

بازشناسي برخط حروف مجزاي دستنويس فارسي بر اساس تشخيص گروه بدنه اصلي با استفاده از ماشين بردار پشتيبان

محمدامين مهرعليان^۱ و كاظم فولادي^۲

^۱ دانشكده ي مهندسي كامپيوتر و فناوري اطلاعات، دانشگاه صنعتي اميركبير، تهران (mehralian@aut.ac.ir)

^۲ دانشكده ي مهندسي برق و كامپيوتر، پرديس دانشكده هاي فني، دانشگاه تهران، تهران (kfouladi@ut.ac.ir)

چكيده

در اين مقاله روشي جديد براي بازشناسي برخط حروف مجزاي فارسي ارائه شده است كه با استفاده از دسته بندي كننده ي ماشين بردار پشتيبان يا SVM عمل مي كند. اين عمليات در دو مرحله صورت مي گيرد. در مرحله ي اول حرف ورود ي با استفاده از دسته بندي كننده در قالب يكي از هجده گروه بدنه ي اصلي حروف، طبقه بندي مي شود و سپس در مرحله ي دوم، موقعيت، تعداد و شكل ساير حركت ها مانند نقطه و سرکش (ريزحركت ها)، حرف نهايي را تعيين مي كند. به عنوان نمونه براي تشخيص حرف «ت» ابتدا گروه بدنه ي «ب، پ، ت، ث» تشخيص داده مي شود و سپس وجود ريزحركت «دونه نقطه» در بالاي آن منجر به انتخاب «ت» از اين گروه مي شود. نتايج تجربي اين كار پژوهشي كه بر اساس مجموعه داده ي Online-TMU صورت گرفته است، متوسط نرخ بازشناسي بدنه ي اصلي را ۹۴٪ نشان مي دهد و با در نظر گرفتن پس پردازش ها بر اساس ريزحركت ها اين نرخ به حدود ۹۸٪ مي رسد.

كلمات كليدي

بازشناسي برخط حروف فارسي، بازشناسي دست نوشته / دستنويس، ريزحركت، ماشين بردار پشتيبان (SVM).

۱- مقدمه

براي حروف دستنويس فارسي ارائه مي كند و خصوصيات آن را بيان مي كند. الگوريتم بازشناسي حروف و جزئيات آن در بخش هاي ۴، ۵ و ۶ تحت عنوان هاي پيش پردازش، استخراج ويژگي و طبقه بندي و پس پردازش ارائه مي شود. در بخش ۷ به مسايل مربوط به پياده سازي و نتايج تجربي پرداخته مي شود. سرانجام بخش ۸ به خلاصه ي مطالب، بحث و نتيجه گيري اختصاص داده شده است.

امروزه ابزارهاي تجاري چون Tablet PC و PDA كه از صفحات حساس به جاي صفحه كليد براي ارتباط با كاربر استفاده مي كنند توسعه ي فراواني يافته است، از اينرو تقاضا براي سيستم هاي بازشناسي برخط (online) افزايش يافته است. اصلي ترين تفاوت ميان دو رويكرد برخط و برون خط (offline) دسترسي به دنباله ي نقاط در هر حركت دست و همچنين ترتيب و موقعيت زير حرركات است.

در اين مقاله قصد داريم به ارائه ي روشي جديد براي بازشناسي برخط حروف مجزاي فارسي پردازيم. در روش ارائه شده سعي شده است كمترين تعداد قيد ممكن بر روي نحوه ي نگارش حروف قرار داده شود. اصلي ترين قيدي كه وجود دارد، اين است كه كاربر بايستي ابتدا «بدنه» ي حروف را بنويسد و سپس نقطه گذاري آنها را انجام دهد.

در مقايسه با زبان هاي لاتين، براي بازشناسي برخط حروف فارسي و عربي، تاكنون كارهاي كمتر انجام شده است. مرور مختصري بر روي مهم ترين كارهاي مرتبط انجام شده در اين زمينه در بخش ۲ ارائه مي شود.

اين مقاله در هشت بخش سازمان ده ي شده است. در بخش ۲ مروري بر كارهاي انجام شده در زمينه ي بازشناسي حروف دستنويس برخط فارسي/ عربي ارائه مي شود. بخش ۳ مدلي را

۲- مروري بر كارهاي انجام شده توسط ديگران

در زمينه ي بازشناسي دست نوشته ي برخط، برخي كارها به بازشناسي ارقام و برخي ديگر به بازشناسي حروف يا كلمات پرداخته اند. از آنجا شرايط آزايش و معيارهاي ارزيايي اين كارها با يكي ديگر متفاوت است، مقايسه ي عددي آنها را بيان نمي كنيم و فقط روش هاي مورد استفاده در آنها را مورد توجه قرار مي دهيم. كارهاي قديمي تر توسط مؤلفين همين مقاله در [۱] بررسي شده است و در اينجا تنها به مرور چند كار مرتبط جديدتر مي پردازيم.

در كار قبلي مؤلفين اين مقاله، [۱]، بازشناسي حروف مجزاي دستنويس فارسي، بر اساس تشخيص بدنه ي اصلي به كمك مدل مخفي ماركوف (HMM) انجام شده است. پس از تشخيص بدنه ي اصلي، ريزحركت ها (مانند نقاط، سرکش و مد) مورد توجه قرار مي گيرند و از تركيب اطلاعات آنها با بدنه ي اصلي

استفاده از رویکرد برخط این امکان را به ما می دهد که اطلاعات ترتیب زمانی در هنگام نوشتن حروف و یا تعداد حرکات قلم را در اختیار داشته باشیم. در این کار فرض می کنیم که بدنه ی اصلی حروف در ابتدا نوشته می شود و سایر حرکات مانند نقاط، سرکش ها و... در مراحل بعدی گذاشته می شود. به عنوان نمونه حرف «چ» با سه حرکت قلم نوشته می شود، به این ترتیب که ابتدا بدنه ی اصلی «ح» و در مراحل بعدی نقاط قرار داده می شوند.

تنها محدودیتی که در نحوه ی نوشتن ریزحرکت ها وجود دارد، این است که «سه نقطه» می بایستی از ترکیب «دو نقطه» و یک «تک نقطه» ساخته شده باشد. همچنین دو سرکش «گ» باید در دو حرکت مجزای قلم نوشته شود. بعلاوه ی دسته ی «ط» و «ظ» و سرکش «ک» بایستی در حرکتی بجز حرکت بدنه ی اصلی نوشته شود. به عبارت دیگر حروف «ط» و «ک» باید در دو حرکت قلم و حروف «گ» و «ظ» باید در سه حرکت قلم نگاشته شوند. شکل (۱)، صورت صحیح نگارش ۳۳ حالت مختلف حروف فارسی (شامل ۳۲ حرف و آ) را نشان می دهد که در آن هر قسمت با یک حرکت قلم نوشته شده است [۳].

آ ا ب پ ت ث ج ح خ د
ر ز س ش ص ض ط ظ ع غ
ف ق ک گ ل م ن و ه ی

شکل (۱): نگارش صحیح برای ۳۳ حالت مختلف از حروف

۴- پیش پردازش

از آنجا که داده های برخط توسط قلم بر روی صفحه ی حساس نوشته می شوند، مختصات نقاط نمونه برداری شده، تعداد و فاصله ی آنها و همچنین ابعاد منحنی های ایجاد شده از یک حرف به حرف دیگر و از یک نویسنده به نویسنده ی دیگر می تواند بشدت تغییر کند. برای اینکه بتوان تاثیر این تغییرات را در فرآیند بازشناسی به حداقل رساند، باید نوعی نرمال سازی روی مجموعه نقاط نمونه برداری شده از حرکات قلم انجام شود.

۴-۱- نرمال سازی فواصل میان نمونه ها

برای نرمال سازی فواصل میان نمونه ها، ابتدا میانگین فاصله ی نقاط متوالی در مجموعه نقاط مربوط به یک حرکت را اندازه گیری می کنیم و نقاطی را که در فاصله ای کمتر از دو برابر این میانگین باشند، حذف می کنیم. با حذف این نقاط اضافی، لرزش های کمتری نیز در مسیر حرکت دیده می شود و از این رو لرزش های ناخواسته در نوشتن حرف نیز تا حدی فیلتر می شود. شکل (۲) دو نمونه از حروف و نتیجه ی نرمال سازی فواصل میان نقاط آنها دیده می شود.

تشخیص داده شده، حرف نوشته شده بازشناسی می شود. در [5] نویسنده با تعریف نقاط بحرانی در شکل حروف (نقاط کمینه و بیشینه محلی) هر یک از آنها را به بخش هایی تقسیم بندی کرده است، سپس با استفاده از یک مجموعه ویژگی های هندسی و با بکارگیری شبکه ی عصبی به دسته بندی حروف پرداخته است.

در [6] برای باشناسی حروف از درخت تصمیم استفاده شده است که در سطح اول از آن، نمونه ها بر اساس تعداد نقاط دسته بندی می شوند. سپس در سطح بعدی موقعیت نقاط مورد بررسی قرار می گیرد که برای این منظور خط کرسی نیز استخراج می گردد. در آخرین مرحله نیز با استفاده از یک الگوریتم تطبیق الگو، بدنه اصلی حروف با بدنه اصلی حروف معیاری که از هر حرف در پایگاه ذخیره شده است مقایسه می گردد. اگرچه نتایج نهایی ارائه شده مطلوب به نظر می رسد اما وجود ویژگی تعداد نقاط در اولین سطح از درخت تصمیم محدودیت هایی را برای کاربران سیستم ایجاد می کند.

در [7] از دو دسته بندی کننده مختلف استفاده شده است که هر یک، ویژگی های متفاوتی را مورد استفاده قرار می دهد. در دسته بندی کننده ی اول، یک ویژگی با عنوان «هیستوگرام همبستگی» به شبکه عصبی پرسپترون داده می شود و در دسته بندی کننده دوم، ویژگی محتوی همبستگی با دسته بندی کننده SVM مورد استفاده قرار داده شده، مشکل اصلی این روش، ابعاد بالای بردارهای ویژگی مورد استفاده است؛ مثلاً در دسته بندی کننده SVM به ازای N نمونه ورودی برای هر حرف، $N(N-1)$ ویژگی استخراج می شود.

۳- مدل حروف فارسی

حروف مجزا در نگارش فارسی در دو بخش اصلی نوشته می شوند. قسمت اصلی حروف «بدنه» نامیده می شود و علائم اضافی، مانند نقطه، سرکش یا کلاه را «ریز حرکت» می گوئیم [۲]. ویژگی مهمی که در بازشناسی حروف فارسی قابل توجه است این است که برخی از حروف بدنه ی یکسان دارند و تنها تفاوت آنها در ریزحرکت های آنهاست. با توجه به بدنه های مشابه، حروف فارسی را می توان در قالب ۱۸ گروه دسته بندی کرد که این گروه ها در جدول (۱) قابل مشاهده هستند.

جدول (۱): گروه بندی انجام شده بر اساس بدنه حروف

گروه ۱	آ، ا	گروه ۷	ص، ض	گروه ۱۳	ل
گروه ۲	ب، پ، ت، ث	گروه ۸	ط، ظ	گروه ۱۴	م
گروه ۳	ج، چ، ح، خ	گروه ۹	ع، غ	گروه ۱۵	ن
گروه ۴	د، ذ	گروه ۱۰	ف	گروه ۱۶	و
گروه ۵	ر، ز، ژ	گروه ۱۱	ق	گروه ۱۷	ه
گروه ۶	س، ش	گروه ۱۲	ک، گ	گروه ۱۸	ی

$$f_1(x_i, y_i) = x_i \quad (1)$$

$$f_2(x_i, y_i) = y_i \quad (2)$$

$$f_3(x_i, y_i) = x_i - x_{i-1} (= \Delta x_i) \quad (3)$$

$$f_4(x_i, y_i) = y_i - y_{i-1} (= \Delta y_i) \quad (4)$$

یکی از بهبودهای این روش در مقایسه با روش قبلی ارائه شده در [۱] کاهش ابعاد بردار ویژگی برای هر نقطه از ۶ به ۴، بدون افت دقت در نتایج نهایی است. دو ویژگی کنار گذاشته شده، معیاری از زاویه‌ی نقطه‌ی فعلی را نسبت به نقطه‌ی قبلی بازنمایی می‌کردند.

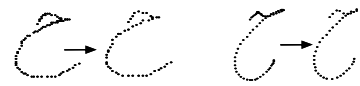
ویژگی (۱) و (۲) همان مختصات دکارتی نقطه است که در مرحله‌ی پیش‌پردازش مقادیر نرمال‌شده‌ی آن محاسبه شده است. ویژگی (۳) و (۴) میزان تغییر مختصات در راستای افقی و عمودی را نشان می‌دهد. از آنجا که در مرحله‌ی یکسان‌سازی تعداد نقاط نمونه در پیش‌پردازش همگی حروف با تعداد ثابتی نقطه بازنمایی شده‌اند، این ویژگی‌ها نیز که متأثر از شکل و طول حروف می‌باشند، می‌توانند یکی از پارامترهای ایجاد تمایز بین بدنه‌ی حروف در نظر گرفته شوند.

مجموعه‌ی این چهار ویژگی برای هر دو نقطه‌ی متوالی استخراج می‌شود و کل آنها به عنوان بردار ویژگی به SVM داده می‌شود.

برای تعیین وضعیت ریزحرکت‌ها، موقعیت هر ریزحرکت با توجه به بدنه‌ی اصلی حرف و در سه دسته‌ی بالا، پایین و وسط تقسیم‌بندی می‌شود. این دسته‌بندی بر اساس مقایسه‌ی نقطه‌ی مرکزی ارتفاع ریزحرکت با حداقل و حداکثر ارتفاع بدنه صورت می‌گیرد. در حالتی که نقطه‌ی مرکزی ارتفاع ریزحرکت بین حداقل و حداکثر ارتفاع بدنه‌ی اصلی قرار داشته باشد (دسته‌ی وسط) بسته به گروه ممکن است به یکی از دو حالت بالا و پایین تبدیل شود. برای مثال برای حرف «ج» وسط به معنی پایین و در حرف «ن» وسط به معنی بالا خواهد بود.

۵-۲- طبقه‌بندی

در این مقاله از «ماشین بردار پشتیبان» یا SVM به عنوان دسته‌بندی‌کننده اصلی در بازشناسی حروف استفاده شده است. بررسی‌های صورت گرفته نشان می‌دهد استفاده از SVM به طور مستقیم در بازشناسی برخاسته از چندین مورد توجه قرار نگرفته است و علت این امر نیز ناشی از متغیر بودن طول ویژگی‌های استخراج شده، بسته به طول منحنی حروف است [9]. به عنوان نمونه در [10] از SVM تنها برای جداسازی (segmentation) حروف استفاده شده است و برای دسته‌بندی نهایی HMM بکارگرفته شده است. اما فرآیند پیش‌پردازش استفاده شده در این مقاله علاوه بر ایجاد یک استانداردسازی اولیه‌ی حروف، قابلیت استفاده از ویژگی‌ها را برای دسته‌بندی‌کننده‌ی SVM فراهم می‌کند.



شکل (۲): دو نمونه از نرمال‌سازی فواصل میان نقاط

۴-۲- یکسان‌سازی تعداد نقاط نمونه برداری

یکسان‌سازی تعداد نقاط نمونه برداری برای سهولت استفاده در SVM و افزایش کارایی آن، ضروری است. در واقع با این کار همه‌ی حروف با یک تعداد مشخص از برش‌های زمانی نمونه برداری مجدد می‌شوند و به داده‌هایی با ابعاد یکسان تبدیل می‌شوند.

برای این منظور، ابتدا با استفاده از اسپیلاین‌های درجه‌ی سوم یک درونیابی بین هر چهار نقطه‌ی متوالی از منحنی حرف انجام می‌شود و به این ترتیب تخمینی از شکل کلی حرف صورت می‌گیرد. سپس طول منحنی به ۲۰ قسمت مساوی تقسیم می‌شود به طوری که طول منحنی بین هر دو نقطه‌ی متوالی، اندازه‌ای برابر دارد و در نتیجه ۲۱ نقطه با فواصل مساوی از روی منحنی حاصل نمونه برداری می‌شود. البته تعداد این نقاط برای ریزحرکت‌ها به ۱۰ نقطه کاهش می‌یابد. شکل (۳) دو مثال از انجام یکسان‌سازی تعداد نقاط نمونه برداری را نشان می‌دهد.



شکل (۳): دو نمونه از یکسان‌سازی تعداد نقاط نمونه برداری

۴-۳- یکسان‌سازی اندازه و مختصات نقطه‌ی شروع

از آنجا که اندازه حرف نوشته شده توسط کاربران می‌تواند متفاوت باشد، بنابراین لازم است پیش از بازشناسی آن، یک همسان‌سازی در ابعاد صورت گیرد. برای این منظور حرف نوشته شده در چارچوب مختصاتی به ابعاد ۱۰۰×۱۰۰ قرار می‌گیرد. بعلاوه از آنجا که نقطه‌ی شروع نگارش حرف می‌تواند هر نقطه‌ای روی صفحه باشد، نقطه‌ی شروع برای همه‌ی حروف به نقطه‌ی (۰, ۰) منتقل می‌شود.

۵- استخراج ویژگی‌ها و طبقه‌بندی

برای انجام بازشناسی بدنه‌ی حروف و زیرحرکت‌ها با استفاده از SVM لازم است که ویژگی‌های مناسبی را جهت بازنمایی هر یک از آنها انتخاب کنیم.

۵-۱- استخراج ویژگی‌ها

برای انتخاب ویژگی‌ها، زیرمجموعه‌ای از ویژگی‌های معرفی شده در [8] را مورد استفاده قرار می‌دهیم. برای هر نقطه‌ی نمونه مجموعه‌ی چهار ویژگی زیر را محاسبه می‌کنیم:

۶- پس پردازش

پس از تشخیص و بازشناسی گروه بدنه‌ی اصلی، نتایج همراه با اطلاعات مربوط به ریزحرکت‌ها، به یک درخت تصمیم داده می‌شود تا با ترکیب اطلاعات این دو بخش بررسی کند که حرف ناشناخته، متعلق به کدام یک از اعضای گروه‌های هجده‌گانه است و بدین ترتیب فرآیند بازشناسی کامل حرف کامل شود [۱].

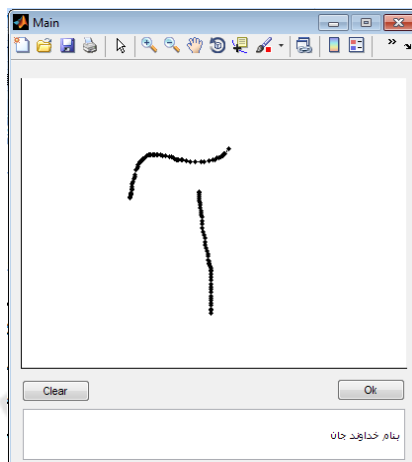
از آنجا که بازشناسی و تفکیک ریزحرکت‌ها مانند نقطه‌ها، سرکش‌ها و... نسبتاً دشوار است و اغلب به لحاظ ظاهری نیز با هم تداخل دارند، سعی ما بر این بوده است که از این اطلاعات در آخرین مراحل تصمیم‌گیری استفاده شود. از این رو در ساختار درخت تصمیم، اولویت تصمیم‌گیری با تعداد و محل ریزحرکت‌هاست، نه شکل واقعی آنها. به عنوان نمونه زمانی که بدنه‌ی اصلی حرفی در قالب گروه ۱، «ا» تشخیص داده می‌شود، تنها وجود یک زیرحرکت در بالای آن کافی است تا سیستم آن را «آ» تشخیص دهد و دیگر نیازی به تشخیص شکل علامت «مد» نمی‌باشد.

همچنین گاهی به خاطر عدم تطابق بین ریزحرکت‌ها و بدنه‌ی اصلی، سیستم قادر به تشخیص خطا در بازشناسی بدنه خواهد بود و می‌تواند بدنه‌ی تشخیص داده شده را تصحیح کند و آن را به بدنه‌ی محتمل دیگری تبدیل نماید. برای مثال ممکن است سیستم بدنه‌ی «ل» را به همراه یک ریزحرکت در بالای آن تشخیص دهد. در چنین حالتی سیستم با توجه عدم تطابق بین

آنچه تشخیص داده شده و آنچه برای حرف «ل» مورد انتظار بوده است، تشخیص بدنه‌ی خود را به «ن» تغییر می‌دهد. به عبارت دیگر بازنگری نهایی بر اساس مشخصات بسیار ساده‌ی ریزحرکت‌ها صورت می‌گیرد.

۷- پیاده‌سازی و نتایج تجربی

برای پیاده‌سازی برنامه‌ی بازشناسی، از نرم‌افزار MATLAB استفاده شده است. برای استفاده از SVM نیز کتابخانه LIBSVM، به کار گرفته شده است. شکل (۴) نمایی از این نرم‌افزار را نشان می‌دهد.



شکل (۴): نمایی از نرم‌افزار بازشناسی برخط حروف فارسی.

جدول (۲): ماتریس پراکنش (Confusion) برای بازشناسی گروه‌ها

	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18	%
1	۷۱	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۱	۰	۰	۰	۰	0.99
2	۰	۱۳۶	۰	۰	۱	۰	۰	۰	۰	۰	۰	۸	۰	۰	۰	۰	۰	۰	0.94
3	۰	۰	۱۴۳	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	1.00
4	۰	۰	۰	۷۰	۱	۰	۰	۱	۰	۰	۰	۱	۰	۰	۰	۰	۰	۰	0.96
5	۰	۰	۰	۱	۱۰۹	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	0.99
6	۰	۰	۰	۰	۱	۶۸	۰	۰	۰	۰	۰	۱	۰	۱	۰	۱	۰	۱	0.93
7	۰	۰	۰	۰	۱	۱	۷۱	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	0.97
8	۰	۰	۰	۰	۰	۰	۰	۶۷	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	1.00
9	۰	۰	۰	۰	۰	۰	۰	۰	۷۲	۰	۰	۰	۰	۰	۰	۰	۰	۰	1.00
10	۰	۲	۰	۰	۰	۰	۰	۰	۰	۳۱	۴	۰	۰	۰	۰	۰	۰	۰	0.84
11	۰	۰	۰	۰	۰	۰	۰	۰	۰	۱	۳۵	۰	۰	۰	۱	۰	۰	۰	0.95
12	۰	۱۲	۰	۰	۱	۰	۰	۰	۰	۱	۱	۵۳	۰	۰	۱	۱	۰	۰	0.76
13	۰	۰	۰	۰	۱	۰	۰	۰	۰	۰	۱	۱	۳۲	۰	۱	۰	۱	۰	0.86
14	۰	۰	۰	۰	۱	۰	۰	۰	۰	۰	۰	۰	۳۶	۰	۰	۰	۰	۰	0.97
15	۰	۰	۰	۰	۰	۰	۰	۰	۰	۱	۲	۰	۰	۳۲	۰	۲	۰	۰	0.86
16	۰	۰	۰	۰	۰	۰	۰	۲	۰	۰	۰	۰	۰	۰	۳۶	۰	۰	۰	0.95
17	۰	۰	۰	۰	۰	۰	۱	۰	۰	۱	۰	۵	۰	۰	۰	۲۸	۰	۰	0.80
18	۰	۰	۰	۰	۱	۱	۰	۰	۰	۰	۰	۰	۰	۰	۰	۰	۳۶	۰	0.95

درصد کل: ۹۴.۳

جدول (۳): مقایسه نتایج ماتریس پراکنش با روش ارائه شده در [۱]

گروه	1	2	3	4	5	6	7	8	9	10	11	12	13	14	15	16	17	18
SVM	۷۱	۱۳۶	۱۴۳	۷۰	۱۰۹	۶۸	۷۱	۶۷	۷۲	۳۱	۳۵	۵۳	۳۲	۳۶	۳۲	۳۶	۲۸	۳۶
HMM	۷۱	۱۲۵	۱۴۱	۶۹	۱۰۲	۷۱	۶۹	۶۴	۶۹	۳۱	۳۲	۵۱	۳۵	۳۶	۳۰	۲۵	۳۳	۳۶
اختلاف	۰	۱۱	۲	۱	۷	-۳	۲	۳	۳	۰	۳	۲	-۳	۰	۲	۱	-۵	۰

در شکل (مانند شباهت مد و دو نقطه) استفاده از SVM برای بازشناسی آن‌ها چندان کارآمد به نظر نمی‌رسد. بنابراین استفاده از سایر اطلاعات مربوط به ریزحرکات مانند موقعیت و تعداد آنها در مرحله پس پردازش، می‌تواند تا حد زیادی مشکل ضعف در بازشناسی حروف را برطرف کند.

در مقایسه‌ی روش ارائه شده در این مقاله با روش قبلی که از HMM برای بازشناسی بدنه‌ی اصلی استفاده می‌کرد، چند بهبود ملاحظه می‌شود. اولاً همان طور که در بخش ۵-۱ اشاره شد، تعداد عناصر بردار ویژگی برای هر نقطه، از شش به چهار کاهش یافته است، بدون اینکه نرخ بازشناسی اولیه کاهش یابد. ثانیاً ساختار SVM نسبت به HMM در هنگام اجرا از سرعت بیشتری برخوردار است که این باعث می‌شود پیاده‌سازی کارآمدتری در عمل ایجاد شود. علاوه بر این بهبود نرخ بازشناسی در حدود یک درصد نیز قابل توجه است.

۹- مراجع

- [۱] مهرعلیان، م.ا.، فولادی، ک.، "بازشناسی برخط حروف مجزای دست نویس فارسی بر اساس تشخیص گروه اصلی بدنه با استفاده از مدل مخفی مارکوف" مجموعه مقالات پانزدهمین کنفرانس مهندسی کامپیوتر ایران، مرکز توسعه فناوری نیرو (متن)، تهران، ۱۳۸۸.
- [۲] سلیمانی باغشاه، م.، باقری شورکی، س.، کسائی، ش.، "بازشناسی و یادگیری حروف مجزای برخط فارسی به روش فازی"، مجموعه مقالات چهاردهمین کنفرانس مهندسی برق ایران، دانشگاه صنعتی امیرکبیر، تهران، ۱۳۸۵.
- [۳] رضوی، س. م.، کبیر، ا.، "بازشناسی برخط حروف مجزای فارسی"، مجموعه مقالات ششمین کنفرانس سیستم‌های هوشمند، دانشگاه شهید باهنر، کرمان، آذر ۱۳۸۳.
- [۴] رضوی، س. م.، کبیر، ا.، "یک پایگاه داده برای بازشناسی دست‌نوشته‌های برخط فارسی"، مجموعه مقالات ششمین کنفرانس سیستم‌های هوشمند، دانشگاه شهید باهنر، کرمان، آذر ۱۳۸۳.

- [5] Harouni, M., Mohamad, D., Rasouli, A., "Deductive method for recognition of on-line handwritten Persian/Arabic characters", *The 2nd International Conference on Computer and Automation Engineering, (ICCAE)*, pp. 791-795, 2010.
- [6] Omer, M.A.H., Ma, S.L., "Online arabic handwriting character recognition using matching algorithm", *The 2nd International Conference on Computer and Automation Engineering, (ICCAE)*, pp. 259-262., 2010.
- [7] Izadi, S., Suen, C.Y., "Integration of contextual information in online handwriting representation", *Lecture Notes in Computer Science, LNCS 5716*, pp. 132-142. 2009.
- [8] Biem, A., "Minimum classification error training for online handwriting recognition", *IEEE Transactions on Pattern Analysis and Machine Intelligence*, Vol. 28, No. 7, 2006.
- [9] Bahlmann, C., Haasdonk, B., Burkhardt, H., "Online handwriting recognition with support vector machines-a kernel approach", *Eighth International Workshop on Frontiers in Handwriting Recognition*, pp. 49-54, 2002.
- [10] Ahmad, A.R., Khalia, M., Viard-Gaudin, C., Poisson, E., "Online handwriting recognition using support vector machine", *IEEE Region 10 Conference (TENCON)*, pp. 311-314, 2004.

برای SVM از کرنل گاوسی استفاده شده که مقدار پارامتر σ در آن برای حروف و ریزحرکات به ترتیب ۰٫۶ و ۱٫۶ در نظر گرفته شده است.

مجموعه داده‌ی Online-TMU که توسط بخش مهندسی برق دانشگاه تربیت مدرس تهیه شده است [۴]، به عنوان مجموعه‌ی داده‌ی آموزشی و آزمایشی مورد استفاده قرار گرفته است. از ۷۰ درصد نمونه‌ها در مرحله آموزش و از ۳۰ درصد مابقی برای آزمایش استفاده شده است. نتیجه‌ی مرحله‌ی آزمون در ماتریس پراکنش موجود در جدول (۲) نشان داده شده است.

همانطور که مشاهده می‌شود، بازشناسی در خروجی SVM با میانگینی برابر با ۹۴٫۳٪ درست صورت می‌گیرد. جدول (۳) مقایسه‌ای بین ماتریس پراکنش حاصل از دسته‌بندی کننده SVM را نسبت به روش ارائه شده در [۱] (با استفاده از HMM) برای هر گروه به صورت جداگانه نشان می‌دهد (اعداد نشان داده شده نماینده تعداد نمونه‌هایی است که درست دسته‌بندی شده‌اند).

نرخ بازشناسی کلی سیستم برای حروف پس انجام پس پردازش به ۹۸٪ درصد می‌رسد. به عنوان نمونه در مورد دو گروه ۲ و ۱۲ (بدنه‌های «ک» و «ب») که تداخل زیادی دیده می‌شود تصمیم‌گیری نهایی با دخیل کردن اطلاعات ریزحرکات صورت می‌گیرد.

جدول (۴) : مقایسه‌ی نرخ بازشناسی صحیح این سیستم با دو کار دیگر بر اساس مجموعه داده‌ی Online-TMU

سیستم ارائه شده در [۲]	۹۰/۲۷٪
سیستم ارائه شده در [۳]	۹۳/۳٪
سیستم پیشنهادی قبلی در [۱]	۹۷٪
سیستم پیشنهادی جدید در این مقاله	۹۸٪

جدول (۴) نرخ بازشناسی کلی این سیستم را با سه سیستم دیگر که آنها نیز برای ارزیابی از مجموعه داده‌ی Online-TMU استفاده کرده‌اند نشان می‌دهد. مشاهده می‌شود که در صورت یکسان بودن شرایط آزمایش‌ها، سیستم پیشنهادی ما نرخ بازشناسی صحیح بالاتری را ارائه می‌کند.

۸- خلاصه، بحث و نتیجه‌گیری

در این مقاله، روشی ساده برای بازشناسی برخط حروف فارسی ارائه شده است که مبنای آن، بازشناسی بدنه‌ی اصلی حروف با استفاده از SVM می‌باشد. بدین صورت که در گام اول بدنه اصلی حروف تشخیص داده می‌شوند سپس با استفاده از موقعیت و شکل ریز حرکات نتیجه نهایی مشخص می‌شود.

آنچه بر مبنای نتایج تجربی به نظر می‌رسد، این است که تشخیص بدنه‌ی اصلی حروف با روش SVM نتایج مطلوبی را در بر دارد، اما بنا به دلایلی چون تعدد شکل ریزحرکات (مانند نوشتن سه نقطه در قالب یک، دو یا سه حرکت) و عدم انحصار