

استفاده از تکنیک‌های داده‌کاوی جهت تشخیص هوشمند شدت اختلال افسردگی

فرشته پارساپور^۱، جاوید پیمانی^{۲*}

• پذیرش مقاله: ۹۸/۱۱/۱۴

• دریافت مقاله: ۹۸/۴/۲۳

مقدمه: پیاده‌سازی روشی که بتواند هر فرد را در تشخیص یا پیشگیری اختلالات روانی یاری رساند می‌تواند گامی مهم در جهت پیشگیری و کنترل این اختلالات به خصوص در مراحل ابتدایی آن‌ها تلقی شود. هدف پژوهش حاضر به‌کارگیری تکنیک‌های داده‌کاوی در تشخیص هوشمند شدت اختلال افسردگی است.

روش: پژوهش کاربردی حاضر با مراجعه به تعدادی کلینیک روانپزشکی در شهر تهران و بررسی پرونده بیماران انجام شد. ۴۲۰ نفر که به پرسشنامه مینه سوتا ۷۱ سؤالی پاسخ داده بودند با استفاده از روش نمونه‌گیری در دسترس به عنوان نمونه انتخاب شدند (۳۰۰ نفر مبتلا به مراتب متفاوتی از افسردگی و ۱۲۰ نفر فاقد آن). پاسخنامه آزمون مینه‌سوتا ۷۱ سؤالی و تشخیص روانپزشک به عنوان داده برای ایجاد مدل توسط الگوریتم‌های K نزدیک‌ترین همسایه، درخت تصمیم و ماشین بردار پشتیبان مورد استفاده قرار گرفت. ۷۰ درصد داده‌ها برای آموزش و ۳۰ درصد داده‌ها برای اعتبارسنجی مدل به کار گرفته شد. جهت تحلیل داده‌ها از نرم‌افزار Matlab استفاده شد.

نتایج: نتایج ارزیابی‌ها نشان داد که الگوریتم درخت تصمیم با دقت ۹۹/۱۶ درصد به دقت بالاتری در مقایسه با دیگر الگوریتم‌ها دست یافت. همچنین با اجرای مدل‌های ایجاد شده بر روی هر سؤال آزمون مینه سوتا ۷۱ سؤالی تأثیر هر سؤال در ارزیابی مشخص شد.

نتیجه‌گیری: تقسیم‌بندی بیماران با رویکرد داده‌کاوی و بر اساس مهم‌ترین ویژگی‌ها، می‌تواند ابزار مفید و مؤثر برای تحلیل و بهبود فرآیند تصمیم‌گیری پزشکان در رابطه با درمان بیماران باشد.

کلید واژه‌ها: افسردگی، تشخیص هوشمند، داده‌کاوی

• **ارجاع:** پارساپور فرشته، پیمانی جاوید. استفاده از تکنیک‌های داده‌کاوی جهت تشخیص هوشمند شدت اختلال افسردگی. مجله انفورماتیک سلامت و زیست پزشکی ۱۳۹۹؛ ۷(۳): ۶۲-۲۵۲.

۱. کارشناسی ارشد روانشناسی بالینی، گروه روانشناسی بالینی، واحد کرج، دانشگاه آزاد اسلامی، البرز، ایران

۲. دکتری تخصصی علوم اعصاب، استادیار، گروه روان‌شناسی بالینی، واحد کرج، دانشگاه آزاد اسلامی، البرز، ایران

* **نویسنده مسئول:** جاوید پیمانی

آدرس: البرز، کرج، دانشگاه آزاد اسلامی، واحد کرج

• **شماره تماس:** ۰۲۶۳۴۲۵۹۵۷۱

• **Email:** Dr.Peymani@yahoo.com

مقدمه

در روانشناسی بالینی، تشخیص‌گذاری با مشاهده، آزمون و مصاحبه انجام می‌شود که در بین آن‌ها مصاحبه از اهمیت بیشتری برخوردار است. هدف از مصاحبه تشخیصی، یافتن دقیق علل و عوامل مؤثر در پیدایش اختلال و مشکل روانی در مراجع می‌باشد. در این رابطه مصاحبه باید بتواند کلیه مسائل و عوامل مؤثر در پیدایش اختلال را از کم اهمیت‌ترین تا پراهمیت‌ترین شناسایی کرده و بتواند تشخیصی مناسب در ارتباط با مشکل درمانجو در اختیار مصاحبه‌گر قرار دهد؛ زیرا تا زمانی که تشخیصی دقیق و درست از نوع اختلال بیمار در دست نباشد نمی‌توان درمان مناسبی برای او عرضه نمود [۱]. در انجام مصاحبه، روانشناس با مشکلاتی مواجه است که از اعتبار مصاحبه به‌عنوان یک وسیله ارزشیابی و تشخیصی می‌کاهد. این مشکلات ناشی از اشتباهاتی هستند که معلول فرآیند مصاحبه‌شونده و مصاحبه‌گر می‌باشند. اشتباهات ناشی از فرآیند مصاحبه را می‌توان معلول شبکه بغرنج و بسیار پیچیده کنش‌هایی دانست که در روابط چهره به چهره بین مصاحبه‌گر و مصاحبه‌شونده صورت می‌گیرد [۲]. موقعیتی که در آن مصاحبه صورت می‌گیرد نیز ممکن است خطاهایی را ایجاد کند. اشتباه دیگر ناشی از ترس‌ها و دلهره‌های مصاحبه‌شونده و نقش‌هایی است که ایفای آن از وی انتظار می‌رود که سازمان‌بخشی دقیق مصاحبه و پیش‌بینی این نقش‌ها تا حدودی از بروز این اشتباه می‌کاهد. اشتباه و مشکل دیگر ناشی از خود مصاحبه‌گر است که بیشتر معلول ویژگی‌های شخصی مصاحبه‌گر، اشتباه در یادداشت کردن جریان مطالعه، استنباط غلط، نداشتن آموزش کافی و مکتبی که مصاحبه‌گر از آن پیروی دارد، می‌باشد [۳]. از طرفی اختلالات روانی به دلیل گستردگی و دارا بودن علائم متعدد، پیچیدگی بالایی در شناسایی دارند به‌ویژه آن که برخی از این اختلالات دارای علائم مشترکی نیز هستند و این مسئله تشخیص دقیق و در نتیجه برنامه‌ریزی درمانی و آموزشی مناسب را برای آن‌ها دشوار می‌کند [۱]. تشخیص بیماری‌های روانی به دلیل نشانه‌های ذهنی و درونی آن پیچیده است و این پیچیدگی با توجه به این حقیقت که نشانه‌ها ممکن است پیش‌رونده باشند و تشخیص اولیه حتی توسط روانپزشک مجرب هم ممکن نباشد، بیشتر می‌شود. به‌علاوه یک ننگ اجتماعی وسیع وجود دارد که ممکن است بیماران روانی طرد شوند. کمبود متخصصان روان آموزش دیده و باتجربه در بسیاری از مناطق جهان باعث حادث شدن مشکل می‌شود. افسردگی یک رنج جدی است که بخش بزرگی از مردم جهان را تحت تأثیر قرار

می‌دهد و به علت علائم گسترده و مختلف، تشخیص مشکل به صورت غیرقطعی داده می‌شود [۴]. یک مطالعه که اخیراً در مجله بین‌المللی پزشکی JAMA منتشر شده است نشان داد که افسردگی بالینی، ضعیف تشخیص داده شده و غالباً بدون درمان سپری می‌شود. در بررسی ۴۶۰۰۰ آمریکایی، نشان داده شده است که ۸/۴ درصد افراد مصاحبه شونده افسردگی داشتند؛ ولی تنها ۲۸/۷ درصد آن‌ها انواع درمانی را دریافت کرده بودند. علاوه بر این از کسانی که برای افسردگی تحت درمان قرار گرفته بودند تنها ۲/۹ درصد از آن‌ها قطعاً افسردگی داشتند. بسیاری از افراد با حالاتی که قطعاً افسردگی نبود، با داروهای ضدافسردگی تحت درمان قرار داشتند [۵]. در مواردی مانند تشخیص افسردگی که عدم اطمینان نقش مهمی در آن ایفا می‌کند، یک الگوریتم هوشمند می‌تواند نسبت به روش‌های سنتی تشخیصی، سطوح بالاتری از اطمینان را فراهم کند. ارتباط و ضرورت وجود چنین سیستم‌های تشخیصی سودمند خودکاری برای کمک به مقابله با این تهدید دیده می‌شود [۴]. همکاری متخصصان در زمینه کامپیوتر و پزشکی راه‌حل جدیدی را در تحلیل داده‌های پزشکی و به‌دست آوردن الگوهای مفید و کاربردی ارائه می‌دهد که همان داده‌کاوی پزشکی است. استخراج دانایی از میان حجم انبوه داده‌های مرتبط با سوابق بیماری و پرونده‌های پزشکی افراد با استفاده از فرآیند داده‌کاوی می‌تواند منجر به شناسایی قوانین حاکم بر ایجاد و رشد بیماری‌ها شده و اطلاعات ارزشمندی را به‌منظور شناسایی علل رخداد بیماری‌ها، تشخیص، پیش‌بینی و درمان بیماری‌ها با توجه به عوامل محیطی حاکم در اختیار متخصصان حوزه سلامت قرار دهد؛ بنابراین پیچیدگی اطلاعات پزشکی و وجود ابزارهای داده‌کاوی باعث می‌شود که داده‌کاوی بر روی داده‌های پزشکی و سلامت مهم تلقی گردد [۶]. پژوهش‌های مختلفی با استفاده از روش‌های داده‌کاوی، انواع مختلفی از اختلالات افسردگی را پیش‌بینی کردند، Tortajada و همکاران [۷] در تحقیقی به پیش‌بینی افسردگی پس از زایمان با استفاده از پرسپترون چندلایه و بر اساس پیشینه روانپزشکی، تغییرات هیجانی در دوران بارداری، عصبی بودن، پرسشنامه افسردگی پس از زایمان ادینبرگ (Edinburgh Postnatal Depression Scale)، حمایت اجتماعی، خطر پری ناتال (Perinatal Risk) و اطلاعات جمعیت شناختی پرداختند و به این نتیجه رسیدند که پرسپترون چندلایه به عنوان یک مدل پیش‌بینی برای افسردگی پس از زایمان عملکرد خیلی خوبی دارد و با دقت ۸۱ درصد به درستی پیش‌بینی می‌کند. ملک‌پور و همکاران [۸] در تحقیقی به

روش

روش پژوهش حاضر کاربردی می‌باشد و پایه اصلی آن بر داده‌کاوی و تحلیل و بررسی روی داده‌های افسردگی موجود در پرونده مراجعه‌کنندگان به تعدادی کلینیک روان‌پزشکی در شهر تهران بنا نهاده شده است. در این مطالعه داده‌های مربوط به بیمارانی مورد بررسی قرار گرفت که آزمون مینه‌سوتا ۷۱ سؤالی فرم ایرانی بر روی آن‌ها اجرا شده بود. از بین پرونده‌های موجود با توجه به نتایج تفسیر آزمون مینه‌سوتا ۷۱ سؤالی توسط متخصص مربوطه، ۴۲۰ پرونده (۳۰۰ نفر (۷۱/۵ درصد) مبتلا به مراتب متفاوتی از افسردگی و ۱۲۰ نفر (۲۸/۵ درصد) فاقد آن) به روش نمونه‌گیری در دسترس انتخاب شد. دامنه سنی ۲۲۶ نفر از افراد بین ۱۶ تا ۳۰ سال و ۱۹۴ نفر بین ۳۰ تا ۶۰ سال قرار داشت. ۲۱۸ نفر (۵۲ درصد) زن و ۲۰۲ نفر (۴۸ درصد) مرد بودند. در مسائل مربوط به مدل‌سازی حجم نمونه مورد نیاز در محدوده $5M < n < 15M$ قرار می‌گیرد که در آن M تعداد متغیرهای مورد بررسی و n حجم نمونه می‌باشد [۱۱]. از آنجا که تعداد متغیرهای مورد مطالعه در این پژوهش، ۷۳ متغیر می‌باشد؛ بنابراین برای پژوهش کنونی بین ۳۶۵ تا ۱۰۹۵ نمونه لازم است. برای رعایت ملاحظات اخلاقی، نوع مطالعه و هدف آن برای کارکنان کلینیک‌ها توضیح و به آن‌ها اطمینان داده شد که اطلاعات گردآوری شده به صورت محرمانه محفوظ خواهد ماند.

آزمون مینه‌سوتا ۷۱ سؤالی شامل ۸ مقیاس بالینی و ۳ مقیاس روانی می‌باشد که در جدول ۱ آورده شده است [۱۲]. پاشا شریفی و نیکخو در تحقیقی که به منظور بررسی پایایی و روایی فرم کوتاه ایرانی این آزمون انجام دادند، آن را واجد شرایط لازم برای کاربرد در پژوهش‌های روانشناختی و فعالیت‌های بالینی دانسته‌اند (آلفای کرونباخ ۰/۹۶) [۱۳].

بررسی عوامل مؤثر بر میزان افسردگی افراد با استفاده از شبکه‌های عصبی پرسپترون چندلایه، درخت تصمیم، درخت چید، الگوریتم رافست و بر اساس متغیرهای جنسیت، تحصیلات، شغل، سن، وضعیت تأهل، تعداد فرزندان، میزان ورزش کردن، میزان گوش دادن به موسیقی، میزان اعتقادات مذهبی، میزان سفر کردن و وضعیت اقتصادی پرداختند. با اعمال هر مدل و الگوریتم، قوانینی تولید شدند که این قوانین کمک بسزایی در شناخت این که چه افرادی و با چه ویژگی‌هایی دارای افسردگی کم، متوسط یا شدید هستند، می‌کنند. Sau و همکاران [۹] با استفاده از شبکه عصبی مصنوعی و بر اساس متغیرهای جمعیت شناختی، مشکلات شنوایی، بینایی، تحرک، خواب و همچنین شاخص‌های پرسشنامه افسردگی سالمندان، وجود و یا عدم وجود افسردگی را در جمعیتی شامل ۱۰۵ سالمند، با دقت ۹۷/۲ درصد به‌درستی پیش‌بینی کردند. تیره و همکاران [۱۰] با استفاده از الگوریتم درخت تصمیم و بر اساس اطلاعات دموگرافیک، پرسشنامه اضطراب بک Beck Anxiety Inventory و پرسشنامه افسردگی بک Beck Depression Inventory به پیش‌بینی وجود و یا عدم وجود افسردگی در میان دانشجویان پرداختند، میزان صحت پیش‌بینی ۷۸ درصد به‌دست آمد. اگرچه مطالعات متعددی در این خصوص انجام شده است و به دقت‌های تشخیصی مطلوبی دست یافته‌اند؛ لیکن در این مطالعه با ارائه یک سیستم پزشکیار تلاش شده است ضمن تشخیص شدت افسردگی، تأثیر هر ویژگی در ارزیابی مشخص شود. این که هر ویژگی دارای چه ارزشی است و چقدر در تشخیص بیماری‌ها نقش دارد، می‌تواند کاربرد گسترده‌ای در زمینه تشخیص اختلالات روانی و آزمون‌های روان‌شناختی داشته باشد؛ لذا پژوهش حاضر با هدف تشخیص هوشمند شدت اختلال افسردگی با استفاده از روش‌های داده‌کاوی و نیز تعیین تأثیر هر ویژگی در ارزیابی انجام شد.

جدول ۱: مقیاس‌های آزمون مینه‌سوتا ۷۱ سؤالی

مقیاس	محتوای مواد
L	کشف تلاش عمده و ساده‌لوحانه آزمودنی برای نشان دادن چهره مطلوب از خود
F	اندازه‌گیری پاسخ‌های انحرافی و غیرعادی
K	تشخیص افرادی که تصویر کاملاً مطلوبی از خود ارائه می‌دهند
خودبیمارانگاری	اشتغال ذهنی با بدن و ترس از بیماری و مریضی
افسردگی	درجه بدبینی و غمگینی فرد را می‌سنجد و به خودپنداره ضعیف، اندوه، بدبینی و ناامیدی مربوط است
هیستری	تشخیص بیماران دچار اختلال حرکتی یا حسی روانزاد
انحراف اجتماعی	بی‌تفاوتی اجتماعی، فقدان سازگاری اجتماعی مانند مشکلات با خانواده، مقام‌های قدرت و بیگانگی اجتماعی
پارانویا	بدگمانی، حساسیت بین فردی و حق به جانب بودن
ضعف روانی	اضطراب، ترس، اعتماد به نفس پایین، حساسیت بی‌مورد و دمدمی بودن
اسکیزوفرنیا	بیگانگی و انزوای اجتماعی، ادراکات عجیب و غریب، شکوه‌های مربوط به بیگانگی با خانواده، دشواری در تمرکز و کنترل تکانه
هیپومانیا	خودمحوری، بزرگ‌منشی و تحریک‌پذیری

آماده سازی داده: در این مرحله داده ها به شکل مناسب جهت پردازش توسط ماشین در آمدند. در مطالعه حاضر پنج برچسب برای متغیر هدف در نظر گرفته شد که عبارت اند از: افراد فاقد افسردگی، افراد مستعد افسردگی، افراد مبتلا به افسردگی خفیف، افراد مبتلا به افسردگی متوسط و افراد مبتلا به افسردگی شدید. برچسب داده های مورد بررسی توسط متخصص مربوطه مشخص و مورد تأیید قرار گرفت. در جهان واقعی، داده همیشه کامل نیست و در مورد اطلاعات پرسشنامه ها این موضوع

همیشه درست است؛ لذا در این مرحله برخی از پاسخنامه ها که ناقص بودند، حذف شدند، تعداد بیماران مورد بررسی به ۴۰۰ مورد رسید. داده ها به صورت فایل اکسل تهیه گردید. در نهایت ۴۲۰ رکورد اولیه از بیماران پس از پالایش و حذف برخی رکوردها به ۴۰۰ رکورد نهایی کاهش یافت. هر رکورد دارای ۷۴ ویژگی می باشد که به همراه مقادیر آن ها در جدول ۲ آورده شده است.

جدول ۲: داده ها و مقادیر آن ها

ویژگی	مجموعه مقادیر
سن	۱۶-۶۰
جنسیت	۰ مرد ۱ زن
پاسخنامه آزمون مینه سوتا ۷۱ سؤالی	۰ خیر ۱ بلی
تشخیص پزشک (شدت افسردگی و عدم افسردگی)	۱ فاقد افسردگی ۲ مستعد افسردگی ۳ افسردگی خفیف ۴ افسردگی متوسط ۵ افسردگی شدید

مدل سازی: در این مرحله، تکنیک های مختلف مدل سازی انتخاب و به کار گرفته شد. در پژوهش حاضر با استفاده از تکنیک طبقه بندی به تولید مدل و الگوی بهینه ای برای تشخیص شدت افسردگی پرداخته شد. برای تولید مدل، داده های مورد بررسی به دو بخش آموزش و آزمایش تقسیم شدند. مدل ابتدا با استفاده از ۷۰ درصد داده ها به عنوان داده های آموزش، آموزش دید و داده های بخش آزمایش که ۳۰ درصد باقی رکوردها بودند، مدل تولید شده را تست و برچسب مربوط به رکوردهای مذکور را تعیین نمودند. با توجه به این که مجموعه داده ها به صورت دودویی و متغیر هدف دارای پنج برچسب است، نیاز به الگوریتم طبقه بندی است که بتواند با برچسب های غیر دودویی کار کند. بدین دلیل برای مدل سازی از الگوریتم های درخت تصمیم، ماشین بردار پشتیبان و K نزدیک ترین همسایه استفاده شد.

درخت تصمیم: درخت تصمیم ساختاری شبیه به درخت دارد که قابلیت یادگیری دارد. الگوریتم های مختلفی برای درخت تصمیم وجود دارد که این الگوریتم ها یا به تولید یک درخت دودویی، یا به تولید یک درخت غیر دودویی می پردازند [۱۴].

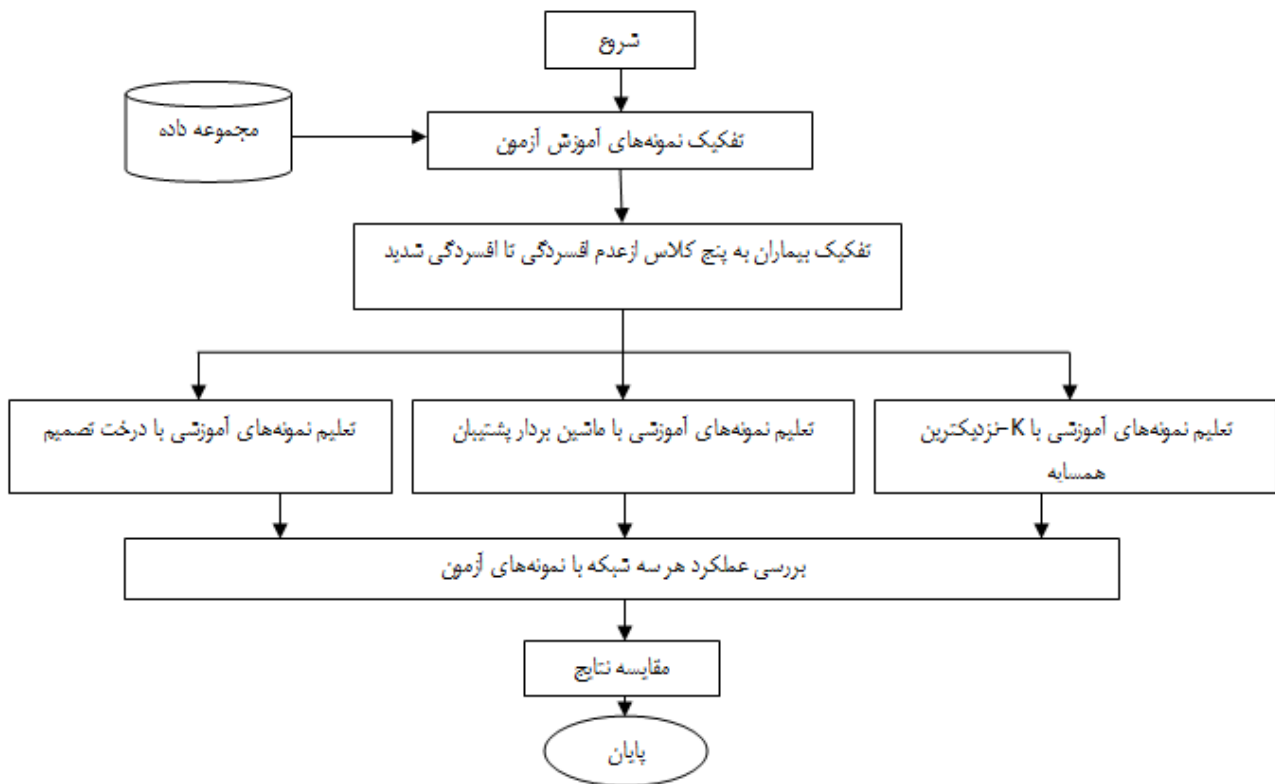
ماشین بردار پشتیبان: الگوریتم ماشین بردار پشتیبان از یک نگاشت غیر خطی برای تبدیل داده های آموزشی اصلی درون یک فضا با ابعاد بزرگ تر استفاده می کنند؛ به عبارت دیگر می توان گفت با استفاده از یک نگاشت غیر خطی داده های آموزشی را از فضای اصلی به فضای ویژگی منتقل می کند. حال در این بعد جدید، الگوریتم به دنبال ابر صفحه جدا کننده بهینه است که بتواند رکوردهای یک کلاس را از کلاس دیگر جدا کند [۱۴].

K-نزدیک ترین همسایه: طبقه بند K-نزدیک ترین همسایه یک طبقه بند یادگیری با ناظر می باشد که نمونه ها را بر اساس شباهت و فاصله از نمونه های آموزشی طبقه بندی می کند. برای یک داده آزمایشی الگوریتم به دنبال K نمونه از نزدیک ترین نمونه ها می گردد. نزدیکی دو نمونه، با به دست آوردن تشابه و یا فاصله میان این دو نمونه محاسبه می شود. هر نمونه می تواند از انواع داده ها تشکیل شده باشد که باید تشابه میان آن ها بررسی شود. پس از یافتن این K داده مشابه با نمونه آزمایشی، رأی اکثریت تعیین کننده برچسب کلاس داده آزمایشی می باشد [۱۵].

در پژوهش حاضر مدل سازی با استفاده از نرم افزار Matlab R2015a انجام شد.

افسردگی متوسط، افسردگی شدید به کمک ۳ نوع الگوریتم طبقه‌بند مختلف به تشخیص شدت افسردگی پرداخته و نتایج با هم مقایسه شده است.

مراحل مدل‌سازی در این پژوهش در شکل ۱ ارائه شده است. همان‌طور که الگوریتم نشان می‌دهد در این پژوهش پس از تفکیک نمونه‌های آزمایش و آموزش و نیز تقسیم آن‌ها به پنج کلاس عدم افسردگی، مستعد افسردگی، افسردگی خفیف،



شکل ۱: مراحل مدل‌سازی با استفاده از الگوریتم‌های پژوهش حاضر

همچنین FN(False Negative) و FP (False Positive) به ترتیب نرخ مثبت کاذب و منفی کاذب را نشان می‌دهند. منحنی مشخصه عملکرد، علاوه بر شاخص‌های فوق یکی از روش‌های مقایسه قدرت مدل‌های مختلف پیش‌بینی کننده است. منحنی ترکیب (شفافیت-۱، حساسیت) منحنی مشخصه عملکرد نامیده می‌شود. مساحت زیر این منحنی یک شاخص ترکیبی است که نشان می‌دهد مدل با چه احتمالی موقعیت مثبت را نسبت به موقعیت منفی انتخاب می‌کند. بیشترین حد رقم این شاخص ۱ و کمترین حد آن برای انتخاب تصادفی ۰/۵ است. مقادیر بیش از ۸۰ درصد مساحت زیر منحنی نشان دهنده توان بالای مدل در پیش‌بینی است [۱۹].

ارزیابی مدل: عملکرد الگوریتم‌های داده‌کاوی اغلب با استفاده از یک ماتریس درهم‌ریختگی اندازه‌گیری می‌شود. ماتریس درهم‌ریختگی شامل اطلاعات ارزشمندی درباره رده‌بندی‌های واقعی و پیش‌بینی شده توسط مدل رده‌بندی است. بر مبنای اطلاعات حاصل از این ماتریس شاخص‌های مقایسه مناسبی محاسبه می‌شود [۱۶]. شاخص‌های مختلفی مانند شفافیت (Specificity)، حساسیت (Sensitivity)، دقت (Accuracy) و صحت (Precision) برای ارزیابی روش‌های خوشه‌بندی وجود دارند که طبق روابط در جدول ۳ محاسبه می‌گردند [۱۷، ۱۸]. TP(True Positive) نرخ مثبت واقعی و TN(True Negative) نرخ منفی واقعی را نشان می‌دهد.

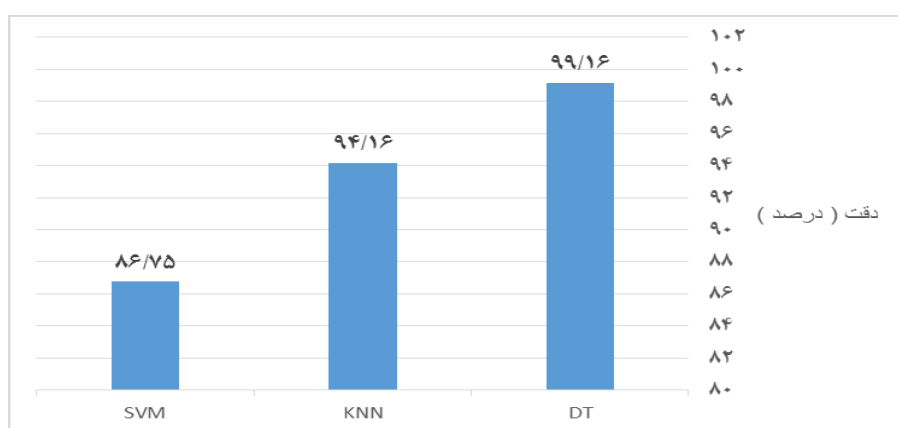
جدول ۳: نحوه محاسبه دقت، صحت، حساسیت و شفافیت

دقت	$\text{Accuracy} = \frac{TP+TN}{TP+FN+FP+TN}$
صحت	$\text{Precision} = \frac{TP}{TP+FP}$
حساسیت	$\text{Sensitivity} = \frac{TP}{TP+FN}$
شفافیت	$\text{Specificity} = \frac{TN}{FP+TN}$

نتایج

دقت طبقه‌بندی‌های درخت تصمیم، ماشین بردار پشتیبان و K نزدیک‌ترین همسایه بر روی مجموعه داده‌ای توصیف شده در این مطالعه مورد آزمایش قرار گرفت. شکل ۲ نشان می‌دهد که

الگوریتم درخت تصمیم با دقت ۹۹/۱۶ و الگوریتم K نزدیک‌ترین همسایه با دقت ۹۴/۱۶ در مقایسه با الگوریتم ماشین بردار پشتیبان به دقت بالاتری در تشخیص شدت افسردگی دست یافته‌اند.



شکل ۲: مقایسه دقت الگوریتم‌های طبقه‌بندی

در جدول ۴، الگوریتم‌های طبقه‌بندی بر اساس معیارهای دقت، صحت، حساسیت، شفافیت و سطح زیر منحنی مشخصه عملکرد مقایسه شده‌اند. الگوریتم درخت تصمیم بر اساس همه معیارها

بهبتر از سایر الگوریتم‌ها عمل کرده است؛ بنابراین الگوریتم درخت تصمیم مدلی مناسب برای پیش‌بینی شدت افسردگی است.

جدول ۴: مقایسه طبقه‌بندی‌های مورد استفاده بر اساس شاخص‌های مختلف

الگوریتم	دقت (درصد)	صحت (درصد)	حساسیت (درصد)	شفافیت (درصد)	سطح زیر منحنی مشخصه عملکرد (درصد)
K- نزدیک‌ترین همسایه	۹۴/۱۶	۸۶/۰۳	۹۴/۷۰	۹۰/۳۱	۹۲/۵۱
درخت تصمیم	۹۹/۱۶	۹۸/۹۴	۹۹/۶۲	۹۸/۰۱	۹۸/۸۱
ماشین بردار پشتیبان	۸۶/۷۵	۸۲/۵۱	۹۴/۴۵	۷۵/۷۹	۸۵/۱۲

جداول ۵ و ۶ ارزش هر ویژگی را به ترتیب با استفاده از الگوریتم‌های درخت تصمیم و k نزدیک‌ترین همسایه نشان می‌دهند. برای محاسبه دقت یک مشخصه مطابق روش توصیف شده در بخش مدل‌سازی عمل شده است. به این صورت که مدل ایجاد شده بر روی هر مشخصه به صورت مجزا اجرا شد و دقت هر مشخصه بر روی داده‌های مجموعه آزمایش محاسبه

شد. دقت یک مشخصه عبارت است از درصد داده‌ای مجموعه آزمایش که برچسب عدم افسردگی، مستعد افسردگی، افسردگی خفیف، افسردگی متوسط و افسردگی شدید آن به درستی توسط آن مشخصه پیش‌بینی شده باشد. با توجه به دقت هر مشخصه در سنجش کلاس‌ها، کلاسی که هر مشخصه با بالاترین دقت می‌سنجد نیز آشکار شد.

جدول ۵: تأثیر هر سؤال در ارزیابی با استفاده از الگوریتم درخت تصمیم

کلاس	دقت (درصد)	مشخصه	کلاس	دقت (درصد)	مشخصه	کلاس	دقت (درصد)	مشخصه	کلاس	دقت (درصد)	مشخصه
۲	۹۰	۵۵	۲	۹۰	۳۷	۳	۹۱/۶۷	۱۹	۲	۹۰	۱
۲	۸۶/۶۷	۵۶	۲	۸۷/۵۰	۳۸	۳	۹۳/۳۳	۲۰	۲	۸۸/۳۳	۲
۱	۹۲/۵۰	۵۷	۲	۸۹/۱۷	۳۹	۲	۹۰	۲۱	۲	۸۸/۳۳	۳
۲	۹۰	۵۸	۲	۸۸/۳۳	۴۰	۳	۹۰	۲۲	۲	۹۰	۴
۱	۸۹/۱۷	۵۹	۱	۸۵/۸۳	۴۱	۲	۸۸/۳۳	۲۳	۲	۸۸/۳۳	۵
۲	۸۸/۳۳	۶۰	۲	۹۰	۴۲	۲	۸۶/۶۷	۲۴	۲	۸۹/۱۷	۶
۲	۹۳/۳۳	۶۱	۱	۸۵/۸۳	۴۳	۲	۸۷/۵۰	۲۵	۲	۸۹/۱۷	۷
۳	۸۷/۵۰	۶۲	۱	۹۲/۵۰	۴۴	۲	۹۰/۸۳	۲۶	۲	۸۷/۵۰	۸
۲	۸۷/۵۰	۶۳	۱	۸۷/۵۰	۴۵	۳	۸۶/۶۷	۲۷	۲	۹۰	۹
۲	۸۶/۶۷	۶۴	۱	۸۶/۶۷	۴۶	۳	۸۵/۸۳	۲۸	۳	۹۰	۱۰
۲	۸۷/۵۰	۶۵	۲	۸۹/۱۷	۴۷	۳	۹۱/۶۷	۲۹	۲	۸۵/۸۳	۱۱
۱	۹۰	۶۶	۲	۸۷/۵۰	۴۸	۲	۹۰	۳۰	۲	۹۱/۶۷	۱۲
۲	۸۸/۳۳	۶۷	۲	۸۶/۶۷	۴۹	۲	۸۷/۵۰	۳۱	۲	۸۵/۸۳	۱۳
۱	۹۰	۶۸	۲	۹۰	۵۰	۲	۹۱/۶۷	۳۲	۲	۸۹/۱۷	۱۴
۱	۹۰	۶۹	۲	۹۰	۵۱	۲	۸۶/۶۷	۳۳	۳	۸۹/۱۷	۱۵
۲	۸۸/۳۳	۷۰	۳	۸۹/۱۷	۵۲	۲	۸۵/۸۳	۳۴	۲	۸۹/۱۷	۱۶
۲	۹۱/۶۷	۷۱	۲	۸۵/۸۳	۵۳	۲	۸۶/۶۷	۳۵	۲	۸۸/۳۳	۱۷
۱	۹۰/۸۳	۷۲	۲	۹۰	۵۴	۲	۹۰/۸۳	۳۶	۲	۹۰	۱۸
۲	۹۰	۷۳									

جدول ۶: تأثیر هر سؤال در ارزیابی با استفاده از الگوریتم K نزدیک ترین همسایه

کلاس	دقت (درصد)	مشخصه	کلاس	دقت (درصد)	مشخصه	کلاس	دقت (درصد)	مشخصه	کلاس	دقت (درصد)	مشخصه
۲	۹۳/۳۳	۵۵	۲	۹۳/۳۳	۳۷	۳	۹۳/۳۳	۱۹	۲	۸۶/۶۷	۱
۲	۸۶/۶۷	۵۶	۲	۹۳/۳۳	۳۸	۳	۹۶/۶۷	۲۰	۲	۹۰	۲
۱	۹۰	۵۷	۲	۹۰	۳۹	۲	۸۶/۶۷	۲۱	۲	۹۳/۳۳	۳
۲	۹۳/۳۳	۵۸	۲	۹۰	۴۰	۳	۹۳/۳۳	۲۲	۲	۹۳/۳۳	۴
۱	۹۳/۳۳	۵۹	۱	۹۰	۴۱	۲	۹۳/۳۳	۲۳	۲	۹۰	۵
۲	۹۰	۶۰	۲	۹۳/۳۳	۴۲	۲	۸۶/۶۷	۲۴	۲	۹۳/۳۳	۶
۲	۹۳/۳۳	۶۱	۱	۸۶/۶۷	۴۳	۲	۸۶/۶۷	۲۵	۲	۹۰	۷
۳	۸۶/۶۷	۶۲	۱	۹۳/۳۳	۴۴	۲	۹۳/۳۳	۲۶	۲	۹۰	۸
۲	۸۶/۶۷	۶۳	۱	۸۶/۶۷	۴۵	۳	۹۰	۲۷	۲	۹۳/۳۳	۹
۲	۸۶/۶۷	۶۴	۱	۸۶/۶۷	۴۶	۳	۹۰	۲۸	۳	۹۰	۱۰
۲	۸۶/۶۷	۶۵	۲	۹۶/۶۷	۴۷	۳	۹۳/۳۳	۲۹	۲	۹۰	۱۱
۱	۸۶/۶۷	۶۶	۲	۹۳/۳۳	۴۸	۲	۹۳/۳۳	۳۰	۲	۹۶/۶۷	۱۲
۲	۹۰	۶۷	۲	۹۰	۴۹	۲	۹۰	۳۱	۲	۹۰	۱۳
۱	۹۶/۶۷	۶۸	۲	۹۳/۳۳	۵۰	۲	۹۳/۳۳	۳۲	۲	۹۳/۳۳	۱۴
۱	۹۰	۶۹	۲	۹۳/۳۳	۵۱	۲	۹۰	۳۳	۳	۹۰	۱۵
۲	۸۶/۶۷	۷۰	۳	۸۶/۶۷	۵۲	۲	۸۶/۶۷	۳۴	۲	۹۳/۳۳	۱۶
۲	۹۶/۶۷	۷۱	۲	۸۶/۶۷	۵۳	۲	۹۰	۳۵	۲	۹۰	۱۷
۱	۹۰	۷۲	۲	۹۳/۳۳	۵۴	۲	۹۰	۳۶	۲	۹۳/۳۳	۱۸
۲	۹۳/۳۳	۷۳									

شد. این که هر یک از ویژگی‌ها دارای چه ارزشی است و چقدر در تشخیص بیماری نقش دارد، مسئله مهمی است. در این مطالعه روشی ارائه شد که دقت هر یک از ویژگی‌ها در سنجش افسردگی و همچنین شدتی از افسردگی که هر ویژگی با بالاترین دقت می‌سنجد به طور دقیق مشخص شد.

با توجه به اهمیت تشخیص اختلال افسردگی، مطالعاتی توسط پژوهشگران در این حوزه انجام گرفته است. بختیاری و همکاران [۱] در مطالعه‌ای با استفاده از الگوریتم‌های درخت تصمیم، رگرسیون خطی، ماشین بردار پشتیبان، K نزدیک‌ترین همسایه بیزین ساده، شبکه عصبی و آدابوست و بر اساس پرسشنامه‌ای محقق ساخته شامل متغیرهایی از جمله ملامت، عدم لذت، خواب، اشتها، وضعیت روانی حرکتی، انرژی، احساس گناه و سرزنش، تمرکز و دودلی، افکار خودکشی، زمان، عملکرد فرد در امور مهم زندگی، مصرف مواد مخدر، بیماری‌های جسمی که علائم شبه افسردگی دارند، هذیان، توهم، صدمه‌های مهم زندگی، اختلالات دیگر، وراثت، سابقه اختلالات روانی و نیز تحریف‌های شناختی وجود و یا عدم وجود افسردگی را در نمونه‌ای با ۲۰۰۰ نفر پیش‌بینی کردند. در این پژوهش شبکه عصبی با دقت ۹۸ درصد به عنوان مدل بهینه انتخاب شد. تعداد قابل ملاحظه نمونه‌ها در مدل‌سازی از مزایای این پژوهش بود؛

دقت تمامی ویژگی‌ها در سنجش شدت افسردگی با استفاده از الگوریتم ماشین بردار پشتیبان ۸۶/۷۵ به دست آمد.

بحث و نتیجه‌گیری

در این پژوهش به تشخیص شدت اختلال افسردگی با استفاده از تکنیک‌های داده کاوی پرداخته شد. تکنیک‌های استفاده شده در این پژوهش، الگوریتم‌های درخت تصمیم، ماشین بردار پشتیبان و K نزدیک‌ترین همسایه است. با توجه به نتایج ارزیابی، دقت الگوریتم‌های درخت تصمیم و K نزدیک‌ترین همسایه به ترتیب ۹۹/۱۶ و ۹۴/۱۶ درصد است که در مقایسه با الگوریتم‌های ماشین بردار پشتیبان به دقت بیشتری رسیده‌اند.

برای هر کدام از طبقه‌بندهای مذکور علاوه بر دقت، معیارهای صحت، حساسیت و ویژگی محاسبه گردید. نتایج ارزیابی نشان داد الگوریتم درخت تصمیم بر اساس اغلب معیارها بهتر از سایر الگوریتم‌ها عمل کرده است و به صحت ۹۸/۹۴، حساسیت ۹۹/۶۲، شفافیت ۹۸/۰۱ و سطح زیر منحنی مشخصه عملکرد ۹۸/۸۱ درصد دست یافته است که برای تشخیص افسردگی قابل قبول به نظر می‌رسد.

در این مطالعه علاوه بر تشخیص شدت افسردگی، روشی برای تشخیص افسردگی با استفاده از مهم‌ترین ویژگی‌ها معرفی

همچنین این مطالعه به دقت قابل ملاحظه‌ای در تشخیص رسیده است. در پژوهش حاضر پنج حالت از نبود افسردگی تا افسردگی شدید پیش‌بینی شد و تأثیر هر ویژگی در تشخیص آشکار شد.

تیره و همکاران [۱۰] با استفاده از الگوریتم درخت تصمیم و بر اساس پرسشنامه اطلاعات دموگرافیک، پرسشنامه‌های اضطراب و افسردگی بک به پیش‌بینی وجود و یا عدم وجود افسردگی در میان دانشجویان پرداختند، میزان صحت پیش‌بینی ۷۸ درصد به دست آمد. مطالعه حاضر با دقت ۹۹/۱۶ با استفاده از الگوریتم درخت تصمیم و با به کارگیری پرسشنامه مینه‌سوتا پنج حالت از نبود افسردگی تا افسردگی شدید را پیش‌بینی نمود.

Mukherjee و همکاران [۲۰] با استفاده از شبکه عصبی مصنوعی و بر اساس ده متغیر از نشانه‌های رایج افسردگی از جمله احساس ناراحتی، فقدان شادی، کاهش وزن، بی‌خوابی و بی‌اشتهایی وجود و یا عدم وجود افسردگی را با خطای ۰/۱۱ تشخیص دادند. تعداد کم متغیرها برای مدل‌سازی از معایب این روش است. در مطالعه حاضر از پرسشنامه مینه‌سوتا که از معتبرترین ابزار سنجش شدت اختلالات روانی می‌باشد جهت جمع‌آوری داده‌ها استفاده شد.

Suhara و همکاران [۲۱] در مطالعه‌ای با استفاده از الگوریتم‌های ماشین بردار پشتیبان و شبکه عصبی حافظه کوتاه‌مدت طویل و بر اساس داده‌های حاصل از خود گزارش‌دهی مجموعه‌ای شامل ۲۳۸۲ نفر در مدت ۳۴۵۱۵۸ روز در رابطه با متغیرهای خلق در ساعات مختلف روز، نوع فعالیت و زمان خواب، خلق افسرده پیش‌بینی کردند. نتایج نشان داد مدل پیشنهادی با سطح زیر منحنی مشخصه عملکرد ۸۸/۶ درصد نسبت به ماشین بردار پشتیبان عملکرد بهتری دارد. تعداد قابل ملاحظه نمونه‌ها در مدل‌سازی از مزایای این پژوهش بود؛ اما پیچیدگی بالای زمانی از معایب این روش بود. چون مطالعه حاضر سعی در طراحی و ارزیابی یک مدل پزشک‌یار در تشخیص شدت افسردگی دارد، پیچیدگی زمانی پایین از اهداف این روش می‌باشد.

Osubor و همکاران [۲۲] در مطالعه‌ای با استفاده از الگوریتم شبکه تطبیقی عصبی‌فازی (Adaptive Neural Fuzzy Inference System) ANFIS و بر اساس داده‌های حاصل از پرسشنامه سنجش افسردگی پس از زایمان (Postpartum Depression Screening Scale) PDSS شدت افسردگی پس از زایمان را در نمونه‌ای شامل ۵۹ نفر با دقت ۹۷ درصد

به‌درستی پیش‌بینی کردند و به این نتیجه رسیدند که مدل پیشنهادی در مقایسه با الگوریتم شبکه عصبی به تنهایی، کارایی بیشتری دارد. استفاده از رویکرد فازی از مزایای این پژوهش بود؛ اما تعداد کم نمونه‌ها از معایب این مطالعه است. در مطالعه حاضر تعداد قابل قبولی از بیماران به کار گرفته شد.

به طور خلاصه نقاط قوت روش پیشنهادی به شرح زیر است:
۱- دقت قابل ملاحظه در تشخیص طیفی از عدم افسردگی تا افسردگی شدید با استفاده از تکنیک‌های پیشنهادی.

۲- تعیین ارزش هر ویژگی در تشخیص افسردگی به منظور ایجاد سیستمی پزشک‌یار در تشخیص به‌موقع و با صرف کمترین زمان و هزینه.

۳- هوشمندسازی آزمون مینه‌سوتا ۷۱ سؤالی به شکل تعیین ارزش هر سؤال در ارزیابی با هدف صرفه‌جویی در زمان، هزینه و جلوگیری از پاسخ‌های تصادفی و بی‌حوصلگی بیماران و در نهایت رسیدن به تشخیص دقیق.

۴- روش پیشنهادی کاملاً خودکار است و نیاز به دخالت انسان ندارد.

با توجه به نتایج پژوهش، در پایگاه داده‌ها که شامل مجموعه‌ای از ارزیابی‌های بالینی بیماران و نهایتاً برچسب اختلال برای هر بیمار می‌باشد، برای نمونه جدید با به کارگیری پژوهش حاضر تنها کافی است مجموعه‌ای از ارزیابی‌هایی که دقت بالایی در تشخیص اختلال دارند، بررسی شوند، به این صورت با تشخیص به‌موقع اختلالات روانی می‌توان از عوارض ناشی از آن‌ها کاست و یا پیشگیری کرد.

از جمله محدودیت‌هایی که طی انجام تحقیق حاضر وجود داشت، مشکلات فراوان در مسیر جمع‌آوری داده بود. ثبت سیستمی اطلاعات بیماران به جای استفاده از پرونده می‌تواند تا حد زیادی از مشکلات جمع‌آوری و آماده‌سازی داده و نیز خطای انسانی ناشی از ثبت دستی اطلاعات بکاهد و باعث صرفه‌جویی در زمان شود.

پیشنهاد می‌شود در تحقیقات آینده از سایر مشخصه‌های بالینی و آزمون‌های روانشناختی جهت تشخیص هوشمند دیگر اختلالات روانی و هوشمندسازی آزمون‌های روانشناختی استفاده شود. همچنین می‌توان دیگر الگوریتم‌های داده‌کاوی را نیز برای رسیدن به این هدف مورد استفاده قرار داد.

تعارض منافع

در پژوهش حاضر هیچ‌گونه تعارض منافی وجود ندارد.

References

1. Bakhtiyari M, Esmaeilpuor M, Ebrahimi ME. Diagnose of Depression Using Artificial Intelligence. 6th Congress of Iranian Psychological Association; 2017 Nov 13-15; Tehran: Iranian Psychological Association; 2017. [In Persian].
2. Babbie ER. The Practice of Social Research. Translated by Faazel R. Tehran: Samt; 2016. [In Persian].
3. Rafiepoor F. Searches and thoughts. Tehran: publication Joint-stock company; 2015. [In Persian].
4. Mukherjee S, Ashish K, Hui N, Chattopadhyay S. Modeling Depression Data: Feed Forward Neural Network vs Radial Basis Function Neural Network. American Journal of Biomedical Science 2014; 6(3): 166-74. doi: 10.5099/aj140300166
5. Rowley R. Can Artificial Intelligence help in the treatment of depression? 2016 [cited 2016 Nov 8]. Available from: <https://flowhealth.com/blog/2016/11>.
6. Jooriyan N, Ashoori M. Predicting the effectiveness of preeclampsia medications based on dose and method of drug consumption using data mining. Iranian Journal of Obstetrics, Gynecology and Infertility 2014; 17(123): 13-22. [In Persian]
7. Tortajada S, García-Gomez JM, Vicente J, Sanjuán J, de Frutos R, Martín-Santos R, et al. Prediction of postpartum depression using multilayer perceptrons and pruning. Methods Inf Med 2009;48(3):291-8. doi: 10.3414/ME0562
8. Malekpour S, Esmaeilpour M. Investigating Effective Factors on Depression in Individuals Using Data Mining. 2nd International Conference on Knowledge-Based Research in Computer Engineering & IT; 2017 Sep 21; Tehran: Kharazmi University; 2017. [In Persian].
9. Sau A, Bhakta I. Artificial Neural Network (ANN) Model to Predict Depression among Geriatric Population at a Slum in Kolkata, India. J Clin Diagn Res 2017;11(5):VC01-VC04. doi: 10.7860/JCDR/2017/23656.9762
10. Tire H, Jafari H, Farbod D. Depression prediction in students using decision tree algorithm. Summary of Articles of the 12th Annual Research Students Congress of the Eastern Medical Sciences Universities; 2017 Nov 29; Gonabad: Gonabad University of Medical Sciences; 2017. [In Persian].
11. Hooman HA. Structural Equation Modeling Using LISREL Software. Tehran: Samt; 2005. [In Persian]
12. Khodayaryfard M, Parand A. Psychological Assessment and Testing. Tehran: Tehran University Press; 2009. [In Persian].
13. Groth-Marnat G, Wright AJ. Handbook of Psychological Assessment. Translated by Pasha Sharif H, Nikkhoo MR. Tehran: Sokhan, Roshd; 2010. [In Persian].
14. Han J, Kamber M. Data Mining. Southeast Asia: Morgan Kaufmann; 2006.
15. Fang X. Are you becoming a diabetic? a data mining approach. Sixth International Conference on Fuzzy Systems and Knowledge Discovery; 2009 Aug 14-16; Tianjin, China: IEEE; 2009. doi: 10.1109/FSKD.2009.807
16. Kohavi R, Provost F. Glossary of terms, editorial for the special issue on applications of machine learning and the knowledge discovery process. Machine Learning 1998; 30(2-3): 271-4. <https://doi.org/10.1023/A:1017181826899>
17. AlizadehSani R. Diagnosis of heart disease using data mining [dissertation]. Tehran: Sharif University of Technology; 2012. [In Persian]
18. Han J, Kamber M, Pei J. Data Mining: Concepts and Techniques. 3rd ed. USA: Morgan Kaufman; 2011.
19. SPSS. Clementine12 User Manual. Chicago: IL; 2007.
20. Mukherjee S, Ashish K, Baran Hui N, Chattopadhyay S. Modeling depression data: feed forward neural network vs. radial basis function neural network. Am J Biomed Sci 2014;6(3):166-74. doi: 10.5099/aj140300166
21. Suhara Y, Xu Y, Pentland A. DeepMood: Forecasting Depressed Mood Based on Self-Reported Histories via Recurrent Neural Networks. 26th International Conference on World Wide Web; 2017 Apr 3-7; Perth, Australia: ACM; 2017. p. 715-24. doi: 10.1145/3038912.3052676
22. Osubor VI, Egwali AO. A neuro fuzzy approach for the diagnosis of postpartum depression disorder. Iran Journal of Computer Science 2018; 1: 217-25.

Using Data Mining Techniques for Intelligent Diagnosis of Severity of Depressive Disorder

Parsapour Fereshteh¹, Peymani Javid^{2*}

• Received: 14 July 2019

• Accepted: 03 Feb 2020

Introduction: Implementing a method that can help individuals diagnose or prevent mental disorders can be an important step in preventing and controlling these disorders especially in the early stages. The objective of this research was to apply data mining techniques for intelligent diagnosis of severity of depressive disorder.

Method: The present applied research was carried out by going to a number of psychiatric clinics in Tehran and investigating patients' medical records. A total of 420 subjects who responded to the Minnesota Multiphasic Personality Inventory (MMPI, 71 questions) were selected through convenience sampling as the sample of the study (300 subjects were diagnosed with different degrees of depression and 120 subjects showed no symptoms of depression). The answer sheet of MMPI and the diagnosis of the psychiatrist were used as data for developing the model by k--Nearest Neighbor (k-NN), Decision Tree, and Support Vector Machine algorithms. About 70 percent of the data were applied for training and 30 percent of the data were used for validating the model. MATLAB software was used for data analysis.

Results: The results of the evaluations showed that Decision Tree algorithm with accuracy of 99.16% had higher accuracy compared to other algorithms. Furthermore, by implementing the developed models on each question of MMPI, the influence of each question on evaluation was determined.

Conclusion: Classifying patients with data mining approach and based on the most important characteristics can be a useful and effective tool for analyzing and improving the decision-making process of physicians regarding the treatment of patients.

Keywords: Depression, Intelligent Diagnosis, Data Mining

• **Citation:** Parsapour F, Peymani J. Using Data Mining Techniques for Intelligent Diagnosis of Severity of Depressive Disorder. Journal of Health and Biomedical Informatics 2020; 7(3): 252-62. [In Persian]

1. M.A. in Clinical Psychology, Clinical Psychology Dept., Karaj Branch, Islamic Azad University, Alborz, Iran

2. Ph.D. in Neuroscience, Assistant Professor, Clinical Psychology Dept., Karaj Branch, Islamic Azad University, Alborz, Iran

• **Corresponding Author:** Javid Peymani

Address: Karaj Branch, Islamic Azad University, Karaj, Alborz, Iran

• **Tel:** 026-34259571

• **Email:** Dr.Peymani@yahoo.com