

مقایسه سامانه‌های پشتیبان تصمیم‌گیری در پیش‌بینی دیابت

میثم جهانی^۱، جلال رضایی‌نور^۲، اسماعیل هداوندی^۳، ایرج صالحی^۴، حبیب‌اله تحسینی^۵

^۱ کارشناسی ارشد مهندسی فناوری اطلاعات، گروه مهندسی کامپیوتر، دانشکده فنی و مهندسی، دانشگاه قم

^۲ استادیار مهندسی صنایع، گروه مهندسی صنایع، دانشکده فنی و مهندسی، دانشگاه قم

^۳ دانشجوی دکتری مهندسی صنایع، گروه مهندسی صنایع، دانشگاه امیرکبیر

^۴ دانشیار فیزیولوژی، گروه فیزیولوژی، دانشکده پیراپزشکی، دانشگاه علوم پزشکی همدان

^۵ دانشجوی کارشناسی ارشد، گروه آمار زیستی، دانشگاه علوم پزشکی همدان

نویسنده رابط: جلال رضایی‌نور، آدرس: قم بلوار امین دانشگاه قم، دانشکده فنی و مهندسی، گروه فناوری اطلاعات، پست الکترونیک: j.rezaee@qom.ac.ir

تاریخ دریافت: ۹۳/۰۳/۲۷؛ پذیرش: ۹۴/۰۲/۰۵

مقدمه و اهداف: در طی سال‌های اخیر از جمله روش‌های پیش‌بینی و تشخیص بیماری در عرصه پزشکی، به کارگیری روش‌های پشتیبان تصمیم با الگوریتم‌های تکاملی و ترکیبی است، که دارای توانمندی بالایی در مدل‌سازی مسائل پزشکی و مهندسی دارند. هدف این مقاله مقایسه‌ی چند سامانه پشتیبان تصمیم‌گیری و بررسی دقت این سامانه‌ها در پیش‌بینی بیماری دیابت است.

روش کار: در مطالعه‌ی حاضر با استفاده روش ترکیب الگوریتم ژنتیک و لونیگ-مارکوارت (Genetic Algorithm and Levenberg-Marquardt; GALM)، به تعیین و بهینه‌سازی اوزان‌های شبکه‌ی عصبی پرداخته شد، و برای بررسی اعتبارسنجی مدل‌ها، از روش‌های اعتبار سنجی سنتی و اعتبار سنجی k باره (K-Fold Cross Validation; K-Fold) استفاده گردید، و در نهایت مدل پیشنهادی (GALM) با مدل‌های رگرسیون لجستیک، الگوریتم ژنتیک از طریق نمودار سطح زیر منحنی (Receiver operating characteristic, (ROC و ماتریس در هم ریختگی (Confusion matrix) مقایسه گردید.

نتایج: پس از انجام بررسی‌ها معلوم شد در بین مدل‌های مقایسه شده، مدل حاصل از الگوریتم GALM دارای حساسیت و ویژگی بالاتری نسبت به مدل‌های رگرسیون لجستیک و الگوریتم ژنتیک می‌باشد. همچنین در بین مدل‌ها، مدل پیشنهادی (الگوریتم GALM) مدلی است که دارای حساسیت، ویژگی، ارزش اخباری منفی (NPV) (Negative Predictive Value; NPV)، ارزش اخباری مثبت (PPV) (Positive Predictive Value; PPV)، بالا و درست‌نمایی منفی (NLR) (Negative Likelihood Ratio; NLR) پایین و نزدیک به صفر می‌باشد، و می‌توان این مدل را به عنوان مدلی مناسب انتخاب کرد.

نتیجه‌گیری: نتایج نشان می‌دهد، مدل GALM با میزان‌های به ترتیب حساسیت، ویژگی، ارزش اخباری مثبت، ارزش اخباری منفی، سطح زیر منحنی ۹۸/۷، ۹۰/۰۱، ۹۱/۸، ۹۸/۳، ۰/۹۷۲، در مقایسه با مدل‌های GA و LR مدلی مناسب برای پیش‌بینی دیابت می‌باشد.

واژگان کلیدی: دیابت، شبکه عصبی، الگوریتم ژنتیک

مقدمه

از دیرباز بشر با بیماری‌های مختلفی روبه‌رو بوده است که در برخی موارد توانسته که بیماری‌ها را تشخیص و راه‌کاری در راستای بهبود آن ارائه دهد، اما متأسفانه در مواقعی بیماری به دلیل عدم تشخیص به صورت خاموش در بیمار تا مدت‌ها باقی مانده و ممکن است جان بیمار را به خطر اندازد. بر همین اساس مطالعه‌ها و بررسی‌های فراوانی در عرصه پیش‌بینی بیماری‌های مختلف انجام گردیده تا به جایی که امروز بشر از مدل‌های پشتیبان از تصمیم و روش‌های هوشمند برای پیش‌بینی، بهره می‌جوید. از جمله کاربردهای مورد استفاده مدل‌های پشتیبان از تصمیم در زمینه پزشکی و تشخیص بیماری‌هایی مانند دیابت

است (۱).

دیابت یک بیماری همه‌گیر می‌باشد که در اثر کاهش یا نبود انسولین در بدن اتفاق می‌افتد. بر همین اساس دیابت بر دو نوع مهم ۱ و ۲ تقسیم می‌شود. دیابت نوع یک یا دیابت وابسته به انسولین، بر اثر تخریب سلول‌های بتای پانکراس که مسئولیت تولید انسولین را به عهده دارند، ایجاد می‌شود، و فرد مبتلا به این بیماری نیازمند تزریق انسولین می‌باشد. این بیماری در هر سنی می‌تواند اتفاق بی‌افتد که در کودکان بسیار متداول بوده است (۲). نوع دیگر دیابت، دیابت نوع دوم می‌باشد. در این نوع دیابت، انسولین (هورمون تنظیم کننده قند خون) حداقل در مراحل اولیه

مورد استفاده قرار گرفت.

۱. غلظت گلوکز پلاسماي خون؛

۲. فشار دیاستولیک (mm/Hg)؛

۳. شاخص توده بدن؛

۴. سن؛ و

۵. تغییر وزن

نرمال‌سازی

در آموزش شبکه‌های عصبی نکته مهمی که باید مورد توجه قرار گیرد، نرمال‌سازی داده‌ها پیش از استفاده در مدل است. در این مطالعه از روش نرمال‌سازی مطابق با فرمول (۱) استفاده شد که داده‌ها را در بازه بین (۱-۰) تبدیل می‌کند.

$$1 - \left(\frac{\text{داده} - \text{کمینه داده}}{\text{حداکثر داده} - \text{کمینه داده}} \right) \times 2 = \text{داده نرمال شده} \quad (1)$$

پس از نرمال‌سازی در ادامه داده‌ها به شبکه عصبی برای آموزش ارسال می‌گردد.

ایجاد شبکه عصبی

شبکه‌ای عصبی یکی از پرکاربردترین روش‌های مدل‌سازی است، که در مسائل مختلفی مانند طبقه‌بندی استفاده می‌گردد. هر شبکه‌ای عصبی از سه لایه اصلی (لایه ورودی، لایه پنهان، لایه خروجی) تشکیل شده است (۸،۹) و با روش‌های مختلفی به هم متصل می‌شوند. در این مطالعه از شبکه عصبی پرسپترون چند لایه (MLP)^۲ استفاده گردید و برای به‌دست آوردن دقت مناسب در پیش‌بینی، پارامترهای مختلفی از توابع انتقال و خطا مورد آزمایش قرار گرفتند و بهترین نتایج مطابق با جدول شماره ۱ به‌دست آمد.

جدول شماره ۱- خلاصه ویژگی‌های انتخاب شده در شبکه عصبی

تابع انتقال لایه اول	سیگموئید
تابع انتقال لایه دوم	تانژانت
تابع خطا	میانگین توان دوم خطا
تعداد نرون ورودی	۵
تعداد نرون میانی	۷
تعداد نرون خروجی	۱

هم‌چنین برای پیدا کردن بهترین میزان یادگیری و تعداد نرون‌های لازم، در شبکه‌ای عصبی یک پارامتر از پارامترهای جدول شماره ۲ در هر مرحله تغییر داده، و بقیه پارامترها ثابت نگه داشته شد. این روند برای تمام پارامترها انجام گرفته، سپس نتایج در جدول شماره ۲ ثبت گردید.

به میزان کافی ترشح می‌شود، اما به نوعی مقاومت در برابر عملکرد طبیعی این هورمون در سطح سلول‌های بدن وجود دارد که نتیجه آن بالا رفتن سطح قند خون می‌باشد (۳).

مرکز کنترل و پیشگیری (CDC)^۱ برآورد نموده که در سال ۲۰۱۱ میلادی نزدیک به ۲۵/۸ میلیون آمریکایی مبتلا به دیابت بوده و ۷۹ میلیون نفر دیگر در خطر ابتلا به این بیماری هستند و همچنین در سال ۲۰۱۲ میلادی بیان شد که میانگین هزینه برای درمان به بیمار دیابتی می‌شود ۸۵۲۰۰ بوده که ۵۳ درصد آن مربوط به عوارض دیابت است (۴).

تأخیر در تشخیص و پیش‌بینی دیابت به دلیل کنترل ناکافی از قند خون عوارض خطر ابتلا به مویرگ‌ها (نفروپاتی، رتینوپاتی، نوروپاتی) و ماکرو واسکولار (سکته مغزی، سکته قلبی، و مرگ قلبی-عروقی)، بیماری‌های چشمی و نارسایی کلیوی را افزایش می‌دهد (۵،۶).

بنابراین ارایه مدلی به منظور پیش‌بینی دیابت برای این‌که پزشکان بتوانند از آن به عنوان مدلی کمکی در پیش‌بینی دیابت استفاده کنند، می‌تواند مفید باشد.

در این مطالعه با استفاده از شبکه عصبی برای پیش‌بینی افراد دیابتی از غیر دیابتی استفاده گردید و از، ترکیب الگوریتم ژنتیک و لونبرگ-مارکوارت (GALM) برای بالا بردن دقت در پیش‌بینی دیابت استفاده شد. روش‌شناسی کار در بخش‌های بعدی توضیح داده شده است.

روش کار

در راستای طراحی مدلی برای پیش‌بینی دیابت، پژوهش حاضر شامل ۴ گام است که در شکل ۱ مشاهده می‌شود.



شکل شماره ۱- گام‌های اجرای پروژه

جمع‌آوری داده‌ها

به منظور انجام مطالعه حاضر با استفاده از روش سرشماری، مجموعه داده‌های مرکز دیابت استان همدان که مشتمل بر ۵۴۵ نمونه انتخاب گردید. با توجه به ویژگی‌های اعلام گردیده از سوی بهداشت جهانی و مقالات مورد بررسی (۷) و مشاوره از پزشک متخصص ۵ ویژگی از پارامترها به شرح زیر در پیش‌بینی دیابت

^۲ Multi Perceptron Layer, MLP

^۱ Centers for Disease Control and Prevention; CDC

جدول شماره ۲- تعیین و مقایسه پارامترهای شبکه عصبی

تعداد نرون لایه مخفی	میزان یادگیری	میزان مومنتم	دقت آموزش	دقت تست
۳	۰/۶۳	۰/۷۳	۹۳	۷۹
۷	۰/۵۷	۰/۷۶	۸۹	۸۸
۱۶	۰/۶۰	۰/۷۷	۸۹	۸۷
۳۶	۰/۶۲	۰/۷۲	۹۱	۸۷

همان‌طور که در جدول شماره ۲ مشخص می‌شود، بهترین نتیجه مربوط لایه مخفی با ۷ نرون می‌باشد. در مرحله بعد به منظور به‌دست آوردن دقت بالاتر و آرایه مدلی بهتر از مدل شبکه عصبی، اوزان شبکه عصبی، از GALM به‌دست آمد.

به‌روزرسانی اوزان

نکته‌ای که در شبکه‌های عصبی باید مورد توجه قرار داد، وزن‌دهی اولیه و بروز رسانی اوزان آن می‌باشد، به طوری که اگر وزن‌های شبکه عصبی به درستی انتخاب شوند، الگوریتم در بهینه‌ی محلی قرار نگرفته و نقطه بهینه سراسری را پیدا خواهد کرد (۱۲-۱۰). در همین راستا در مطالعه حاضر از الگوریتم ژنتیک و ترکیبی از الگوریتم ژنتیک با روش لونبرگ-مارکوارت (LM)^۱ که نخستین بار توسط جان هلند معرفی شد، برای وزن‌های شبکه‌ی عصبی استفاده شد، که در ادامه به توضیح هر یک پرداخته می‌شود.

الگوریتم ژنتیک

الگوریتم ژنتیک^۲ یکی از الگوریتم‌های جستجوی تصادفی است، که ایده‌ی آن برگرفته از طبیعت می‌باشد. الگوریتم ژنتیک علاوه بر این‌که در روش‌های مسائل خطی، محدب و برخی مشکلات مشابه بسیار موفق بوده است، برای حل مسایل گسسته و غیرخطی نیز بسیار کارا می‌باشد. ایده‌ی ژنتیک طبیعی در جانداران، بدین صورت است که از ترکیب کروموزوم‌های بهتر، نسل‌های بهتری پدید می‌آیند. هم‌چنین در این بین گاهی اوقات جهش‌هایی در کروموزوم‌ها رخ می‌دهد که ممکن است باعث بهتر شدن نسل بعدی شود. الگوریتم ژنتیک نیز با استفاده از این ایده به حل مسائل می‌پردازد. روند استفاده از الگوریتم‌های ژنتیک به صورت زیر می‌باشد (۱۳-۱۵):

۱. انتخاب جمعیت اولیه از یک فضای راه حل؛

۲. ارزیابی جمعیت: مهم‌ترین قسمت الگوریتم ژنتیک بخش

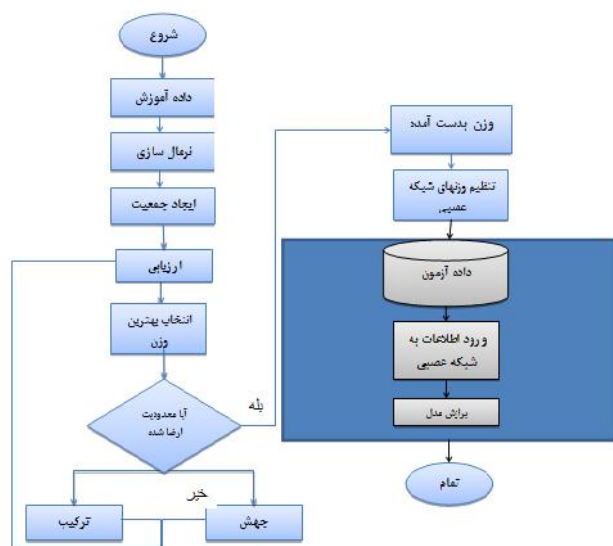
ارزیابی می‌باشد که اگر به درستی انتخاب شود، تأثیر به سزایی در بالا رفتن دقت تشخیص خواهد داشت. در این قسمت معمولاً از روش‌هایی مانند: میانگین توان دوم (MSE)^۳، برای ارزیابی استفاده می‌شود.

۳. انتخاب بهترین کروموزوم از جمعیت موجود: در این مرحله از بین جمعیت موجود مطابق با الگوریتم‌های انتخاب، والدها برای ترکیب و جهش انتخاب می‌شوند. از متداول‌ترین روش‌های انتخاب می‌توان به چرخ رولت اشاره کرد.

۴. ترکیب^۴: در این مرحله از ترکیب دو والد منتخب که در مرحله قبل به‌دست می‌آید، فرزندان جدید ایجاد می‌شود.

۵. جهش^۵: در قسمت جهش والدی انتخاب گردیده و با تغییر در برخی ژن‌های آن فرزند جدید ایجاد می‌گردد.

فلوچارت کلی الگوریتم ژنتیک را می‌توان در شکل شماره ۲ مشاهده کرد.



شکل شماره ۲- فلوچارت الگوریتم ژنتیک

در این مرحله از مطالعه، جمعیت و نرخ‌گذار و جهش‌های متفاوت در الگوریتم ژنتیک (GA) مطابق با فلوچارت شکل ۲ آزمایش شد و پارامترهای الگوریتم ژنتیک مناسب انتخاب شد. در همین راستا به دلیل وجود حالات متفاوت در انتخاب پارامترها، جمعیت، میزان گذار، و میزان جهش، از روش تاگوچی برای

^۲ Mean Square Error, MSE

^۳ Cross

^۵ Mutation

^۱ Levenberg-Marquardt, LM

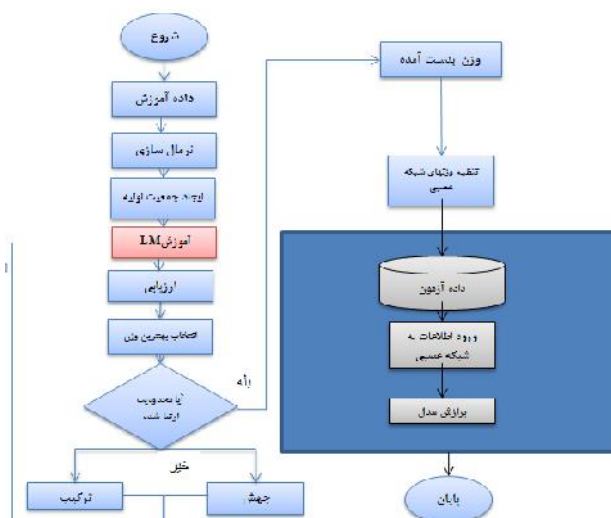
^۲ Genetic Algorithm, (GA)

الگوریتم GALM

مدلی که در به‌روز رسانی اوزان استفاده گردید، استفاده از روش ترکیب الگوریتم ژنتیک و لونیگ-مارکوارت (GALM) است، به طوری که داده‌ها پس از هر عمل ترکیب و جهش با روش لونیگ-مارکوارت آموزش داده شده و از طریق روش اعتبار سنجی k باره اعتبار سنجی گردید (شکل شماره ۱) که برای مطالعه بیشتر در مورد الگوریتم می‌توان به منبع شماره ۱۸ مراجعه شود. نتایج پس از ورود اطلاعات و اعمال پارامترها به شبکه GALM در جدول شماره ۵ ثبت گردید.

جدول شماره ۵- نتایج به‌دست آمده از ترکیب LM و ژنتیک

نسل	k	میزان ترکیب	میزان جهش	دقت آموزش	دقت تست
۲۰	۵	۰/۹	۰/۰۵	۹۶/۷	۹۵/۵



شکل شماره ۳- فلوچارت ترکیب الگوریتم ژنتیک و LM

میانگین توان دوم خطا

میانگین توان دوم خطا (MSE) یک معیار برای به‌دست آوردن بهترین برآورد می‌باشد، که در بین متخصصان آمار از مطلوبیت خاصی برخوردار است، این معیار از تفاوت بین مقدار واقعی مدل و مقدار برآورد شده از خروجی مدل به‌دست می‌آید (۱۹). با استفاده از این معیار، برآوردی که دارای کم‌ترین MSE باشد، انتخاب می‌شود. فرمول مورد استفاده در MSE به صورت زیر می‌باشد.

$$MSE = \frac{1}{n} \sum_{i=1}^n (Y_i^A - Y_i)^2 \quad (2)$$

بررسی حالت‌ها استفاده شد که نتایج بررسی در جدول شماره ۳ آمده است. تاگوچی روشی در طراحی آزمایش‌ها می‌باشد. در روش طراحی آزمایش‌ها تغییراتی را به روی متغیرهای ورودی ایجاد کرده تا تأثیر هر پارامتر روی خروجی مشخص شود. تغییر این پارامترها و استفاده از روش آزمایش‌ها زمانی مفید می‌باشد که متغیرهای مورد بررسی زیاد بوده و نتوان همه‌ی حالت‌ها را بررسی کرد، که در این حالت روش یاد شده چند مقدار را برای هر پارامتر برای بررسی، پیشنهاد داده و پس از آزمایش، حالت‌های پیشنهادی و ثبت نتایج، تاگوچی بهترین پارامترها را معرفی می‌کند (۱۶).

جدول شماره ۳- نتایج به‌دست آمده از الگوریتم ژنتیک و آزمایش تاگوچی

تعداد نسل	جمعیت اولیه	میزان گذار	میزان جهش	دقت آموزش	دقت تست
۵۹	۲۰	۰/۶	۰/۰۵	۹۴	۸۲
۷۵	۲۰	۰/۸	۰/۱	۹۳	۸۶
۵۱	۲۰	۰/۹	۰/۲	۹۴	۷۱
۶۸	۳۰	۰/۶	۰/۱	۹۴	۸۶
۷۹	۳۰	۰/۸	۰/۲	۹۳	۸۶
۷۴	۵۰	۰/۶	۰/۲	۹۴	۸۷
۶۰	۵۰	۰/۸	۰/۰۵	۹۳	۵۰
۸۵	۵۰	۰/۹	۰/۱	۹۴	۸۸
۵۸	۳۰	۰/۹	۰/۱	۹۳	۴۸

پس از اجرای پیشنهادی طراحی آزمایش‌ها، تاگوچی میزان گذار ۰/۹، میزان جهش ۰/۰۵ و جمعیت اولیه ۳۰ را به عنوان بهترین پارامترها پیشنهاد داد، که پس از به‌دست آوردن این پارامترهای مناسب، نمونه‌ها با الگوریتم ژنتیک، و روش اعتبار سنجی k باره (K-Fold) اعتبارسنجی گردید که نتایج آن در جدول شماره ۴ مشاهده می‌شود. در روش اعتبارسنجی k باره، داده‌ها به k زیر مجموعه و با اندازه برابر تقسیم می‌شوند که در آن یکی از k مجموعه داده به عنوان اعتبار نگه داشته شده و $k-1$ زیر مجموعه دیگر برای آموزش استفاده می‌گردد. این کار تا زمانی که هر زیر مجموعه دقیقاً یک بار آموزش داده شود، ادامه می‌یابد (۱۷). نتایج حاصل از اجرای الگوریتم با روش اعتبار سنجی k باره در جدول شماره ۴ مشاهده می‌شود.

جدول شماره ۴- نتایج به‌دست آمده از الگوریتم ژنتیک به روش k-Fold

نسل	K	میزان ترکیب	میزان جهش	دقت آموزش	دقت تست
۲۰	۵	۰/۹	۰/۰۵	۹۰	۹۲

همان‌طور که در جدول شماره ۴ مشخص است، دقت آزمایش نسبت به حالتی که از روش اعتبار سنجی سنتی استفاده شد بیشتر است.

تحلیل سطح زیر منحنی و ماتریس در هم ریختگی

نکاتی که در انتخاب یک مدل مطرح است این است که آیا مدلی مناسب است که دارای حساسیت بالا باشد، یا مدلی که ویژگی آن بالاست؟

هنگامی که یک مدل تحلیلی برای بیماری ارایه می‌گردد، اغلب گروهی از مردم را مدل مثبت تشخیص می‌دهد. این افراد هم شامل افراد واقعاً بیمار (مثبت حقیقی) و هم افراد سالم (مثبت کاذب) هستند. مثبت کاذب و ویژگی، از این نظر دارای اهمیت است که تمام افراد غربال شده در این گروه باید آزمون‌های پیچیده‌تر و پرهزینه‌تری را انجام دهند. یکی از مشکلات عدیده‌ای که به وجود می‌آید، تحمیل بار اضافه به سامانه مراقبت‌های بهداشتی است. مشکل دیگر القای اضطراب و نگرانی در افرادی است که به آن‌ها گفته شده نتیجه آزمایش‌های آن‌ها مثبت است.

متقابلاً اگر فردی بیمار باشد، اما به اشتباه نتیجه‌ی آزمایش در وی منفی گردد، به‌ویژه هنگامی که پای یک بیماری جدی و به دنبال آن لزوم اقدام‌های درمانی اساسی درمیان باشد، مشکل واقعاً بحرانی است. به‌عنوان مثال اگر بیماری مورد نظر نوعی سرطان و یا دیابت باشد که فقط در مراحل اولیه قابل علاج باشد، نتیجه منفی کاذب و حساسیت پایین می‌تواند فرد را به مصرف قند زیادی و عوارض پس از آن دچار نماید. بنابراین در انتخاب مدلی که حساسیت بالا و یا مدلی که ویژگی بالایی داشته باشد، باید برحسب نوع بیماری مدل انتخاب گردد. با توجه به توضیحات یاد شده اگر مدلی انتخاب گردد، که هم حساسیت و هم ویژگی بالایی داشته باشد و سطح زیر منحنی آن در نمودار ROC بالا باشد، می‌توان این مدل را مناسب برای آن بیماری تشخیص داد. نمودار ROC نموداری است که در آن محور x آن ویژگی و محور y آن حساسیت می‌باشد که هر چه نمودار به سمت مساحت ۱ نزدیک باشد و سطح زیر منحنی آن بیشتر باشد بهتر است (۲۰، ۲۱).

یافته‌ها

جدول شماره ۶- ماتریس در هم ریختگی

نسبت درست‌نمایی منفی	نسبت درست‌نمایی مثبت	ارزش پیشگویی منفی	ارزش پیشگویی مثبت	ویژگی	حساسیت	
۰/۱۴	۹/۳	۹۸/۳	۹۱/۸	۹۰/۰۱	۹۸/۷	GALM
۰/۱۴	۶/۵۰	۹۸/۲۰	۸۸/۶۰	۸۴/۸۰	۹۸/۷۰	GA
۰/۰۵	۷/۸۰	۹۳/۵	۹۰/۴	۸۷/۹۰	۹۴/۹	LR

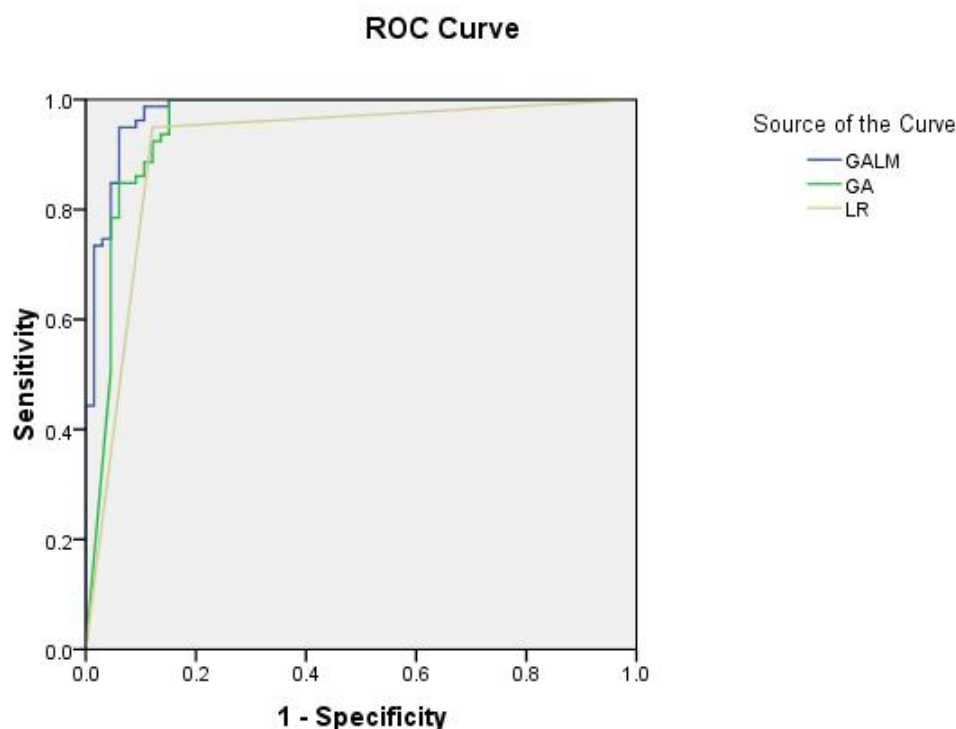
به منظور بررسی و مشخص کردن مدل مناسب از بین مدل‌های رگرسیون لجستیک (LR)، الگوریتم ژنتیک (GA)، الگوریتم ژنتیک و لونبرگ-مارکوارت (GALM) در این مطالعه از ماتریس آشفتگی (Confusion matrix) و نمودار سطح زیر منحنی ROC استفاده گردید.

ماتریس آشفتگی جدول ۶، نتایج حاصل از مدل‌های به‌دست آمده از الگوریتم‌های LR، GALM، GA را با یک دیگر از نظر حساسیت، ویژگی، ارزش اخباری مثبت، ارزش اخباری منفی درست‌نمایی مثبت و درست‌نمایی منفی مقایسه کرده است.

همان‌طور که در جدول شماره ۶ مشخص است در بین مدل‌های مقایسه شده، مدل‌های حاصل از الگوریتم GALM و GA دارای حساسیت بالاتری نسبت به مدل رگرسیون لجستیک می‌باشند، اما در بین این دو مدل، ویژگی مدل حاصل از الگوریتم GALM در مقایسه با سایر مدل‌ها بالاتر است و همان‌طور که در قسمت روش تحقیق توضیح داده شد، مدلی مناسب است که دارای حساسیت و ویژگی، ارزش اخباری منفی، ارزش اخباری مثبت بالاتر و درست‌نمایی منفی پایین داشته باشد. پس در بین مدل‌ها، مدل پیشنهادی (الگوریتم GALM) مدلی است که دارای حساسیت و ویژگی ارزش اخباری منفی، ارزش اخباری مثبت بالاتر و درست‌نمایی منفی پایین و نزدیک به صفر است. پس می‌توان این مدل را به عنوان مدلی مناسب انتخاب کرد.

نمودار ROC

مناطق زیر منحنی نمودار AUC که در جدول شماره ۷ و شکل شماره ۴ مشاهده می‌شود، نشان دهنده‌ی این است که الگوریتم ترکیبی GALM از مساحت زیر منحنی بالاتری نسبت به الگوریتم‌های GA و LR برخوردار است. بنابراین مدل GALM مطابق با نمودار ROC نیز دارای سطح زیر منحنی بالاتری می‌باشد و نسبت به مدل‌های دیگر مناسب‌تر است.



شکل شماره ۴- نمودار ROC

جدول شماره ۷- مناطق زیر منحنی نمودار ROC

سطح	نتایج آزمون متغیر
۰/۹۷۲	GALM
۰/۹۵۲	GA
۰/۹۱۴	LR

بالتر است (۲۲).

بحث

نتایج تحلیل مطالعه حاضر نیز با نمودار ROC انجام گردیده و در صورت انتخاب وزن مناسب نشان از دقت بالای شبکه ANN دارد که با نتیجه‌ی این پژوهش هم‌خوانی دارد. همچنین در این مطالعه نیز مدل پیشنهادی با رگرسیون چند متغیره مقایسه گردید که نشان می‌دهد مدل GALM از رگرسیون دقت بالاتری دارد.

تمورتر^۱ و همکاران، مقایسه مطالعه‌هایی روی مدل‌های تشخیص دیابت داشته و با طراحی مدلی از شبکه عصبی نشان دادند در بین شبکه‌ی عصبی و مدل‌هایی هم‌چون رگرسیون مدل شبکه‌ی عصبی عملکرد بهتری دارد (۱).

امیرها با روش‌های سیستم استنتاج فازی- عصبی تطبیقی^۲ و مجموعه راف (Rough Set^۳) روی دوزهای دارویی کارکرده و این

یکی از مسائل مهم در عرصه پزشکی، تشخیص انواع بیماری‌ها، از جمله دیابت است. در ایران، دیابت یکی از شایع‌ترین بیماری‌ها بوده که در ۷۰ درصد موارد افراد از بیماری خود آگاه نمی‌باشند. بدین منظور امروزه استفاده از دانش نرم‌افزاری در عرصه پزشکی رشد چشم‌گیری داشته و با توسعه این فناوری‌ها، مدل‌ها و الگوریتم‌هایی معرفی گردیدند که به کمک حوزه‌ی بهداشت آمده، و استفاده از سامانه‌های پشتیبان از تصمیم به عنوان یکی از نتایج این پیشرفت‌ها در کاهش زمان پیش‌بینی و بهبود دقت تشخیص است.

چونجیان و همکاران، مطالعه موردی روی ۸۶۴۰ بیمار برای شناسایی بیماران DM۲T و آزمایش کردن شبکه مصنوعی (ANN) و لجستیک چند متغیره داشته و تحلیل‌های خود را در نمودار ROC مقایسه کردند و به این نتیجه رسیدند دقت ANN از MLR

^۱ Temurtas

^۲ Adaptive Neuro Fuzzy Inference System, (ANFIS)

به ترتیب با نرون، میزان یادگیری، میزان مومنتم، و دقت ۵۷/۷، ۰/۷۶ و ۸۸ درصد مناسب‌ترین شبکه عصبی انتخاب شد (جدول شماره ۲). پس از پیدا کردن شبکه عصبی مناسب، به تنظیم اوزان شبکه عصبی با الگوریتم GALM با جمعیت اولیه، میزان ترکیب، و میزان جهش ۳۰، ۰/۹ و ۰/۰۵ و روش اعتبار سنجی kباره پرداخته شد که پس از اجرا الگوریتمی با دقت پیش‌بینی ۹۵/۵ به‌دست آمد.

تشکر و قدردانی

نگارندگان مقاله بر خود لازم می‌دانند تا از همکاران دانشگاه علوم پزشکی همدان و دکترمهدی مهدوی و تمامی کسانی که نگارندگان را در این مقاله یاری کردند، تشکر و قدردانی نمایند.

دو روش را با هم مقایسه کردند و به این نتیجه رسیدند، که در روش نتایج حساسیت روش ANFIS دقت بالاتری دارد، اما در نتایج RMSE، دقت مجموعه راف بالاتر است (۲۳)، که در این مطالعه نیز از روش نتایج حساسیت استفاده شده است.

نتیجه‌گیری

در مطالعه حاضر سعی شده تا با مقایسه روش‌های مبتنی بر هوش مصنوعی به پیش‌بینی بیماری دیابت پرداخته شود، تا از این طریق به حذف یا کاهش خطای انسانی در پیش‌بینی بیماری و هوشمند نمودن روند پیش‌بینی کمکی شود. در این مطالعه در ابتدا شبکه عصبی را برای پیدا کردن میزان و تعداد نرون مناسب آموزش داده، و در بین مدل‌ها، شبکه‌ی عصبی

منابع

- 1 Temurtas H, Yumusak N, Temurtas F. A comparative study on diabetes disease diagnosis using neural networks, *Expert Syst. Appl.* 2009; 36: 8610–15.
- 2 Nielsen D, Krych L, Buschard K. Beyond genetics. Influence of dietary factors and gut microbiota on type 1 diabetes, *FEBS Lett.* 2014; 588: 4234–43.
- 3 Pei E, Li J, Lu C, Xu J, Tang T, Ye M, et al. Effects of lipids and lipoproteins on diabetic foot in people with type 2 diabetes mellitus: a meta-analysis. *J Diabetes Complications.* 2014; 28: 559–64.
- 4 CDC. National diabetes fact sheet: national estimates and general information on diabetes and prediabetes in the United States, 2011, Atlanta, GA US Dep. Heal. Hum. Serv. Centers Dis. Control Prev. 2011 201. Available at: <http://scholar.google.com/scholar?hl=en&btnG=Search&q=i ntitle:National+diabetes+fact+sheet,+2011.+2011#25> (Accessed 2011)
- 5 Chavey A, Kioon M, Bailbé D. Maternal diabetes, programming of beta-cell disorders and intergenerational risk of type 2 diabetes *Diabetes.* 2014; 40 (5): 323–30.
- 6 Manzella D, Grella R, Abbatecola AM, Paolisso G, Repaglinide Administration Improves Brachial Reactivity in Type 2 Diabetic Patients. *Diabetes Care.* 2005; 28: 366–71.
- 7 Smith JW, Everhart JE, Dickson WC, Knowler WC, Johannes RS. Using the ADAP Learning Algorithm to Forecast the Onset of Diabetes Mellitus. *Proc. Annu. Symp. Comput. Appl. Med. Care.* 1998: 261–65.
- 8 Igor H, Bohuslava J, Martin J, Martin N. Application of Neural Networks in Computer Security, *Procedia Eng.* 2014; 69:1209–15.
- 9 Narasinga Rao MR, Sridhar GR, Madhu K, Rao AA. A clinical decision support system using multi-layer perceptron neural network to predict quality of life in diabetes, *Diabetes Metab. Syndr. Clin. Res. Rev.* 2010; 4(1): 57–59.
- 10 Gershenson C. Artificial Neural Networks for Beginners, *Networks*, vol. cs.NE/0308,, 2003, p:8.
- 11 Tu JV. Advantages and disadvantages of using artificial neural networks versus logistic regression for predicting medical outcomes. *J. Clin. Epidemiol.* 1996; 49: 1225–31.
- 12 Landau L. DM Concepts & Techniques Han&Kamber," *Zhurnal Eksp. i Teor. Fiz.* 1937, 40: 9823.
- 13 Knowles JD, Corne DW. M-PAES: A memetic algorithm for multiobjective optimization," in *Evolutionary Computation, 2000. Proceedings of the 2000 Congress on*, 2000; 1: 325–32.
- 14 Guan X, Zhang X, Han D, Zhu Y, Lv J, Su J. A strategic flight conflict avoidance approach based on a memetic algorithm, *Chinese J.* 2014; 27: 93–101.
- 15 Kin J. Feed-Forward Neural Networks and Genetic Algorithms for Automated Financial Time Series Modelling, London, October, 1995.
- 16 W.H.Yang YST. Design optimization of cutting parameters for tuning operation based on taguchi method., *J. Mater. Process. Technol.* 1997: 123–29.
- 17 Rodríguez JD, Pérez A, Lozano JA. Sensitivity analysis of kappa-fold cross validation in prediction error estimation. *IEEE Trans. Pattern Anal. Mach. Intell.* 2010; 32: 569–75.
- 18 Alba E, Chicano J. Training neural networks with GA hybrid algorithms. *Genet. Evol. Comput.* 2004: 852–63.
- 19 Kleffe J, Rao JN. Estimation of mean square error of empirical best linear unbiased predictors under a random error variance linear model. *J. Multivar. Anal.* 1992; 43: 1–15.
- 20 Hanley A, Mcneil J. The meaning and use of the area under a receiver operating characteristic (ROC) curve." *Radiology.* 1982; 143: 29–36.
- 21 Bradley AP. The use of the area under the ROC curve in the evaluation of machine learning algorithms. *Pattern Recognit.* 1997; 30: 1145–59.
- 22 Wang C, Li L, Wang L, Ping Z, Flory M. Evaluating the risk of type 2 diabetes mellitus using artificial neural network: An effective classification approach," *Diabetes Res.* 2013; 100: 111–8..
- 23 Gülçin Yıldırım E, Karahoca A, Uçar T. Dosage planning for diabetes patients using data mining methods. *Procedia Computer Science, Elsevier.* 2011; 3:1374–80.

Original Article

Comparison of Decision Support Systems for Diabetes Prediction

Jahani M¹, Rezaenoor J¹, Hadavani E², Tahsini H³, Iraj Salehi⁴

1- Department of Information Technology, University of Qom, Qom, Iran

2- Department of Textile Engineering, Amirabad University of Technology, Tehran

3- Department of Biostatistics & Epidemiology, University Hamadan of Medical Sciences, Hamadan, Iran

4- Department of Physiology, Hamadan University of Medical Sciences, Hamadan, Iran

Corresponding author: Rezaenoor J, rezaenoor@hotmail.com

Background & Objectives: In recent years, different decision support systems (DSS) have been used to predict and diagnose diseases. The purpose of this paper was to compare some DSSs and to evaluate their accuracy in predicting diabetes.

Methods: In this research, determination and optimization of the weights of the neural network were undertaken using genetic algorithm and Levenberg-Marquardt (GALM). Traditional and K-Fold Cross Validation were used to verify the models. Finally, the proposed model (i.e. GALM) was compared using logistic regression and genetic algorithm based on area under curve (AUC), and Confusion Matrix.

Results: After evaluating the results, the model based on the GALM algorithm showed better sensitivity and specificity in comparison with models based on the logistic regression (LR) and genetic algorithm (GA). Furthermore, among other models, the proposed model had a high sensitivity, specificity, negative predictive value (NPV), positive predictive value (PPV), and a small negative likelihood.

Conclusion: The results showed that the GALM model with a sensitivity, specificity, PPV, NPV, and AUC of 98.7, 90.01, 91.8, 98.3 and 0.979 respectively was an appropriate model for predicting diabetes in comparison with models of GA and LR.

Keywords: Diabetes, Neural Network, Genetic Algorithm, Levenberg-Marquardt