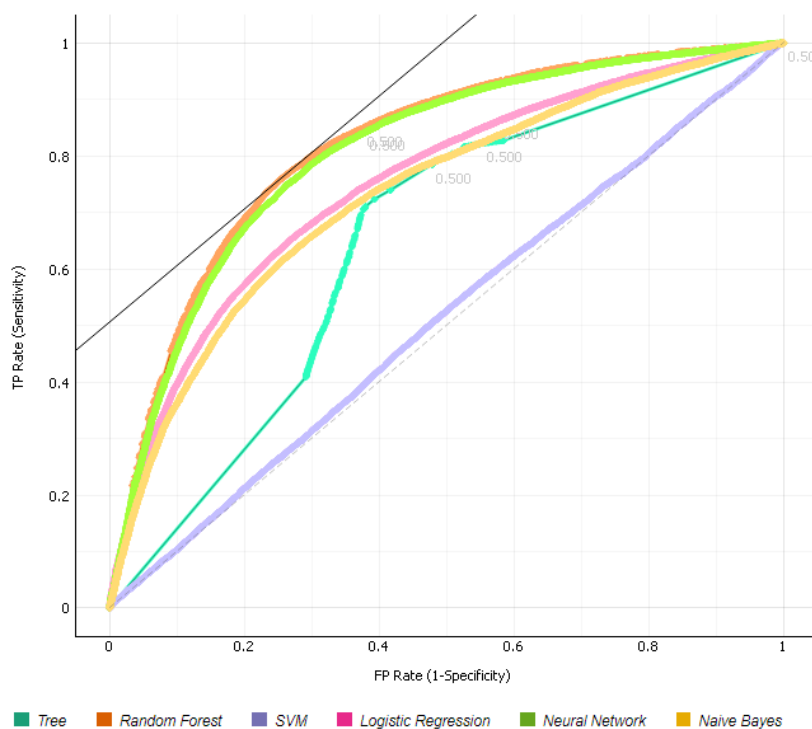
جدول ۲-۴ نتایج مدل‌سازی جنسیت، سن و جنسیت، سن و درنهایت سن، جنسیت و داده ژنومی را با روش‌های مختلف یادگیری ماشین، نشان می‌دهد. معیارهای اصلی بررسی مدل‌های پیش‌بینی کننده شامل مساحت زیر نمودار ROC، دقت، معیار F ، صحت و بازخوانی هست که مهم‌ترین معیار AUC است.

جدول ۲-۴ نتایج مدل‌سازی روش‌های یادگیری ماشین ستی

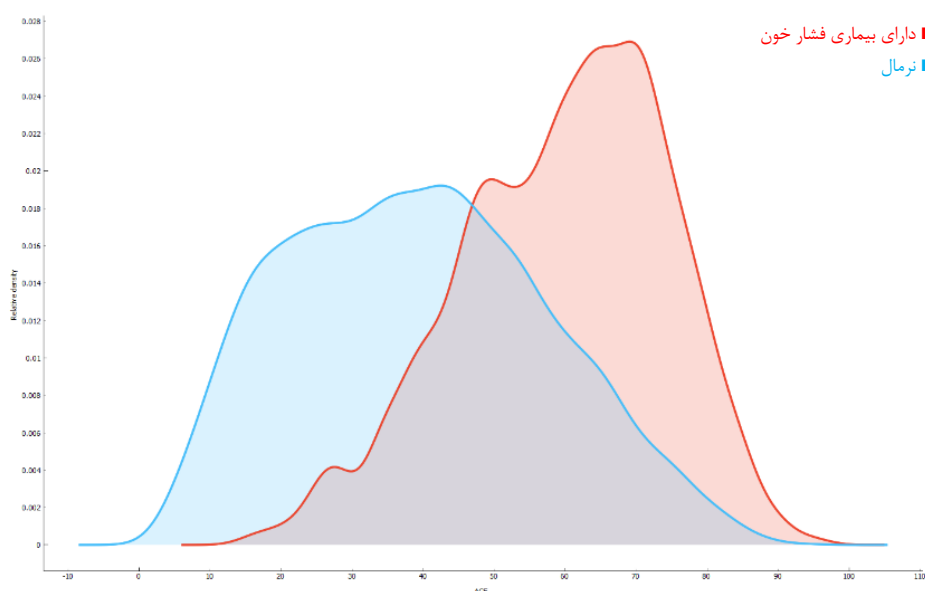
روش	مدل	AUC	صحت	مقدار F	دقت	بازخوانی
Random Forest	جنسیت	0.507	0.507	0.505	0.507	0.507
	سن	0.806	0.741	0.741	0.742	0.741
	سن و جنسیت	0.809	0.74	0.74	0.743	0.74
	جنسیت، سن و داده ژنومی	0.868	0.888	0.881	0.878	0.888
Neural Network	جنسیت	0.504	0.505	0.505	0.505	0.505
	سن	0.808	0.743	0.742	0.746	0.743
	سن و جنسیت	0.814	0.742	0.741	0.745	0.742
	جنسیت، سن و داده ژنومی	0.885	0.891	0.889	0.887	0.891
Logistic Regression	جنسیت	0.507	0.507	0.505	0.507	0.507
	سن	0.808	0.743	0.743	0.743	0.743
	سن و جنسیت	0.808	0.743	0.743	0.744	0.743
	جنسیت، سن و داده ژنومی	0.811	0.853	0.819	0.817	0.853
Tree	جنسیت	0.507	0.507	0.505	0.507	0.507
	سن	0.806	0.741	0.74	0.742	0.741
	سن و جنسیت	0.81	0.741	0.741	0.743	0.741
	جنسیت، سن و داده ژنومی	0.610	0.852	0.844	0.838	0.852

SVM	جنسیت	0.5	0.503	0.498	0.503	0.503
	سن	0.397	0.303	0.303	0.303	0.303
	سن و جنسیت	0.549	0.53	0.498	0.54	0.53
	جنسیت، سن و داده ژنومی	0.565	0.810	0.780	0.759	0.810

نمودار ROC حاصل از مدل سازی جنسیت، سن و داده ژنومی در شکل ۳-۴ و توزیع سنی بر اساس فشارخون در شکل ۴-۴ آمده است.



شکل ۳-۴ نمودار ROC مدل سازی جنسیت، سن و داده ژنومی



شکل ۴-۴ نمودار توزیع سنی بر اساس برچسب فشارخون

نتایج حاصل از بخش بندی سن در تصویر چندگانه ۴-۵ آمده است.

Scores						Scores					
Method	AUC	CA	F1	Precision	Recall	Method	AUC	CA	F1	Precision	Recall
Neural Network	0.848	0.868	0.861	0.858	0.868	Neural Network	0.850	0.867	0.862	0.858	0.867
Random Forest	0.844	0.869	0.860	0.856	0.869	Random Forest	0.842	0.870	0.861	0.857	0.870
Logistic Regression	0.755	0.842	0.791	0.789	0.842	Logistic Regression	0.755	0.842	0.792	0.790	0.842
Naive Bayes	0.672	0.756	0.760	0.765	0.756	Naive Bayes	0.672	0.756	0.760	0.765	0.756
Tree	0.606	0.833	0.825	0.819	0.833	Tree	0.606	0.834	0.826	0.820	0.834
SVM	0.505	0.784	0.762	0.745	0.784	SVM	0.505	0.784	0.762	0.745	0.784

Scores						Scores					
Method	AUC	CA	F1	Precision	Recall	Method	AUC	CA	F1	Precision	Recall
Neural Network	0.672	0.990	0.986	0.981	0.990	Neural Network	0.769	0.845	0.775	0.792	0.845
Logistic Regression	0.665	0.990	0.986	0.981	0.990	Logistic Regression	0.767	0.841	0.787	0.781	0.841
Naive Bayes	0.652	0.990	0.986	0.981	0.990	Random Forest	0.765	0.844	0.779	0.778	0.844
SVM	0.632	0.990	0.986	0.981	0.990	Tree	0.757	0.844	0.777	0.776	0.844
Random Forest	0.607	0.990	0.986	0.981	0.990	Naive Bayes	0.757	0.845	0.774	0.714	0.845
Tree	0.500	0.990	0.986	0.981	0.990	SVM	0.414	0.596	0.650	0.753	0.596

Scores						Scores					
Method	AUC	CA	F1	Precision	Recall	Method	AUC	CA	F1	Precision	Recall
Neural Network	0.863	0.886	0.880	0.877	0.886	Neural Network	0.863	0.886	0.880	0.877	0.886
Random Forest	0.857	0.884	0.876	0.873	0.884	Random Forest	0.857	0.884	0.876	0.873	0.884
Logistic Regression	0.785	0.860	0.815	0.811	0.860	Logistic Regression	0.785	0.860	0.815	0.811	0.860
Naive Bayes	0.691	0.771	0.780	0.789	0.771	Naive Bayes	0.691	0.771	0.780	0.789	0.771
Tree	0.594	0.855	0.846	0.840	0.855	Tree	0.594	0.855	0.846	0.840	0.855
SVM	0.512	0.761	0.764	0.768	0.761	SVM	0.512	0.761	0.764	0.768	0.761

شکل ۴-۵ تصویر جداول مقایسه‌ای روش‌های با نظارت مبتنی بر بخش‌بندی سن

با توجه به نتایج به‌دست‌آمده، بخش‌بندی بر روی سن باعث کاهش دقت نهایی مدل می‌گردد.

نتایج این بخش در (S Ali Lajevardi et al., 2022) منتشر گردید.^۱

۴-۳- محاسبه امتیاز خطر ژنتیکی مبتنی بر GRS

برای مقایسه نتایج به‌دست‌آمده از روش‌های یادگیری ماشین با روش‌های آماری مرسوم، محاسبه امتیاز خطر ژنتیکی با ابزارهای GCTA و PLINK انجام شده است.

برای این کار ابتدا، داده‌ها به ۱۰ قسمت (بخش) تصادفی تقسیم شدند. در مرحله بعدی با ۹ قسمت محاسبه تأثیر فردی^۲ و پس از آن محاسبه تأثیر SNP توسط بسته GCTA انجام می‌شود. سپس با اجرای دستورات لازم در PLINK محاسبه امتیاز انجام می‌شود. در نهایت قسمت کنار گذاشته توسط یک دسته‌بند در هر بخش محاسبه و نتایج آن ادغام می‌گردد. برای این کار از نرم‌افزار Orange استفاده شده است. لیست دستورات در شکل ۴-۶ آمده است. پرونجا ژنوم ورودی به صورت genome.ped هست.

^۱ <https://github.com/18026070/Hypertension-CNN>

^۲ Individual

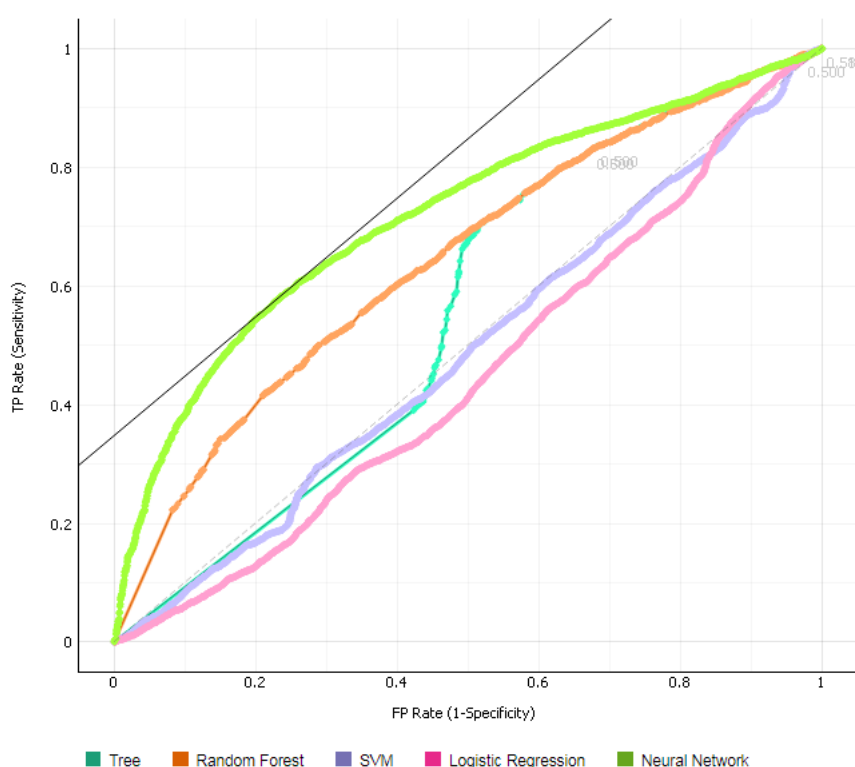
```
plink --noweb --ped genome.ped --map genome.map --make-bed --out genome
gcta64 --bfile genome --make-grm --reml --thread-num 10 --out genome
gcta64 --reml --grm genome --pheno genome.phen --cvblup --out genome
plink --noweb --score genome.tr_n.snp.blp --bfile genome --pheno
genome.ts_n.phen --covar genome.cov --out genome.t_n
```

شکل ۴-۶ نمونه کد محاسبه امتیاز خطر ژنتیکی

نتایج به دست آمده در جدول ۳-۴ آمده است و نمودار ROC نیز در شکل ۴-۷ آمده است. نتایج اولیه در مقاله (لاجوردی 1399، *et al.*) منتشر گردید.

جدول ۳-۴ نتایج دسته‌بندی‌های مختلف برای امتیاز خطر ژنتیکی

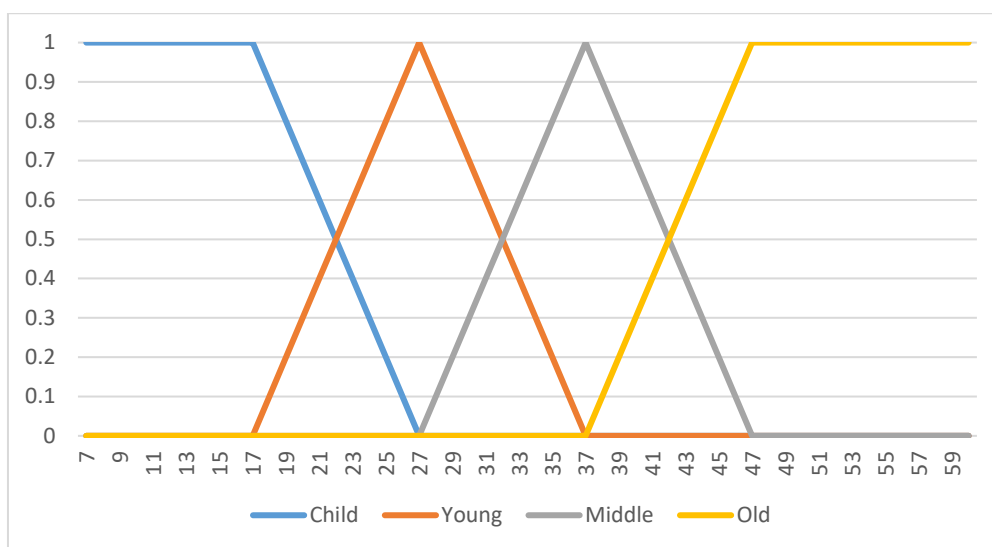
Method	AUC	CA	F1	Precision	Recall
Neural Network	0.714	0.742	0.648	0.652	0.742
Random Forest	0.641	0.702	0.692	0.685	0.702
Tree	0.550	0.704	0.693	0.685	0.704
SVM	0.464	0.675	0.647	0.630	0.675
Logistic Regression	0.425	0.747	0.639	0.558	0.747



شکل ۴-۷ نمودار ROC مدل‌سازی امتیاز خطر ژنومی

۴-۴- مدل‌سازی فازی

با توجه به متغیرهای در نظر گرفته شده در مدل‌سازی، تنها متغیری که امکان فازی سازی بر روی آن وجود داشت، سن است. به همین دلیل بر اساس ادبیات موضوع فازی سازی سن بر اساس شکل ۴-۸ انجام گردید (Satyanarayana *et al.*, 2020; Singh, Verma and Singla, 2020). در نهایت با در نظر گرفتن سن فازی و SNP های در نظر گرفته شده در مدل‌سازی یادگیری ماشین، مدل‌سازی دیگری نیز برای خطر پرفشاری خون انجام گردید.



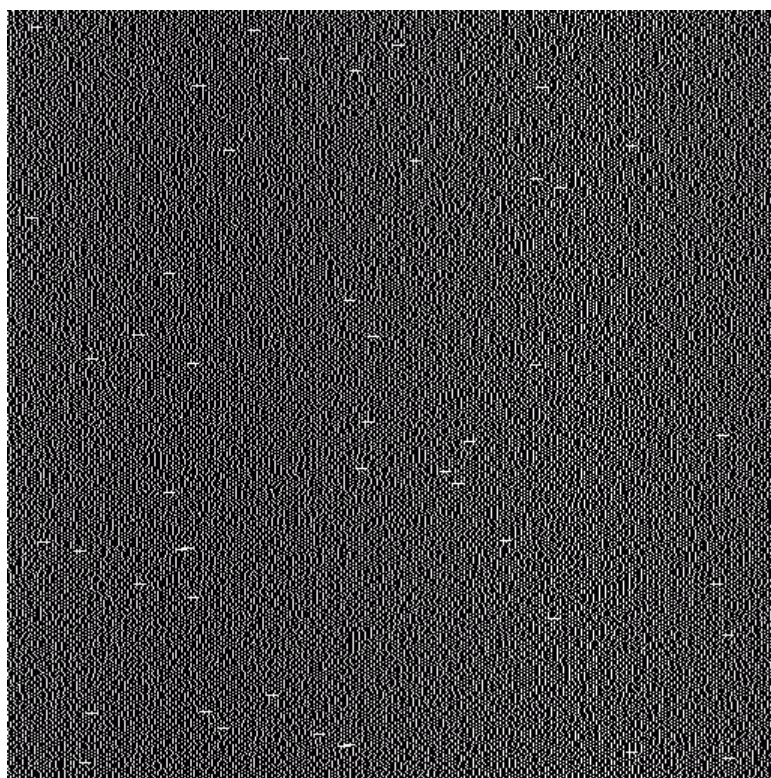
شکل ۸-۴ فازی سازی سن

۴-۵- اجرای شبکه عمیق CNN

همان‌طور که در فصل قبل ذکر شد برای شروع شبکه CNN ابتدا داده‌ها به صورت تصویر درآمدن است.

۴-۵-۱- تبدیل داده ژنومی به تصویر

برای اجرای الگوریتم CNN بر روی داده‌های ژنومی، اولین کار تبدیل داده ژنومی شامل ۶۵۲۹۱۹ SNP، به یک تصویر هست. برای این کار داده‌های ژنوم که شامل رشته‌ای از T, G, C, A هست تبدیل به بیت شده و به صورت تصویر سیاه و سفید در یک پرونجا تصویری با فرمت GIF ذخیره سازی انجام می‌شود. نتیجه حاصل شده به صورت شکل ۴-۹ هست.



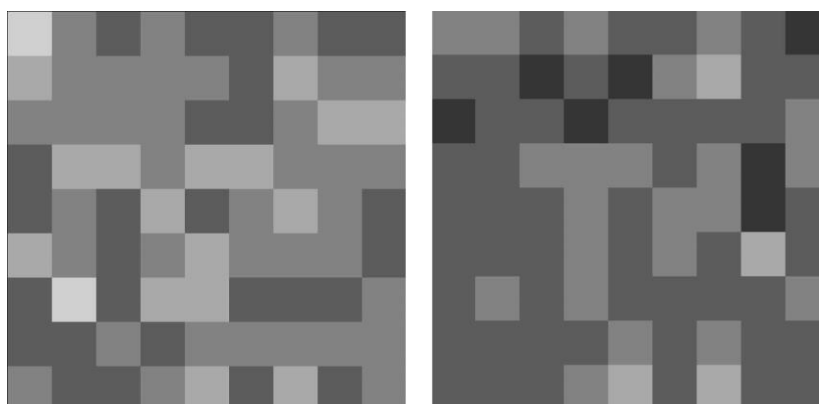
شکل ۹-۴ تصویر ساخته شده از داده ژنومی

با این کار علاوه بر اینکه تصویر آماده طراحی شبکه CNN هست، داده ژنومی به شکل مطلوبی فشرده‌سازی شده است که نوع فشرده‌سازی به حالت بدون از دست رفتگی^۱ و بدون نیاز به الگوریتم بازیابی جهت خارج کردن از حالت فشرده هست.

۴-۵-۲- اجرای شبکه CNN

طراحی شبکه عمیق مبتنی بر CNN بر اساس ۲۴ پالایه به صورت $4*4$ که در هر پالایه در هر ردیف تنها یک ستون دارای مقدار یک و مابقی صفر هست و پس از اعمال پالایه، تمامی نتایج در همدیگر ادغام می‌شود (اجرای پیچشی CNN). تعداد لایه‌های CNN وابسته به پیچیدگی تغییرات هست که با توجه به داده‌های آزمون، یک لایه مناسب است.

در مرحله بعد MaxPool استفاده شده است که به صورت $2*2$ هست که در هر ۴ خانه بیشینه را به لایه بعدی منتقل می‌کند و باعث کاهش بعد $1/2$ در طول و در ارتفاع تصویر می‌شود که در نتیجه در هر مرحله تصویر $1/4$ می‌شود که با انجام تعداد لایه معین می‌توان به مقدار دلخواه برای ورودی به لایه بعدی FCN رسید. با توجه به داده‌های آزمون، ۶ لایه استفاده شده است. تابع فعالیت FCN به صورت Relu هست. در نهایت برای یکدست‌سازی اطلاعات تصاویر به دست آمده مبتنی بر max-min نرمال‌سازی شده و برای هر ورودی، مقیاس ۱۰۰ برای نمایش خروجی در نظر گرفته شده است. تصویر $4-10$ دو نمونه از این پردازش هست.



شکل ۱۰-۴ خروجی Normalize، MaxP، CNN

ساخت هر تصویر نهایی در حدود ۳ دقیقه زمان می‌برد که این مرحله بالغ بر ۲۳ روز به طول انجامید.

۳-۵-۴- آماده‌سازی داده‌های طولی رخ‌نمود

با توجه به اینکه تمامی داده‌های رخ‌نمودی در پژوهشکده بر روی یک پروتکتور SPSS هست، در مرحله بعدی داده‌های SPSS رخ‌نمودهای ثبت شده بر مبنای رخ‌نمودهای مورد نیاز و پردازش اولیه اطلاعات بر اساس توزیع رخ‌نمودها و داده‌های ثبت شده در ۶ دوره TLGS، دسته‌بندی گردید. سپس داده‌های رخ‌نمود به ژن‌نمود مبتنی بر کدهای صادره از TLGS متصل شد که منجر به استخراج صحیح ۱۱۰۸۸ ردیف اطلاعاتی شامل داده‌های ژن‌نمود و رخ‌نمود گردید.

¹ Lossless

برای اجرای FCN بر مبنای (Basile et al., 2019) برچسب‌های زیر مدنظر قرار گرفت:

۱- عادی: فشار دیاستولی کمتر از ۸۰ و فشار سیستولی کمتر از ۱۲۰

۲- فشار دیاستولی بالاتر از ۸۰

۳- فشار سیستولی بالاتر از ۱۲۰

۴- فشار دیاستولی بالاتر از ۸۰ یا فشار سیستولی بالاتر از ۱۲۰

لازم به ذکر است با توجه به اینکه داده سن و جنسیت تأثیر مهمی در فشارخون دارند (Basile et al., 2019)، این دو داده نیز علاوه بر داده ژن‌نمود به‌عنوان داده مستقل وارد مدل می‌شوند که بر این اساس می‌توان داده‌های دوره‌های مختلف را نیز برای یک فرد پردازش نمود.

۴-۵-۴- اجرای شبکه تمام متصل^۱

برای اجرای شبکه تمام‌اتصال انتهایی شبکه عمیق، از روش‌های شبکه عصبی ساده، جنگل تصادفی، درخت تصمیم، رگرسیون لجستیک و SVM استفاده گردید. نتایج حاصل در شکل‌های ۴-۱۱ الی ۴-۱۳ آمده است.

Method	AUC	CA	F1	Precision	Recall
Neural Network	0.809	0.939	0.932	0.929	0.939
Random Forest	0.785	0.936	0.929	0.925	0.936
Logistic Regression	0.729	0.932	0.901	0.878	0.932
Tree	0.661	0.928	0.918	0.913	0.928
SVM	0.633	0.920	0.909	0.900	0.920

شکل ۴-۱۱ جدول مقایسه‌ای روش‌های با نظارت در فشار دیاستولی بالا

Method	AUC	CA	F1	Precision	Recall
Neural Network	0.918	0.925	0.922	0.919	0.925
Random Forest	0.885	0.924	0.918	0.915	0.924
Logistic Regression	0.866	0.911	0.893	0.887	0.911
SVM	0.657	0.914	0.888	0.887	0.914
Tree	0.650	0.915	0.905	0.900	0.915

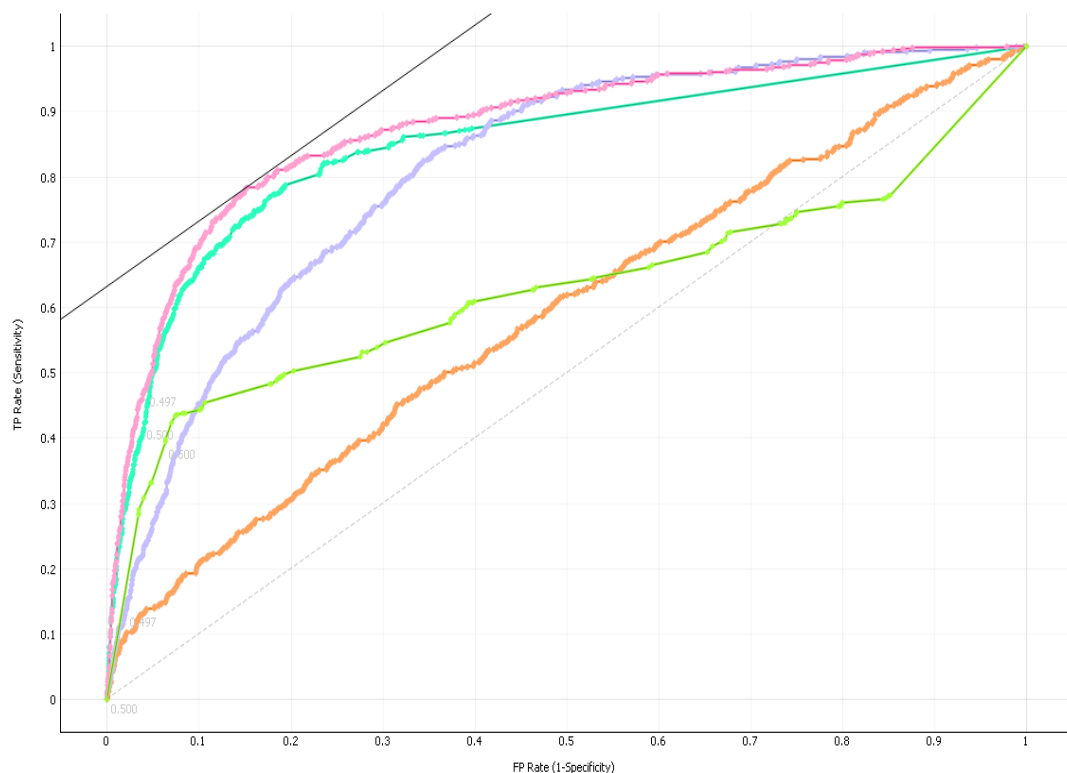
شکل ۴-۱۲ جدول مقایسه‌ای روش‌های با نظارت در فشار سیستولی بالا

Method	AUC	CA	F1	Precision	Recall
Neural Network	0.877	0.901	0.896	0.893	0.901
Random Forest	0.849	0.896	0.889	0.885	0.896
Logistic Regression	0.811	0.881	0.851	0.845	0.881
Tree	0.621	0.879	0.870	0.864	0.879
SVM	0.608	0.857	0.838	0.825	0.857

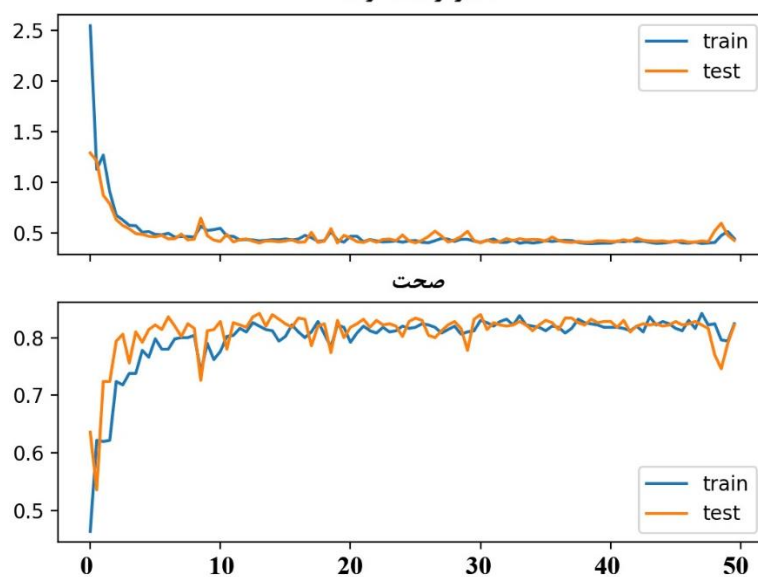
شکل ۴-۱۳ جدول مقایسه‌ای روش‌های با نظارت در فشار سیستولی یا دیاستولی بالا

نمودار ROC مدل‌سازی CNN با توابع مختلف در شکل ۴-۱۴ آمده است.

¹ Fully Connected Network (FCN)

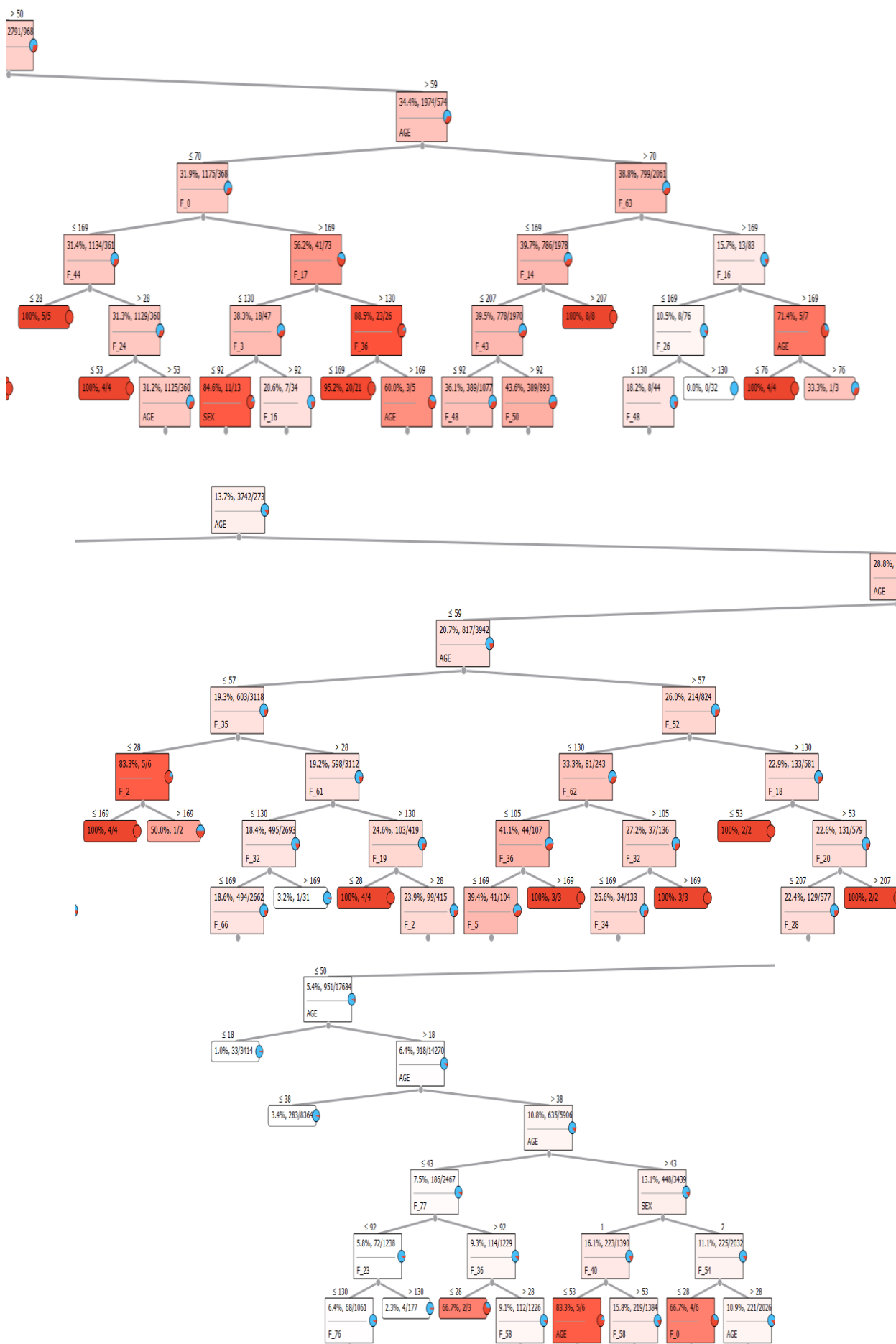


شکل ۴-۱۴ نمودار ROC مدل سازی CNN
مقدار از دست رفته



شکل ۴-۱۵ نمودار صحت و مقدار از دست رفته

با استفاده از الگوریتم درخت تصمیم بر روی داده‌ها، درختی حاصل می‌شود که در تصویر ۴-۱۵ آمده است. در این درخت تفکیک سنی به دست آمده کاملاً متناسب با تحقیقات زیستی در حوزه فشارخون هست و نشان می‌دهد که مهم‌ترین عامل در بیماری فشارخون سن هست و در سطوح اهمیت بعدی SNP های متفاوتی برای جنسیت زن و مرد نقش دارند. در این تصویر هر چه قدر کادری قرمز رنگ‌تر باشد احتمال وجود فشارخون در آن دسته افراد بالاتر است که به ترتیب برای بالای ۶۰ سال، بالای ۵۰ سال و بالای ۱۸ سال هست.



شکل ۱۶-۴ درخت تصمیم ناشی از مدل سازی فشارخون

۶-۴- جمع‌بندی نتایج

با توجه به اینکه روش‌های مختلفی در پیش‌بینی فشارخون بر مبنای داده‌های ژنومی در این تحقیق انجام شده است، خلاصه‌ی نتایج در جدول ۴-۴ آمده است.

جدول ۴-۴ خلاصه نتایج مدل‌سازی‌های انجام‌شده در پیش‌بینی بیماری پرفشاری خون

نوع	داده	تعریف	CNN	GRS	Fuzzy	Top10	ML (NN Best)				
متغیر	جنسیت	مرد/زن	*	*	*	*	*	*	*	*	*
	سن	سال دوره‌ای	*			*	*	*	*	*	*
		سن فازی		*							
		سن تعیین‌شده		*					*		
	ژنوم	تصویر 600K	*								
		600K GRM		*							
		GWAS SNP			*	*	*	*	*	*	*
		Farhood SNP									
	برچسب	فشارخون	*	*	*	*	*	*	*	*	*
	نتایج	AUC	0.877	0.71	0.866	0.83	0.542	0.817	0.885	0.81	0.804
		صحت	0.901	0.74	0.88	0.861	0.884	0.771	0.891	0.884	0.883
		مقدار F	0.896	0.65	0.878	0.832	0.829	0.768	0.889	0.834	0.838
		دقت	0.893	0.65	0.876	0.834	0.781	0.767	0.887	0.841	0.839
		بازخوانی	0.901	0.74	0.88	0.861	0.884	0.771	0.891	0.884	0.883

۷-۴- شناسایی SNP های مهم

برای شناسایی SNP های مهم روش‌های مختلف رتبه‌سنجی وجود دارد. این روش‌ها امکان شناسایی را با توجه به توزیع متغیر مربوطه در داده‌ها بر مبنای برچسب و دیگر متغیرهای ورودی انجام می‌دهد.

به جهت اهمیت استفاده از روش‌های مختلف، روش‌های بهره اطلاعات^۱، نسبت بهره^۲، شاخص جینی^۳، کای دو^۴، رلیف^۵ و FCBF استفاده گردید. درنهایت با توجه به رتبه‌ای که هر SNP در هر روش به دست آورد، اهمیت میانگین SNP محاسبه گردید و بر مبنای آن لیست مرتب مهم‌ترین SNP ها در جدول ۴-۵ آمده است. در این جدول ستون‌های TCGS بر اساس اطلاعات بومی در داده‌ها اندازه‌گیری شده است و اطلاعات GWAS نیز در ستون‌های مربوطه قرار گرفته است.

برای مقایسه با دیگر مدل‌ها، خطر پرفشاری خون بر اساس ۱۰ SNP برتر نیز محاسبه گردید. نتایج کامل در مقاله (S. Ali Lajevardi *et al.*, 2022) منتشر گردید^۶.

جدول ۵-۴ لیست مرتب شناسایی ۱۰ SNP مهم

RS	TCGS				GWAS				
		MAF	Ref	Alt	TRAIT	POS	Gene	RISK ALLELE	P-VALUE
rs6506537	1	36.79	T	C	Hypertension	18p11.32	NR	rs6506537-T	4.00E-09
rs10021303	2	38.53	G	A	Hypertension	4q22.3	NR	rs10021303-A	7.00E-06
rs380914	3	44.07	C	T					
rs3768939	4	29.6	G	A	Cardiovascular disease in hypertension (calcium channel blocker interaction)	2q37.2	CENTG2	rs3768939-A	9.00E-07
rs4150161	5	19.71	T	C	Systolic blood pressure response to hydrochlorothiazide in hypertension	16q24.1	TAF1C	rs4150161-T	1.00E-06
rs31864	6	35.05	A	G	Hypertension	5q33.3	EBF1	rs31864-A	3.00E-06
rs1925458	7	17.34	G	T	Alzheimer's disease in hypertension	6p22.3	NRSN1	rs1925458-?	1.00E-06
rs991316	8	48.38	C	T	Hypertension	4q23	ADH7	rs991316-T	5.00E-06
rs1799945	9	11.15	C	G	Hypertension	6p22.2	HFE	rs1799945-G	2.00E-10
rs12509878	10	44.69	T	G	Diastolic blood pressure night-to-day ratio in hypertension	4q22.3	lncRNA	rs12509878-C	6.00E-06

¹ Information gain

² gain ratio

³ gini index

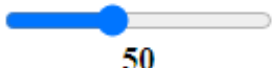
⁴ χ^2

⁵ relief

⁶ <https://github.com/18026070/Hypertension-ML>

۴-۸- ابزار سنجش خطر بیماری فشارخون بالا

بر مبنای SNP های مهم شناسایی شده در روش رتبه‌بندی و تأثیر عامل سن در محاسبه خطر، ابزاری به زبان PHP در آدرس <https://work.feham.ir/CalH> منتشر شده است. تصویر ابزار به صورت شکل ۴-۱۶ هست.

RS	Status	Age
rs6506537	☒ Ref ☐ Het ☐ Hom	 50 Calculated HTN Risk 61.3%
rs1021303	☐ Ref ☒ Het ☐ Hom	
rs380914	☐ Ref ☒ Het ☐ Hom	
rs3768939	☐ Ref ☒ Het ☐ Hom	
rs4150161	☐ Ref ☐ Het ☒ Hom	
rs31864	☐ Ref ☒ Het ☐ Hom	
rs1925458	☒ Ref ☐ Het ☐ Hom	
rs991316	☐ Ref ☐ Het ☒ Hom	
rs1799945	☐ Ref ☒ Het ☐ Hom	
rs12509878	☒ Ref ☐ Het ☐ Hom	

شکل ۴-۱۷ تصویر ابزار ساخته شده جهت پیش‌بینی خطر HTN

۴-۹- تحلیل ساختار علمی موضوع امتیاز خطر ژنتیکی

محققین در یک حوزه علمی به منظور دیدن فراسوهای دانش در حوزه تخصصی خود، معمولاً آثار پژوهشگران پیشین خود را مرور می‌کنند؛ به بیان دیگر، با اتکا به گذشته علم، آینده علمی حوزه تخصصی خود را ترسیم می‌کنند. داشتن درک و نمایی کلی از چارچوب علمی حوزه موردنظر، یکی از راه‌هایی است که پژوهشگران را برای رسیدن به اهداف پژوهشی در یک حوزه تخصصی کمک می‌کند. در این راستا ترسیم نقشه یا ساختار علمی آن حوزه، امری ضروری به نظر می‌رسد. در یک نقشه علمی که بر اساس برون‌دادهای علمی-پژوهشی محققین یک حوزه علمی ترسیم می‌شود، نویسندگان تأثیرگذار، خوشه‌ها و زیر حوزه‌های موضوعی شکل گرفته در طول زمان، روند حرکتی آن حوزه و موضوعات پژوهشی آینده حول آن حوزه تخصصی تعیین و معرفی می‌شوند. در نقشه‌های علمی ظهور حیطه‌های جدید و توقف برخی حوزه‌های علمی اشباع شده به وضوح نمایان است. به بیانی ساده نقشه علمی مصورسازی نتایج برآمده از تجزیه و تحلیل انتشارات یک حوزه علمی از زوایای گوناگون و ارائه چارچوبی کلی از آن حوزه است.

در تعریف نقشه‌های علمی اعتقاد بر این است که نقشه علمی، نمایی از حوزه‌های علمی است که با تجزیه و تحلیل کمی اطلاعات کتابشناختی تهیه می‌شود. عناصر تشکیل دهنده نقشه‌های علمی، برون‌دادهای حوزه‌های پژوهشی هستند. در این نقشه‌ها، حوزه‌های علمی که دارای ارتباط مفهومی قوی‌تری هستند، بر کنار هم و حوزه‌هایی که ارتباط ضعیف‌تری دارند در فاصله دورتری از هم قرار می‌گیرند (Noyons, 1999).

با توجه به رشد روزافزون حجم اطلاعات و سرعت فزاینده تولید علم، رصد روندهای علمی و شناسایی دانش پنهان شده در متون کاری دشوار است. لذا امروزه برای کشف لایه‌های پنهان دانش در علوم مختلف از روش‌های مختلفی از جمله متن کاوی استفاده می‌شود. محققین توانسته‌اند با بهره‌گیری از این روی‌ها به کشف دانش پنهان موجود در متون علمی دست پیدا کنند. درحالی‌که دستیابی به چنین دانشی از طریق روش‌های سنتی امکان‌پذیر نیست. یکی از این فنون که طی سال‌های اخیر خیلی مورد استقبال و اقبال متخصصان رشته‌های مختلف قرار گرفته تحلیل هم رخدادی واژگان^۱ است که به‌عنوان شاخصی نوین به جامعه علمی عرضه شد (Rip and Courtial, 1984). هم رخدادی واژگان یک فن تحلیل محتوا هست که هم فراوانی موضوعات و هم ارتباط بین آن‌ها را بیان می‌کند. تحلیل هم رخدادی واژگان، ابزاری برای کشف الگوهای پنهان و رویدادهای نوظهور مفهومی است که از طریق آن می‌توان مفاهیم اصلی یک زمینه یا حوزه علمی را شناخت و به‌واسطه این شناخت، الگوها و رویدادهای مفهومی، ساختار علمی، شبکه مفهومی، روابط سلسله مراتبی مفاهیم، و مقولات مفهومی آن حوزه را کشف، ترسیم و مدیریت کرد. روش تحلیل هم رخدادی واژگان با تکیه بر واژه‌های پربسامد می‌تواند مهم‌ترین موضوعات پژوهشی در هر حوزه علمی را مشخص کند؛ یعنی میزان رخداد یک واژه، نشان‌دهنده میزان اهمیت آن در حوزه علمی موردنظر است. کالون^۲ در سال ۱۹۸۳ اولین کسی بود که این روش را توصیف کرد و از آن در پژوهش‌های خود استفاده کرد (He, 1999). در این تحلیل، هم رخدادی کلیدواژه‌ها در عنوان، چکیده یا متن مقالات بررسی می‌شود. هم رخدادی کلیدواژه‌ها میزان ارتباط شناختی میان یک مجموعه مدارک را نشان می‌دهد. با مقایسه نقشه‌های حاصل در دوره‌های زمانی مختلف، پویایی علم ردیابی می‌شود. با به‌کارگیری روش تجزیه و تحلیل هم رخدادی واژگان می‌توان موضوعات علمی را استخراج و ارتباط میان آن‌ها را به‌صورت مستقیم از محتوای موضوعی کشف کرد. تحقیقات نشان داده که مصورسازی دانش از طریق تحلیل هم رخدادی واژگان می‌تواند تصویر روشن‌تری از آن حوزه نمایان سازد چراکه تحلیل واژگان موضوعات اساسی آن حوزه را در یک شمای کلی در اختیار پژوهشگر قرار می‌دهد. از آنجاکه نقشه‌های علمی دارای ساختاری مشابه ساختار شبکه‌های اجتماعی^۳ هستند، برای مصورسازی و تحلیل و تفسیر آن‌ها از فنون تحلیل شبکه‌های اجتماعی استفاده می‌شود (Guns, Liu and Mahbuba, 2011).

در این تحقیق از روش تحلیل هم رخدادی واژگان که یکی از روش‌های علم‌سنجی است، استفاده شده تا از حوزه امتیاز خطر ژنتیکی اطلاعات کمی قابل تحلیل به دست آید. جامعه پژوهش ۱۱۷۵ مقاله در حوزه امتیاز خطر ژنتیکی است که در پایگاه داده ScienceDirect در بازه سال‌های ۲۰۱۶ تا ۲۰۲۰ نمایه شده است. جهت انجام این تحقیق، عبارت امتیاز خطر ژنتیکی به‌عنوان هسته اصلی در نظر گرفته شده است. تمام کلمات کلیدی مربوط به ۱۱۷۵ مقاله از پایگاه داده ScienceDirect که شامل این عنوان در مجموعه کلیدواژه‌های خود بودند،

^۱Co-word analysis

^۲Callon

^۳Social networks

جمع‌آوری شد. سپس هر یک از کلمات کلیدی به‌عنوان یک گره در شبکه در نظر گرفته شده است. یال‌های شبکه نیز به‌صورت ارتباط بین کلمات تعریف شده‌اند. به‌طوری‌که بین دو گره یال وجود خواهد داشت در صورتی‌که آن دو گره (کلمه) به‌صورت مشترک در لیست کلمات کلیدی یک مقاله آمده باشند. پس از جمع‌آوری تمام کلمات، ارتباط دوبه‌دو بین تمام کلمات بررسی شد و در صورتی‌که دو کلمه دارای مقاله مشترک بودند، بین آن‌ها یک یال تعریف شده است. با توجه به این‌که محققین برای انتخاب کلمات کلیدی، استاندارد معینی را در نظر نمی‌گیرند و انتخاب این کلمات به‌صورت سلیقه‌ای انجام می‌شود. در دادگان جمع‌آوری شده، بسیاری از کلمات کلیدی علی‌رغم یکسان بودن از نظر مفهوم و معنا، در مقالات متفاوت در نظر گرفته شده است که همه یک مفهوم را متبادر می‌سازد. این موضوع باعث می‌شود در شبکه حاصل از ارتباط کلمات کلیدی مقالات، به ازای هر کلمه کلیدی یک گره در نظر گرفته شود و در نتیجه شبکه دارای گره‌های تکراری از لحاظ مفهوم باشد. این امر علاوه بر ایجاد گره‌های ایزوله در شبکه و تشکیل اجتماعات بسیار کوچک با حداکثر ۳ یا ۴ گره، باعث می‌شود نتایج تحلیل منعکس‌کننده واقعیت مسئله نباشد. برای حل این مشکل، تمام کلمات از نظر ظاهر و مفهومی بررسی شده و کلمات هم‌معنا به‌عنوان یک گره واحد در نظر گرفته شده‌اند. این کار به‌صورت ماشینی و با استفاده از الگوریتم فاصله لون‌اشتاین انجام شده است.

۴-۹-۲- نتایج و ارزیابی آن

در جدول ۴-۶ ده مجله که بیشترین ارتباط و همسویی با موضوع مورد بررسی را داشته‌اند مشاهده می‌شود. که می‌تواند جهت تصمیم‌گیری برای ارسال مقالات مرتبط با موضوع به مجلات جهت چاپ مورد استفاده قرار گیرد.

جدول ۴-۶ نام مجلات با بیشترین ارتباط با موضوع

ردیف	نام مجله	تعداد مقالات چاپ شده	IF ^۱	IF ساله ۵
۱	Gastroenterology	۸۵	۱۹.۲۳۳	۱۸.۸۴۴
۲	European Neuropsychopharmacology	۵۱	۴.۴۶۸	۴.۳۹۴
۳	Journal of the American College of Cardiology	۴۳	۱۸.۶۳۹	۱۹.۰۶۸
۴	Atherosclerosis	۴۱	۳.۹۱۹	۴.۶۳۹
۵	Ophthalmology	۲۱	۷.۷۳۲	۷.۸۴
۶	Alzheimer's & Dementia	۲۱	۱۴.۴۲۳	۱۳.۵۳۳
۷	The American Journal of Human Genetics	۱۷	۱۰.۵۰۲	۱۰.۳۴۴
۸	The Lancet Diabetes & Endocrinology	۱۶	۲۴.۵۴۰	۲۳.۰۱۳
۹	Journal of Clinical Lipidology	۱۶	۳.۸۶۰	۴.۰۸۰
۱۰	Journal of Allergy and Clinical Immunology	۱۵	۱۰.۲۲۸	۱۱.۵۹۲

برای اجتماع یابی از الگوریتم اجتماع یابی^۲ بکار گرفته شده در نرم‌افزار Gephi 0.9.2 موسوم به الگوریتم لووین^۳ که توسط بلاندل^۴ و همکاران در سال ۲۰۰۸ ارائه شد، استفاده شده است (Blondel et al., 2008).

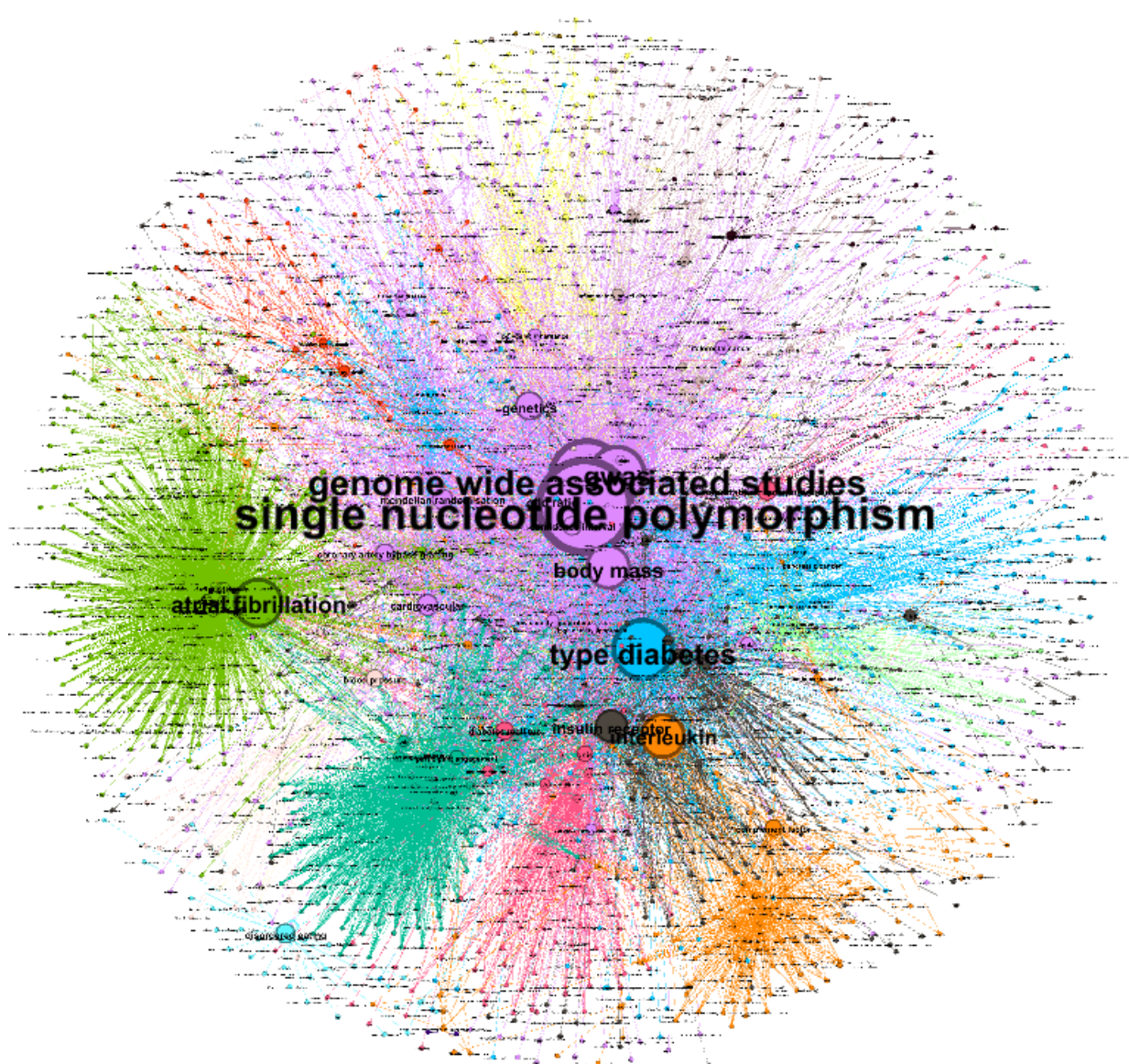
^۱ Impact Factor

^۲ Community detection

^۳ Louvain

^۴ Blondel

در شکل ۴-۱۷ خروجی اجتماع یابی شبکه بر اساس الگوریتم لووین مشاهده می‌شود. در شبکه به‌دست‌آمده مقدار ماژولاریتی ۰.۵۹ هست. در این شکل هر اجتماع با یک رنگ مجزا مشخص شده است. بر اساس (Xie and Szymanski, 2011) ماژولاریتی بین ۰.۳ تا ۰.۷ نشان‌دهنده ساختار قوی خوشه‌هاست. نتایج اجتماع یابی نشان می‌دهد در حوزه پژوهش‌های این موضوع به ترتیب زیر حوزه‌های بیماری‌های مرتبط با دیابت (۲۶.۳۲٪)، بیماری‌های قلبی (۱۳.۲۹٪)، بیماری‌های خونی (۶.۳۱٪) و فشارخون (۵.۳۴٪) از منظر مسائل خطر ژنتیکی بیشترین پژوهش‌ها را به خود اختصاص داده است. نتایج به‌دست‌آمده در مقاله (لاجوردی، علوی و کارگری، ۱۳۹۹) منتشر شده است.



شکل ۴-۱۸ نمایی از شبکه ارتباط کلمات کلیدی در مقالات

۴-۱۰- جمع‌بندی

نتایج کامل تحقیق در این فصل به ترتیب گفته‌شده در فصل قبل ارائه گردید و درنهایت نقشه علمی موضوع تحقیق ارائه شد. ضمناً در این فصل مقالاتی مستخرج از تز نیز ارائه شد و درنهایت در فصل بعدی به بیان نقاط قوت و ضعف تحقیق و همچنین کاربرد آن پرداخته می‌شود.

۵- بحث و نتیجه‌گیری

در این فصل به تحلیل نتایج به‌دست‌آمده در فصل قبل پرداخته می‌شود و درنهایت با جمع‌بندی مطالب مستند تر به اتمام می‌رسد.

۵-۱- بررسی مدل پیش‌بینی کننده خطر HTN مبتنی بر SNP

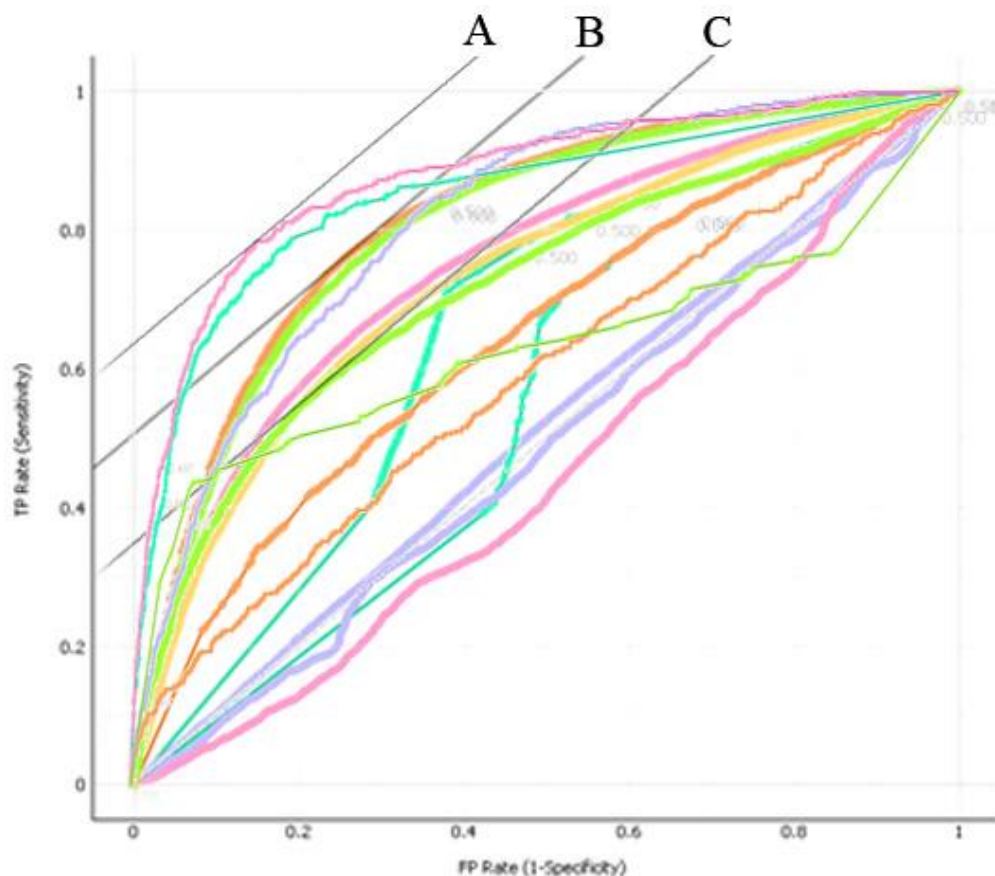
در مطالعه حاضر، بیش از ۴۸۰۰۰ برچسب فشارخون در شرکت‌کنندگان TCGS توسط پنج مدل مختلف یادگیری ماشینی و یادگیری عمیق مبتنی بر CNN ارزیابی و پیش‌بینی شد که نتیجه آن مقدار AUC ۰.۸۸۵ برای بهترین مدل یادگیری ماشین و مقدار ۰.۸۷۷ برای CNN را بدون در نظر گرفتن عوامل خطر اضافی نشان می‌دهد. در مقایسه با سایرین، علاوه بر اینکه این مدل بر اساس AUC از کیفیت بهتری برخوردار است، به دلیل عدم استفاده از سایر پارامترهای بیولوژیکی، ابزار ساده‌تری برای پیش‌بینی خطر است.

بالاترین AUC گزارش‌شده در ادبیات موضوع برای تشخیص بالینی فشارخون بالا بر اساس عوامل بالینی فشارخون ۰.۷۵۷ است. مطالعات دیگر شامل عوامل خطر ژنتیکی و بالینی فشارخون بر اساس GRS است که خطر فشارخون را با AUC ۰.۷۸ و ۰.۸۵۴ پیش‌بینی می‌کند (Li et al., 2020; Wu et al., 2020; Niu et al., 2021). بنابراین، یافته‌های این تحقیق عملکرد بهتری را به‌ویژه پس از در نظر گرفتن نشانگرهای ژنومی، جنسیت و سن نشان می‌دهد. همچنین با مقایسه نتایج مدل‌ها نشان داده می‌شود که SVM برای مدل‌سازی این داده‌ها نامناسب است و جنسیت تأثیر قابل‌توجهی بر فشارخون بالا ندارد.

امروزه، مهم‌ترین دستاورد عملی در عصر ژنومی، ایجاد ابزاری برای تشخیص زودهنگام خطر فشارخون بر اساس نشانگرهای ژنومی برای بهبود پیشگیری و درمان در افراد پرخطر است. در این تحقیق ابزاری اصلاح‌شده برای پیش‌بینی خطر فشارخون بالا در سنین پایین‌تر ارائه شده است.

۵-۲- مقایسه نتایج و معنی‌داری بهبود

بر اساس (Pepe et al., 2013) برای نشان دادن معنی‌دار بودن بهبود در مدل‌های پیش‌بینی، بر اساس نمودار ROC مقایسه انجام می‌گردد و اگر در تمامی نقاط مساحت زیر نمودار بهتر باشد، نشان‌دهنده معنی‌داری بهبود دقت برای آن مدل است. در شکل ۵-۱ نمودارهای ROC مدل‌های CNN (A) و یادگیری ماشین مبتنی بر کاندید (B) و مبتنی بر امتیاز خطر ژنتیکی (C) آمده است. همان‌طور که در شکل نشان داده شده، مدل CNN در تمامی نقاط دارای عملکرد بهتری هست که نشان‌دهنده معنی‌داری بهبود هست.



شکل ۱-۵ مقایسه نمودارهای ROC

۳-۵- بررسی نتایج مدل سازی درخت تصمیم

با توجه به مدل سازی درخت تصمیم (شکل ۴-۱۵) بر روی داده ها نکات ذیل به دست می آید.

۱- اولین انشعاب در سن ۵۰ سالگی هست که نشان‌دهنده مرز مشخص ۵۰ سال برای افزایش خطر بیماری پرفشاری خون هست.

۲- در شاخه کم‌خطر، انشعاب بعدی در سن ۱۸ سالگی هست که مرز اولیه خطر پرفشاری خون را نشان می‌دهند.

۳- در شاخه پرخطر، انشعاب بعدی در سن ۵۹ سال است که مرز نهایی خطر پرفشاری خون است.

۴- تا سطح چهارم انشعاب بر روی سن هست که نشان می‌دهد سن مهم‌ترین عامل در بروز بیماری پرفشاری خون است.

۵- پس از سن عوامل ژنومی تأثیر بالاتری از جنسیت دارند و می‌توان گفت جنسیت تأثیر مستقیمی در بروز بیماری پرفشاری خون ندارد. این مسئله که جنسیت تأثیر مستقیمی ندارد در مقاله (Hengel, Benitah and Wenzel, 2022) نیز بیان شده است.

۶- تفاوت معنی‌داری در ژنوم‌های مؤثر بیماری پرفشاری خون در بین زنان و مردان نیست. این مسئله در جدول ۴-۵ آمده است که ۱۰ SNP با اهمیت بر روی کروموزوم‌های غیر جنسیتی هستند.

۵-۴- بررسی نتایج حاصل از رتبه‌بندی SNP

در میان نشانگرهای قبلی مرتبط با فشارخون، ده SNP برتر در این مطالعه رتبه‌بندی شده و نتایج نشان داد که SNP های برتر بالاترین p-value را ندارند. یافته‌های مبتنی بر رتبه‌بندی نشان می‌دهد که برخلاف نتایج قبلی بر اساس p-value، باید نشانگرها را بر اساس ارتباط رتبه‌بندی کنیم (Sullivan and Feinn, 2012). البته لازم به ذکر است که SNP هایی را باید در نظر گرفت که دارای سطح معنی‌داری خوبی با فشارخون بالا هستند.

۵-۵- خطاهای مؤثر در تحقیق

با توجه به اینکه تحقیق و مدل‌سازی انجام‌شده بر مبنای داده هست، ممکن است خطاهای مختلفی در حین مدل‌سازی وجود داشته باشد.

۱- خطای ناشی از اندازه‌گیری رخ‌نمود فشارخون: این خطا در هنگام اندازه‌گیری فشارخون ممکن است رخ دهد. با توجه به توضیحات نحوه اندازه‌گیری این رخ‌نمود، اندازه فشارخون توسط دستگاه فشارسنج حیوهای و پس از استراحت کافی هست که در دو مرحله این مقدار سنجیده شده و درنهایت میانگین آن گزارش‌شده است که دارای حداقل خطا هست.

۲- خطای ناشی از شناسایی SNP درروش NGS: با توجه به تحقیق (Petrackova et al., 2019) خطای فرایند NGS کمتر از ۱ درصد هست (مابین ۰.۱ درصد الی ۱ درصد). این خطا را می‌توان با افزایش حداقل عمق پوشش در فرایند NGS کاهش داد.

۳- خطای ناشی از جنس داده و نوع مدل: در مدل‌سازی‌های انجام‌شده مدل SVM برخلاف تحقیقات مختلف دارای رتبه پایین‌تری نسبت به شبکه عصبی است. دلایل این کاهش کیفیت شامل: بزرگ بودن داده و نداشتن تابع فعال‌سازی مناسب، نویزی بودن داده. در این تحقیق توابع فعال‌سازی مختلفی برای SVM پیاده گردید اما نتیجه مناسب نبود و البته وجود نویز بالا در داده‌ها نیز مزید بر علت بود. این نویز در داده‌های ژنومی با توجه به نرخ خطای فرایند NGS ایجادشده است.

۵-۶- پاسخ به سؤالات تحقیق

سؤالات اصلی تحقیق به‌صورت زیر هست:

✓ آیا می‌توان از روش‌های یادگیری ماشین در پیش‌بینی خطر پرفشاری خون مبتنی بر SNP به شکل مؤثرتر از روش‌های موجود بهره برد؟ بله. با توجه به روش‌های استفاده‌شده نشان داده‌شده است که یادگیری ماشین می‌تواند باعث افزایش دقت و کیفیت پیش‌بینی خطر پرفشاری خون باشد.

سؤال فرعی تحقیق نیز به‌صورت زیر هست:

✓ آیا می‌توان اطلاعات جمعیت‌شناختی و داده‌های طولی رخ‌نمود را در بهبود مدل پیش‌بینی کننده استفاده کرد؟ بر اساس داده‌های استفاده‌شده که در شش دوره هست، این امکان ایجاد شد که بتوانیم مقایسه

مناسبی از لحاظ تأثیر سنی و همچنین وضعیت فرد در سال‌های بعد داشته باشیم که این امر در مدل‌سازی استفاده شده است.

✓ با توجه به مدل ارائه شده، آیا می‌توان SNP‌های مهم ژنومی را استخراج کرد؟ روش یادگیری ماشین به صورت جعبه سیاه هست و تشخیص SNP با اهمیت دچار مشکل است و بر این اساس از روش‌های رتبه‌بندی به جهت سنجش اهمیت برای SNP‌هایی که معنی‌داری ارتباط با HTN آن‌ها توسط GWAS بیان شده بود، استفاده گردید که نتایج آن به صورت یک ابزار پیش‌بینی خطر پیاده‌سازی گردید.

۵-۷- نوآوری‌های تحقیق

در این تحقیق با توجه به مدل‌سازی‌های انجام شده می‌توان نوآوری‌های به دست آمده را به صورت ذیل بیان کرد:

۱- با توجه به اهمیت بیماری پرفشاری خون که در حال حاضر بیش از یک میلیارد نفر را در جهان درگیر کرده است، ایجاد یک ابزار که بتواند در سنین جوانی خطر بیماری را پیش‌بینی کند و پیش‌آگهی بدهد ضروری است که این امر باعث افزایش مراقبت‌ها و سبک زندگی افرادی است که دارای خطر بالایی می‌باشند. در این تحقیق بر اساس مدل‌های یادگیری ماشین و تنها داده ژنومی، قبل از بروز عوامل دیگر مانند چاقی، پیش‌آگهی مناسبی داده می‌شود که می‌تواند توسط پزشکان استفاده گردد.

۲- فرد می‌تواند قبل از آنکه دچار بیماری پرفشاری خون شود با سنجش خطر در سنین پایین، سبک زندگی مناسب‌تری را انتخاب کند که مانع از درگیری با بیماری پرفشاری خون می‌شود.

۳- نوع مدل‌سازی انجام شده که حاصل ادغام داده‌های طولی در یک مدل یادگیری ماشین هست نوآوری مهم این تحقیق است. البته ایجاد تصویر از داده‌های ژنومی جهت اجرا در مدل CNN نیز یکی از نوآوری‌هاست که باعث کاهش حجم پروتینا ژنوم نیز می‌گردد.

۴- به جهت استفاده، ابزاری اولیه ایجاد شده است که جزو نوآوری‌های تحقیق است که در بخش ۴-۸ اشاره گردید. این ابزار به سادگی می‌تواند با تعداد محدود SNP مؤثر، خطر بیماری را در سنین مختلف محاسبه نماید.

۵- یکی از نتایج تحقیق عدم تأثیرگذاری جنسیت در خطر بیماری پرفشاری خون است که در تحقیقات جدید نیز به این مسئله اشاره شده است (Hengel, Benitah and Wenzel, 2022). آنچه باعث افزایش خطر می‌شود، افزایش استرس است که معمولاً در مردان بیشتر است و افزایش استرس نیز عوامل مختلفی دارد من جمله عوامل ژنومی.

۶- یکی از خطاهای مهم در روش امتیاز خطر چندژنی، خطای تأثیر کوچک^۱ هست. در این خطا به دلیل اینکه تک تک نوکلئوتیدها دارای p-value کمتر از حد تعریف شده هستند از مدل‌سازی خارج می‌شود. این مسئله در صورتی است که تأثیر چند تک نوکلئوتید باهم دیگر می‌تواند عامل مؤثر در بروز رخ نمود باشد. در این تحقیق با توجه به در نظر گرفتن تمامی ژنوم، این مسئله مرتفع شده است.

¹ Small Effect

۵-۸- محدودیت‌ها و پیشنهادهای آتی

این تحقیق به دلیل استفاده از اعتبارسنجی داده‌های داخلی دارای محدودیت‌هایی است. نتایج این مطالعه خارج از داده‌های TCGS برای اعتبارسنجی خارجی استفاده نشده است و همچنین داده‌های مورد استفاده برای ارزیابی محدود به داده‌های جمع‌آوری شده در TCGS است که ترکیب قومی کمی دارد.

برای پیشنهادهای آتی، می‌توان متغیرهای بالینی مؤثر مانند نژاد را وارد مدل کرد. البته باید دقت کرد که تنها متغیرهایی وارد مدل شوند که وابستگی به سن نداشته باشند.

همچنین می‌توان برای مدل‌سازی بهتر فشارخون به صورت مجزا مدل‌سازی IDH و ISH را نیز انجام داد. برای مدل‌سازی بهتر می‌توان از الگوریتم‌های LightGBM و XGBoost و Ensemble‌های دیگر نیز استفاده و نتایج آن را با نتایج مدل‌های ارائه شده مقایسه کرد. همچنین می‌توان وضعیت فرد در دوره قبلی را به عنوان یکی از ورودی‌های دوره جدید قرار داد و بر اساس آن با پنجره زمانی مدل‌سازی برای پیش‌آگهی انجام داد. یکی از پیشنهادهای مهم می‌تواند بررسی الگوی ژنومی فشارخون بر اساس داده‌های بومی باشد که در این تحقیق نسخه اولیه آن انجام شد و در ادامه می‌توان آن را تکمیل و گزارش کرد.

۵-۹- خلاصه فصل

در این فصل به بررسی یافته‌های تحقیق، نتیجه‌گیری و جمع‌بندی نتایج پرداختیم. علاوه بر این محدودیت‌ها و پیشنهادهای پژوهشی آینده بیان گردید.

فهرست منابع

- Adams, F. (1922) 'The Genuine Works of Hippocrates', *Southern Medical Journal*. Southern Medical Association, 15(6), p. 519. doi: 10.1097/00007611-192206000-00025.
- Alizadehsani, R. *et al.* (2019) 'Machine learning-based coronary artery disease diagnosis: A comprehensive review', *Computers in Biology and Medicine*. Elsevier Ltd, 111(May), p. 103346. doi: 10.1016/j.compbiomed.2019.103346.
- Anis, A., Fahiem, M. A. and Tauseef, H. (2016) 'A Classification Approach Based on Genetic-Data-Structuring for the Prediction of Hypertension', 4(4), pp. 236–242.
- Aulchenko, Y. S., De Koning, D.-J. J. and Haley, C. (2007) 'Genomewide rapid association using mixed model and regression: A fast and simple method for genomewide pedigree-based quantitative trait loci association analysis', *Genetics*. Oxford University Press (), 177(1), pp. 577–585. doi: 10.1534/genetics.107.075614.
- Azizi, F. *et al.* (2009) 'Prevention of non-communicable disease in a population in nutrition transition: Tehran Lipid and Glucose Study phase {II}', *Trials*. Springer Science and Business Media , 10(1). doi: 10.1186/1745-6215-10-5.
- de Bakker, P. I. W. *et al.* (2008) 'Practical aspects of imputation-driven meta-analysis of genome-wide association studies', *Human Molecular Genetics*. Oxford University Press (), 17(R2), pp. R122–R128. doi: 10.1093/hmg/ddn288.
- Basile, J. *et al.* (2019) 'Overview of hypertension in adults - UpToDate', *Up To Date*, pp. 1–56. Available at: https://www.uptodate.com/contents/overview-of-hypertension-in-adults?search=hypertension&source=search_result&selectedTitle=1~150&usage_type=default&display_rank=1%0Ahttps://www-uptodate-com.puce.idm.oclc.org/contents/overview-of-hypertension-in-adults?se.
- Bauer, C. *et al.* (2015) '{BioMiner}: Paving the Way for Personalized Medicine', *Cancer Informatics*. {SAGE} Publications, 14, p. CIN.S20910. doi: 10.4137/cin.s20910.
- Blondel, V. D. *et al.* (2008) 'Fast unfolding of communities in large networks', *Journal of statistical mechanics: theory and experiment*. IOP Publishing, 2008(10), p. P10008.
- Breiman, L. (2000) 'Randomizing outputs to increase prediction accuracy', *Machine Learning*. Kluwer Academic Publishers, 40(3), pp. 229–242. doi: 10.1023/A:1007682208299.
- Bush, W. S. (2019) 'Genome-Wide Association Studies', in *Encyclopedia of Bioinformatics and Computational Biology*. Elsevier, pp. 235–241. doi: 10.1016/B978-0-12-809633-8.20232-X.
- Çalışkan, M., Erol, O. and Öz, G. C. (2020) *The Recent Topics in Genetic Polymorphisms*. Edited by M. Çalışkan, O. Erol, and G. Cevahir Öz. Rijeka: IntechOpen. doi: 10.5772/intechopen.77777.
- Cardon, L. R. and Bell, J. I. (2001) 'Association study designs for complex diseases', *Nature Reviews Genetics*. Springer Science and Business Media , 2(2), pp. 91–99. doi: 10.1038/35052543.
- Cebamano, L. *et al.* (2014) 'Regional heritability advanced complex trait analysis for {GPU} and traditional parallel architectures', *Bioinformatics*. Oxford University Press (), 30(8), pp. 1177–1179. doi: 10.1093/bioinformatics/btt754.
- Chowdhury, B. and Garai, G. (2017) 'A review on multiple sequence alignment from the perspective of genetic algorithm', *Genomics*. Academic Press, 109(5–6), pp. 419–431. doi: 10.1016/J.YGENO.2017.06.007.
- Chowdhury, M. Z. I. *et al.* (2022) 'Prediction of hypertension using traditional regression and machine learning models: A systematic review and meta-analysis', *PLOS ONE*. Public Library of Science, 17(4), p. e0266334. doi: 10.1371/JOURNAL.PONE.0266334.
- Cios, K. J. and William Moore, G. (2002) 'Uniqueness of medical data mining', *Artificial Intelligence in Medicine*, 26(1–2), pp. 1–24. doi: 10.1016/S0933-3657(02)00049-0.
- Daneshpour, M. S. *et al.* (2017) 'Rationale and Design of a Genetic Study on Cardiometabolic Risk Factors: Protocol for the Tehran Cardiometabolic Genetic Study (TCGS)', *JMIR Research Protocols*. {JMIR} Publications Inc., 6(2), p. e28. doi: 10.2196/resprot.6050.

- Davies, G. *et al.* (2011) 'Genome-wide association studies establish that human intelligence is highly heritable and polygenic', *Molecular Psychiatry*. Springer Science and Business Media , 16(10), pp. 996–1005. doi: 10.1038/mp.2011.85.
- Diamandis, E. P. (2012) 'Biomarker validation is still the bottleneck in biomarker research', *Journal of Internal Medicine*. Wiley, 272(6), p. 620. doi: 10.1111/j.1365-2796.2012.02579.x.
- Eshmurodova Dilrabo Khudayberdi kizi, Ergasheva Iroda Akbarali kizi, K. F. N. and Siab (2023) 'ASSESSMENT OF POLYGENIC RISK OF ARTERIAL HYPERTENSION', in *SCIENTIFIC APPROACH TO THE MODERN EDUCATION SYSTEM*, pp. 39–43. Available at: www.interonconf.com.
- Esteva, A. *et al.* (2019) 'A guide to deep learning in healthcare', *Nature Medicine*. Springer Science and Business Media , 25(1), pp. 24–29. doi: 10.1038/s41591-018-0316-z.
- Evangelou, E. and Ioannidis, J. P. A. A. (2013) 'Meta-analysis methods for genome-wide association studies and beyond', *Nature Reviews Genetics*. Springer Science and Business Media , 14(6), pp. 379–389. doi: 10.1038/nrg3472.
- Farris, J. *et al.* (2021) 'Genomic analyses of glycine decarboxylase neurogenic mutations yield a large-scale prediction model for prenatal disease', *PLOS Genetics*. Edited by G. S. Barsh, 17(2), p. e1009307. doi: 10.1371/journal.pgen.1009307.
- Fawcett, T. (2006) 'An introduction to {ROC} analysis', *Pattern Recognition Letters*. Elsevier , 27(8), pp. 861–874. doi: 10.1016/j.patrec.2005.10.010.
- Filshstein, T. J. *et al.* (2019) 'Reserve and Alzheimer's disease genetic risk: Effects on hospitalization and mortality', *Alzheimer's and Dementia*. Elsevier Inc., 15(7), pp. 907–916. doi: 10.1016/j.jalz.2019.04.005.
- Ganesalingam, J. and Bowser, R. (2010) 'The application of biomarkers in clinical trials for motor neuron disease', *Biomarkers in Medicine*. Future Medicine Ltd, 4(2), pp. 281–297. doi: 10.2217/bmm.09.71.
- Garimella, P. S. *et al.* (2023) 'A genomic deep field view of hypertension', *Kidney International*, 103(1), pp. 42–52. doi: 10.1016/j.kint.2022.09.029.
- Gibson, G. (2010) 'Hints of hidden heritability in {GWAS}', *Nature Genetics*. Springer Science and Business Media , 42(7), pp. 558–560. doi: 10.1038/ng0710-558.
- Gilmour, A. R., Thompson, R. and Cullis, B. R. (1995) 'Average Information {REML}: An Efficient Algorithm for Variance Parameter Estimation in Linear Mixed Models', *Biometrics*. JSTOR, 51(4), p. 1440. doi: 10.2307/2533274.
- Ginsburg, G. S., Donahue, M. P. and Newby, L. K. (2005) 'Prospects for Personalized Cardiovascular Medicine', *Journal of the American College of Cardiology*. Elsevier , 46(9), pp. 1615–1627. doi: 10.1016/j.jacc.2005.06.075.
- Ginsburg, G. S. and Willard, H. F. (2009) 'Genomic and personalized medicine: foundations and applications', *Translational Research*. Elsevier , 154(6), pp. 277–287. doi: 10.1016/j.trsl.2009.09.005.
- Giontella, A. *et al.* (2021) 'Clinical Evaluation of the Polygenetic Background of Blood Pressure in the Population-Based Setting', *Hypertension*, 77(1), pp. 169–177. doi: 10.1161/HYPERTENSIONAHA.120.15449.
- Gray, A., Stewart, I. and Tenesa, A. (2012) 'Advanced Complex Trait Analysis', *Bioinformatics*. Oxford University Press (), 28(23), pp. 3134–3136. doi: 10.1093/bioinformatics/bts571.
- Guns, R., Liu, Y. X. and Mahbuba, D. (2011) 'Q-measures and betweenness centrality in a collaboration network: a case study of the field of informetrics', *Scientometrics*. Springer, 87(1), pp. 133–147.
- He, Q. (1999) 'Knowledge discovery through co-word analysis'. Graduate School of Library and Information Science. University of Illinois~....
- Hebbring, S. (2019) 'Genomic and Phenomic Research in the 21st Century', *Trends in Genetics*. Elsevier , 35(1), pp. 29–41. doi: 10.1016/j.tig.2018.09.007.
- Hengel, F. E., Benitah, J.-P. and Wenzel, U. O. (2022) 'Mosaic theory revised: inflammation and salt play central roles in arterial hypertension', *Cellular & Molecular Immunology*. Springer

- US, 19(5), pp. 561–576. doi: 10.1038/s41423-022-00851-8.
- Heo, B. M. (2018) ‘Prediction of Prehypertension and Hypertension Based on Anthropometry , Blood Parameters , and Spirometry’. doi: 10.3390/ijerph15112571.
- Hernández-Rodríguez, S. *et al.* (2016) ‘A recommender system applied to the indirect materials selection process (RS-IMSP) for producing automobile spare parts’, *Computers in Industry*. Elsevier B.V., 82, pp. 233–244. doi: 10.1016/j.compind.2016.07.004.
- Hirschhorn, J. N. and Daly, M. J. (2005) ‘Genome-wide association studies for common diseases and complex traits’, *Nature Reviews Genetics*. Springer Science and Business Media , 6(2), pp. 95–108. doi: 10.1038/nrg1521.
- Ho, D. S. W. *et al.* (2019) ‘Machine learning SNP based prediction for precision medicine’, *Frontiers in Genetics*, 10(MAR), pp. 1–10. doi: 10.3389/fgene.2019.00267.
- Hordri, N. F. *et al.* (2017) ‘A systematic literature review on features of deep learning in big data analytics’, *International Journal of Advances in Soft Computing and its Applications*, 9(1), pp. 32–49.
- Jain, K. K. (2006) ‘A critical review of the Royal Society’s report on personalized medicine’, *Drug Discovery Today*. Elsevier , 11(13–14), pp. 573–575. doi: 10.1016/j.drudis.2006.05.005.
- Khan, A. *et al.* (2020) ‘A survey of the recent architectures of deep convolutional neural networks’, *Artificial Intelligence Review*. Springer Science and Business Media , 53(8), pp. 5455–5516. doi: 10.1007/s10462-020-09825-6.
- Kheradmand, M. *et al.* (2015) *Population Based Cohort Studies in Iran: A Review Article*, *Journal of Mazandaran University of Medical Sciences*. Journal of Mazandaran University of Medical Sciences. Available at: <http://jmums.mazums.ac.ir/article-1-5714-en.html> (Accessed: 14 February 2021).
- Khorraminezhad, L. *et al.* (2020) ‘Statistical and Machine-Learning Analyses in Nutritional Genomics Studies’, *Nutrients*, 12(10), p. 3140. doi: 10.3390/nu12103140.
- Kim, H. and Andrade, F. C. D. (2019) ‘Diagnostic status and age at diagnosis of hypertension on adherence to lifestyle recommendations’, *Preventive Medicine Reports*. Elsevier, 13(October 2018), pp. 52–56. doi: 10.1016/j.pmedr.2018.11.005.
- Kolifarhood, G., Daneshpour, M. S., *et al.* (2019) ‘Generality of genomic findings on blood pressure traits and its usefulness in precision medicine in diverse populations: A systematic review’, *Clinical Genetics*, 96(1), pp. 17–27. doi: 10.1111/cge.13527.
- Kolifarhood, G., Daneshpour, M., *et al.* (2019) ‘Heritability of blood pressure traits in diverse populations: a systematic review and meta-analysis’, *Journal of Human Hypertension*. Springer US, 33(11), pp. 775–785. doi: 10.1038/s41371-019-0253-4.
- Kolifarhood, G. *et al.* (2021) ‘Genome-wide association study on blood pressure traits in the Iranian population suggests ZBED9 as a new locus for hypertension’, *Scientific Reports*. Nature Publishing Group UK, 11(1), pp. 1–13. doi: 10.1038/s41598-021-90925-w.
- Krizhevsky, A., Sutskever, I. and Hinton, G. E. (2017) ‘Imagenet classification with deep convolutional neural networks’, *Communications of the ACM*. AcM New York, NY, USA, 60(6), pp. 84–90.
- Kurniansyah, N. *et al.* (2022) ‘A multi-ethnic polygenic risk score is associated with hypertension prevalence and progression throughout adulthood’, *Nature Communications*. Springer US, 13(1), p. 3549. doi: 10.1038/s41467-022-31080-2.
- Kurtz, T. W. and Spence, M. A. A. (1993) ‘Genetics of essential hypertension’, *The American Journal of Medicine*. Elsevier , 94(1), pp. 77–84. doi: 10.1016/0002-9343(93)90124-8.
- Lajevardi, S. Ali *et al.* (2022) ‘A CNN deep learning model to improve SNP-based hypertension risk prediction accuracy’.
- Lajevardi, S. Ali *et al.* (2022) ‘Hypertension risk prediction based on SNPs by machine learning models’, *Current Bioinformatics*. Bentham Science Publishers Ltd., 17. doi: 10.2174/1574893617666221011093322.
- Lajevardy, S. A. and Kargari, M. (2022) ‘Developing new genetic algorithm based on integer

- programming for multiple sequence alignment', *Soft Computing* 2022 26:8. Springer, 26(8), pp. 3863–3870. doi: 10.1007/S00500-022-06790-W.
- Lecun, Y., Bengio, Y. and Hinton, G. (2015) 'Deep learning', *Nature*. Springer Science and Business Media , 521(7553), pp. 436–444. doi: 10.1038/nature14539.
- Li, C. *et al.* (2019) 'A prediction model of essential hypertension based on genetic and environmental risk factors in northern han chinese', *International Journal of Medical Sciences*, 16(6), pp. 793–799. doi: 10.7150/ijms.33967.
- Li, S. *et al.* (2020) 'Genetic risk scores to predict the prognosis of chronic heart failure patients in Chinese Han', *Journal of Cellular and Molecular Medicine*, 24(1), pp. 285–293. doi: 10.1111/jcmm.14722.
- Libbrecht, M. W. and Noble, W. S. (2015) 'Machine learning applications in genetics and genomics', *Nature Reviews Genetics*. Springer Science and Business Media , 16(6), pp. 321–332. doi: 10.1038/nrg3920.
- Ligon, A. H. *et al.* (2000) 'Identification of female carriers for Duchenne and Becker muscular dystrophies using a {FISH}-based approach', *European Journal of Human Genetics*. Springer Science and Business Media , 8(4), pp. 293–298. doi: 10.1038/sj.ejhg.5200450.
- Loganathan, L. *et al.* (2019) 'Computational and Pharmacogenomic Insights on Hypertension Treatment: Rational Drug Design and Optimization Strategies', *Current Drug Targets*, 21(1), pp. 18–33. doi: 10.2174/1389450120666190808101356.
- Mahajan, S. *et al.* (2019) 'Prevalence, Awareness, and Treatment of Isolated Diastolic Hypertension: Insights From the China PEACE Million Persons Project', *Journal of the American Heart Association*, 8(19), pp. 1–17. doi: 10.1161/JAHA.119.012954.
- Maher, B. (2008) 'Personal genomes: The case of the missing heritability', *Nature*. Springer Science and Business Media , 456(7218), pp. 18–21. doi: 10.1038/456018a.
- Mahrokh, L. *et al.* (2017) 'An efficient microbial fuel cell using a {CNT} {\textendash} {RTIL} based nanocomposite', *Journal of Materials Chemistry A*. Royal Society of Chemistry ({RSC}), 5(17), pp. 7979–7991. doi: 10.1039/c6ta09766a.
- Maj, C. *et al.* (2022) 'Dissecting the Polygenic Basis of Primary Hypertension: Identification of Key Pathway-Specific Components', *Frontiers in Cardiovascular Medicine*, 9(February), pp. 1–10. doi: 10.3389/fcvm.2022.814502.
- Maltby, D. (2011) 'Big Data Analytics', *ASIST*.
- Marcos-Zambrano, L. J. *et al.* (2021) 'Applications of Machine Learning in Human Microbiome Studies: A Review on Feature Selection, Biomarker Identification, Disease Prediction and Treatment', *Frontiers in Microbiology*, 12(February). doi: 10.3389/fmicb.2021.634511.
- Martinez-Ríos, E. *et al.* (2021) 'A review of machine learning in hypertension detection and blood pressure estimation based on clinical and physiological data', *Biomedical Signal Processing and Control*. Elsevier Ltd, 68(March), p. 102813. doi: 10.1016/j.bspc.2021.102813.
- Mattoo, T. K. (2019) 'Definition and diagnosis of hypertension in children and adolescents - UpToDate', *UpToDate*, (Cv), pp. 1–34. Available at: [https://www.uptodate.com/contents/definition-and-diagnosis-of-hypertension-in-children-and-adolescents?search=tension arterial&source=search_result&selectedTitle=1~150&usage_type=default&display_rank=1#H12](https://www.uptodate.com/contents/definition-and-diagnosis-of-hypertension-in-children-and-adolescents?search=tension%20arterial&source=search_result&selectedTitle=1~150&usage_type=default&display_rank=1#H12) (Accessed: 14 February 2021).
- McCarthy, M. I. *et al.* (2008) 'Genome-wide association studies for complex traits: Consensus, uncertainty and challenges', *Nature Reviews Genetics*. Springer Science and Business Media , 9(5), pp. 356–369. doi: 10.1038/nrg2344.
- McKusick, V. A. *et al.* (1960) 'Medical genetics 1959', *Journal of Chronic Diseases*. Elsevier, 12(1), pp. 1–202. doi: 10.1016/0021-9681(60)90134-X.
- McKusick, V. A. (1960) 'Genetics and the Nature of Essential Hypertension', *Circulation*, XXII.
- Montañez, C. A. C., Fergus, P. and Chalmers, C. (2018) 'Deep Learning Classification of Polygenic Obesity using Genome Wide Association Study SNPs', in *2018 International Joint*

Conference on Neural Networks (IJCNN).

- Mosley, J. D. *et al.* (2020) 'Predictive Accuracy of a Polygenic Risk Score Compared with a Clinical Risk Score for Incident Coronary Heart Disease', *JAMA - Journal of the American Medical Association*, 323(7), pp. 627–635. doi: 10.1001/jama.2019.21782.
- Ng, H. T., Goh, W. B. and Low, K. L. (1997) 'Feature selection, perceptron learning, and a usability case study for text categorization', *{ACM} {SIGIR} Forum*. Association for Computing Machinery ({ACM}), 31(SI), pp. 67–73. doi: 10.1145/278459.258537.
- Niu, M. *et al.* (2021) 'Identifying the predictive effectiveness of a genetic risk score for incident hypertension using machine learning methods among populations in rural China', *Hypertension Research*. Springer US. doi: 10.1038/s41440-021-00738-7.
- Nour, M. and Polat, K. (2020) 'Automatic Classification of Hypertension Types Based On Personal Features By Machine Learning Algorithm', *Mathematical Problems in Engineering*. Hindawi Limited, 2020, pp. 1–13. doi: 10.1155/2020/2742781.
- Noyons, E. C. M. (1999) 'Bibliometric mapping as a science policy and research management tool'. DSWO Press.
- Padmanabhan, S. and Dominiczak, A. F. (2021) 'Genomics of hypertension: the road to precision medicine', *Nature Reviews Cardiology*. Springer US, 18(4), pp. 235–250. doi: 10.1038/s41569-020-00466-4.
- Pattin, K. A. and Moore, J. H. (2009) 'Role for protein\textendashprotein interaction databases in human genetics', *Expert Review of Proteomics*. Informa {UK} Limited, 6(6), pp. 647–659. doi: 10.1586/epr.09.86.
- Pepe, M. S. *et al.* (2013) 'Testing for improvement in prediction model performance', *Statistics in Medicine*, 32(9), pp. 1467–1482. doi: 10.1002/sim.5727.
- Pérez-Enciso, M. *et al.* (2019) 'A Guide for Using Deep Learning for Complex Trait Genomic Prediction', *Genes*. MDPI AG, 10(7), p. 553. doi: 10.3390/genes10070553.
- Petrackova, A. *et al.* (2019) 'Standardization of Sequencing Coverage Depth in NGS: Recommendation for Detection of Clonal and Subclonal Mutations in Cancer Diagnostics', *Frontiers in Oncology*, 9(September), pp. 1–6. doi: 10.3389/fonc.2019.00851.
- Phulka, J. S. *et al.* (2023) 'Current State and Future of Polygenic Risk Scores in Cardiometabolic Disease: A Scoping Review', *Circulation: Genomic and Precision Medicine*. doi: 10.1161/CIRCGEN.122.003834.
- Powers, D. M. W. and Ailab (2011) 'EVALUATION: FROM PRECISION, RECALL AND F-MEASURE TO ROC, INFORMEDNESS, MARKEDNESS & CORRELATION', *International Journal of Machine Learning Technology*, pp. 37–63.
- Quinlan, J. R. (1986) 'Induction of decision trees', *Machine Learning*. Springer Science and Business Media , 1(1), pp. 81–106. doi: 10.1007/bf00116251.
- Rafiei, A. and Amjadi, O. (2013) 'Personalized medicine; a bridge between current medicine and the future healthcare', *Journal of Clinical Excellence*, 1(2), pp. 47–68.
- Ramezankhani, A. *et al.* (2016) 'Classification-based data mining for identification of risk patterns associated with hypertension in Middle Eastern population', 0(June).
- Reddy, B. and Fields, R. (2022) 'Multiple Sequence Alignment Algorithms in Bioinformatics', *Lecture Notes in Networks and Systems*. Springer Science and Business Media Deutschland GmbH, 286, pp. 89–98. doi: 10.1007/978-981-16-4016-2_9/COVER.
- Rip, A. and Courtial, J. (1984) 'Co-word maps of biotechnology: An example of cognitive scientometrics', *Scientometrics*. Akadémiai Kiadó, co-published with Springer Science+ Business Media BV~..., 6(6), pp. 381–400.
- Risch, N. and Merikangas, K. (1996) 'The Future of Genetic Studies of Complex Human Diseases', *Science*. American Association for the Advancement of Science ({AAAS}), 273(5281), pp. 1516–1517. doi: 10.1126/science.273.5281.1516.
- Ruiz, M. E. and Srinivasan, P. (1999) 'Hierarchical neural networks for text categorization (poster abstract)', in *Proceedings of the 22nd annual international {ACM} {SIGIR} conference on Research and development in information retrieval - {SIGIR} {\textquotesingle}99*. {ACM}

- Press. doi: 10.1145/312624.312700.
- Runcie, D. E. and Crawford, L. (2018) 'Fast and flexible linear mixed models for genome-wide genetics', *PLoS Genetics*. Cold Spring Harbor Laboratory, 15(2), pp. 1–24. doi: 10.1101/373902.
- Russakovsky, O. *et al.* (2015) '{ImageNet} Large Scale Visual Recognition Challenge', *International Journal of Computer Vision*. Springer Science and Business Media , 115(3), pp. 211–252. doi: 10.1007/s11263-015-0816-y.
- Russom, P. (2018) 'Introduction to Big Data Analytics', in *Smart Grid Technology*. Cambridge University Press, pp. 38–48. doi: 10.1017/9781108566506.005.
- Satyanarayana, N. *et al.* (2020) 'An intelli AFM: An intelligent association based fuzzy rule miner to predict high blood pressure using bio-psychological factors', *Intelligent Decision Technologies*. Edited by K. Chrysafiadi, 14(2), pp. 227–237. doi: 10.3233/IDT-190156.
- Shah, S. *et al.* (2015) 'Improving Phenotypic Prediction by Combining Genetic and Epigenetic Associations', *American Journal of Human Genetics*. Cell Press, 97(1), pp. 75–85. doi: 10.1016/j.ajhg.2015.05.014.
- Shrestha, A. and Mahmood, A. (2019) 'Review of deep learning algorithms and architectures', *IEEE access*. IEEE, 7, pp. 53040–53065.
- Silver, D. *et al.* (2016) 'Mastering the game of Go with deep neural networks and tree search', *Nature*. Springer Science and Business Media , 529(7587), pp. 484–489. doi: 10.1038/nature16961.
- Singh, D., Verma, S. and Singla, J. (2020) 'A Comprehensive Review of Intelligent Medical Diagnostic Systems', in *2020 4th International Conference on Trends in Electronics and Informatics (ICOEI)*(48184). IEEE, pp. 977–981. doi: 10.1109/ICOEI48184.2020.9143043.
- Singh, M. *et al.* (2016) 'Molecular genetics of essential hypertension', 1963(April). doi: 10.3109/10641963.2015.1116543.
- Skol, A. D. *et al.* (2007) 'Optimal designs for two-stage genome-wide association studies', *Genetic Epidemiology*. Wiley, 31(7), pp. 776–788. doi: 10.1002/gepi.20240.
- Stehman, S. V (1997) 'Selecting and interpreting measures of thematic classification accuracy', *Remote Sensing of Environment*. Elsevier , 62(1), pp. 77–89. doi: 10.1016/s0034-4257(97)00083-7.
- Strachan, T. and Read, A. P. (2018) 'Human population genetics', in *Human Molecular Genetics*. Garland Science, pp. 397–418. doi: 10.1201/9780429448362-12.
- Sullivan, G. M. and Feinn, R. (2012) 'Using Effect Size—or Why the P Value Is Not Enough', *Journal of Graduate Medical Education*, 4(3), pp. 279–282. doi: 10.4300/JGME-D-12-00156.1.
- Sun, Y. *et al.* (2019) 'Identification of 12 cancer types through genome deep learning', *Scientific Reports*. Cold Spring Harbor Laboratory, 9(1). doi: 10.1101/528216.
- Szymczak, S. *et al.* (2009) 'Machine learning in genome-wide association studies', *Genetic Epidemiology*. Wiley, 33(S1), pp. S51--S57. doi: 10.1002/gepi.20473.
- Tohidi, M. *et al.* (2012) 'Triglycerides and triglycerides to high-density lipoprotein cholesterol ratio are strong predictors of incident hypertension in Middle Eastern women', *Journal of Human Hypertension*. Nature Publishing Group, 26(9), pp. 525–532. doi: 10.1038/jhh.2011.70.
- Torkamani, A., Wineinger, N. E. and Topol, E. J. (2018) 'The personal and clinical utility of polygenic risk scores', *Nature Reviews Genetics*. Springer US, 19(9), pp. 581–590. doi: 10.1038/s41576-018-0018-x.
- Trerotola, M. *et al.* (2015) 'Epigenetic inheritance and the missing heritability', *Human Genomics*. Springer Science and Business Media , 9(1). doi: 10.1186/s40246-015-0041-3.
- Visscher, P. *et al.* (2012) 'Five Years of {GWAS} Discovery', *The American Journal of Human Genetics*. Elsevier , 90(1), pp. 7–24. doi: 10.1016/j.ajhg.2011.11.029.
- Visscher, P. M. *et al.* (2014) 'Statistical Power to Detect Genetic (Co)Variance of Complex Traits Using {SNP} Data in Unrelated Samples', *{PLoS} Genetics*. Edited by G. S. Barsh. Public

- Library of Science ({PLoS}), 10(4), p. e1004269. doi: 10.1371/journal.pgen.1004269.
- Waezizadeh, T. *et al.* (2018) 'Mathematical models for the effects of hypertension and stress on kidney and their uncertainty', *Mathematical Biosciences*. Elsevier, 305, pp. 77–95. doi: 10.1016/j.mbs.2018.08.013.
- Wang, B. *et al.* (2020) 'Prediction model and assessment of probability of incident hypertension: the Rural Chinese Cohort Study', *Journal of Human Hypertension*. Springer US. doi: 10.1038/s41371-020-0314-8.
- WATSON, J. D. and CRICK, F. H. C. (1953) 'Molecular Structure of Nucleic Acids: A Structure for Deoxyribose Nucleic Acid', *Nature*. Springer Science and Business Media, 171(4356), pp. 737–738. doi: 10.1038/171737a0.
- Willer, C. J., Li, Y. and Abecasis, G. R. (2010) '{METAL}: fast and efficient meta-analysis of genomewide association scans', *Bioinformatics*. Oxford University Press (), 26(17), pp. 2190–2191. doi: 10.1093/bioinformatics/btq340.
- Wu, X. *et al.* (2020) 'Value of a Machine Learning Approach for Predicting Clinical Outcomes in Young Patients With Hypertension', *Hypertension*. Ovid Technologies (Wolters Kluwer Health), 75(5), pp. 1271–1278. doi: 10.1161/hypertensionaha.119.13404.
- Xie, J. and Szymanski, B. K. (2011) 'Community detection using a neighborhood strength driven label propagation algorithm', in *2011 IEEE Network Science Workshop*, pp. 188–195.
- Xu, J. *et al.* (2019) 'Translating cancer genomics into precision medicine with artificial intelligence: applications, challenges and future perspectives', *Human Genetics*. Springer Science and Business Media, 138(2), pp. 109–124. doi: 10.1007/s00439-019-01970-5.
- Yang, J. *et al.* (2010) 'Common {SNPs} explain a large proportion of the heritability for human height', *Nature Genetics*. Springer Science and Business Media, 42(7), pp. 565–569. doi: 10.1038/ng.608.
- Yang, J. *et al.* (2011) 'Genome partitioning of genetic variation for complex traits using common {SNPs}', *Nature Genetics*. Springer Science and Business Media, 43(6), pp. 519–525. doi: 10.1038/ng.823.
- Yang, J. *et al.* (2013) 'Genome-Wide Complex Trait Analysis ({GCTA}): Methods, Data Analyses, and Interpretations', in *Methods in Molecular Biology*. Humana Press, pp. 215–236. doi: 10.1007/978-1-62703-447-0_9.
- Yang, J. *et al.* (2017) 'Concepts, estimation and interpretation of {SNP}-based heritability', *Nature Genetics*. Springer Science and Business Media, 49(9), pp. 1304–1310. doi: 10.1038/ng.3941.
- Yin, B. *et al.* (2019) 'Using the structure of genome data in the design of deep neural networks for predicting amyotrophic lateral sclerosis from genotype', *Bioinformatics*. Oxford University Press (), 35(14), pp. i538–i547. doi: 10.1093/bioinformatics/btz369.
- Zhang, S. *et al.* (2018) 'Learning for Personalized Medicine: A Comprehensive Review from a Deep Learning Perspective', *IEEE Reviews in Biomedical Engineering*. IEEE, 12(August), pp. 194–208. doi: 10.1109/RBME.2018.2864254.
- Zhao, J. *et al.* (2018) 'Learning from Longitudinal Data in Electronic Health Record and Genetic Data to Improve Cardiovascular Event Prediction', *Scientific Reports*. Cold Spring Harbor Laboratory, 9(1), pp. 1–10. doi: 10.1101/366682.
- Zhu, D. *et al.* (2018) 'A machine learning approach for air quality prediction: Model regularization and optimization', *Big Data and Cognitive Computing*, 2(1), pp. 1–15. doi: 10.3390/bdcc2010005.
- Ziegler, A. *et al.* (2012) 'Personalized medicine using {DNA} biomarkers: a review', *Human Genetics*. Springer Science and Business Media, 131(10), pp. 1627–1638. doi: 10.1007/s00439-012-1188-9.
- Zou, J. *et al.* (2019) 'A primer on deep learning in genomics', *Nature Genetics*. Springer US, 51(1), pp. 12–18. doi: 10.1038/s41588-018-0295-5.

تیموری، م.، ابراهیمی، ا. و علوی نیا، س. م. (۱۳۹۴) 'مقایسه روش‌های مختلف یادگیری ماشین در تشخیص پرفشاری خون در بیماران دیابتی با و بدون در نظر گرفتن هزینه‌ها'، *مجله تخصصی اپیدمیولوژی ایران*.

دانشپور، م. ا. (1395). مطالعه گسترده ژنومی در طرح ژنتیک کاردیومتابولیک تهران بستر مناسبی جهت گسترش پزشکی فرادقیق در ایران، پژوهنده، ۱۱۲.

عزیزی، ف.، هدایتی، م. و دانشپور، م. ا. (۱۳۹۶) 'طراحی و تکمیل نخستین نسخه ژنوم مرجع ایرانیان، عدد درون ریز و متابولیسم / ایران، ۵.

لاجوردی، س. (1399). بهبود دقت پیش‌بینی خطر بیماری پرفشارخون بر مبنای تغییرات ژنتیکی با استفاده از روش‌های یادگیری ماشین in 'مفدهمین کنفرانس بین‌المللی مهندسی صنایع.

لاجوردی، س.، علوی، س. و کارگری، م. (۱۳۹۹) 'تحلیل ساختار علمی موضوع امتیاز خطر ژنتیکی با استفاده از روش تحلیل هم‌رخدادی واژگان مبتنی بر مقالات پایگاه in 'ScienceDirect ششمین کنفرانس بین‌المللی مهندسی صنایع و سیستم‌ها ICISE ۲۰۲۰.

واژه نامه فارسی به انگلیسی

Pooling	ادغام
Genome wide Rapid Association Using Mixed Model and Regression	ارتباط سریع ژنومی با استفاده از مدل ترکیبی و رگرسیون
Positive Prediction Value	ارزش پیش‌بینی مثبت
Negative Predictive Value	ارزش های پیشگویانه منفی
Missing	از دست رفته
Cross Validation	اعتبار سنجی متقابل
Risk Scroe	امتیاز خطر
Polygenic Risk Score	امتیاز خطر چند ژنی
Effect size	اندازه اثر
Resequencing	باز توالی
Binding	بخش بندی مقید
Lossless	بدون خطا
Dynamic Programming	برنامه نویسی پویا
Iteratively	بصورت مکرر
Overfitting	بیش از حد برازش
Overfit	بیش برازش
Hypertension	پرفشاری خون
Human Genome Project	پروژه ژنوم انسانی
Iranian Reference Genome Project	پروژه ژنوم مرجع ایران
Electronic Health Record	پرونده الکترونیک سلامت
Personalized Medicine	پزشکی شخصی
back propagation	پس انتشار
Single Nucleotide Polymorphism	چند ریختی تک نوکلئوتیدی
Restriction fragment length polymorphisms	چند ریختی طول قطعه محدود
Wrapper	پوشانه
Convolution	پیچشی
Genetic linkage	پیوند ژنتیکی
loss function	تابع خطا
Advanced Complex Trait Analysis	تجزیه و تحلیل صفات پیچیده پیشرفته
Genome-wide Complex Trait Analysis	تجزیه و تحلیل صفات پیچیده در سطح ژنوم
omic-data-based complex trait analysis	تجزیه و تحلیل صفات پیچیده مبتنی بر داده‌های امیک
Principal Component Analysis	تجزیه و تحلیل مؤلفه های اصلی
Co-word analysis	تحلیل هم واژه
Sequence Alignment	تراز توالی
Pairwise Sequence Alignment	تراز توالی دوبه‌دو
Chip	تراشه
relief	تسکین
ascertainment	تشخیص
Community detection	تشخیص جوامع

Variable number of tandem repeats	تعداد متغیر تکرارهای پشت سر هم
Copy number variation	تغییرات تعداد کپی
short tandem repeats	تکرارهای پشت سر هم کوتاه
Variety	تنوع
Single nucleotide variation	تنوع تک نوکلئوتیدی
Whole genome sequencing	توالی یابی کل ژنوم
feed forward	جلوبرنده
Founder populations	جمعیت های بنیانگذار
hesitation margin	حاشیه تردید
Restricted Maximum Likelihood	حداکثر احتمال محدود
Deletion	حذف
indel	حذف - اضافه
Sensitivity	حساسیت
Relative risk	خطر نسبی
Insertion	درج
Decision Tree	درخت تصمیم
Accuracy	صحت
bimodality	دو وجهی
Phenotype	رخ نمود
Longitudinal Phenotype	رخ نمود طولی
Classification	دسته بندی
Auto encoder	رمزگذار خودکار
Genotype	ژن نمود
historical control genotypes	ژن نمودهای کنترل تاریخی
Velocity	سرعت
Body Mass Index	شاخص توده بدن
gini index	شاخص جینی
Residual network	شبکه باقی مانده
Convolutional Neural Network	شبکه عصبی پیچشی
radial basis function neural network	شبکه عصبی تابع پایه شعاعی
Feed forward neural network	شبکه عصبی جلوبرنده
Kohonen self organizing neural network	شبکه عصبی خود سازماندهی کوهونن
Deep Neural Network	شبکه عصبی عمیق
Recurrent neural network	شبکه عصبی گردشی
Fully Connected Network	شبکه کاملاً متصل
Social networks	شبکه های اجتماعی
Adversarial networks	شبکه های متخاصم
Individual	شخصی
Precision	دقت
Complex Trait	صفت پیچیده
Information gain	بهره اطلاعات
Impact Factor	ضرب تأثیر
Meta-Analysis	فرا تحلیل

Recall	بازخوانی
Blood Pressure	فشارخون
Diastolic Blood Pressure	فشارخون دیاستولیک
Systolic Blood Pressure	فشارخون سیستولیک
propensity	گرایش
Genetic Relation Matrix	ماتریس رابطه ژنتیکی
Confusion Matrix	ماتریس مختلط
Support Vector Machine	ماشین بردار پشتیبانی
Mini satellites	ماهواره های کوچک
Positive	مثبت / صحیح
Data Set	دادگان
fuzzy set	مجموعه فازی
Mixed Linear Model	مدل خطی مختلط
Fuzzy Boundry	مرز فازی
Phenome-wide association study	مطالعه ارتباط گسترده پدیده
Genome Wide Association Study	مطالعه ارتباط گسترده ژنوم
Missed Value	مقادیر از دست رفته
Regularization	قاعده مند سازی
Negative	منفی / غلط
Area Under Curve	ناحیه زیر منحنی
True Positive Rate	نرخ مثبت واقعی
True Negative Rate	نرخ منفی واقعی
gain ratio	نسبت بهره
Biomarker	نشانگر زیستی
Evidence Theory	نظریه شواهد
Linkage Mapping	نقشه پیوند
Random Sub-sampling	نمونه گیری تصادفی فرعی
Polymerase chain reactions	واکنش های زنجیره ای پلیمرز
Inheritance	وراثت
Missing Heritability	وراثت پذیری گمشده
Regional heritability advanced complex trait analysis	وراثت پذیری منطقه ای تجزیه و تحلیل صفات پیچیده پیشرفته
Specificity	ویژگی
Receiver Operating Characteristic	ویژگی عملکرد گیرنده
Multiple Sequence Alignment	هم ترازای چندگانه
Deep learning	یادگیری عمیق
Genome Deep Learning	یادگیری عمیق ژنوم
Machine Learning	یادگیری ماشین

واژه نامه انگلیسی به فارسی

Accuracy	صحت
Advanced Complex Trait Analysis	تجزیه و تحلیل صفات پیچیده پیشرفته
Adversarial networks	شبکه‌های متخاصم
angiotensinogen	آنژیوتانسینوژن
Area Under Curve	ناحیه زیر منحنی
ascertainment	تشخیص
Auto encoder	رمزگذار خودکار
back propagation	پس انتشار
bimodality	دو وجهی
Binding	بخش بندی مقید
Biomarker	نشانگر زیستی
Blood Pressure	فشارخون
Body Mass Index	شاخص توده بدن
Chip	تراشه
Classification	دسته‌بندی
Community detection	تشخیص جوامع
Complex Trait	صفت پیچیده
Confusion Matrix	ماتریس مختلط
Convolution	پیچش
Convolutional Neural Network	شبکه عصبی پیچیده
Copy number variation	تغییرات تعداد کپی
Co-word analysis	تحلیل هم واژه
Cross Validation	اعتبار سنجی متقابل
Data Set	دادگان
Decision Tree	درخت تصمیم
Deep learning	یادگیری عمیق
Deep Neural Network	شبکه عصبی عمیق
Deletion	حذف
Diastolic Blood Pressure	فشارخون دیاستولیک
Dynamic Programming	برنامه نویسی پویا
Effect size	اندازه اثر
Electronic Health Record	پرونده الکترونیک سلامت
Evidence Theory	نظریه شواهد
feed forward	جلوبرنده
Feed forward neural network	شبکه عصبی جلوبرنده
Founder populations	جمعیت های بنیانگذار
Fully Connected Network	شبکه کاملاً متصل
Fuzzy Boundry	مرز فازی
fuzzy set	مجموعه فازی
gain ratio	نسبت بهره
Genetic linkage	پیوند ژنتیکی

Genetic Relation Matrix	ماتریس رابطه ژنتیکی
Genome Deep Learning	یادگیری عمیق ژنوم
Genome Wide Association Study	مطالعه ارتباط گسترده ژنوم
Genome wide Rapid Association Using Mixed Model and Regression	ارتباط سریع ژنومی با استفاده از مدل ترکیبی و رگرسیون
Genome-wide Complex Trait Analysis	تجزیه و تحلیل صفات پیچیده در سطح ژنوم
Genotype	ژننمود
gini index	شاخص جینی
hesitation margin	حاشیه تردید
historical control genotypes	ژننمودهای کنترل تاریخی
Human Genome Project	پروژه ژنوم انسانی
Hypertension	پرفشاری خون
Impact Factor	ضریب تأثیر
indel	حذف - اضافه
Individual	شخصی
Information gain	بهره اطلاعات
Inheritance	وراثت
Insertion	درج
Iranian Reference Genome Project	پروژه ژنوم مرجع ایران
Iteratively	بصورت مکرر
Kohonen self organizing neural network	شبکه عصبی خود سازماندهی کوهونن
Linkage Mapping	نقشه پیوند
Longitudinal Phenotype	رخنمود طولی
loss function	تابع خطا
Lossless	بدون خطا
Machine Learning	یادگیری ماشین
Meta-Analysis	فراتحلیل
Mini satellites	ماهواره های کوچک
Missed Value	مقادیر از دست رفته
Missing	از دست رفته
Missing Heritability	وراثت پذیری گمشده
Mixed Linear Model	مدل خطی مختلط
Multiple Sequence Alignment	هم ترازای چندگانه
Negative	منفی / غلط
Negative Predictive Value	ارزش های پیشگویانه منفی
omic-data-based complex trait analysis	تجزیه و تحلیل صفات پیچیده مبتنی بر داده های امیک
Overfit	بیش برازش
Overfitting	بیش از حد برازش
Pairwise Sequence Alignment	تراز توالی دوبه دو
Personalized Medicine	پزشکی شخصی
Phenome-wide association study	مطالعه ارتباط گسترده پدیده
Phenotype	رخنمود

Polygenic Risk Score	امتیاز خطر چند ژنی
Polymerase chain reactions	واکنش های زنجیره ای پلیمرز
Pooling	ادغام
Positive	مثبت / صحیح
Positive Prediction Value	ارزش پیش بینی مثبت
Precision	دقت
Principal Component Analysis	تجزیه و تحلیل مؤلفه های اصلی
propensity	گرایش
radial basis function neural network	شبکه عصبی تابع پایه شعاعی
Random Sub-sampling	نمونه گیری تصادفی فرعی
Recall	بازخوانی
Receiver Operating Characteristic	ویژگی عملکرد گیرنده
Recurrent neural network	شبکه عصبی گردشی
Regional heritability advanced complex trait analysis	وراثت پذیری منطقه ای تجزیه و تحلیل صفات پیچیده پیشرفته
Regularization	قاعده مند سازی
Relative risk	خطر نسبی
relief	تسکین
Resequencing	باز توالی
Residual network	شبکه باقی مانده
Restricted Maximum Likelihood	حداکثر احتمال محدود
Restriction fragment length polymorphisms	پلی مورفیسم های طول قطعه محدود
Risk Score	امتیاز خطر
Sensitivity	حساسیت
Sequence Alignment	تراز توالی
short tandem repeats	تکرارهای پشت سر هم کوتاه
Single Nucleotide Polymorphism	پلی مورفیسم تک نوکلئوتیدی
Single nucleotide variation	تنوع تک نوکلئوتیدی
Social networks	شبکه های اجتماعی
Specificity	ویژگی
Support Vector Machine	ماشین بردار پشتیبانی
Systolic Blood Pressure	فشار خون سیستولیک
True Negative Rate	نرخ منفی واقعی
True Positive Rate	نرخ مثبت واقعی
Variable number of tandem repeats	تعداد متغیر تکرارهای پشت سر هم
Variety	تنوع
Velocity	سرعت
Whole genome sequencing	توالی یابی کل ژنوم
Wrapper	پوشانه

Abstract

Hypertension is one of the most significant underlying ailments of cardiovascular disease; hence, methods that can accurately reveal the risk of hypertension at an early age are essential. Also, one of the most critical personal health objectives is to improve disease prediction accuracy by examining genetic variants.

Therefore, various clinical and genetically based methods are used to predict the disease; however, the critical issue with these methods is the high number of input variables as genetic markers with small samples. One approach that can be used to solve this problem is machine learning.

This study was conducted on the participants' genetic markers in the 20-year research of cardiometabolic genetics in Tehran (TCGS). Various machine learning methods were used, including linear regression, neural network, random forest, decision tree, and support vector machine. The top ten genetic markers were identified using importance-based ranking methods, including information gain, gain ratio, Gini index, χ^2 , relief, and FCBF. A model based on a neural network with AUC of 89% was presented. This model has an accuracy and an f-measure of 0.89, which shows the quality. The final results indicate the success of the machine learning approach.

The initial processing method and data modeling have innovation, especially in the deep learning network. The research results in the two machine learning and deep learning models have AUC of 0.87 and 0.89, respectively. This modeling provides the possibility of prognosis at young ages with appropriate accuracy.

Keywords: Deep Learning, Genetic Markers, Hypertension, Disease Risk.



**Department of Information Technology
Faculty of Industrial and Systems Engineering**

**Dissertation Submitted in Partial Fulfillment of the
Requirements for the Degree of Doctor of Philosophy (Ph.D)**

**Design and presenting an intelligent
model for predicting SNP-based
hypertension in uncertainty**

**By:
Seyed Ali Lajevardy**

**Supervisor:
Dr. Kargari**

**Co-Supervisor:
Dr. Daneshpour**

**Advisor:
Dr. Akbarzadeh**

June 2023