

پیش بینی صحیح دُز " انسولین و گلیبن کلامید " در بیماران دیابتی مبتنی بر ترکیب سیستم‌های هوشمند و سابقه بیماری

محمد فیوضی^{۱،*}، جواد حدادینیا^۱، نسرين ملانیا^۲، مریم هاشمیان^۳، کاظم حسن پور^۴

چکیده

مقدمه: یکی از عوارض خطرناک بیماری دیابت نوع ۱، افزایش و یا کاهش ناگهانی سطح غلظت قند خون می‌باشد که باعث بروز خطراتی مانند اغما و بیهوشی خواهد شد. انسولین و گلیبن کلامید، دو دارویی هستند که در اثر تجویز دُز صحیح، سطح غلظت قند خون را در نهایت، به درستی تنظیم می‌کنند تا از بروز چنین عوارضی پیشگیری شود. بنابراین استفاده از روشی مناسب به منظور پیش‌بینی یا تجویز صحیح این داروها و در نهایت پیشگیری از این عوارض، گام مهمی در جهت کنترل بهینه بیماری محسوب می‌شود. در این تحقیق ما بر آن هستیم، تا با استفاده از ترکیب الگوریتم‌های داده‌کاوی و هوش مصنوعی، پیش‌بینی صحیحی از دُز انسولین و گلیبن کلامید برای بیماران انجام دهیم.

روش‌ها: مراحل ایجاد سیستمی هوشمند به منظور تعیین و پیش‌بینی میزان صحیح انسولین و گلیبن کلامید بدین صورت است؛ که ابتدا لازم است تا شرایط و پارامترهایی که در میزان و دُز دارو (انسولین و گلیبن کلامید) برای بیماران دخیل هستند را شناسایی نماییم، سپس بانک جامعی مبتنی بر این ویژگی‌ها تشکیل دهیم، با همکاری تیم پزشکی مرکز تحقیقات دیابت شهرستان سبزوار، بانک جامعی در ۳ مرحله براساس اطلاعات ۱۱۰ بیمار و ۲۵۸ مورد مشکوک آن مرکز طراحی شد، سپس طبق قاعده سیستم‌های هوشمند و شبکه‌های عصبی مصنوعی، مدلی به منظور پیش‌بینی میزان دُز انسولین و گلیبن کلامید برای بیماران طراحی شد.

یافته‌ها: سیستم پیشنهادی با استفاده از ترکیب روش‌های مذکور موفق شد با تکیه بر ویژگی‌های پایگاه داده در قالب ترکیب و تعامل به دقت پیش‌بینی بیش از ۹۵٪ دست یابد. در مقایسه با روش‌های رایج از یکطرف و روش‌های مصنوعی از طرف دیگر (غالباً در بهترین حالت دقت و صحت قدرت پیش‌بینی آن‌ها در حدود ۸۵٪ برای مدت زمان کمتر از ۳ ساعت بوده است) به خوبی مشخص شد که، عملکرد سیستم پیشنهادی، مناسب‌تر، سریع‌تر و تا حدودی مطمئن‌تر از سایر روش‌های ترکیبی هوشمند است.

نتیجه‌گیری: این تحقیق با ارائه یک روش و مدل برگرفته از واقعیت و شرایط بیماران؛ میزان صحیح و تجویز شده دارو برای آنها را مشخص می‌کند، این در حالی است که نسبت به روند طبیعی تجویز توسط پزشکان، علاوه بر افزایش سرعت در تعیین میزان دارو، دارای صحت و مقبولیت و دقت قابل قبولی می‌باشد.

واژگان کلیدی: الگوریتم ژنتیکی، دیابت، دُز صحیح انسولین و گلیبن کلامید، شبکه‌های هوشمند مصنوعی

۱- بخش مهندسی پزشکی، دانشکده مهندسی برق و کامپیوتر، دانشگاه حکیم سبزواری، سبزوار، ایران

۲- مرکز تحقیقات فناوری‌های نوین پزشکی، دانشگاه علوم پزشکی سبزوار، سبزوار، ایران

۳- گروه زیست شناسی، دانشکده علوم پایه، دانشگاه حکیم سبزواری، سبزوار، ایران

۴- گروه علوم بالینی، دانشکده پزشکی، دانشگاه علوم پزشکی سبزوار، سبزوار، ایران

***فشنانی:** خراسان رضوی، شهرستان سبزوار، توحیدشهر، دانشگاه حکیم سبزواری، دانشکده مهندسی برق و کامپیوتر، بخش مهندسی

پزشکی، کدپستی: ۹۶۱۷۹۷۶۴۸۷، تلفن: ۰۵۷۱۴۴۱۰۱۰۴، نمابر: ۰۵۷۱۴۴۱۰۱۰۴، پست الکترونیک: mohammad.fiuzy@yahoo.com

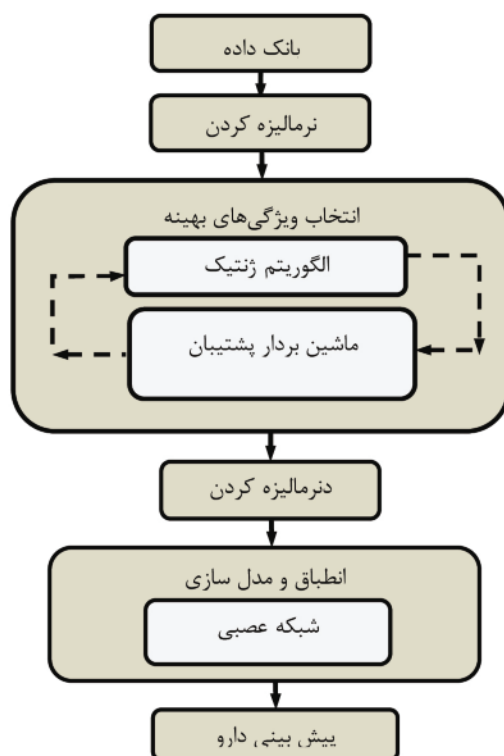
مقدمه

بیماری دیابت یا مرض قند یکی از شایع ترین بیماری های شناخته شده در دنیا می باشد. براساس گزارش ها حدود ۲۲۰ میلیون نفر در جهان از این بیماری رنج می برند [۱-۳]. دیابت عمدتاً از طریق آزمایش های قند خون (Plasma یا Glucose یا Fasting) تشخیص داده می شود که به ۳ دسته تقسیم می شود. دیابت نوع ۱ یا دیابت وابسته به انسولین (IDDM) بیشتر در سن پایین و کودکان دیده می شود. دیابت نوع ۲ یا دیابت غیر وابسته به انسولین (NIDDM) در ۹۰٪ تا ۹۵٪ بیماران دیابتی مشاهده می شود [۳]. در این دسته انسولین ترشح می شود، ولی بدن در برابر مصرف انسولین خود مقاومت نشان می دهد. دیابت نوع ۳ بیشتر در زنان باردار دیده می شود که عمدتاً پس از بارداری به دیابت نوع ۲ تبدیل می شود [۴]. البته در زمینه تشخیص بیماری دیابت توسط همین مولفین تحقیقاتی [۵،۶] انجام شده است، اما در زمینه پیش بینی و تعیین صحیح دژ انسولین، تحقیقات چندانی انجام نشده است (تا آنجا که مولفین می دانند)، مشکل عمده ای که در حال حاضر در رابطه با این بیماری مخرب و خطرناک و به خصوص در نوع ۱ آن وجود دارد، افزایش سطح قند خون (Hyperglycemia) و یا کاهش بیش از حد و ناگهانی آن است که از جمله عوارض بسیار خطرناک آن محسوب می گردد، که می تواند به بیهوشی، اغماء و گاهی حتی مرگ بیمار منجر شود [۷]. بنابراین استفاده از روشی مناسب به منظور پیش بینی و تعیین میزان صحیح دز انسولین و حتی داروهایی دیگر همچون "متفورمین"، "گلیبن کلامید" و در جهت پیشگیری از این عوارض، گام مهمی برای کنترل بهتر این بیماری محسوب می شود.

با توجه به مطالب بیان شده و بررسی کارهای گذشته لزوم به نتیجه رسیدن این گونه پژوهش ها، بسیار مهم و سرنوشت ساز خواهد بود. از آنجا که در تحقیقات علوم پزشکی و اپیدمیولوژی اغلب مسئله سلامت انسان مطرح است، پیش بینی درست نتایج اهمیت بیشتری می یابد، بنابراین لازم است تا جایی که امکان دارد کمترین خطا و بیشترین اطمینان حاصل شود. با توجه به این که مسائل مرتبط با پاسخ های آمیخته و به وفور در مطالعات پزشکی مشاهده می شود و با در نظر گرفتن این که روش های موجود در آمار کلاسیک برای مدل بندی و پیش بینی، به دلیل محدودیت در عمل کارایی چندانی ندارد، ارائه روش هایی که بتواند راه گشای این گونه مسائل باشد، بسیار مفید و ارزنده به نظر می رسد. در این تحقیق برآن هستیم تا با ارائه مدلی دقیق و البته با کمترین خطا، میزان صحیح دژ دارو برای بیماران را بدست آوریم. ابتدا روش کلی را بیان می کنیم سپس با توجه به مطالب فلوچارت ۱ به بحث و بررسی رئوس مطالب می پردازیم. امید است تا این تحقیق گامی هرچند اندک ولی موثر در جهت بهبود روش ها و معالجه این بیماری داشته باشد.

روش ها

پس از مطالعه مختصر روی سیستم ها و تحقیقات انجام شده، سیستم پیشنهادی به گونه ایی که در فلوچارت ۱ مشخص شده است، طراحی و معرفی شد. بانک داده ایی از بیماران تشکیل شد، سپس براساس داده ها، که گونه ایی دنباله یا سری ایی در واحد زمان با مقادیر متفاوت هستند، توسط شبکه های عصبی مصنوعی، پیش بینی براساس مدل تشکیل شده صورت گرفت.



فلو چارت ۱- روند الگوریتم پیشنهادی به منظور تعیین دُز صحیح داده‌های مورد استفاده

الف: بانک داده

استفاده شد. در جدول ۱ نمونه‌ای از ویژگی‌ها به همراه نحوه کدگذاری پیشنهادی مشخص شده است.

جدول ۱- اطلاعات و داده‌های بانک اطلاعاتی

ویژگی‌ها	کد	[کیمینه / بیشینه]
جنسیت	۱+ یا ۱-	[۱+ / ۱-]
سن	مقدار مربوطه	[۱۲۰ / ۰]
تاهل	۱ یا ۰	[۱+ / ۱]
وزن	مقدار مربوطه	[۲۰۰ / ۰]
قد	مقدار مربوطه	[۲۰۰ / ۱۴۰]
تاریخچه بیماری دیابت	۵ یا ۲ یا ۱ یا ۰	[۵ / ۰]
فشار خون دیاستول	مقدار مربوطه	[۱۰۰ / ۱۰]
فشار خون سیستول	مقدار مربوطه	[۱۸۰ / ۴۰]
قند خون سریع	مقدار مربوطه	[۵۰۰ / ۵۰]
قند خون	مقدار مربوطه	[۵۰۰ / ۵۰]
تری گلیسرید	مقدار مربوطه	[۵۰۰ / ۵۰]
کلسترول	مقدار مربوطه	[۴۰۰ / ۱۰۰]
شاخص جرم توده	مقدار مربوطه	[۵۰ / ۱۴]
نبض	مقدار مربوطه	[۲۰۰ / ۲۰]
نوع تشخیص	۱ یا ۲ یا ۳	[۳ / ۱]
نوع درمان	۱ یا ۲ یا ۳	[۳ / ۱]
سابقه بیماری	۱ یا ۲ یا ۱۱	[۱۱ / ۱]
سابقه دارو	۱+ یا ۱-	[۱+ / ۱-]

ادامه جدول ۱ در صفحه بعد

به منظور تعیین و پیش‌بینی دُز دارو برای بیماران، از مرکز تحقیقات دیابت شهرستان سبزوار، وابسته به دانشگاه علوم پزشکی شهرستان سبزوار یاری گرفته شد. به این ترتیب که با هماهنگی و همکاری تیم پزشکی و متخصص غدد، پرسش‌نامه‌هایی (به منظور مشخص کردن ویژگی‌های موثر) طراحی گردید. این پرسش‌نامه‌ها توسط ۳۶۸ فرد (۱۱۰ بیمار به همراه ۲۵۸ مورد مشکوک)، که به صورت کاملاً تصادفی (سیستماتیک) انتخاب شده بودند، در طی ۳ مرحله تکمیل شد (از هر فرد ۳ بار نمونه‌گیری به عمل آمد). براساس نظر پزشکان سعی شد تا ویژگی‌هایی تعریف شوند که در روند درمان، مدل‌سازی رفتار بیمار و تصمیمات پزشکان بیشترین تاثیر را داشته باشند. بعد از مشخص شدن ویژگی‌ها، کدگذاری آن‌ها مهم است، که در پاره‌ای از موارد مقدار دقیق عددی آزمایش را قرار می‌دهیم و در قسمت‌هایی دیگر لازم است تا مقدار کد شده قرار داده شود. به منظور کدگذاری از نظر تیم پزشکی براساس راهنمای‌های بیماران،

را Wrapper یا حلقه بسته می‌گویند. این روش، جستجو در فضای زیر مجموعه‌ها را براساس تخمین دقت ناشی از انتخاب یک زیر مجموعه خاص تحت شرایط الگوریتم طبقه‌بندی مورد استفاده (به عنوان معیاری از بهینگی آن زیر مجموعه) انجام می‌شود [۱۲]. الگوریتم‌های دسته دوم معمولاً نتایج بهتری به دست می‌آورند. مهمترین بخش هر روش انتخاب ویژگی حلقه بسته، الگوریتم جستجویی است، که در آن به کار رفته است. طی دهه گذشته محققان روی الگوریتم‌های جستجوی تکاملی مثل الگوریتم ژنتیک [۱۳، ۱۴]، ACO^۳، [۱۲، ۱۵، ۱۶] و ترکیب ACO/PSO^۴ [۱۷] تمرکز داشته‌اند.

ج: فرآیند اجرای مراحل در روش پیشنهادی

در الگوریتم پیشنهادی ابتدا داده‌ها در بانک اطلاعاتی تهیه شده قرار گرفتند و سپس توسط رابطه (۱) نرمالیزه شدند.

$$y_i = \frac{x_i - m}{\sigma_i} \quad (1)$$

در این رابطه x_i اعداد و یا داده‌ها؛ m میانگین داده‌های ویژگی i ام و σ_i واریانس داده‌های ویژگی i ام (در هر ویژگی) هستند. توسط الگوریتم ژنتیک باینری موثرترین ویژگی‌هایی که در الگوریتم طبقه‌بندی کننده ماشین بردار پشتیبان (SVM) باعث بهترین و بیشترین طبقه‌بندی می‌شوند، انتخاب شدند و در نهایت براساس بهترین ویژگی‌هایی که انتخاب شدند، عمل مدل‌سازی و انطباق فضای ورودی (ویژگی‌های بیماران) با فضای خروجی (تصمیمات پزشک مبنی بر تجویز دارو) توسط شبکه‌های عصبی مصنوعی صورت گرفت.

د: الگوریتم ژنتیکی دودویی

در این تحقیق از یک نوع الگوریتم ژنتیک ساده توام با الگوریتم طبقه‌بندی ماشین بردار پشتیبان (SVM) برای انتخاب پارامترهای (ویژگی‌های) بهینه استفاده شد. فلوجارت ۲، نمودار بلوکی سیستم مورد استفاده را نمایش می‌دهد؛ البته جایگاه قرارگیری مجموعه پیشنهادی در تابع هزینه الگوریتم ژنتیک می‌باشد. ابتدا یک بردار دودویی به

ادامه جدول ۱		
کلیست‌رول خوب	مقدار مربوطه	[۱۰ ۴۰۰]
کلیست‌رول بد	مقدار مربوطه	[۱۰ ۲۰۰]
اسید اوریک	مقدار مربوطه	[۰ ۱۰۰]
کراتین	مقدار مربوطه	[۰ ۱۰]
هموگلوبین A1c	مقدار مربوطه	[۵ ۲۰]

ب: اهمیت انتخاب ویژگی و بازشناسی الگو

ابتدا باید به مساله انتخاب ویژگی‌های مناسب در فضای تبدیل پرداخت. انتخاب ویژگی در واقع برگزیدن ویژگی‌هایی است که حداکثر توان را در پیشگویی خروجی دارا باشند [۸]. تعریف اینکه زیر مجموعه بهینه چه می‌تواند باشد، به مسئله‌هایی که قرار است حل شوند وابسته است [۹]. در روش‌هایی که بحث شده‌اند مقادیر ویژه بردارها اهمیت آنها را نشان می‌دهد. اهمیت دادن به تمام ویژگی‌ها بر این اساس همیشه مناسب نیست و علاوه بر این بردارهای ویژگی نیز همیشه شامل اطلاعات مفید نیستند. انتظار می‌رود با داشتن اطلاعات بیشتر، نتایج بهتر شود؛ اما در عمل مشاهده می‌شود که این مساله باعث افت دقت می‌شود [۱۰، ۱۱]. بنابراین برای رسیدن به یک بازنمایی مناسب، انتخاب بهترین بردار ویژگی یا ساده سازی بردار ویژگی امری ضروری است. یک سیستم شناسایی الگوی متداول شامل ۴ بخش استخراج ویژگی، انتخاب ویژگی، طراحی و آموزش طبقه‌بندی کننده و سرانجام آزمایش است. در این تحقیق از الگوریتم ژنتیک دودویی^۱ (GA) برای کاوش و انتخاب ویژگی‌های موثر، الگوریتم ماشین بردار پشتیبان (SVM)^۲ برای طبقه‌بندی نمونه‌ها و همین طور از شبکه‌های عصبی مصنوعی برای طراحی طبقه‌بندی کننده و آزمایش نهایی استفاده شده است. الگوریتم‌های انتخاب ویژگی بسته به روند ارزیابی آنها به دو دسته عمده تقسیم می‌شوند. اگر انتخاب ویژگی مستقل از هرگونه الگوریتم یادگیری انجام شود (یعنی به صورت یک پیش پردازنده کاملاً مجزا)، آن روش، فیلتر یا حلقه باز می‌باشد. در این مورد ویژگی‌های نامطلوب قبل از استنتاج دور ریخته می‌شوند. اما، اگر روند ارزیابی با یک الگوریتم طبقه‌بندی در ارتباط باشد، روش انتخاب ویژگی

3. Ant Colony Optimization

4. Ant Colony Optimization/Particle Swarm Optimization

1. Binary Genetic Algorithm

2. Support Vector Machine

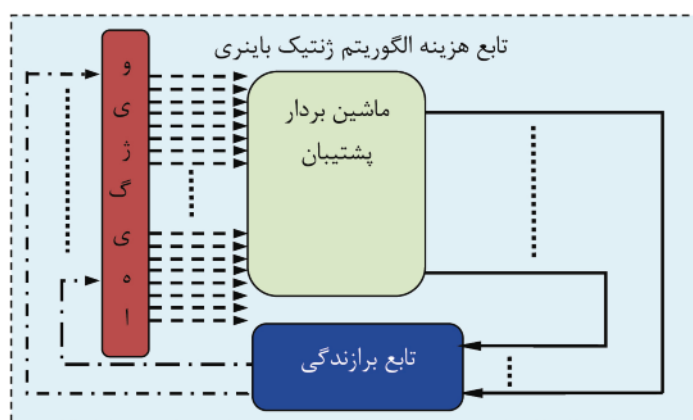
۵۰۰ نسل تنظیم شدند. سایر پارامترهای الگوریتم ژنتیک نیز مطابق جدول ۲، براساس اقتباس از [۲۰-۱۸] هستند.

جدول ۲- پارامترهای الگوریتم ژنتیک پیشنهادی

تعداد نسل ها	۵۰۰
اندازه جمعیت	۲۰
احتمال تقاطع	۰٫۸
احتمال جهش	۰٫۱
نوع باز ترکیب	دو نقطه ای
نوع انتخاب	چرخ رولت

در این تحقیق به عنوان تابع برازندگی الگوریتم ژنتیک، از حاصل ضرب حساسیت طبقه بندی در ویژگی طبقه بندی ماشین بردار پشتیبان (SVM) استفاده شدند، و متغیرهای ورودی این شبکه یک رشته دودویی با طول ۶۹ بیت (واحد) می باشد که، هریک از بیت های این رشته معرف انتخاب یا عدم انتخاب یک پارامتر تشخیصی می باشد.

طول ۶۹ بیت (هر بیمار ۳ بار داده گیری ۲۳ نمونه ای) واحد (به تعداد پارامترها یا ویژگی ها) بصورت تصادفی انتخاب می شود که هر کدام از بیت های این رشته دودویی متناظر با یک پارامتر در ماتریس آموزش شبکه است. اگر بیت متناظر با هر پارامتر صفر باشد، آن پارامتر حذف شده و اگر یک باشد، آن پارامتر حذف نشده و در آموزش دخالت داده می شود. پس از اجرای یک نسل از الگوریتم ژنتیک با پارامترهای تصادفی مقدار برازندگی هر رشته دودویی (کروموزوم) برای ۳۰ نمونه تست محاسبه شد و از بهترین کروموزوم های نسل قبل، ۲ کروموزوم به همراه ۱۸ کروموزوم جدید که با استفاده از عملگرهای ژنتیک تولید شدند، نسل جدید (دوم) الگوریتم ژنتیک تولید می شود. این چرخه اینقدر ادامه پیدا می کند، تا به دقت مورد تنظیم دست یابیم و یا تعداد نسل ها به تعداد حداکثر تنظیم شده رسیده باشد. در اجرای الگوریتم ژنتیک حداکثر تعداد نسل ها را



فلوچارت ۲- تابع هزینه الگوریتم ژنتیک

هزینه ماشین بردار پشتیبان

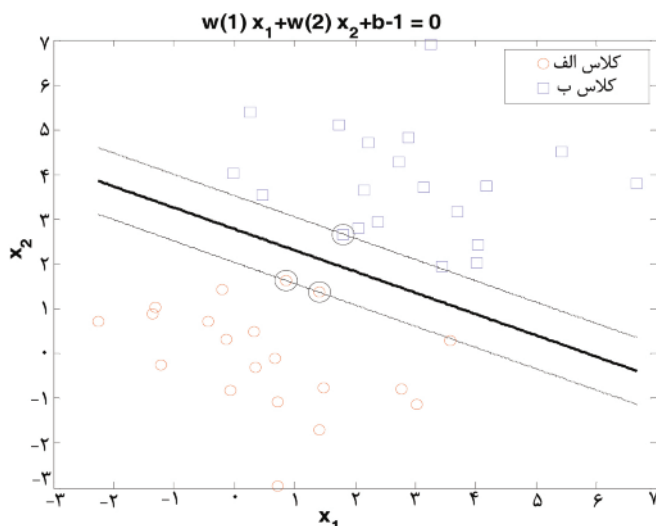
SVM یک الگوریتم یادگیری با سرپرستی است که می تواند در کاربردهایی نظیر جداسازی و طبقه بندی داده مورد استفاده قرار گیرد [۲۱]. این الگوریتم می تواند داده های آموزش را به دو دسته تقسیم کند. اصولاً داده ها توسط یک خط قابل جداسازی هستند، و در غیر این صورت الگوریتم با نگاشت ویژگی ها به وسیله تابع ϕ به ابعاد بالاتر سعی می کند توسط یک فرا صفحه با دقت بالاتری داده ها را تفکیک نماید (شکل ۱).

علت استفاده از تابع هزینه (رابطه ۱) در الگوریتم ژنتیک این است که در این حالت مقدار حساسیت (رابطه ۲) و ویژگی (رابطه ۳) را همزمان افزایش می یابند. وقتی مقدار حساسیت و ویژگی تشخیصی هر دو افزایش یابند، مقدار دقت و سطح زیر منحنی ROC^۱ (منحنی جهت مقایسه عملکرد و مقایسه ۲ سیستم) نیز افزایش می یابند.

$$Fitness = Se \times Sp \quad (1)$$

$$Se = \frac{Tp}{Tp + Fn} \quad (2)$$

$$Sp = \frac{Tn}{Tn + Fp} \quad (3)$$



شکل ۱- جداسازی داده‌ها در ۲ کلاس توسط الگوریتم SVM

ژنتیک را دارد، هدف مشخص کردن ویژگی‌هایی است که بیشترین تاثیر در تمایز بین ویژگی‌ها به منظور پیش‌بینی را داشته باشند. به منظور تسریع و کاهش محاسبات در این کار از تابع خطی SVM به منظور جداسازی ویژگی‌های موثر از غیر موثر استفاده شده است. تا این مرحله بهترین نتایج یعنی موثرترین ویژگی‌هایی که در امر پیش‌بینی نقش دارند مشخص می‌شوند، یعنی ویژگی‌های شماره (۵، ۶، ۱۱، ۱۲، ۱۵، ۱۶، ۱۷، ۱۸، ۱۹، ۲۰ و ۲۲) یا همان ویژگی‌هایی که در (تصویر ۲) مشخص شده است.

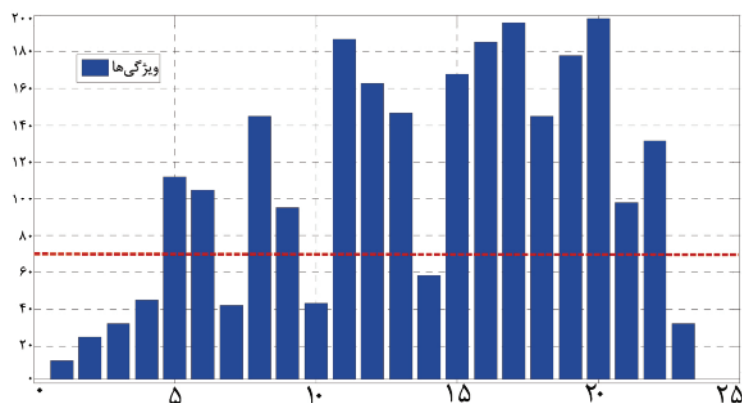
معادله (۴) رابطه فراصفحه طبقه‌بندی کننده داده‌ها است. معادلات (۵ و ۶) روابط فراصفحه‌های موازی براساس شرایط حداکثر حاشیه می‌باشند. در صورت ترسیم این توابع ساده فاصله بین دو صفحه حاشیه برابر می‌باشد.

$$w \times x - b = 0 \quad (4)$$

$$w \times x - b = 1 \quad (5)$$

$$w \times x - b = -1 \quad (6)$$

در این روابط x متغیر ورودی، w بردار نرمال خط جدا کننده و b عرض از مبدا خط جدا کننده است. SVM در اینجا نقش کلاس‌بندی نمونه‌های خروجی از الگوریتم

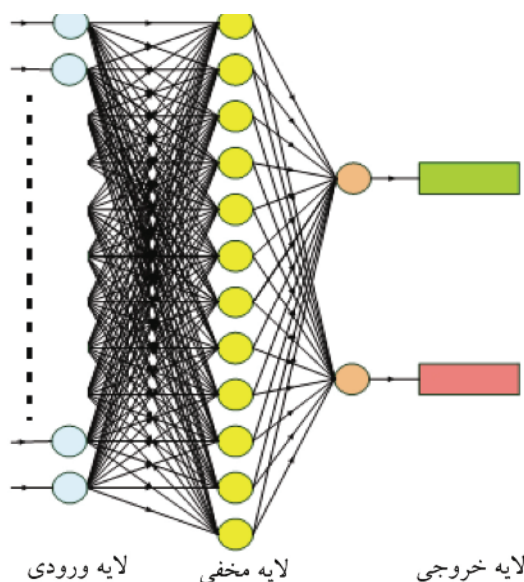


شکل ۲- بهترین ویژگی‌های انتخاب شده توسط الگوریتم ژنتیک و طبقه‌بندی کننده SVM

مشخص شد. در مرحله بعدی باید تنها ویژگی‌های مشخص شده از افراد را براساس اصل سری‌های زمانی (در

در اثر فقط ۲۰۰ تکرار، تمایزی فاحش بین ویژگی‌های مهم و سرنوشت‌ساز در تعیین میزان دژ دارو برای بیماران

است. در این روش خطا به روش خاص، لایه به لایه، به عقب برگشت داده می‌شود و در هر لایه اصلاحات لازم روی وزن‌ها انجام می‌گیرد. این روند آنقدر ادامه می‌یابد تا خطای خروجی به سمت مقدار کمینه‌ای به نام Error Goal (خطای هدف) همگرا شود و با تعداد دفعات تکرار حداکثر از قبل تعیین شده برابر شود. این مقادیر باید طوری تعیین شوند که از یادگیری بیش از اندازه شبکه جلوگیری شود. تعداد دفعات تکرار در این تحقیق ۱۰۰۰ بار در نظر گرفته شده است و مقدار Error Goal نیز برابر ۰/۰۲ تعیین شده است. مدل کامل یک نوع شبکه عصبی در شکل ۳ آمده است.



شکل ۳- شمای نمونه از مدل کامل شبکه عصبی

برای بهینه سازی ضریب یادگیری، از تعداد کل پارامترهای ورودی به عنوان تعداد لایه ورودی، از ۵ نرون به عنوان تعداد نرون‌های لایه میانی و یک نرون نیز به عنوان لایه خروجی استفاده شد. با اقتباس از [۲۳] ما این تعداد نرون‌ها را به همراه، پارامترهای دیگر شبکه از قبیل دفعات تکرار و Error Goal (خطای هدف) را ثابت نگه داشتیم و ضریب یادگیری را از ۰/۰۰۱ تا ۱ تغییر دادیم و هر بار شبکه را آموزش دادیم و مقدار MSE (میانگین مربعات خطا) را در هر مرحله مشخص نمودیم. به این ترتیب نرخ یادگیری که منجر به حداقل MSE (میانگین مربعات خطا)

ذیل توضیح داده می‌شود)، به منظور انطباق، مدل‌سازی و تصمیم گیری با توجه به مدل و داده‌های ورودی به سیستم یا شبکه عصبی (مشخصات هر فرد)، تحویل شود تا تصمیمی صحیح در مورد میزان داروی تجویزی برای هر فرد بدست آید. دلیل استفاده از سری زمانی بیماران این است که از هر بیمار در ۳ نوبت با فواصل یک هفته نمونه گیری شد. در این تحقیق بدست آوردن بهترین عملکرد و یا خلاصه‌ای از گذشته سیستم فاکتور مهمی است که توسط طبقه‌بندی کننده الگوریتم ژنتیک و الگوریتم SVM بهترین کارایی یا ویژگی سیستم در گذشته بدست می‌آید تا بوسیله شبکه‌های عصبی مصنوعی، رفتار آینده سیستم پیش بینی شود (با اقتباس از [۲۲]).

و: شبکه‌های عصبی مصنوعی

در این پروژه از یک شبکه عصبی ۳ لایه پیش خور با تابع فعالیت سیگموئید در لایه میانی و لایه خروجی و تابع خطی برای لایه ورودی استفاده گردید. تعداد نرون‌های لایه ورودی ۱۲ نرون در نظر گرفته شد. لایه میانی این شبکه برای انتخاب بهترین ساختار، از ۳ تا ۱۷ نرون یعنی به صورت (۱۵، ۱۳، ۱۱، ۹، ۷، ۵، ۳ و ۱۷) تغییر داده شد، تا با توجه به بهترین نتایج بدست آمده، تعداد نرون‌های لایه میانی انتخاب شود که در نهایت بهترین نتایج با ۹ نرون در لایه میانی بدست آمد. انتخاب تعداد نرون‌های لایه میانی بسیار مهم است چرا که اگر تعداد آنها کم باشد، شبکه برای حل مسائل غیر خطی و پیچیده با کمبود منابع یادگیری مواجه می‌شود و اگر زیاد باشد، خود باعث ایجاد دو مشکل می‌شود. اول آنکه زمان آموزش شبکه افزایش می‌یابد و دوم آنکه ممکن است شبکه نظام بی‌اهمیت داده‌های آموزش را یاد بگیرد و در حل مسائل ضعیف عمل نماید. در این پروژه، برای آموزش شبکه از الگوریتم پس انتشار خطا استفاده شده است. در این روش ماتریس‌های وزنی شبکه به گونه‌ای تغییر می‌یابند، تا مقدار متوسط مربعات خطای شبکه (MSE^1) یا میانگین مربعات خطا از مقدار خاصی که در برنامه مشخص می‌شود، (۰/۰۲) کمتر شود. خروجی مطلوب شبکه در واقع همان نتایجی است که توسط پزشک تجویز شده

می‌شود، حدود ۰/۲ بدست آمد. مشخصات شبکه عصبی به کار رفته در این پروژه در جدول ۳ آمده است.

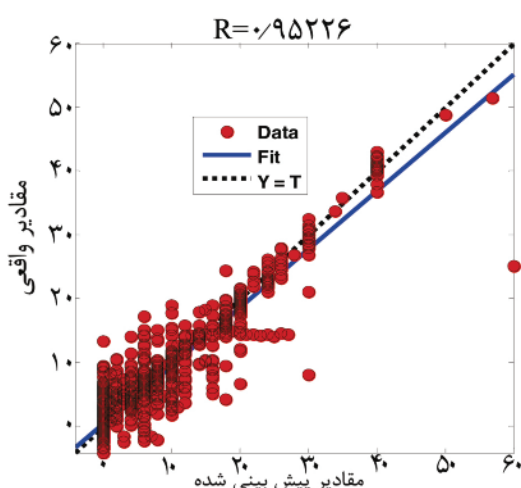
جدول ۳- پارامترهای شبکه عصبی پیشنهادی

مشخصات و پارامترهای شبکه عصبی	
تعداد نرون‌های لایه ورودی	۱۲ عدد
تعداد لایه‌های پنهان	۱ عدد
تعداد نرون‌های لایه پنهان	۵ عدد
تعداد نرون‌های لایه خروجی	۱ عدد
نوع تابع کار شبکه	سیگموئید
نوع آموزش شبکه	پس انتشار خطا
میانگین مربعات خطا	۰/۰۲
نرخ یادگیری	۰/۰۲
تعداد تکرار	۱۰۰۰

ز: صحت آزمایش و ارزیابی نهایی

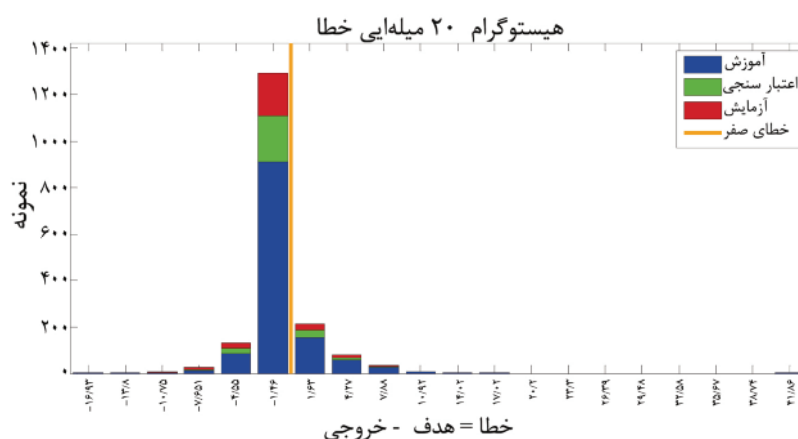
در حالت کلی می‌توان تمام فرآیند انتخاب ویژگی که منجر به انتخاب ویژگی به منظور پیش‌بینی دژ صحیح دارو شد را به این صورت بیان کرد، ابتدا بانک داده تشکیل شد، سپس توسط الگوریتم باینری ژنتیک با توجه به تعداد ویژگی‌های بانک داده، بردارهایی به ازای هر نمونه سرشار از "۰" و "۱" تولید گردید، رشته در مشخصات هر نمونه (گزارش) ضرب شده و با توجه به "۰" یا "۱" شدن ویژگی‌های هر نمونه، آرایه حاصل وارد تابع هزینه (در تابع برازندگی) شدند، سپس خروجی تشخیصی هر نمونه بدست آمد و متناسب با خروجی‌های تشخیصی، به منظور طبقه‌بندی صحیح، نتیجه وارد الگوریتم SVM گردید و مجدداً به الگوریتم GA باز گردید. این فرآیند در تابع هزینه الگوریتم GA و تکرار الگوریتم GA منجر به انتخاب بهترین ویژگی‌های ممکن به منظور پیش‌بینی شد. براساس بهترین ویژگی‌ها، بانک داده‌ای شامل فقط ویژگی‌های مهم وارد شبکه عصبی می‌شود. شبکه عصبی براساس مدل مشخص شده که متناسب با نظرات پزشک مبتنی بر بانک داده است، تصمیم‌گیری بر طبق پیش‌بینی را انجام می‌دهد. برای اثبات قابلیت‌های سیستم پیشنهادی، فرایند اعتبار سنجی عرضی ۱۰- دسته‌ای (توسط K-fold) بر روی پایگاه داده (مجموعاً ۳۶۸ نمونه حاوی ویژگی‌های نمونه) پیاده شد که به ۱۰ دسته تقسیم شدند. به طوری که هر

قسمت دارای ترکیبی از تمام نمونه‌ها می‌باشد. ۸ دسته ۳۶ داده‌ای و ۲ دسته ۳۷ داده‌ای. در این نوع اعتبار سنجی، ۹ قسمت از داده‌ها برای آموزش و قسمت باقیمانده برای آزمایش به کار رفتند. این فرایند آموزش و آزمایش ۱۰ بار به صورت متوالی چرخشی انجام شد به طوری که هر بار یک قسمت متفاوت برای آزمایش کنار گذاشته شد، تا بهترین ساختار شبکه مشخص شود، پس از مشخص شدن بهترین ساختار شبکه، می‌توان تمام آزمایشات را با این ساختار از شبکه انجام داد. بهترین نتیجه بدست آمده براساس نتایج واقعی در حالت رگرسیونی در شکل ۴ ترسیم شد، که مقدار انحراف و دقت بدست آمده قابل قبول می‌باشد (در حد بیش از ۰/۹۵ درصد).



شکل ۴- نمایش رگرسیونی مقادیر پیش‌بینی شده بر حسب نتایج واقعی

نمودار رگرسیونی در شکل ۴ نشان می‌دهد که نتایج بدست آمده در مقایسه با مقادیر واقعی در حالت کلی از دقت بالایی برخوردار است. به طوری که می‌توان گفت تقریباً مقادیر پیش‌بینی شده، بیش از ۰/۹۵ مقادیر واقعی هست. در عملکرد نهایی سیستم پیشنهادی، میزان خطایی در مقایسه بین مقادیر واقعی و مقادیر پیش‌بینی شده بوجود آمد؛ که برای نمایش این خطا یا اصطلاحاً، نمودار خطای هیستوگرام شبکه از شکل ۵ استفاده گردید.



شکل ۵- نمایش عملکرد خطای شبکه در محاسبه مقادیر واقعی و مقادیر پیش‌بینی شده

یافته‌ها

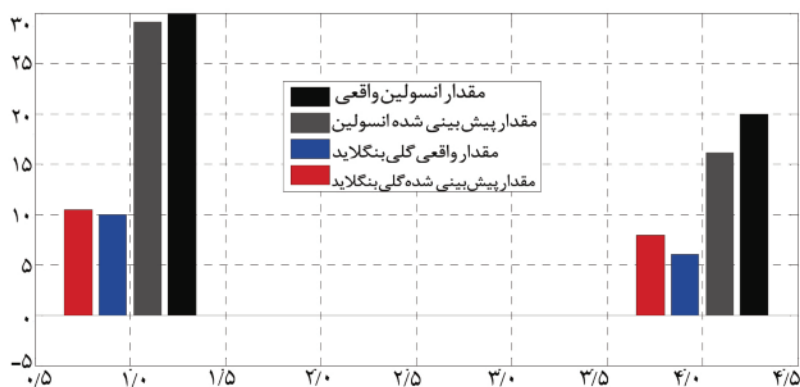
نتایج شبیه‌سازی‌ها و ارزیابی‌های انجام شده

به منظور شبیه‌سازی، از نرم‌افزار MATLAB® و همین‌طور برای ترسیم نمودارها از نرم‌افزار ORIGIN® استفاده شده است. نتایج شبیه‌سازی‌ها نشان می‌دهد که انتخاب صحیح ورودی‌ها، به شبکه‌های عصبی بسیار مهم است، چرا که ورودی‌ها می‌توانند دقت و کارایی شبکه را افزایش و زمان همگرایی را به طور چشمگیری کاهش دهد. ویژگی‌های اصلی، سرنوشت‌ساز و تاثیرگذار در تعیین میزان دُز انسولین و دارو در بیماران تعیین شدند. سپس شبیه‌سازی نهایی توسط شبکه عصبی مصنوعی مطابق شکل ۶ برای ۲۴ ساعت آینده بیمار تهیه شد. البته مقادیر پس از این که توسط شبکه عصبی، پیش‌بینی شدند، باید به صورت قبلی بازگردند، به همین خاطر آنها طبق رابطه (۱۲) درمالایز گردیدند، سپس مقادیر واقعی مشخص شدند.

$$x_i = y_i \times \sigma_i + m \quad (12)$$

در این رابطه x_i مقدار واقعی پیش‌بینی شده است، y_i مقدار پیش‌بینی شده توسط شبکه است، m مقدار میانگین داده‌ها و مقدار انحراف از معیار داده‌ها (منظور در هر ویژگی) است. در شکل ۶ میزان داروی، پیش‌بینی شده برای یک فرد که به صورت تصادفی انتخاب شده است (فرد شماره ۹)، مشخص شده است.

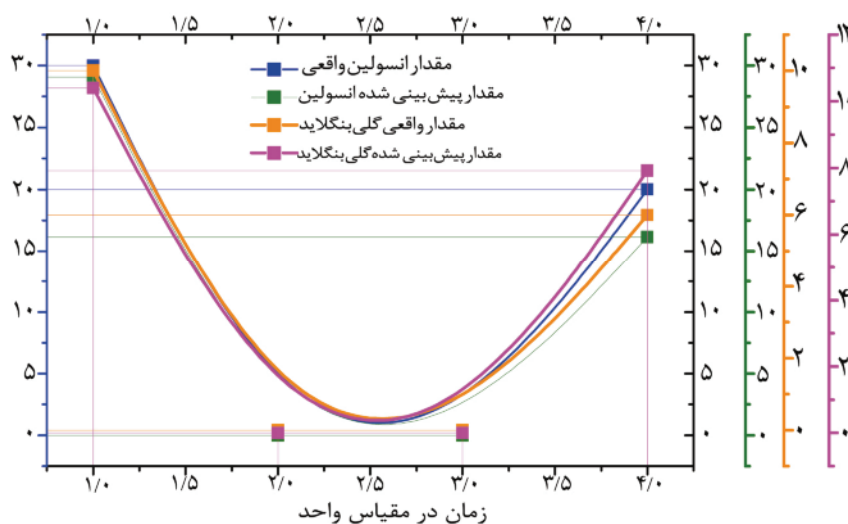
خطای سیستم که توسط تابعی در ساختار انطباقی شبکه عصبی تعریف شده است، (اختلاف مقادیر تابع هدف و تابع پیش‌بینی شده) توسط شکل ۵ بیان گردید. البته هر چقدر نمودار شکل ۵ به حالت گوسی نزدیک‌تر باشد مطمئناً شبکه عملکرد بهتری داشته است. در این شکل نمودار دارای ۴ قسمت است، در محور عمودی مقدار یا تعداد نمونه‌ها، مشخص شده است، در محور افقی، مقادیر عددی خطا (حول نقطه صفر یا خط عمودی نارنجی رنگ)، مشخص شده است، نمودارهای میله‌ای سه رنگ متفاوت دارند، اولین رنگ قرمز است، که بیان‌کننده خطا در نمونه‌های تست است، رنگ سبز بیانگر مقدار خطا در نمونه‌های اعتبار سنجی است و همین‌طور قسمت آبی رنگ که بیانگر مقدار خطا در نمونه‌های آموزشی هستند. تمام مقادیر خطا که در نمودار بالا ترسیم شده است، براساس نتایج نهایی و کلی است که در مقدار خطای "۱/۴۶۱-۱" شبکه به اشباع رسید.



شکل ۶- مقادیر پیش‌بینی شده و واقعی برای انسولین و گلیبن‌کلامید

گلیبن‌کلامید در ساعات اولیه (۰ بامداد تا ۶ صبح)، ۲۰ واحد انسولین و ۶ واحد گلیبن‌کلامید در ساعت (۱۸ عصر تا ۲۴ نیمه شب)، در صورتی که مقادیر پیش‌بینی شده برای این بیمار ۲۹ واحد انسولین و ۱۰ واحد گلیبن‌کلامید در ساعات اولیه (۰ بامداد تا ۶ صبح) و مقادیر ۱۶ واحد انسولین و ۷ واحد گلیبن‌کلامید در ساعات (۱۸ عصر تا ۲۴ نیمه شب) پیش‌بینی شده است. در شکل ۷ مقادیر پیش‌بینی شده بر حسب مقدار و زمان (بر حسب واحد) برای بیمار شماره ۹ ترسیم شده است.

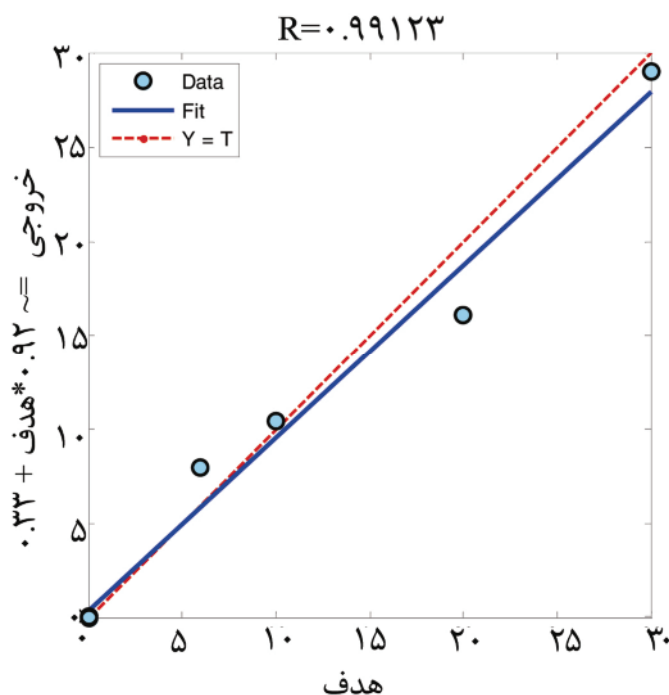
در شکل ۶ میزان تجویز دارو برای بیمار شماره ۹ مشخص شد. محور افقی این نمودار بیان‌کننده زمان واحد است، هر واحد ۶ ساعت آینده را بیان می‌کند که در مجموع ۲۴ ساعت آینده را بیان می‌شود، شماره ۱ بیان‌کننده ساعات (۰ بامداد تا ۶ صبح)، شماره ۲ بیان‌کننده ساعات (۶ صبح تا ۱۲ ظهر)، شماره ۳ بیان‌کننده ساعات (۱۲ ظهر تا ۱۸ عصر) و شماره ۴ بیان‌کننده ساعات (۱۸ عصر تا ۲۴ نیمه شب) می‌باشد، برای این بیمار مقادیر واقعی که مقرر شده بودند به این صورت است که، ۳۰ واحد انسولین، ۱۰ واحد



شکل ۷- مقادیر عددی پیش‌بینی شده و واقعی انسولین و گلیبن‌کلامید

که انحراف و ضریب R آن (بیش ۰/۹۹) بیان‌کننده مناسب و دقیق بودن پیش‌بینی صورت گرفته است.

اگر مقادیر پیش‌بینی شده را بر حسب مقادیر واقعی به وسیله نمایش رگرسیون بیان شود، شکل ۸ حاصل می‌شود



شکل ۸- نمایش مقادیر پیش‌بینی شده و مقادیر واقعی شده انسولین و گلیکین‌کلامید برای بیمار شماره ۹

ممکن است نتوان به درستی و به صورت دقیق کارایی شبکه و تصمیمات پزشک را با یکدیگر مقایسه کرد، لذا از تحلیل منحنی‌های ROC برای ارزیابی کارایی مدل در مقایسه با آنچه که از پزشک (متخصص غدد) بدست آمده است، استفاده شد.

جدول ۴- پارمترهای مورد نیاز در مبحث مقایسه

مقادیر واقعی	مقادیر پیش‌بینی شده
Tn (منفی درست)	Fp (مثبت نادرست)
Fn (منفی نادرست)	Tp (مثبت درست)
$Sen = \frac{Tp}{Tp + Fn}$	$Spe = \frac{Tn}{Tn + Fp}$
کسر حساسیت	کسر ویژگی
$Positive\ Predictive\ Value = \frac{Tp}{Tp + Fp}$	$Accuracy = \frac{Tp + Tn}{Tp + Fn + Tn + Fp}$
کسر مقادیر پیش‌بینی شده مثبت	کسر دقت پیش‌بینی شده منفی
$Negative\ Predictive\ Value = FPF = \frac{Tn}{Tn + Fn}$	کسر مقادیر پیش‌بینی شده منفی

براساس رابطه‌های جدول ۴، مقادیر دقت، حساسیت و ویژگی را برای داده‌های نهایی یعنی مقادیر پیش‌بینی شده بیماران را در ۲ قسمت دژ تجویز شده انسولین و دژ گلیکین‌کلامید، در جدول ۵ بررسی و مشخص شدند.

شکل ۸ بیان‌کننده این اصل است که برای فرد بیمار شماره ۹، صحت مقادیر پیش‌بینی شده ۰/۹۹ صحت مقادیر واقعی تجویز شده از سوی پزشک برای آن بیمار می‌باشد. تا این قسمت از مباحثی که در مورد پیش‌بینی و طراحی سیستم پیش‌بینی بحث شد، عملکرد سیستم در یک فرد تصادفی مورد بررسی قرار گرفت. که البته مسلماً با یک نفر نمی‌توان در مورد یک مجموعه صحبت، تصمیم و یا حتی اعتبار آن را بررسی نمود. به همین خاطر از نمودار ROC، پارامترهای حساسیت، ویژگی و دقت استفاده گردید.

ارزیابی مدل پیشنهادی

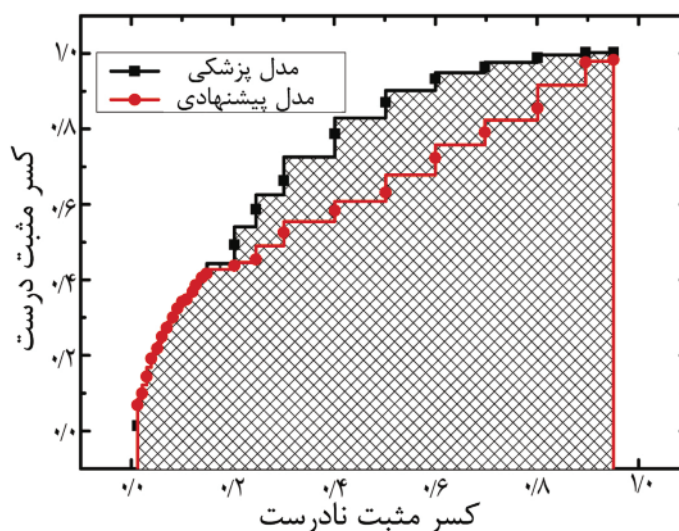
برای ارزیابی سیستم پیشنهادی، ۵۰ نمونه از بانک اطلاعاتی به طور تصادفی انتخاب شدند که شامل نمونه‌هایی با دژ داروهای تجویز شده در دو نوع انسولین و داروی گلیکین‌کلامید می‌باشند. معیارهای ما برای ارزیابی سیستم پیشنهادی، میزان دقت، حساسیت و ویژگی‌های تعیین شده و همچنین سطح زیر منحنی ROC^۹ برای این ۵۰ نمونه می‌باشد. نحوه محاسبه این پارامترها طبق جدول ۴ می‌باشد. از آنجایی که حساسیت، ویژگی و دقت، تا حد زیادی وابسته به جمعیت بیماران غالب بوده و از این رو

جدول ۵- مقادیر دقت نهایی، حساسیت و ویژگی پیش‌بینی شده برای ۲ مقدار انسولین و گلیبن‌کلامید

دقت الگوریتم پیشنهادی براساس درصد	گلیبن‌کلامید پیش‌بینی شده	دقت الگوریتم پیش‌بینی شده براساس درصد	انسولین پیش‌بینی شده
۸۳	حساسیت	۹۲/۷	حساسیت
۶۱	ویژگی	۸۴/۱	ویژگی
۹۲/۴۵	دقت	۹۵/۵۲	دقت

پس از تحلیل کامل و بررسی چند گانه در مورد دقت و قدرت پیشگویی سیستم پیشنهادی؛ به مناسب بودن این سیستم و اطمینان این سیستم می‌توان پی برد. در مرحله بعد، به تحلیل نمودار ROC در سیستم پیشنهادی پرداخته می‌شود. تحلیل ROC، شکلی از یک اندازه‌گیری است که عمدتاً برای مقایسه کارایی روش‌های تشخیصی بکار می‌رود. در تحلیل ROC، کارایی تشخیص و یا تصمیم‌گیری بر حسب دو شاخص ارزیابی می‌شود: ۱- کسر مثبت درست (TPF)، ۲- کسر مثبت نادرست (FPF). کسر اول به عنوان کسری از تصمیم‌گیری‌هایی است که به درستی تشخیص داده شده‌اند (مثل حساسیت) و کسر دوم به کسری از تصمیم‌گیری‌هایی است که به درستی تشخیص داده نشده‌اند،

از مقادیر فوق می‌توان برای رسم منحنی‌های ROC استفاده نمود. سطح زیر منحنی (ROC) معیاری برای میزان کارایی روش بکار رفته است. همان‌طور که قبلاً گفته شد، هر چقدر که سطح زیر این نمودار بیشتر باشد نشان دهنده عملکرد بهتر آن سیستم خواهد بود، البته اگر در حالت مقایسه با سیستمی مثل تصمیمات پزشک باشد، سیستم پیشنهادی زمانی در حد سیستم واقعی یا اصلی عمل می‌کند، که سطح زیر منحنی برابر سطح زیر منحنی تصمیمات پزشک باشد. در اثر تعریف پارامترها و مشخص شدن بازه‌ها، نمودار ROC برای بیماران براساس نتایج سیستم پیشنهادی و تصمیمات پزشک ارائه شده ترسیم شدند.



شکل ۹- نمودار ROC در سیستم پیشنهادی و سیستم مبتنی بر نظریات پزشک

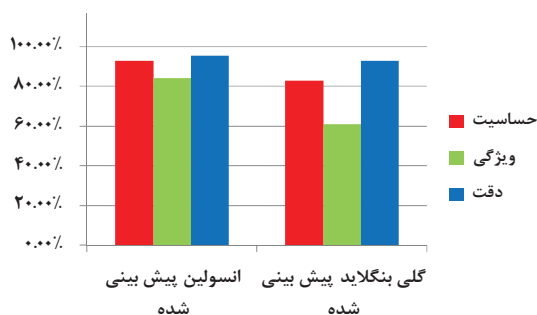
مساحت‌های سطح زیر ۲ نمودار در سیستم پیشنهادی و سیستمی که مبتنی بر نظریات و تصمیمات پزشک است، در شکل ۹ مشخص شده است، به راحتی می‌توان در مورد عملکرد مناسب سیستم پیشنهادی به خاطر برابر بودن نسبی

مساحت‌ها و محدوده عملکرد دو سیستم با توجه به خطوط قرمز و مشکی در شکل ۹ پی برد.

بحث

دیابت از شایع‌ترین عامل‌های مرگ و میر در سراسر دنیا، در بین تقریباً تمامی اقوام جهان می‌باشد، با توجه به پیشرفت‌های بدست آمده در روش‌های تشخیص، اما هنوز هم تعیین میزان صحیح دُز داروی این بیماران خصوصاً انسولین و گلیکین کلامید، در مراحل اولیه کار دشوار، مهم و سرنوشت‌ساز است و در بسیاری از موارد تعیین دُز دارو، با قطعیت و صحت کامل انجام نمی‌شود. در همین راستا تحقیقات جدیدی انجام گرفته است؛ از آن جمله می‌توان به مدل‌های ریاضی و مینی‌مال براساس ویژگی‌های حساسیت بدن بیمار در پاسخ به گلوکز تزریقی، دُز انسولین مصرفی و سطح فعلی قند خون که عمدتاً به منظور تعیین حساسیت بدن بیمار در پاسخ به گلوکز تزریقی و نیز پیش‌بینی کوتاه مدت سطح قند و انسولین موجود در بدن، مورد استفاده قرار گرفته‌اند، اشاره کرد [۲۶-۲۴]. Huang به همراه همکاران، با استفاده از شبکه‌های عصبی در سال ۲۰۱۰ پیش‌بینی سطح گلوکز خون بیماران را براساس ویژگی‌ها و اطلاعات هر بیمار شامل؛ وزن، زمان اولیه تزریق، زمان انتهایی تزریق و کربوهیدرات مصرفی برای بیماران را پیش‌بینی کرده‌اند [۲۷]. Yu در سال ۲۰۱۲ توسط روش‌های محاسباتی و برنامه‌ریزی LSSVM1 پیش‌بینی دُز انسولین را انجام داده است [۲۸]. البته باید گفت که از گذشته تحقیق و تامل بر روی موضوع پیش‌بینی دُز انسولین برای بیماران رایج بوده است. Canete و همکاران، یک سیستم حلقه بسته عملی آزمایشگاهی به صورت آنلاین برای کنترل گلوکز خون براساس میانگین وزن؛ گلوکز پلاسما و انسولین پیشنهادی تهیه کرده است [۲۹]. Pappada و همکاران در سال ۲۰۰۸ تنها برای ۵۰-۱۸۰ دقیقه توانستند براساس ویژگی‌های؛ نحوه زندگی، وضعیت روحی، دُز انسولین، وعده غذایی و کربوهیدرات مصرفی پیش‌بینی سطح انسولین برای بیماران را انجام دهند [۳۰]. Stavroula در سال ۲۰۰۶ با استفاده از ۴ پارامتر عملکرد کوتاه، متوسط، سریع، بلند مدت و کربوهیدرات مصرفی و غلظت قند خون که تابعی از زمان بعد از تزریق انسولین در بدن هستند (بیماران در اینجا کودک می‌باشند) توسط شبکه‌های عصبی مصنوعی میزان

انسولین را برای بیماران مشخص کرده‌اند [۳۱]. Wahab و همکاران با استفاده از کنترل کننده تطبیقی، طراحی یک کنترل کننده PID2 براساس ۲ معادله دیفرانسل مرتبه اول برای انسولین و تنظیم گلوکز در ۲ زیر سیستم به منظور تعیین دُز صحیح انسولین را انجام داده‌اند [۳۲]. لزوم استفاده از روشی صحیح و منطقی مبتنی بر واقعیت در این فرآیند نسبتاً پیچیده و حیاتی لازم است، سیستم‌های هوشمند مصنوعی که به نوعی تلفیقی از دانش فرد خبره یا بشری با مدلی از سیستم مورد بررسی است، می‌تواند به عنوان ابزاری مناسب در جهت تعیین این فرایند مهم استفاده شود. در این تحقیق سعی کردیم مدلی ارائه شود تا با استفاده از سیستم‌های هوشمند مصنوعی و تصمیمات پزشکان از یک طرف و رفتار بیماران از طرف دیگر در قالب بانک داده، فرآیند تعیین صحیح دُز انسولین و دُز گلیکین کلامید در بیماران را انجام دهیم. سیستم پیشنهادی همان طور که در شکل ۱۰ می‌بینیم، به دقت (و نتایج) بیش از ۹۵٪ دست یافت که در بررسی‌های به عمل آمده از نتایج حاصل، می‌توان به دقت و صحت این سیستم پی برد.



شکل ۱۰- نتایج نهایی سیستم پیشنهادی

به طوری که نتایج موجود و به دست آمده (در مقالات و تحقیقات ارائه شده قبلی سیستمی مشابه و یا همانند نبود؛ تا آنجا که مولفین اطلاع دارند) در مقایسه کلی نتایج، دقت و قدرت پیش‌بینی برای بیماران، به میزان قابل توجهی بهبود یافت. همچنین مقایسه نتایج به دست آمده در این تحقیق با نتایج به دست آمده از تحقیقات مشابه قبلی نشان دهنده افزایش قابل ملاحظه میزان صحت پیش‌بینی سری زمانی در سطح غلظت قند خون است، چرا که سیستم پیشنهادی موفق شد با در دست داشتن چند ویژگی از

داشت و بعد عملی به آنها بخشید. میل و رسیدن به چنین هدفی تنها در صورت دانش بالا، فهم درست و تسلط کافی از یک طرف به سیستم‌های هوشمند و از طرفی دیگر به بیماری و داده‌کاوی امکان‌پذیر است، که در آینده‌ای نه چندان دور انشاله نایل خواهد شد.

سپاسگزاری

در انتها از تمام همکاری‌ها و حمایت‌های بی‌شائبه و صمیمانه ریاست مرکز تحقیقات دیابت شهرستان سبزوار و همین طور کارکنان دلسوز و مهربان آن مرکز که ما را در تهیه این تحقیق یاری رساندند بسیار ممنون و متشکریم.

بیماران مبتلا به دیابت، سطح قند خون تا ۲۴ ساعت آینده بیمار را پیش‌بینی و کنترل نماید. از شبکه‌های مصنوعی عصبی به منظور انطباق، یادگیری، الگوریتم بردار ماشین پشتیبان به منظور دسته‌بندی و الگوریتم ژنتیک به منظور انتخاب بهترین ویژگی‌ها استفاده شد. بسیاری از محققین تمایل زیادی در استفاده از این ابزار را دارند، اما چالش آموزش شبکه‌های عصبی مهمترین دغدغه آنها می‌باشد. ترکیب الگوریتم‌های تکاملی و سیستم‌های داده کاوی می‌تواند یک ایده مناسب برای افزایش کارایی این شبکه‌ها در جهت پیش‌بینی صحیح باشند. البته هدف اصلی در این گونه تحقیقات بدست آوردن صحت و دقت ۱۰۰ درصدی می‌باشد، تا بتوان به سیستم‌های هوشمند مصنوعی اطمینان

مأخذ

- World Health Organization, Diabetes Center, Fact Sheet N 312, www.who.int/mediacentre/factsheets/fs312/en, (08/29/2012).
- Centers for Disease Control and Prevention, National Diabetes, Fact Sheet 2011 www.cdc.gov/diabetes, (08/29/2012).
- نیم آبادی، محمد رضا؛ امیر احمدی چماچار، نوشاز؛ تهامی، احسان؛ ربانی، حسین. تشخیص دیابت با استفاده از SVM. چکیده‌نامه چهاردهمین کنفرانس دانشجویی مهندسی برق ایران؛ کرمانشاه؛ ایران؛ کرمانشاه، دانشگاه صنعتی کرمانشاه؛ ص ۲۵۱-۲۵۸.
- American Diabetes Association, Diabetes Basics www.diabetes.org/diabetes-basics, (08/29/2012).
- فیوضی، محمد؛ قره‌خانی، اعظم؛ حدادینا، جواد. تشخیص بیماری دیابت براساس ترکیب و تعامل سیستم‌های فازی، درخت تصمیم‌گیری و سیستم‌های فازی عصبی. چکیده‌نامه چهارمین کنفرانس ملی فناوری اطلاعات و دانش؛ اردیبهشت ۱۳۹۱؛ بابل؛ ایران؛ بابل: دانشگاه صنعتی بابل؛ ص ۱۵۰-۱۵۶.
- فیوضی، محمد؛ حدادینا، جواد؛ ملانیا، نسرین. تشخیص بیماری دیابت براساس روش ترکیبی هوشمند. مجله علمی و ترویجی هوش مصنوعی و ابزار دقیق ایران؛ تابستان ۱۳۹۱، دوره اول (شماره ۳۲): ۸-۱.
- قوچانی، سعید راحتی؛ تهامی، سیداحسان. پیش‌بینی نوسانات سطح قند خون در بیماران مبتلا به دیابت نوع ۱ با استفاده از شبکه‌های عصبی خود بازگشتی المن. مجله علمی و پژوهشی فنی و مهندسی دانشگاه آزاد اسلامی مشهد؛ پاییز ۱۳۸۷، دوره دوم، (شماره ۱): ۹۴-۸۱.
- Jensen R. Ph.D. Thesis: Combining rough and fuzzy sets for feature selection, [School of informatics], University of Edinburgh; 2005 .
- علیپور، محمد، حدادینا جواد. معرفی یک سیستم دقیق تشخیص سرطان پستان. فصلنامه بیماری‌های پستان ایران، تابستان ۱۳۸۸، شماره دوم (دوره ۴)، ص ۱۸-۲۹.
- Xiang T, and Gong S. Spectral Clustering with eigenvector selection. *Journal of Pattern Recognition*, 2008; 41(3): p 1012-1029.
- اشک زری طوسی سهیلا، صدوقی یزدی هادی. خوشه بندی طیفی با انتخاب ویژگی Kernel PCA. چکیده‌نامه شانزدهمین کنفرانس ملی سالانه انجمن کامپیوتر ایران؛ ایران؛ دانشگاه صنعتی شریف، تهران، ایران؛ اسفند ۱۳۸۹، ص ۱۱۸.
- Kanan H, Faez K, Hosseinzadeh M. Face recognition system using ant colony optimization-based selected features. *Proceeding of IEEE Symp and Computational Intelligence in Security and Defense Applications*, 1-5 April 2007; p 57-62.
- Siedlecki W, Sklansky J. A note on genetic Algorithms for large-scale feature selection. *Journal of Pattern Recognition Letters*, 1989. 10(5): p. 335-47.
- Bon A, Ogier J, Razali A, Yasin A. Feature selection in beltline Moulding process using genetic algorithm. *Journal of Application Science and Research*, 2008; 4(7):p 783-92.
- Dorigo M, Maniezzo V, and Colona A. The ant system optimization by a colony of cooperating agents. *IEEE Transaction on Systems, Man, and Cybernetics-part B*, 1996: 26 (1): p 1-13.

16. Liu B, Abbass H, McKay B. Classification rule discovery with ant colony optimization. *Proceeding of IEEE Computational Intelligence Bulletin*, 25-29 April 1986; p 152 .
17. Meng Y. A swarm intelligence based algorithm for proteomic pattern detection of ovarian cancer. *Proceeding of IEEE Symp. Computational Intelligence and Bioinformatics and Computational Biology (CIBCB)*, 18-20 April 2006; p 7.
18. Lawrence D. *Handbook of genetic algorithms*. Chapman and Hall Press, New York, 1991, pp. 7-12 .
19. Michalewicz Z. *Genetic Algorithms + Data Structures = Evolution Programs*. Springer Press, New York, Third Edition, 1992; pp.3-15.
20. Zhang P, Verma B, Kumar K. Statistical classifier in conjunction with genetic algorithm based feature selection. *Journal of Pattern Recognition Letters* 2005. 26(2): p. 909-918.
۲۱. نعیم آبادی، محمد رضا؛ امیراحمدی چماچار، نوشاز، تهامی، احسان و ربانی، حسین. تشخیص دیابت با استفاده از SVM و MLP، چکیده نامه چهاردهمین کنفرانس دانشجویی مهندسی برق ایران/ کرمانشاه؛ ایران؛ کرمانشاه، دانشگاه صنعتی کرمانشاه؛ ص ۲۵۸-۲۵۱.
۲۲. اصغری اسکوئی، محمد رضا . کاربرد شبکه های عصبی در پیش بینی سری های زمانی. فصلنامه پژوهش های اقتصادی ایران، دوره ۴، شماره ۱۲، پاییز ۱۳۸۱: ۵۹-۶۸.
۲۳. نیروئی مهیار، عبدالمالکی پرویز، گیتی معصومه. شبیه سازی یک مدل ترکیبی به کمک الگوریتم ژنتیکی و شبکه عصبی مصنوعی برای تفکیک الگوهای خوش خیم و بدخیم سرطان سینه در ماموگرافی. مجله علمی و پژوهشی فیزیک پزشکی ایران، دوره ۳ (شماره ۱۳)، زمستان ۱۳۸۵: ۲۵-۳۶.
24. Li J, Kuang Y, Mason C. Modelling the glucose-insulin regulatory system and ultradian insulin secretory oscillations with two time delays. *The journal of Mathematics and computer Science* 2012; 8(4), 25-33.
25. Leaning M.S. and Boroujerdi M.A. A system for compartmental modelling and simulation. Computer. *Methods programs biomed Journal* 1991; 35(4), 24-36.
26. Tresp V, Moody J and Delong W.R. Neural modelling of physiological processes. *Journal of Computational Learning Theory and Natural Learning Systems* 1994; 2(5): 254-268.
27. Huang H.P, Shih-Wei Liu, I-Lung Chien, Chia-Hung Lin. A Dynamic Model with Structured Recurrent Neural Network to Predict Glucose-Insulin Regulation of Type 1 Diabetes Mellitus. *Proceeding of 9th International Symposium on Dynamics and Control of Process Systems (DYCOPS 2010)*; Leuven; Belgium; 2010: 5-7 July 2010, p 105 .
28. Lean Yu. An evolutionary programming based asymmetric weighted least squares support vector machine ensemble learning methodology for software repository mining. *Journal of Information Sciences* 2012; 191(4): 31-46.
29. Fernandez de Canete J, Gonzalez-Perez S, Ramos-Diaz J.C. Artificial neural networks for closed loop control of in silico and ad hoc type 1 diabetes. *Journal of Computer methods and programs in biomedicine* 2010; 4(2): 128-135.
30. Scott M, Pappada B.S, Brent D, Cameron, Paul M. Development of a Neural Network for Prediction of Glucose Concentration in Type 1 Diabetes Patients. *Journal of Diabetes Science and Technology* 2008; 2(5): 105-118.
31. Stavroula G, Mougiakakou , Aikaterini P, Dimitra I, Konstantina S, Nikita A. Neural Network based Glucose – Insulin Metabolism Models for Children with Type 1 Diabetes. *Journal of biosim* 2011; 25(8): 15-25.
32. Wahab A, Kong Y. K, and Quek C. Model Reference Adaptive Control on Glucose for the Treatment of Diabetes Mellitus. *Proceedings of the 19th IEEE Symposium on Computer-Based Medical Systems (CBMS'06)*; India, New Delhi; India; 2006; p 158.