

## پیش‌بینی بیماری دیابت با استفاده از ماشین بردار پشتیبان کلاسیکی و کوانتومی

شقایق شاه‌محمدی<sup>۱</sup>، حمیدرضا باغشاهی<sup>۱\*</sup> و حسین داودی‌یگانه<sup>۲</sup>

۱. گروه فیزیک، دانشکده علوم، دانشگاه ولی عصر (عج) رفسنجان، رفسنجان

۲. مرکز محاسبات کوانتومی آریاکوانتا، تهران

پست الکترونیکی: baghshahi@vru.ac.ir

(دریافت مقاله: ۱۴۰۴/۰۷/۰۵؛ دریافت نسخه نهایی: ۱۴۰۴/۱۰/۰۷)

### چکیده:

قندخون یکی از شاخص‌های تأثیرگذار بر سازوکار سوخت‌وساز بدن است. پیش‌بینی دقیق و پیشگیری این شاخص، از اهمیت زیادی برخوردار است. در این راستا، تحلیل و بررسی الگوریتم‌های یادگیری ماشین کوانتومی برای پردازش داده‌های پزشکی و افزایش دقت الگوهای پیش‌بینی در حوزه سلامت، از زمینه‌های پیشرو در مطالعات اخیر به شمار می‌آید. الگوریتم‌های یادگیری ماشین کوانتومی، با تکیه بر اصول فیزیک و محاسبات کوانتومی، جایگزینی مؤثر برای روش‌های کلاسیکی در پیش‌بینی بیماری‌های مزمن محسوب می‌شوند. این الگوریتم‌ها با نگاشت داده‌های کلاسیکی به فضای کوانتومی و استفاده از ویژگی‌هایی مانند برهم‌نهی و درهم‌تنیدگی، قابلیت شناسایی الگوهای پنهان و پیچیده در داده‌ها را به‌طور چشم‌گیری افزایش می‌دهند. در این پژوهش، ابتدا داده‌ها پس از پیش‌پردازش با روش تحلیل مولفه‌های اصلی (PCA) پردازش شده‌اند. سپس عملکرد ماشین بردار پشتیبان کوانتومی (QSVM) و مدار کوانتومی وردشی (VQC) در پیش‌بینی قند خون نسبت به روش کلاسیک SVM، بررسی و مقایسه شده است. در نتیجه الگوی QSVM با دقت ۹۰ درصد، عملکرد بهتری نسبت به SVM کلاسیکی با دقت ۷۴/۴ درصد نشان داده است، درحالی‌که الگوی آمیخته کلاسیکی-کوانتومی VQC با ثبت دقت ۹۴/۲ درصد بالاترین سطح دقت و کارایی را در پیش‌بینی سطح قندخون ارائه کرده است.

**واژه‌های کلیدی:** الگوریتم یادگیری ماشین کوانتومی، قندخون، ماشین بردار پشتیبان کوانتومی، مدار کوانتومی وردشی

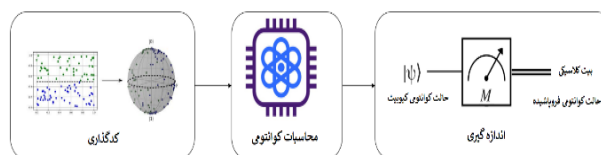
### ۱. مقدمه

است و انتظار می‌رود این روند ادامه‌دار باشد. بنابراین، روش‌های پیش‌بینی و مدیریت قندخون به اولویتی اساسی در علم پزشکی تبدیل شده است [۴]. محاسبات کوانتومی و روش‌های یادگیری ماشین کوانتومی باعث بهبود عملکرد الگوها و همچنین افزایش سرعت تشخیص بیماری و کاهش زمان پیش‌بینی از چندین سال به چند دقیقه شده است [۵-۷]. در روش یادگیری ماشین کلاسیک<sup>۱</sup>، الگوهای بین داده‌ها شناسایی شده و بدون نیاز به برنامه‌ریزی صریح، فرآیند تصمیم‌گیری هوشمندانه‌ای ارائه می‌شود. یکی از الگوریتم‌های رایج در این حوزه، ماشین بردار پشتیبان با هسته<sup>۲</sup> است، که با

دیابت یکی از چالش‌های اصلی سلامت جهانی است که ناشی از نقص در تولید یا استفاده مؤثر از انسولین است [۱]. این بیماری، هفتمین عامل مرگ و میر جهانی به شمار می‌رود و سالانه ۱/۶ میلیون نفر را به کام مرگ می‌کشاند [۲]. از چهار نوع اصلی دیابت، دیابت نوع ۱ (T1D) و نوع ۲ (T2D) شایع‌تر هستند و می‌توانند عوارض جدی همچون سندرم پای دیابتی، نارسایی کلیه و سکتة مغزی ایجاد کنند [۳]. هزینه‌های درمان دیابت بین سال‌های ۲۰۰۷ تا ۲۰۱۹، ۲۲۸ درصد افزایش یافته

۱. Classic Machine Learning (CML)

۲. Support Vector Machine with Kernel (Kernel SVM)



شکل ۱. کاربرد کامپیوتر کوانتومی بر روی یک مجموعه داده کلاسیکی [۱۴].

با توجه به شکل ۱ مشاهده می‌شود که الگوریتم‌های کوانتومی با کدگذاری داده‌ها در حالت‌های کوانتومی و پردازش آن‌ها از طریق مدارهای کوانتومی، قابلیت تحلیل داده‌ها را امکان‌پذیر می‌سازند.

در این پژوهش، با استفاده از روش کلاسیکی، پیش‌بینی قندخون الگوسازی شده است. سپس با استفاده از الگوهای کوانتومی، دقت پیش‌بینی، مقایسه و ارتقا داده شده است. در ادامه، در بخش سوم، به بررسی مزایا و محدودیت‌های این روش‌ها و همچنین تحلیل و مقایسه نتایج الگوها پرداخته شده است تا برتری و چالش‌های هر یک تحلیل شود. در نهایت، در بخش چهارم جمع‌بندی و نتیجه‌گیری پژوهش ارائه شده و عملکرد الگوهای کلاسیکی با الگوهای کوانتومی ارزیابی و مقایسه شده است. همچنین، پیشنهادهایی برای پژوهش‌های آینده با هدف بهبود کارایی الگوهای کوانتومی در کاربردهای عملی و مسائل دنیای واقعی ارائه شده است. این مطالعه با بهره‌گیری از فناوری‌های نوین و محاسبات کوانتومی، گامی در راستای ارتقای کیفیت زندگی و سلامت عمومی جامعه برداشته است.

## ۲. روش کار

در ابتدا، داده‌ها از نمونه‌های پزشکی افراد از طریق آزمایش‌های استاندارد در یک آزمایشگاه آسیب‌شناسی به عنوان مرجع بیمارستانی جمع‌آوری شده‌اند. این داده‌ها شامل ویژگی‌های زیستی مرتبط و عوامل مؤثر بر سطح قندخون و دیابت هستند. مجموعاً ۱۰۰ نمونه واقعی که به صورت تصادفی به دو بخش ۷۰ نمونه (هفتاد درصد) برای آموزش الگو و ۳۰ نمونه (سی درصد) برای ارزیابی عملکرد الگو تقسیم شده‌اند.

ویژگی‌های مورد استفاده در این پژوهش شامل جنسیت، سن، قد، وزن، فعالیت بدنی، رژیم غذایی، سطح استرس، کیفیت خواب و وراثت خانوادگی است.

نگاشت داده‌های غیرخطی به فضایی با ابعاد بالاتر، امکان تفکیک خطی آن‌ها را فراهم می‌کند [۸]. در مقابل، روش یادگیری ماشین کوانتومی<sup>۱</sup> با استفاده از کدگذاری داده‌ها<sup>۲</sup> در حالت‌های کوانتومی انجام شده و پردازش از طریق مدارهای کوانتومی صورت می‌گیرد که در نهایت، تحلیل داده‌های پیچیده را امکان‌پذیر می‌سازد.

الگوریتم‌های کوانتومی با پیچیدگی محاسباتی چندجمله‌ای نسبت به تعداد کیوبیت‌ها، قادرند داده‌ها را در فضایی با ابعاد بالا پردازش کنند. این موضوع موجب افزایش بهره‌وری و کاهش زمان پردازش می‌شود [۹]. روش‌های آمیخته‌ای مانند الگوریتم‌های وردشی<sup>۳</sup>، با کدگذاری داده‌ها در کیوبیت‌ها و استفاده از درهم‌تنیدگی، تعداد پارامترها را کاهش داده و باعث بهبود الگوسازی می‌شود [۱۰ و ۱۱]. این رویکرد که شامل ترکیب روش‌های کلاسیکی و کوانتومی است می‌تواند دقت پیش‌بینی دیابت را افزایش دهد [۱۲]. الگوریتم‌هایی از جمله ماشین بردار پشتیبان کوانتومی<sup>۴</sup> و مدارهای کوانتومی وردشی<sup>۵</sup>، با استفاده از ویژگی‌های منحصربه‌فرد فضای هیلبرت و درهم‌تنیدگی کوانتومی، امکان شناسایی الگوهای پیچیده در داده‌هایی با حجم زیاد را نیز فراهم می‌کند. یادگیری ماشین کوانتومی و تحلیل ماتریس‌های هم‌بستگی، می‌توانند الگوهای مرتبط با قندخون را شناسایی کنند تا پزشکان و بیماران را در مدیریت بهتر آن یاری کند [۱۳].

در این راستا، شرکت‌های بزرگی مانند Google، IBM و D-Wave با ارائه دسترسی ابری به ماشین‌های کوانتومی، نقش مهمی در توسعه فناوری‌های جدید و ارتقاء کیفیت یادگیری ماشین دارند.

۱. Quantum Machine Learning (QML)

۲. Encoding

۳. Variational Algorithms

۴. Quantum Support Vector Machine (QSVM)

۵. Variational Quantum Circuits (VQCs)

جدول ۱. آمار ویژگی‌های زیستی بیماران.

| شماره | ویژگی‌ها       | محدوده  | میانگین | انحراف معیار |
|-------|----------------|---------|---------|--------------|
| ۱     | جنسیت          | ۰-۱     | ۰/۲۹    | ۰/۴۶         |
| ۲     | سن             | ۶-۸۳    | ۴۴/۳۵   | ۱۶/۵۵        |
| ۳     | قد             | ۱۰۷-۱۸۷ | ۱۶۳/۶۰  | ۱۰/۴۷        |
| ۴     | وزن            | ۷-۹۷    | ۶۸/۴۷   | ۱۲/۶۹        |
| ۵     | فعالیت بدنی    | ۰-۱     | ۰/۳۵    | ۰/۴۸         |
| ۶     | رژیم غذایی     | ۰-۱     | ۰/۲۸    | ۰/۴۵         |
| ۷     | سطح استرس      | ۰-۱     | ۰       | ۰            |
| ۸     | کیفیت خواب     | ۰-۱     | ۰/۶۳    | ۰/۴۸         |
| ۹     | وراثت خانوادگی | ۰-۱     | ۰/۴۲    | ۰/۴۹         |

در جدول ۱، آمار توصیفی این ویژگی‌های زیستی (میانگین و انحراف معیار) جهت تحلیل و تشخیص الگوهای بیماری، سازمان‌دهی و ارائه شده است. ویژگی‌های فعالیت بدنی، رژیم غذایی، سطح استرس، کیفیت خواب و وراثت خانوادگی به صورت مقادیر دودویی (۰-۱) کدگذاری شده‌اند.

این ساده‌سازی برای هم‌خوانی داده‌ها با الگوهای کوانتومی که تعداد کیوبیت محدود دارند، انجام شده است. برای هر ویژگی استانداردهای مشخصی تعیین شده است. به عنوان مثال، برای کیفیت خواب، در صورت داشتن ۷ ساعت خواب منظم، بدون بیداری‌های مکرر با مقدار ۱، یعنی حالت مطلوب این ویژگی در نظر گرفته شده است. در مورد سطح استرس، عدم اضطراب مداوم و توانایی در کنترل هیجان و عواطف با مقدار ۰ تعیین شده است. به دلیل این که روابط بین ویژگی‌های زیستی پیچیده و غیرخطی هستند، پیش‌بینی قندخون افراد بر اساس ۹ ویژگی زیستی، یک مسئله غیرخطی با داده‌های غیرخطی است.

این داده‌ها و ویژگی‌ها، نقش مهمی در تشخیص دیابت دارند و انتخاب مؤثرترین ویژگی‌ها، می‌تواند نقش بسزایی در نتیجه پیش‌بینی داشته باشد.

در این پژوهش، مدارهای کوانتومی با استفاده از زبان برنامه‌نویسی پایتون و کتابخانه کیس‌کیت<sup>۱</sup>، بر روی شبیه‌ساز

qasm-simulator اجرا شده‌اند. لازم به ذکر است که به دلیل محدودیت دسترسی به پلتفرم ابری IBM Quantum، استفاده از شبیه‌سازهای ابری مانند ibm-qasm-simulator یا سخت‌افزارهای کوانتومی امکان‌پذیر نبوده است.

با این حال، دستورهای استفاده شده در صورت دسترسی به محیط ابری، قابل انتقال و اجرا بر روی ibm-qasm-simulator یا دستگاه‌های واقعی هستند. همچنین پیش از اجرای الگوهای یادگیری ماشین کوانتومی و کلاسیکی، فرآیند پیش‌پردازش داده‌ها شامل مراحل استانداردسازی و کاهش ابعاد انجام شده است.

داده‌های نوع CQ<sup>۲</sup>، امکان استفاده از قابلیت‌های کوانتومی برای پردازش داده‌های کلاسیکی را فراهم می‌کنند. به منظور سازگاری با محدودیت تعداد کیوبیت‌های قابل دسترسی در سخت‌افزارهای کوانتومی فعلی و کاهش پیچیدگی محاسباتی، از روش تحلیل مؤلفه‌های اصلی<sup>۳</sup> برای کاهش ابعاد داده‌ها استفاده شده است. یکی از مراحل مهم در شناسایی الگوها، مقایسه ویژگی‌های زیستی در افراد است. بدین منظور، توزیع ویژگی‌های زیستی بیماران، در دو گروه دیابتی و غیردیابتی در شکل ۲ مقایسه شده است. با توجه به این شکل، وراثت خانوادگی و سن به عنوان مهم‌ترین عوامل مؤثر در تشخیص دیابت شناسایی شدند. سن به عنوان یک شاخص مهم، نشان می‌دهد که افراد دیابتی اکثراً در گروه سنی بالاتر قرار دارند. همچنین وراثت خانوادگی نیز اهمیت زیادی دارد، زیرا توزیع چگالی این ویژگی در افراد دیابتی به‌طور واضح نمایانگر ابتلا به دیابت است. تحلیل این مشاهدات می‌تواند در بهبود الگوهای پیش‌بینی و مدیریت بهتر این بیماری نقش مهمی داشته باشد.

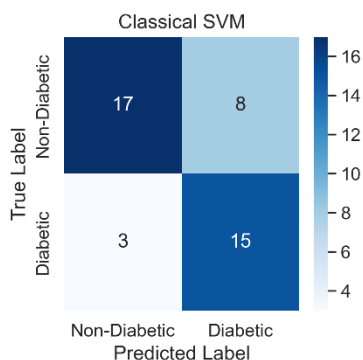
## ۲. ۱. ماشین بردار پشتیبان کلاسیک

ماشین بردار پشتیبان کلاسیک یکی از روش‌های یادگیری ماشین کلاسیک با ناظر است که برای حل مسائل رگرسیون و طبقه‌بندی استفاده می‌شود. هدف اصلی این روش، تعیین ابرصفحه‌ای بهینه با بیشترین فاصله ممکن، برای تفکیک داده‌ها

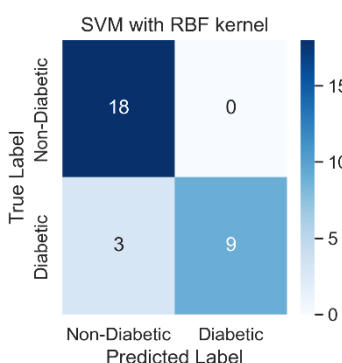
۱. Qiskit

۲. Classical Quantum

۳. Principal Component Analysis (PCA)

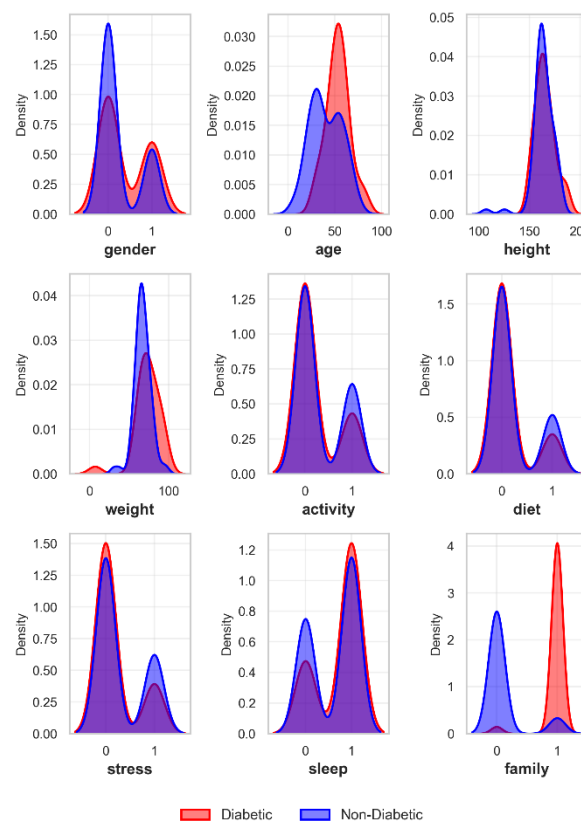


شکل ۳. ماتریس درهم‌ریختگی برای ماشین بردار پشتیبان کلاسیکی با هسته خطی.



شکل ۴. ماتریس درهم‌ریختگی SVM کلاسیکی با هسته غیرخطی (RBF).

پیچیدگی محاسباتی الگوی SVM با افزایش حجم داده‌ها، افزایش می‌یابد و برای داده‌ها با حجم زیاد استفاده از روش‌های جایگزین، ضرورت دارد [۱۵]. با توجه به شکل ۳، الگوی SVM کلاسیکی در تشخیص دیابت عملکرد قابل قبولی داشته است. در محاسبات این الگو، ۷۰ درصد داده‌ها برای آموزش الگو و ۳۰ درصد باقیمانده برای ارزیابی عملکرد آن اختصاص داده شده است. این تقسیم‌بندی به صورت تصادفی انجام شده تا قابلیت تعمیم الگو به داده‌های جدید حفظ شود. از میان افراد دیابتی تعداد ۱۵ نمونه و از میان افراد غیر دیابتی، ۱۷ نمونه به درستی شناسایی شده‌اند که این موضوع نشان‌دهنده حساسیت بالای الگو در تشخیص دیابت است. لازم به ذکر است که ابتدا از ماشین بردار پشتیبان SVM با هسته خطی به عنوان یک الگوی پایه کلاسیک استفاده شده است. به منظور بهبود عملکرد الگو و فراهم کردن مقایسه‌ای منصفانه‌تر با الگوی کوانتومی QSVM، که ماهیت غیرخطی آن از نگاشت ویژگی‌های کوانتومی در فضای هیلبرت ناشی می‌شود، در



شکل ۲. مقایسه توزیع چگالی ویژگی‌های زیستی بین افراد دیابتی و غیردیابتی.

به دسته‌های جداگانه است. از بردارهای پشتیبان به عنوان نقاطی که بیشترین فاصله را از مرز تصمیم‌گیری دارند برای تعیین این ابرصفحه استفاده می‌شود. این روش، ساختار داده‌ها را به خوبی تحلیل کرده و مرزهای تصمیم‌گیری بهینه را تعیین می‌کند. از آنجایی که در دنیای واقعی داده‌ها به ندرت به صورت خطی قابل تفکیک هستند، یادگیری ماشین کلاسیک (CML) برای حل مسائل طبقه‌بندی غیرخطی، از توابع هسته استفاده می‌کند. این توابع با نگاشت داده‌ها از طریق تابع هسته به فضای ویژگی با ابعاد بالاتر، امکان جداسازی آن‌ها را به صورت غیرخطی فراهم می‌کند و موجب ارتقاء عملکرد الگو در شناسایی الگوهای پیچیده می‌شود. همچنین، این توابع برای محاسبه ضمنی حاصل ضرب داخلی داده‌ها در فضای ویژگی با ابعاد بالا استفاده می‌شود که معمولاً با استفاده از رابطه (۱) نمایش داده می‌شود:

$$k(x_i, x_j) = \langle \phi(x_i), \phi(x_j) \rangle \quad (1)$$

به طوری که  $\phi$  نگاشت غیرخطی ویژگی و  $k$  نمایانگر ماتریس هسته است.

کوانتومی (ZZFeatureMap) و همچنین استفاده از ساختار درهم‌تنیدگی به حالت‌های کوانتومی تبدیل می‌شوند. نگاشت کوانتومی ZZFeatureMap برای کدگذاری داده‌های کلاسیکی به حالت‌های کوانتومی استفاده می‌شود. این نگاشت، هر ویژگی عددی را به صورت زاویه‌ای در گیت‌های چرخشی  $R_z$  و  $R_y$  کدگذاری می‌کند. به علاوه، با اعمال گیت CNOT بین تمام زوج کیوبیت‌ها، درهم‌تنیدگی ایجاد می‌کند. این نگاشت به دلیل ایجاد هسته‌های غیرخطی، برای مسئله طبقه‌بندی در یادگیری ماشین کوانتومی استفاده می‌شود [۱۷].

در این نگاشت، با توجه به رابطه (۲)، هر ویژگی به عنوان یک پارامتر در گیت‌های دورانی تک‌کیوبیتی اعمال می‌شود:

$$|\phi(x)\rangle = U_{\text{feature}}(x)|0\rangle^{\otimes n} \quad (2)$$

این عمل، داده‌های کلاسیک را در فضای هیلبرت به ابعاد بالا می‌برد و همین امر، امکان الگوسازی روابط غیرخطی پیچیده را ممکن می‌سازد. در مرحله بعد، باتوجه به شکل ۵، مدار کوانتومی وردشی سه کیوبیتی با پارامترهای قابل تنظیم طراحی شده است. این مدار پس از کدگذاری ویژگی‌ها، برای محاسبه هسته کوانتومی استفاده می‌شود. در الگوی QSVM استاندارد، مقدار هسته به عنوان هم‌پوشانی بین حالت‌های کوانتومی تعریف می‌شود. این رویکرد، با بهره‌گیری از اصول مکانیک کوانتومی مانند درهم‌تنیدگی و تداخل، امکان تحلیل پیچیده‌تری را نسبت به الگوهای کلاسیکی فراهم می‌کند. این برهم‌نهی اجازه می‌دهد هر کیوبیت هم‌زمان در ترکیبی از حالت‌های  $|0\rangle$  و  $|1\rangle$  قرار گیرد. این ویژگی، امکان کدگذاری هم‌زمان چندین ویژگی را فراهم کرده و فضای هیلبرت را به گونه‌ای گسترش می‌دهد که روابط غیرخطی پیچیده‌تری بین داده‌ها قابل الگوسازی باشد. همچنین درهم‌تنیدگی نیز با ایجاد هم‌بستگی کوانتومی بین کیوبیت‌ها، وابستگی‌های غیرکلاسیک بین ویژگی‌ها را در فضای ویژگی کوانتومی نشان می‌دهد. در نگاشت ZZ هر زوج کیوبیت با یکدیگر درهم‌تنیده می‌شوند، این امر به الگو اجازه می‌دهد تعاملات پیچیده بین ویژگی‌های داده (مانند تأثیر استرس یا خواب بر میزان قندخون) را شناسایی کند. این همان چیزی است که در الگوهای کلاسیک به راحتی قابل دستیابی نیست.

الگوی SVM از هسته شعاعی  $RBF^1$  استفاده شده است. این انتخاب به معنای رسیدن به بهترین نتیجه نهایی نیست، زیرا ممکن است هسته‌های دیگر ارزیابی بهتری داشته باشند. اما به طور کلی در مسائل غیرخطی برای مقایسه با الگوهای یادگیری ماشین کوانتومی از این نوع هسته استفاده می‌شود. به منظور ارزیابی تأثیر نوع هسته بر دقت طبقه‌بندی، عملکرد SVM با هسته خطی و غیرخطی RBF در بخش نتایج و بحث، تجزیه و تحلیل شده است.

با توجه به شکل ۴ می‌توان عملکرد الگوی کلاسیکی را با معیارهای مشخص ارزیابی کرد. نتایج حاصل از ماتریس درهم‌ریختگی، نشان می‌دهد که الگوی SVM با هسته غیرخطی، در طبقه‌بندی داده‌ها عملکرد نسبتاً خوبی داشته است. از مجموع داده‌های آزمایشی، ۱۸ نمونه غیر دیابتی به درستی شناسایی شده‌اند. از طرفی، این الگو موفق شده است ۹ نمونه دیابتی را نیز به درستی تشخیص دهد. این نتایج نشان‌دهنده تعادل نسبی دقت الگو در طبقه‌بندی است. اگرچه عملکرد الگو در تشخیص نمونه‌های غیردیابتی دقت پایین‌تری نسبت به شناسایی نمونه‌های دیابتی دارد، اما این اطلاعات می‌تواند در تحلیل بهتر نقاط قوت و ضعف این الگو مؤثر باشد.

## ۲.۲. ماشین بردار پشتیبان کوانتومی

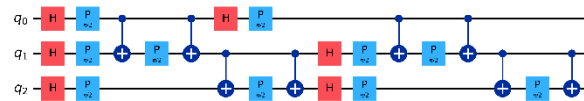
ماشین بردار پشتیبان کوانتومی روش یادگیری با ناظر است و برای دسته‌بندی داده‌ها در مسائل دودویی و چنددسته‌ای کاربردی است. در این روش، محاسبات کوانتومی با بهره‌گیری و استفاده از مدارهای پارامتریک، موازی‌سازی و درهم‌تنیدگی کوانتومی، امکان نگاشت داده‌ها به فضای هیلبرت با ابعاد بالا را فراهم می‌کند که در الگوهای کلاسیک قابل دسترس نیستند. این الگوریتم، با ترکیب فناوری‌های پیشرفته کلاسیکی و کوانتومی، دقت و کارایی در دسته‌بندی داده‌ها را بهبود می‌بخشد. همچنین، با تعریف مرزهای تصمیم‌گیری غیرخطی، قابلیت تحلیل داده‌های پیچیده و استخراج الگوهای پنهان در آن‌ها را به شکل قابل توجهی افزایش می‌دهد [۱۶]. پس از پیش‌پردازش داده‌ها و کاهش ابعاد با استفاده از PCA به ۳ ویژگی، داده‌های کلاسیکی با استفاده از نگاشت ویژگی

حالت کیوبیت‌ها، و گیت  $P$  به عنوان بخشی از ساختار وردشی برای اعمال فاز کنترل شده، پیاده‌سازی شده است. همچنین، پارامترهای قابل تنظیم این گیت‌ها امکان پیش‌بینی و طبقه‌بندی داده‌ها در چارچوب الگوی موردنظر را فراهم می‌کنند.

پس از انجام اندازه‌گیری روی حالت‌های کوانتومی، نتایج حاصل به شکل کلاسیکی استخراج و مورد تحلیل قرار می‌گیرند. نقاط داده جدید بر اساس فاصله خود از بهینه‌ترین ابرصفحه جداکننده طبقه‌بندی می‌شوند. این فرآیند امکان استخراج ویژگی‌های مؤثر از داده‌ها را فراهم می‌کند و فضای تصمیم‌گیری بهینه‌ای برای الگوریتم QSVM ایجاد می‌شود. همچنین، این رویکرد نشان‌دهنده کاربرد فضاهای کوانتومی در بهبود دقت محاسباتی و بهره‌گیری از قابلیت‌های محاسبات کوانتومی در الگوهای طبقه‌بندی و پیش‌بینی است. نتایج به‌دست آمده نشان می‌دهند که این الگو توانسته است عملکرد بهتری نسبت به روش‌های مرسوم کلاسیکی ارائه دهد.

پس از نگاشت ویژگی کوانتومی، همان‌طور که در شکل ۶ مشاهده می‌شود، ماتریس هسته مربوط به دست می‌آید که بیانگر میزان شباهت و ارتباط میان داده‌ها در فضای کوانتومی است. این ماتریس نه تنها شباهت بین داده‌ها را در فضای کوانتومی نشان می‌دهد، بلکه تعامل ویژگی‌های را نیز در فضای هیلبرت تحلیل می‌کند. همچنین ماتریس هسته، هم‌پوشانی حالت‌ها را بین تمامی داده‌های آموزشی در فضای هیلبرت نمایش می‌دهد. مقدار هر مولفه  $K_{ij}$ ، برابر با رابطه (۲) است و نشان‌دهنده درجه شباهت بین ورودی‌های  $X_i$  و  $X_j$  از طریق نگاشت کوانتومی  $\phi$  است.

تحلیل این ماتریس نشان می‌دهد نمونه‌های مربوط به طبقه‌بندی (دیابتی - غیردیابتی) روی قطر هستند و شباهت بیشتری به هم دارند. توانایی نگاشت کوانتومی در افزایش تفکیک‌پذیری داده‌ها در فضای ویژگی کوانتومی دارای اهمیت است. وجود این ساختار منظم در ماتریس هسته بیانگر آن است که هسته کوانتومی می‌تواند الگوهای معنادار موجود در داده‌ها را کشف کند و به الگوی QSVM کمک کند تا مرز تصمیم‌گیری دقیق‌تری را یاد بگیرد.



شکل ۵. مدار کوانتومی با ۳ کیوبیت و نگاشت ZZFeatureMap، شامل گیت‌های  $R_z$ ،  $R_y$  و CNOT و گیت فاز  $P$  در لایه برگشت پذیر طراحی شده است.

#### Kernel Matrix

```
[[1.03142844 0.14741287 0.11693693 ... 0.50324084 0.29497751 0.10166008]
 [0.14741287 1.03809719 0.12826009 ... 0.14716816 0.29171495 0.1144895 ]
 [0.50324084 0.12826009 1.0274556 ... 0.10780368 0.06090624 1.00082107]
 ...
 [0.50324084 0.14716816 0.10780368 ... 1.032281318 0.27772061 0.09395596]
 [0.29497751 0.29171495 0.06090624 ... 0.27772061 1.03242273 0.07128574]
 [0.10166008 0.11448905 1.00082107 ... 0.09395596 0.07128574 1.02518419]]
```

شکل ۶. ماتریس هسته کوانتومی.

در نهایت این دو خاصیت باهم، منجر به ایجاد یک هسته کوانتومی با قابلیت تفکیک‌پذیری بالاتر نسبت به هسته‌های کلاسیک می‌شوند و همین عامل عملکرد الگوی QSVM را در تشخیص دیابت افزایش می‌دهد.

در یک QSVM استاندارد، هم‌پوشانی بین حالت‌های کوانتومی<sup>۱</sup> به عنوان مقدار هسته به صورت معادله زیر تعریف می‌شود:

$$k(x_i, x_j) = |\langle \phi(x_i), \phi(x_j) \rangle|^2 \quad (3)$$

در این پژوهش، داده‌ها با استفاده از یک عملگر یکانی به نام  $U_{\phi(x)}$  و از طریق کدگذاری زاویه‌ای، کدگذاری شده‌اند. سپس هر ویژگی ورودی، به وسیله گیت‌های دورانی تک‌کیوبیتی، به یک زاویه چرخشی بر روی یک کیوبیت نگاشته شده است. با داشتن  $n$  ویژگی، بردار ورودی به صورت رابطه زیر کدگذاری می‌شود:

$$|X\rangle = U_{\phi(x)} |0\rangle^{\otimes n} \quad (4)$$

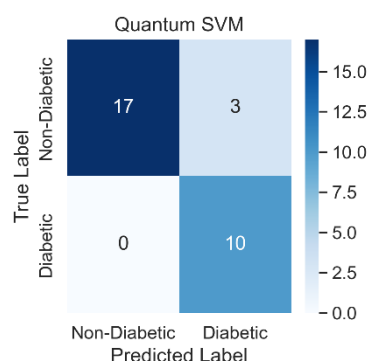
داده‌ها ابتدا پس از فرآیند کاهش ابعاد و تبدیل به فضای ویژگی، با استفاده از نقشه ویژگی کوانتومی آماده‌سازی شده‌اند. در مرحله بعد، مدار کوانتومی بر اساس شکل ۵ با بهره‌گیری از سه کیوبیت و کدگذاری داده‌ها در فضای هیلبرت طراحی شده است. این مدار کوانتومی با استفاده از سه کیوبیت و اعمال گیت‌های چرخشی و CNOT برای ایجاد درهم‌تندگی بین کیوبیت‌ها، به کارگیری گیت هادامارد برای ایجاد و بازیابی

مهمی ایفا می‌کنند. این بهینه‌سازی توسط رایانه‌های کلاسیکی انجام می‌شود، در حالی که محاسبات اصلی در محیط کوانتومی صورت می‌گیرد. این فرآیند با کاهش خطاها از طریق اندازه‌گیری‌های پی‌درپی و کنترل نوفه، دقت طبقه‌بندی و پیش‌بینی را افزایش می‌دهد [۱۹ و ۲۰].

در این رویکرد برخلاف QSVM که به ایجاد مرزها در قالب یک الگوریتم مداری ترکیبی وابسته است، مدارهای کوانتومی وردشی (VQCs)، مقادیر چشم‌داشتی حاصل از اندازه‌گیری‌های کوانتومی را جمع‌آوری کرده و پس از اعمال یک آستانه مشخص، برای پیش‌بینی نتایج مورد استفاده قرار می‌گیرند. در نتیجه، عملکرد خود را با استفاده از روش بهینه‌سازی وردشی بهبود می‌بخشد.

همان‌طور که در شکل ۸ نشان داده شده است، برای پیاده‌سازی روش VQC در زمینه پیش‌بینی قندخون و طبقه‌بندی افراد به دو دسته دیابتی و غیردیابتی، ابتدا فرآیند پیش‌پردازش داده‌ها انجام می‌گیرد. در این راستا، ۹ ویژگی اولیه به سه مؤلفه اصلی کاهش یافته، به طوری که هر کیوبیت نمایانگر یکی از ویژگی‌های زیستی داده‌ها است و بیش از ۹۰ درصد از واریانس کل داده‌های اصلی حفظ شده است. این مقدار واریانس نشان می‌دهد که اطلاعات مهم برای تشخیص الگوی دیابت تا حد قابل قبولی حفظ شده است.

پس از این مرحله، داده‌های کاهش‌یافته به‌عنوان ورودی به الگوریتم کوانتومی VQC که از یک مدار کوانتومی وردشی با عمق کم استفاده می‌کند، داده شده است. همچنین، برای مقایسه بهتر، از همین داده‌های پیش‌پردازش‌شده برای آموزش الگوهای کلاسیک مانند SVM استفاده شده است. با ترکیب نگاشت ویژگی کوانتومی و یک مدار پارامتری، الگو قادر به یادگیری الگوهای پیچیده در داده‌های غیرخطی است. با تنظیم عمق ثابت مدار کوانتومی، پیچیدگی محاسباتی مسئله به‌طور قابل توجهی کاهش می‌یابد. در این ساختار، ابتدا ورودی‌ها از طریق یک نگاشت ویژگی به حالت‌های کوانتومی کدگذاری می‌شوند و سپس توسط یک مدار وردشی پردازش می‌شوند.



شکل ۷. ماتریس درهم‌ریختگی برای ماشین بردار پشتیبان کوانتومی.

این ویژگی، عملکرد بهتر الگوی پیشنهادی را در مقایسه با روش‌های کلاسیک توجیه می‌کند. بنابراین هسته کوانتومی به عنوان یک مؤلفه مهم عمل می‌کند و با بهره‌گیری از اصول مکانیک کوانتومی، ساختار و الگوهای پیچیده در داده‌ها را به خوبی آشکار می‌کند. استفاده از این ماتریس نقش مهمی در بهبود دقت طبقه‌بندی و درک عمیق‌تر داده‌ها دارد.

بر اساس نتایج نمایش داده شده در شکل ۷، خروجی الگوی QSVM نشان‌دهنده عملکرد نسبتاً رضایت‌بخش آن در طبقه‌بندی داده‌هاست. از مجموع داده‌های آزمایشی، ۱۰ نمونه مربوط به افراد دیابتی و ۱۷ نمونه مربوط به افراد غیردیابتی به درستی تشخیص داده شده‌اند، در حالی که الگو در تشخیص سه نمونه دچار خطا شده است. این نتایج نشان می‌دهد الگو در شناسایی و طبقه‌بندی دو گروه موفق بوده است.

### ۲.۳. مدارهای کوانتومی وردشی

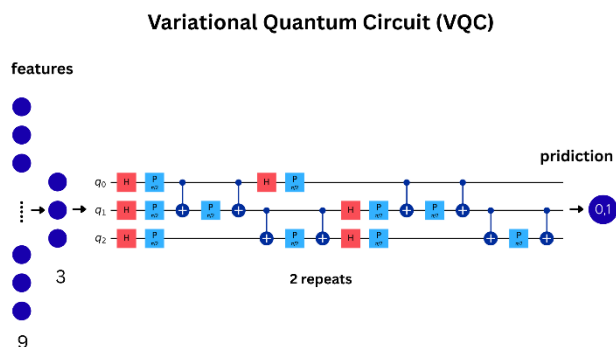
VQC از پیشرفته‌ترین الگوریتم‌های یادگیری ماشین کوانتومی به‌شمار می‌رود که توسط ویچک هاولیچک<sup>۱</sup> و همکاران توسعه یافته است [۱۸]. در این روش، داده‌های کلاسیکی، شامل اعداد و ویژگی‌ها، با فرآیند نگاشت<sup>۲</sup> به داده‌های کوانتومی تبدیل می‌شوند تا امکان پردازش آن‌ها در سامانه‌های کوانتومی فراهم شود. سپس مدارهای کوانتومی وردشی که پارامترهای آن‌ها به تعداد اندازه‌گیری‌ها و ابعاد داده‌ها وابسته است، طراحی می‌شوند. نتایج حاصل از محاسبات کوانتومی به‌عنوان بازخورد به سامانه کلاسیکی بازمی‌گردند و در بهینه‌سازی مدار نقش

۱. Vojtech Havlicek

۲. Mapping

کوانتومی QSVM از یک نگاشت غیرخطی در فضای هیلبرت استفاده شده است. بدین منظور، الگوی کلاسیکی SVM علاوه بر هسته خطی، با هسته غیرخطی RBF نیز ارزیابی شده است. مقایسه عملکرد این الگو بر اساس نتایج حاصل از جدول ۲، قابل مشاهده است. الگو با هسته RBF (غیرخطی) به طور قابل توجهی عملکرد بهتری نسبت به الگوی خطی دارد. دقت الگوی غیرخطی ۹۰ درصد است، در حالی که الگوی خطی ۷۴/۴ درصد دقت دارد. در مقابل، الگوی خطی با نرخ بازیابی پایین ۶۵/۲ درصد و  $F_1$  برابر با ۷۳/۲ درصد، عملکرد ضعیف‌تری در شناسایی نمونه‌های مثبت دارد.

در نهایت، نتایج نشان می‌دهند SVM با هسته RBF عملکرد بهتری نسبت به نسخه خطی خودش دارد، و استفاده از هسته غیرخطی در مسائل با ساختار پیچیده‌تر می‌تواند منجر به بهبود در عملکرد الگو شود. همچنین با توجه به جدول ۳، الگوی کوانتومی QSVM عملکرد بهتری نسبت به SVM کلاسیکی در زمینه طبقه‌بندی و پیش‌بینی داده‌های پزشکی دارد. این برتری به‌ویژه در معیارهای دقت و امتیاز  $F_1$  مشهود است. به‌طوری‌که الگوی QSVM به دقت ۹۰ درصد و امتیاز  $F_1$  به میزان ۸۶/۹ درصد دست یافته است. در مقابل، الگوی کلاسیکی تنها به دقت ۷۴/۴ درصد و امتیاز  $F_1$  برابر با ۷۳/۲ درصد رسیده است. این بهبود قابل توجه در عملکرد الگوی کوانتومی را می‌توان به دو عامل اصلی نسبت داد. نخست، استفاده از فضای هیلبرت کوانتومی امکان نمایش داده‌ها و استخراج ویژگی‌های پنهان و مؤثر را با کارایی بیشتری فراهم می‌کند. دوم، پردازش موازی و بهینه در محیط کوانتومی امکان‌پذیر است. این دو ویژگی به الگوی QSVM کمک می‌کنند تا الگوهای پیچیده‌تر موجود در داده‌ها را با دقت بیشتری شناسایی کرده و در نتیجه، عملکرد بهتری در طبقه‌بندی ارائه دهد. الگوی VQC با دقت ۹۴/۲ درصد و امتیاز  $F_1$  ۹۲/۱ درصد، نسبت به همه معیارهای ارزیابی، عملکرد برجسته و حتی بهتر از الگوی QSVM را نشان می‌دهد. این بهبود در عملکرد VQC را می‌توان به دلیل طراحی مدار وردشی با پارامترهای قابل تنظیم، همراه با بهینه‌سازی مستمر در فضای هیلبرت دانست. این رویکرد اجازه می‌دهد الگوریتم با توجه به ساختار داده‌ها، بهترین راهکار



شکل ۸. مدار کوانتومی وردشی.

جدول ۲. مقایسه هسته‌های خطی و غیرخطی در الگوی SVM.

| Model | Kernel        | Accuracy | Precision | Recall | F1-Score |
|-------|---------------|----------|-----------|--------|----------|
| SVM   | Linear kernel | ٪۷۴/۴    | ٪۸۳/۳     | ٪۶۵/۲  | ٪۷۳/۲    |
| SVM   | RBF kernel    | ٪۹۰      | ٪۱۰۰      | ٪۷۵    | ٪۸۵/۷۱   |

جدول ۳. عملکرد الگوهای SVM، QSVM، VQC.

| Classifiers | Accuracy | Precision | Recall | F1-score |
|-------------|----------|-----------|--------|----------|
| SVM         | ٪۹۰      | ٪۱۰۰      | ٪۷۵    | ٪۸۵/۷۱   |
| QSVM        | ٪۹۰      | ٪۱۰۰      | ٪۷۶/۹  | ٪۸۶/۹    |
| VQC         | ٪۹۴/۲    | ٪۱۰۰      | ٪۸۵/۴  | ٪۹۲/۱    |

این مدار دارای پارامترهای قابل تنظیم است که در فرآیند آموزش، با استفاده از یک الگوریتم بهینه‌سازی تکرارشونده، به‌گونه‌ای تنظیم می‌شوند که خروجی الگو به مقدار واقعی نزدیک‌تر شود. در این کار، از بهینه‌ساز COBYLA استفاده شده است که بدون نیاز به محاسبه گرادیان، پارامترها را بهینه می‌کند. همچنین تابع هزینه نیز به صورت تفاوت بین پیش‌بینی الگو و مقدار واقعی تعریف شده است و در هر تکرار آموزش، مقدار آن کاهش می‌یابد. این فرآیند تا رسیدن به حداکثر تعداد، ادامه می‌یابد. با اجرای این مدار و استخراج نتایج اندازه‌گیری، معیارهای ارزیابی مانند دقت و امتیاز  $F_1$  محاسبه و عملکرد الگو ارزیابی شده است.

این الگوریتم در چارچوب یادگیری با ناظر طراحی شده است و با بهینه‌سازی می‌تواند پیش‌بینی‌هایی با دقت بالا ارائه دهد. این رویکرد امکان استفاده از قابلیت‌های محاسبات کوانتومی را در حل مسائل تشخیصی پزشکی فراهم می‌کند.

### ۳. نتایج و بحث

برای اطمینان از مقایسه بهتر الگوی کلاسیکی با الگوی

طبقه‌بندی را یاد بگیرد.

در الگوهای یادگیری کوانتومی، پدیدهٔ برهم‌نهی به کیوبیت‌ها این امکان را می‌دهد تا چندین حالت را به طور همزمان پردازش کنند. این قابلیت، انجام محاسبات موازی و به‌طور کلی افزایش سرعت پردازش در الگوریتم‌های کوانتومی را فراهم می‌کند. استفاده از برهم‌نهی، امکان رمزگذاری مؤثر داده‌های کلاسیکی به حالت‌های کوانتومی، پردازش هم‌زمان نقاط داده‌های مختلف و محاسبهٔ کارآمد هسته‌های کوانتومی را فراهم می‌کند.

اگرچه نتایج شبیه‌سازی نشان‌دهندهٔ عملکرد قابل قبول الگوی کوانتومی در تشخیص دیابت است، اما اجرای آن بر روی سخت‌افزارهای کوانتومی واقعی با محدودیت مواجه است.

در دستگاه‌های NISQ، وجود نوفه، خطاهای کوانتومی، زمان کوتاه همدوسی، دقت محاسبات را تحت تأثیر قرار می‌دهد. درحالی‌که در شبیه‌سازهای آرمانی این عوامل نادیده گرفته می‌شوند. بنابراین، در دنیای واقعی این موارد می‌توانند منجر به ناپایداری نتایج شوند. همچنین، عمق مدار کوانتومی در این پژوهش با محدود کردن تعداد تکرار و ساده‌سازی ساختار کنترل شده است، تا با دستگاه‌های فعلی سازگار باشد. با این حال، افزایش عمق مدار به دلیل تجمع خطاها، می‌تواند چالش‌برانگیز باشد.

تعداد کیوبیت‌های مورد نیاز نیز برابر با ابعاد ویژگی‌های ورودی است. در این پژوهش، با استفاده از روش PCA و کاهش ابعاد از ۹ به سه ویژگی، تعداد کیوبیت‌ها محدود شده‌اند تا با دستگاه‌های کوانتومی با تعداد کیوبیت محدود سازگار باشند. علاوه بر این، روش کدگذاری زاویه‌ای مورد استفاده، برای داده‌ها با ابعاد بالا کارایی کمتری دارد و نیازمند روش‌های پیشرفته‌تری، از جمله کدگذاری دامنه است.

در کارهای آینده، می‌توان با بهره‌گیری از روش‌های کاهش خطا، بهینه‌سازی مدار و استفاده از دستگاه‌های پیشرفته‌تر، به سمت اجرای عملی‌تر و مقیاس‌پذیرتر این الگو رفت. لازم به ذکر است که بهبود جزئی عملکرد الگوی QSVM در مقایسه با SVM با هستهٔ غیرخطی لزوماً نشان‌دهندهٔ برتری ذاتی محاسبات کوانتومی نیست، بلکه ممکن است ناشی از تفاوت

در ساختار هسته و توانایی نگاشت کوانتومی در الگوسازی با تعاملات پیچیدهٔ ویژگی‌ها باشد. بنابراین، تمایز بین توانایی هسته و توانایی کوانتومی، یکی از چالش‌های اصلی در این حوزه محسوب می‌شود و نیازمند طراحی آزمایش‌های دقیق‌تری است و می‌تواند در پژوهش‌های آتی مورد بررسی قرار گیرد.

همان‌طور که قبلاً گفته شد، تمرکز اصلی این پژوهش بر الگوهای مبتنی بر ماشین بردار پشتیبان کلاسیکی و کوانتومی است. درحالی‌که، مقایسهٔ VQC با الگوهای یادگیری عمیق مانند CNN یا Deep CNN به‌خاطر قابلیت‌های محاسبات کوانتومی، می‌تواند نتایج خوبی را ارائه دهد. نتایج به‌دست‌آمده، علاوه بر تأیید برتری محاسبات کوانتومی در کشف ویژگی‌های پنهان در داده‌ها، قادر است یک فضای تصمیم‌گیری بهینه‌تری نسبت به روش‌های کلاسیکی ایجاد کند. علاوه بر این، الگوهای کوانتومی از توانایی تعمیم‌پذیری بالاتری نسبت به الگوهای کلاسیک برخوردار هستند، که این ویژگی می‌تواند در تشخیص دقیق و مطمئن بیماری‌هایی مثل دیابت، مؤثر واقع شود.

در نهایت، نتایج به دست آمده نشان می‌دهند، ادغام اصول مکانیک کوانتومی در الگوریتم‌های یادگیری ماشین می‌تواند در حل مسائل تشخیصی پیچیده، به‌ویژه در حوزهٔ سلامت، به بهبود قابل ملاحظه‌ای در عملکرد الگوها منجر شود. این دستاورد، پتانسیل قابل توجه روش‌های کوانتومی در کاربردهای عملی هوش مصنوعی را برجسته می‌کند.

#### ۴. نتیجه‌گیری

با توجه به اهمیت بیماری دیابت، پیشگیری از آن برای سلامت انسان بسیار لازم و ضروری است. امروزه استفاده از الگوهای یادگیری ماشین به عنوان راهکاری مؤثر برای پیش‌بینی و تشخیص زودهنگام دیابت، مورد توجه قرار گرفته است. نتایج این پژوهش نشان می‌دهد که روش‌های کوانتومی قابلیت پردازش داده‌ها با دقت بالاتر و تعداد پارامترهای کمتر، نسبت به روش‌های کلاسیک را دارند.

در شرایطی که حجم داده‌های آموزشی محدود باشد، QSVM عملکرد قابل قبولی را ارائه می‌دهد، در حالی که VQC هم توانست در مجموعهٔ داده‌های محدود، دقت بالایی به‌دست

آورد. یکی از مهمترین مزایای روش‌های کوانتومی نسبت به روش‌های کلاسیک، توانایی کشف الگوها و ویژگی‌های پیچیده در داده‌های بالینی و استفاده از خواص مکانیک کوانتومی مانند برهم‌نهی و درهم‌تنیدگی است. مقایسه عملکرد این الگوها نشان می‌دهد که استفاده از روش‌های یادگیری ماشین کوانتومی می‌تواند در پیش‌بینی بیماری‌ها، به‌ویژه دیابت، نتایج بهتری را ارائه دهد. اگرچه بهبود عملکرد الگوی QSVM لزوماً نشان‌دهنده برتری ذاتی محاسبات کوانتومی نیست، اما تحلیل نداشت ویژگی کوانتومی نشان می‌دهد که خواصی مانند برهم‌نهی و درهم‌تنیدگی، سازوکاری فراهم می‌کنند که می‌تواند به ایجاد هسته‌هایی با قابلیت تفکیک‌پذیری بالاتر منجر شود. این چیزی است که در الگوهای کلاسیکی به‌راحتی قابل دستیابی نیست. به علاوه، الگوهای آمیخته مانند VQC که ترکیبی از پردازش کوانتومی و بهینه‌سازی کلاسیک را به‌کار می‌گیرند، قابلیت اجرا روی دستگاه‌ها و رایانه‌های کوانتومی فعلی را دارند.

یکی از محدودیت‌های این پژوهش، عدم دسترسی به بسترهای ابری کوانتومی مانند IBM Quantum است که استفاده از آن در حال حاضر امکان‌پذیر نیست. بنابراین، تمام شبیه‌سازی‌ها در محیطی آرمانی و بدون نوفه انجام شده است. امیدواریم کارهای آینده، با رفع محدودیت‌ها و دسترسی به محیط‌های

## ۵. مراجع

ابری، بتوان عملکرد الگو را در شرایط واقعی‌تر همراه با نوفه و خطا ارزیابی کرد.

از آنجایی که الگوهای کلاسیکی با حجم بالا با چالش محاسباتی مواجه می‌شوند در این پژوهش، مقایسه زمانی دقیقی با سخت‌افزار کوانتومی واقعی، انجام نشده است. در نتیجه به صورت کلی، سرعت پیش‌بینی و تحلیل داده‌ها، به پیچیدگی الگوریتم وابسته است و مزیت الگوریتم‌های کوانتومی با اجرا شدن بر روی سخت‌افزار کوانتومی معنا پیدا می‌کند. بنابراین، در سامانه‌های کلاسیکی و سامانه‌های کوانتومی، طراحی و بهبود الگوریتم‌های دقیق و کاربردی، نقش مهمی دارد.

در مجموع، یافته‌های این پژوهش توانایی بالای یادگیری ماشین کوانتومی را در حل مسائل زیستی و پزشکی برجسته می‌کند. بنابراین افزایش دقت پیش‌بینی، کاهش پیچیدگی پردازش و امکان بهره‌برداری از دستگاه‌های کوانتومی نسل فعلی، از جمله ویژگی‌هایی هستند که الگوها و مدارهای کوانتومی را به گزینه امیدوارکننده‌ای برای مطالعات آینده در زمینه تحلیل و طبقه‌بندی داده‌های زیستی تبدیل می‌کنند. بنابراین در تحقیقات پیش‌رو می‌توان با افزایش حجم داده‌ها و بهینه‌سازی پارامترهای مختلف الگو و با بررسی تأثیر ویژگی‌های زیستی خاص برای الگو، به درک بهتر و قابلیت روش‌های کوانتومی جدید کمک کرد.

1. M Fisher, J. *Diabetes Endocr. Pract.* **5** (2022) 135.
2. H Naz and S Ahuja, *J. Diabetes Metab. Disord.* **19** (2020) 391.
3. H Zhou, R Myrzasheva and R Zheng, *EURASIP J. Wirel. Commun. Netw.* **2020** (2020) 148.
4. P L Panuganti, S N Hinkle, S Rawal, L G Grunnet, Y Lin, A Liu and C Zhang, *Nutrients* **10** (2018) 938.
5. H Gupta, H Varshney, T K Sharma, N Pachauri and O P Verma, *Complex Intell. Syst.* **8** (2022) 3073.
6. J Biamonte, P Wittek, N Pancotti, P Rebentrost, N Wiebe and S Lloyd, *Nature* **549** (2017) 195.
7. V Kecman, *Support Vector Machines: Theory and Applications*, Springer, Berlin (2005) 1.
8. C J Burges, *Data Min. Knowl. Discov.* **2** (1998) 955.
9. O Zang, G Barru  and T Quertier, in: *Proc. Int. Conf. Quantum Eng. Sci. Technol. Ind. Serv.*, Springer, Cham (2025) 270.
10. J Romero, J P Olson and A Aspuru-Guzik, *Quantum Sci. Technol.* **2** (2017) 045001.
11. A Macaluso, L Clissa, S Lodi and C Sartori, in: *Proc. Int. Conf. Comput. Sci.*, Springer, Cham (2020) 591.
12. W R Khan, M A Kamran, M U Khan, M M Ibrahim, K S Kim and M U Ali, *Int. J. Intell. Syst.* **2025** (2025) 1351522.
13. J R McClean, J Romero, R Babbush and A Aspuru-Guzik, *New J. Phys.* **18** (2016) 023023.
14. O Farooq, M Shahid, S Arshad, A Altaf, F Iqbal, Y A M Vera and I Ashraf, *Sci. Rep.* **14** (2024) 19521.
15. D Maheshwari, B Garcia-Zapirain and D Sierra-Sosa, in: *Proc. IEEE Int. Symp. Signal Process. Inf. Technol.*, IEEE, Piscataway (2020) 1.
16. M Schulz, arXiv:2101.11020 (2021).

17. N Singh and S R Pokhrel, arXiv:2501.08205 (2025).
18. V Havlíček, A D Córcope, K Temme, A W Harrow, A Kandala, J M Chow and J M Gambetta, *Nature* **567** (2019) 209.
19. N Innan, M A Z Khan, B Panda and M Bennai, *Quantum Inf. Process.* **22** (2023) 374.
20. D Sierra-Sosa, J D Arcila-Moreno, B Garcia-Zapirain and A Elmaghraby, *Comput. Mater. Contin.* **67** (2021) 1849.