

تشخیص هویت از دست نوشته با رویکرد هوش مصنوعی

کاظم فلاح یخدانی

دکتری مهندسی کامپیوتر دانشگاه آزاد اسلامی واحد میبد

چکیده

دست خط افراد ممکن است شباهت‌های زیادی به هم داشته باشند، شباهت‌هایی که به صورت اتفاقی یا با تمرین ایجاد می‌شود. دست خط‌های مشابه احتمال کلاهبرداری و سرقت‌های مادی و معنوی را افزایش می‌دهد. خالی شدن حساب‌های بانکی یا منتشر شدن اطلاعات شخصی و محرمانه تنها نمونه‌های کوچکی از پیامدهای جعل دست خط است. روش تحقیق در این پژوهش ابتدا شبکه‌های عصبی پیچشی یا همگشتی (CNN) و پس از آن، شبکه‌های حافظه طولانی کوتاه‌مدت (LSTM) می‌باشد. نتایج نشان دادند روش پیشنهادی قادر به یادگیری ویژگی‌ها و به دست آوردن نتایج شناسایی قانع‌کننده‌ای است. در مقایسه با ویژگی‌های مهندسی، روش پیشنهادی صحت شناسایی را حدوداً ۱۳ درصد افزایش می‌دهد و همچنین وابستگی کمتری به دانش تخصص دارد. با توجه به نتایج ارائه‌شده، می‌توان بیان کرد روش پیشنهادی برای شناسایی خودکار نویسنده بسیار رضایت‌بخش و مناسب است و می‌تواند با ورود به حوزه کاربردی، دستیار خوبی برای متخصصان شناسایی دستخط باشد.

کلمات کلیدی: جعل دست نوشته، هوش مصنوعی، شبکه‌های عصبی کانولوشنال

مقدمه

با گذشت سال‌ها، افراد، صنایع و دولت‌ها در سراسر جهان به‌طور فزاینده‌ای با روش‌های جدیدتر و مدرن‌تری از کلاهبرداری و سرقت روبه‌رو هستند. در این میان جعل امضا و هویت در کنار تحمیل خسارت‌های سنگین به دولت‌ها و جوامع، امنیت افراد و سازمان‌ها را هم به‌طور جدی تهدید کند؛ اما تکنولوژی همان‌طور که مسئله و چالشی را ایجاد کرده، راهکاری هم برای این مسائل پیدا می‌کند. هوش مصنوعی همان راهکاری است که می‌تواند دست‌خط و امضاهای مختلف را تجزیه و تحلیل کند و جعلی بودن آن‌ها را تشخیص دهد. در طول زمان‌های گذشته، از دست‌خط برای ایجاد ارتباط میان مردم استفاده شده که شامل نمادهایی برای نمایش زبان‌های مختلف است. دست‌نوشته‌ها ویژگی‌های قابل اندازه‌گیری دارند که می‌توانند نویسندگان را توصیف کنند. شناسایی دست‌نوشته‌ها یکی از فعال‌ترین حوزه‌های تحقیقاتی در زمینهٔ بنیایی ماشین و پردازش الگو است. از کاربردهای تشخیص دست‌نوشته‌ها می‌توان به دسته‌بندی نامه‌های پستی، خواندن خودکار مبالغ چک‌های بانکی، شناسایی خودکار اطلاعات ثبت‌شده در فرم‌ها، شناسایی اسناد دست‌نویس جعلی و غیره اشاره کرد (زرنندی و همکاران، ۱۴۰۳). در تمامی این کاربردها، همواره مسئله دقت و سرعت سیستم اهمیت بسیار زیادی داشته است. جرم جعل در اسناد دست‌نویس به سرعت در حال رشد است؛ بنابراین، شناسایی اسناد دست‌نویس جعلی به‌عنوان یک مشکل بسیار حیاتی برای سازمان‌های دولتی (مانند دادگستری و دفاتر ثبت اسناد) در نظر گرفته می‌شود. به‌طور کلی، شناسایی اسناد جعلی پیچیدگی زیادی دارد. روش عملی در کشف اسناد جعلی شامل دو مرحله است؛ به دست آوردن نمونه‌های دست‌خط از مظنونین و استخراج تفاوت بین ویژگی‌های دست‌خط. تمامی روش‌ها بر پایهٔ استخراج ویژگی‌ها از سند اصلی و مقایسهٔ آنها با سند جعلی استوار است. ویژگی‌های دست‌خط به سبک‌های نوشتاری زبان بستگی دارد. این به‌عنوان یک نشانهٔ اساسی در تجزیه و تحلیل اسناد دست‌نویس در نظر گرفته می‌شود (خسروشاهی و همکاران، ۱۴۰۱).

بیشتر محققان زبان انگلیسی را برای تجزیه و تحلیل اسناد دست‌نویس انتخاب کرده‌اند. با این حال، تعداد کمی از محققان روی زبان‌های راست به چپ مانند عربی، فارسی، اردو و غیره تمرکز کرده‌اند؛ بنابراین، مجموعه داده‌های استاندارد بسیار کمی در زبان‌های راست به چپ وجود دارد. زبان‌های راست به چپ به‌طور گسترده در میان اقشار مختلفی از مردم جهان استفاده می‌شوند و از نظر چگونگی پیوند چند حرف با یکدیگر سبک نوشتاری خاصی دارند. مشکل اصلی در تحلیل دست‌نوشته‌های راست به چپ، کمبود مجموعه داده‌های جامع است. مجموعه داده‌ها شامل انواع حروف، کلمات، جملات، اعداد و انواع مختلف اتصالات است. یک دلیل و توجیه خلاً تحقیق دربارهٔ زبان‌های راست به چپ این است که آنها از روش‌های پیچیده‌ای برای اتصال حروف استفاده می‌کنند؛ از این‌رو، آنها به‌عنوان سبک‌های نوشتاری ناشناخته و کم مطالعه باقی مانده‌اند (شریف زاده و همکاران، ۱۴۰۳).

به‌طور کلی، مطالعات تحقیقاتی در زمینهٔ اسناد دست‌نویس، به سه گروه عمده طبقه‌بندی می‌شوند: تشخیص دست‌خط، شناسایی نویسنده و تأیید نویسنده. نیاز سازمان‌های دولتی (مانند دادگستری و دفاتر ثبت اسناد) عمدتاً مربوط به شناسایی و تأیید نویسنده است. شناسایی و تأیید نویسنده معمولاً توسط متخصصان براساس شناسایی بصری انجام می‌شوند که این کار زمان‌بر، خسته‌کننده و نادرست است. مطالعات تحقیقاتی بسیاری برای حل این مشکلات انجام شده‌اند. برخی از مطالعات بر پایهٔ رویکردهای سنتی یادگیری ماشین و برخی دیگر بر پایهٔ رویکردهای جدید یادگیری ماشین، مانند یادگیری عمیق استوارند. بسیاری از محققان روش‌های طبقه‌بندی را برای مسئله شناسایی نویسنده پیشنهاد داده‌اند. استخراج ویژگی در یادگیری ماشین نقشی اساسی دارد و یک زمینه مهم در ادبیات فن محسوب می‌شود. در رویکردهای سنتی یادگیری ماشین، مراحل استخراج و طبقه‌بندی ویژگی‌ها از یکدیگر جدا هستند؛ در حالی که در رویکردهای جدید یادگیری ماشین، این مراحل با هم ادغام شده‌اند. این رویکردها، برخلاف رویکردهای سنتی، نیازی به دانش قبلی از مسئله ندارند و می‌توانند به‌صورت سلسله‌مراتبی ویژگی‌های مطلوب را از داده‌های خام استخراج کنند (زارع چاهوکی و کارگرشورکی، ۱۳۹۵). به‌طور کلی، تشخیص دست‌نوشته‌ها به‌صورت مرسوم به دو دسته آنلاین (برخط) و آفلاین (غیر برخط) تقسیم می‌شوند. در تشخیص آنلاین ترتیب زمانی از مختصات دریافت می‌شود که بیان‌کنندهٔ حرکات نوک قلم شخص است؛ در حالی که در روش‌های آفلاین تنها تصویر متن در دسترس است.

آوایدا و همکاران یک روش شناسایی نویسنده بر پایه متون عربی با استفاده از ویژگی‌های آماری و ساختاری ارائه دادند. در این روش الگوریتم نزدیک‌ترین همسایه به همراه معیارهای فاصله اقلیدسی استفاده شد. همچنین، از الگوریتم‌های کاهش داده، برای کاهش ابعاد داده استفاده شد. دویست و پنجاه نویسنده، یک پایگاه داده ۵۰۰ پاراگرافی را به زبان عربی نوشته‌اند. نتایج نشان‌دهنده عملکرد مطلوب روش بیان‌شده در شناسایی نویسنده بود. شهابی و همکاران روشی را برای شناسایی آفلاین اسناد فارسی دست‌نویس با استفاده از فیلتر چندکاناله گابور^۱ ارائه دادند. این روش می‌توانست ویژگی‌ها را استخراج کند و یک مجموعه داده محدود را با توجه به معیارهای فاصله اقلیدسی طبقه‌بندی کند. با روش بیان‌شده عملکرد خوبی درباره اسناد دست‌نویس فارسی حاصل شد. باغساز و همکاران روش جدیدی را برای شناسایی آنلاین (برخط) دست‌نوشته‌های فارسی ارائه دادند (زرندی، ۱۴۰۰). در روش بیان‌شده، ویژگی‌ها با توجه به معیارهای فاصله اقلیدسی استخراج شدند و پس از آن یک طبقه‌بند فازی برای شناسایی اسناد دست‌نویس استفاده شد. این روش به دست‌نوشته‌های زبان فارسی اعمال شد و در مقایسه با سایر رویکردهای موجود به صحت بالاتری از شناسایی دست‌یافت. احمد و همکاران یک طبقه‌بند مبتنی بر فاصله را با مجموعه‌ای از ویژگی‌های استخراج‌شده با الگوریتم مور^۲ به‌منظور شناسایی آفلاین نویسنده بررسی کردند. این روش به چهار مجموعه داده جامع اعمال شد و به میزان دقت چشمگیری در شناسایی اسناد دست‌نویس دست‌یافت (مسیبی اطاقسرا و همکاران، ۱۴۰۳).

بررسی مطالعات شناسایی نویسنده نشان می‌دهد اگرچه تاکنون مطالعات زیادی در این زمینه انجام شده است، محدودیت‌هایی در این مطالعات وجود دارد. در بیشتر این مطالعات، از روش‌های سنتی مبتنی بر استخراج و انتخاب ویژگی‌ها برای شناسایی نویسنده استفاده شده است. علاوه بر این، در بیشتر این مطالعات شرایط محیطی مختلف در تهیه پایگاه داده‌های مختلف در نظر گرفته نشده است؛ در صورتی که برای ورود به حوزه کاربردی، در نظر گرفتن کلیه شرایط محیطی ضروری است. براساس این، اولین هدف این مقاله ارائه یک سیستم شناسایی آفلاین نویسنده در شرایط آزمایشی مختلف است که مستقل از هر زبانی باشد. همچنین، با بررسی مطالعات پیشین مشاهده می‌شود مجموعه داده جامع در رابطه با دست‌نوشته‌های راست به چپ وجود ندارد که بتواند به‌عنوان یک مجموعه پایگاه داده مرجع برای بررسی زبان‌های راست به چپ استفاده شود. درواقع دست‌نوشته‌های راست به چپ به‌طور عمده در مطالعات مرتبط نادیده گرفته شده‌اند (جوادی مقدم و همکاران، ۱۴۰۳). براساس این، در دومین هدف این مطالعه، تلاش شده است با تمرکز بر شناسایی نویسنده براساس دست‌نوشته‌های راست به چپ، این شکاف تحقیقاتی بررسی شود که به‌عنوان یک موضوع بسیار مهم و بحث‌برانگیز در سازمان‌های دولتی (مانند دادگستری و دفاتر ثبت اسناد) شناخته می‌شود. برای این منظور، یک مجموعه داده راست به چپ شامل کلمات، جملات و اعداد جمع‌آوری شده است. این مجموعه داده شامل ۸۶۳۰۴ نمونه از افراد مختلف با جنسیت، گروه سنی، شغل و سطح تحصیلات مختلف است. این مجموعه داده در فواصل زمانی مختلف در شرایط آزمایشی مختلف براساس استانداردهای آمریکایی (ASTM) جمع‌آوری شده است. همچنین، یادگیری عمیق به‌طور گسترده و با موفقیت زیادی در تجزیه و تحلیل تصاویر و سیگنال‌ها استفاده شده است. در سومین هدف این مقاله، یک مدل شبکه عصبی کانولوشن عمیق بهبودیافته بر پایه شبکه از پیش آموزش‌دیده^۳ طراحی شده است تا ویژگی‌ها را به‌صورت سلسله‌مراتبی از داده‌های خام دست‌خط یاد بگیرد. مهم‌ترین جنبه مدل پیشنهادی قابلیت آن در طبقه‌بندی مجموعه داده‌های نامتجانس است؛ به این معنا که در هر دوره، اگرچه نمونه‌های تصادفی برای مراحل آموزش و ارزیابی، به یک شخص خاص تعلق داشته است، ممکن است لزوماً یکسان نباشند؛ حتی ممکن است هیچ شباهتی نداشته باشند. استفاده از نمونه‌های نامتجانس چهارمین هدف این مقاله است که به‌طور عمده در مطالعات قبلی نادیده گرفته شده است. درواقع، این نوآوری در روش شناسایی، برجسته‌ترین جنبه مطالعه حاضر است.

¹ Multi-Channel Gabor Filter (MGF)

² Moor

³ pre-trained network

چالش‌های تشخیص دست‌خط جعلی

- ممکن است دست‌خط افراد مختلف به‌صورت کاملاً تصادفی به هم شباهت زیادی داشته باشد.
- دست‌خط یک فرد طی زمان خواه ناخواه تغییر می‌کند و همیشه فرم و الگوی ثابتی ندارد.
- اسناد یا تصاویر قدیمی ممکن است دچار فرسودگی شده باشند و کیفیت پایینی داشته باشند.
- متن اسناد چاپ‌شده در یک خط مستقیم قرار دارند؛ در حالی که دست‌خط افراد روی یک کاغذ سفید لزوماً در یک مسیر مستقیم و منظم شکل نمی‌گیرد.
- دست‌خط‌های شکسته و پیوسته تشخیص حروف را بسیار دشوار می‌کند.
- برای یادگیری عمیق نرم‌افزارهای تشخیص جعل دست‌خط، جمع‌آوری داده‌های اصلی و موثق نسبت به داده‌های مصنوعی دشوارتر و گران‌تر است.

روش تحقیق

در این بخش ابتدا شبکه‌های عصبی پیچشی یا همگشتی (CNN)^۴ و پس از آن، شبکه‌های حافظه طولانی کوتاه‌مدت (LSTM)^۵ بررسی می‌شوند که زیرمجموعه‌ای از شبکه‌های عصبی بازگشتی‌اند.

شبکه‌های عصبی کانولوشنال

شبکه عصبی کانولوشنال، درواقع یک شبکه عصبی بهبودیافته است. در این شبکه، چندین لایه با روشی قدرتمند در کنار هم آموزش می‌بینند. این روش، بسیار کارآمد بوده و یکی از رایج‌ترین روش‌ها در کاربردهای مختلف بینایی ماشین است (دینگ و همکاران،^۶ ۲۰۲۵). همانند شبکه‌های عصبی مصنوعی^۷، تصمیم خروجی نهایی شبکه عصبی کانولوشنال براساس وزن و بایاس لایه‌های قبلی در ساختار شبکه است (زرندی و جیوان، ۲۰۱۵).

شبکه‌های عصبی بازگشتی

شبکه‌های عصبی بازگشتی شاخه‌ای مهم از شبکه‌های عصبی عمیق‌اند که به‌منظور تحلیل سیستم‌های پیچیده استفاده می‌شوند. این شبکه‌ها می‌توانند با کاهش ابعاد داده ورودی، بار محاسباتی را کاهش دهند و همچنین باعث بهبود عملکرد آموزش شوند. علاوه بر این، این شبکه‌ها امکان تلفیق اطلاعات بین ورودی‌های مختلف را به‌منظور دستیابی به ویژگی‌هایی فراهم می‌کنند که نمی‌توان با استفاده از روش‌های سنتی استخراج کرد (طاهری حاجی وند و همکاران،^۶ ۱۴۰۳). شبکه حافظه طولانی کوتاه‌مدت ازجمله شبکه‌های عصبی بازگشتی‌اند که به‌منظور رفع ضعف‌های شبکه‌های بازگشتی ازجمله حل مشکل پراکندگی گرادیان یا مشکلات انفجاری گرادیان به کار برده می‌شوند. برخلاف شبکه عصبی بازگشتی که صرفاً جمع متوازن سیگنال‌های ورودی را محاسبه می‌کند و سپس از یک تابع فعال‌ساز عبور می‌دهد، هر واحد LSTM از یک حافظه C_t در زمان t بهره می‌برد (زارع و همکاران،^۶ ۱۴۰۳). یک سلول حافظه از چهار عنصر اصلی تشکیل شده است: یک دروازه ورودی یا دروازه به‌روزرسانی، یک نورو با اتصال خودبازگشتی، یک دروازه فراموشی و یک دروازه خروجی فعال‌سازی واحد LSTM به‌صورت رابطه زیر تعریف می‌شود:

$$h_t = \Gamma_o \cdot \tanh(C_t)$$

دروازه خروجی با رابطه زیر محاسبه می‌شود:

$$\Gamma_o = \sigma(W_o \cdot [h_{t-1}, X_t] + b_o)$$

^۴ Convolutional Neural Network (CNN)

^۵ Long Short Term Memory (LSTM)

^۶ Ding, H., Hou, H., Wang, L., Cui, X., Yu, W., & Wilson, D. I.

^۷ Artificial neural networks (ANNs)

سلول حافظه نیز با فراموشی نسبی حافظه فعلی و اضافه کردن محتوای حافظه جدید به صورت از رابطه زیر به روزرسانی می شود که در آن، محتوای حافظه جدید از رابطه زیر به دست می آید:

$$\tilde{C}_t = \tanh(W_C \cdot [h_{t-1}, X_t] + b_C)$$

$$C_t = \Gamma_f \cdot C_{t-1} + \Gamma_u \cdot \hat{C}_t$$

میزان حافظه فعلی که باید فراموش شود، توسط دروازه فراموشی کنترل می شود و مقدار حافظه جدید که باید به سلول حافظه اضافه شود، توسط دروازه به روزرسانی (دروازه ورودی) انجام می گیرد. این عملیات در رابطه های زیر نشان داده شده است:

$$\Gamma_f = \sigma(W_f \cdot [h_{t-1}, X_t] + b_f)$$

$$\Gamma_u = \sigma(W_u \cdot [h_{t-1}, X_t] + b_u)$$

شکل ۱ ساختار یک شبکه عصبی بازگشتی LSTM را نشان می دهد. در این شبکه که یک ورودی یا همان دارد، دو خروجی تولید می شود: یک خروجی و خروجی دیگر؛ به دو بخش تقسیم می شود؛ بخشی به گام زمانی بعد، منتقل و بخشی نیز در صورت نیاز به تولید خروجی در گام زمانی فعلی استفاده می شود. دروازه فراموشی وظیفه کنترل جریان اطلاعات از گام زمانی قبلی را دارد. این دروازه مشخص می کند آیا اطلاعات حافظه از گام زمانی قبل استفاده شود یا خیر و اگر باید از گام زمانی قبل چیزی وارد شود، به چه میزان باشد. دروازه به روزرسانی وظیفه کنترل جریان اطلاعات جدید را بر عهده دارد. این دروازه مشخص می کند آیا در گام زمانی فعلی باید از اطلاعات جدید استفاده شود یا خیر و اگر بلی به چه میزان. دروازه خروجی نیز مشخص می کند چه میزان از اطلاعات گام زمانی قبل با اطلاعات گام زمانی فعلی به گام زمانی بعد منتقل شود.

روش پیشنهادی

در این بخش، روش پیشنهادی مقاله ارائه می شود. شکل ۲ بلوک دیاگرام الگوریتم پیشنهادی را نشان می دهد.

جمع آوری داده ها

در روش پیشنهادی یک پایگاه داده جامع طراحی و توسعه داده شده است. همچنین، در این مطالعه علاوه بر پایگاه داده طراحی شده از پایگاه داده های IAM، CVL، KHATT و IFN/ENIT نیز استفاده شده است تا روش پیشنهادی بهتر ارزیابی شود. پایگاه داده IAM شامل ۴۸۹۹ نمونه از ۱۵۰ نویسنده و پایگاه داده CVL شامل ۱۸۵۴ نمونه از ۳۰۹ نویسنده به زبان انگلیسی است. پایگاه داده KHATT شامل ۱۰۸۹۸ نمونه از ۸۲۸ نویسنده و پایگاه داده IFN/ENIT شامل ۲۶۴۵۹ نمونه از ۴۱۱ نویسنده به زبان عربی است. این پایگاه داده ها در مطالعات اخیر بسیار استفاده شده اند و در زمینه شناسایی نویسنده از پایگاه داده های قابل اعتماد و پر کاربرد هستند.

در پایگاه داده پیشنهادی، بر پایه استانداردهای ASTM، نمونه های دستخط از ۶۲ شرکت کننده در بازه های زمانی مختلف در شرایط محیطی مختلف جمع آوری شدند (شکل ۳ را ببینید). از ۶۲ شرکت کننده در این آزمایش، ۳۴ نفر مرد و ۲۸ نفر زن، با میانگین سنی ۲۲ تا ۵۴ سال اند. همچنین ۶ نفر از این شرکت کنندگان چپ دست و ۵۶ نفر راست دست بودند. در نهایت، براساس استانداردهای از پیش تعریف شده، ویژگی های بافتی و ساختاری دستخط به دست آمدند. ویژگی های بافتی مربوط به یک سند دست نویس به عواملی مانند کاغذ، جوهر، ابزارهای نوشتن و غیره بستگی دارد؛ در صورتی که ویژگی های ساختاری به سبک های نوشتاری، نحوه پیوستن حروف و غیره اشاره دارد (حری، ۱۴۰۳). گفتنی است زبان فارسی به عنوان زبان مقصد انتخاب شده که مجموعه داده مربوط به آن به دست آمده است. انتخاب این زبان به این علت است که محققان به نمونه های این زبان دسترسی آسان و کافی داشتند.

هر حرف از زبان فارسی چهار فرم نوشتاری مختلف دارد؛ فرم جداگانه^۸ (S)، فرم شروع کلمه^۹ (BOW)، فرم میانی کلمه^{۱۰} (MOW) و فرم انتهایی کلمه^{۱۱} (EOW). انتخاب فرم به موقعیت آن در یک کلمه بستگی دارد (مشایخی و همکاران، ۱۴۰۳). از شرکت‌کنندگان چهار جمله مربوط به چهار فرم بیان شده گرفته شد و از آنها خواسته شد تا دستورالعمل‌های زیر را دنبال کنند: در مرحله اول، هر جمله دوازده بار براساس استاندارد ASTM روی یک برگ کاغذ جداگانه نوشته شد. در مرحله دوم، هر چهار جمله در یک برگ جداگانه دیگر نوشته شدند. از دو نوع کاغذ استاندارد متفاوت «PaperOne» و «Double-A» استفاده شد که مشخصات آنها در جدول ۱ نشان داده شده است.

در این مطالعه از دو نوع مختلف از خودکارهای استاندارد، به مارک «Schneider» و «Faber-Castell» با رنگ‌های «آبی» و «سیاه» استفاده شد. مشخصات این خودکارها در جدول ۲ آورده شده است. نمونه‌هایی از فرم‌های نوشتاری شرکت‌کنندگان در شکل ۳ نشان داده شده‌اند. گفتنی است شرکت‌کنندگان فرم اصلی را به زبان فارسی پر کردند و به‌ازای هر فرد، مجموعه کاملی از نمونه‌ها گرفته شد.

همه نمونه‌ها روی دو پد نوشتاری مختلف نوشته شدند که به‌عنوان پدهای «سخت» و «نرم» هستند. استفاده از این دو نوع پد، برای نشان دادن میزان فشار قلم در نظر گرفته شده است. نمونه‌های جمع‌آوری شده توسط RICOH Aficio MP ۶۰۰۱ با رزولوشن ۳۰۰ dpi در حالت رنگی اسکن شدند. ترتیب نمونه‌های جمع‌آوری شده، اطلاعات و جزئیات موردنیاز به‌عنوان کتاب کد در مجموعه داده‌ها ذخیره شده است. مجموعه داده جمع‌آوری شده از ۶۲ شرکت‌کننده شامل ۴۴۱ صفحه و ۴۲۰۳ جمله از هم جدا شده است. همان‌طور که در شکل ۴ نشان داده شده است، ارتفاع نمونه‌های جمله‌ای ۲۳۶ پیکسل و عرض آنها متغیر است. اندازه صفحات نمونه‌ها برابر با ۱۶۵۶×۲۳۳۹ پیکسل است.

شبکه عمیق پیشنهادی

شبکه عمیق پیشنهادی در این مطالعه از ترکیب یک شبکه از پیش آموزش‌دیده کانولوشنال-Resnet-۱۵۲ با شبکه LSTM ایجاد شده است. با ترکیب شبکه-Resnet-۱۵۲ با شبکه LSTM می‌توان از مزایای هر دو شبکه به‌طور هم‌زمان استفاده کرد. در بسیاری از مطالعات، از ترکیب شبکه‌های LSTM با شبکه‌های کانولوشنال عمیق به‌منظور کاهش ابعاد ویژگی، افزایش پایداری، کاهش نوسانات، بهبود فرایند آموزش و افزایش صحت شناسایی استفاده شده است. شبکه پیشنهادی بر پایه شبکه از پیش آموزش‌دیده-Resnet-۱۵۲ با یک بلوک پیشنهادی ترکیب می‌شود که شامل دو لایه LSTM، سه لایه نرمال‌سازی دسته‌ای، سه لایه حذف تصادفی و دو لایه FC است (شکل ۲ را ببینید). شبکه‌های از پیش آموزش‌دیده از چند لایه تشکیل شده‌اند که هر لایه ویژگی‌های خاصی را یاد می‌گیرد (زرنی، ۱۴۰۲). لایه‌های اولیه ویژگی‌های پایه‌ای و سطح پایین و لایه‌های بعدی ویژگی‌های پیچیده و سطح بالا را یاد می‌گیرند. در این فرایند ماتریس وزن با روند آموزش تشکیل و تنظیم می‌شود. معماری بلوک پیشنهادی به‌صورت زیر انتخاب شده است: (۱) یک لایه FC با تابع خطی^{۱۲} به همراه یک لایه نرمال‌ساز دسته‌ای با تابع Relu که پس از آن، یک لایه حذف تصادفی قرار می‌گیرد. (۲) یک لایه LSTM با تابع Relu که پس از آن، لایه‌های نرمال‌ساز دسته‌ای و حذف تصادفی قرار می‌گیرند. (۳) معماری مرحله قبل، یک‌بار دیگر تکرار می‌شود. (۴) یک لایه FC با تابع غیرخطی سافت‌مکس برای دسترسی به لایه خروجی استفاده می‌شود. در شبکه پیشنهادی، خروجی شبکه از پیش آموزش‌دیده یک بردار ویژگی با اندازه ۵۱۲×۲۵۶ است. در اولین لایه بلوک پیشنهادی، یعنی FC، تابع خطی روی وزن‌های قابل یادگیری

⁸ Separated form (S)

⁹ Beginning Of Word form (BOW)

¹⁰ Middle Of Word form (MOW)

¹¹ End Of Word form (EOW)

¹² Linear

ویژگی‌های به‌دست‌آمده (W) اعمال می‌شود. مقادیر پیش‌بینی‌شده بایاس در نظر گرفته می‌شود تا ابعاد بردار ویژگی را به 1×256 تغییر دهد. همان‌طور که ملاحظه می‌شود، کاهش ابعاد در لایه‌های پنهان از 112×112 (اندازه ورودی) به 128 (بردار ویژگی انتخاب‌شده) ادامه یافته است که در نهایت، بردار ویژگی انتخاب‌شده به یک لایه FC با تابع غیرخطی سافت‌مکس متصل می‌شود (شکل ۲ را ببینید).

در این مطالعه همه ابر پارامترهای شبکه پیشنهادی به دقت تنظیم شده‌اند تا بهترین نرخ همگرایی را به دست آورند و در نهایت، تابع خطای کراس آن‌تروپی^{۱۳} و بهینه‌ساز SGD ^{۱۴} با نرخ یادگیری 0.01 انتخاب شده‌اند. روش مرسوم پس‌انتشار خطا با اندازه دسته ۱۰۰ برای آموزش شبکه استفاده شده است. ابر پارامترهای بهینه انتخاب‌شده برای مدل پیشنهادی در جدول ۳ نشان داده شده است.

براساس مطالعات صورت‌گرفته در ادبیات پیشین، از تکنیک‌های افزایش داده برای بهبود صحت، جلوگیری از پدیده بیش‌برازش و بهبود روند آموزش در شبکه‌های عصبی استفاده می‌شود؛ با وجود این، تکنیک‌های افزایش داده باید با دقت انتخاب شوند و از هر تکنیک افزایش داده نمی‌توان برای داده‌های دست‌نویس استفاده کرد. در این مطالعه تکنیک‌های مقیاس خاکستری تصادفی^{۱۵}، تغییر رنگ^{۱۶} و چرخش تصادفی^{۱۷} به‌عنوان تکنیک‌های افزایش داده انتخاب شده‌اند. نمونه‌هایی از داده‌های افزوده‌شده در شکل ۷ نشان داده شده‌اند. پس از استفاده از تکنیک‌های افزایش داده، مجموعه داده‌های آموزش ۵۰ درصد افزایش می‌یابند.

به‌عنوان آخرین مرحله برای افزایش و به حداکثر رساندن دقت پیش‌بینی، از تکنیک TTA ^{۱۸} استفاده می‌شود. همانند داده‌افزایی روی مجموعه داده آموزشی برای بهبود عملکرد مدل، هدف از TTA انجام تغییرات تصادفی روی مجموعه داده (تصاویر) آزمون است؛ بنابراین، به جای نشان دادن فقط یک‌بار تصاویر معمولی به مدل آموزش‌دیده، چندین‌بار تصاویر تقویت‌شده به آن نشان داده می‌شوند. پس از آن، میانگین پیش‌بینی‌های هر تصویر، محاسبه و به‌عنوان پیش‌بینی نهایی در نظر گرفته می‌شود. این روش، تکنیکی برای بهبود پیش‌بینی مدل است که به‌طور متداول برای محاسبه پیش‌بینی‌های میانگین در بسیاری از مطالعات طبقه‌بندی استفاده شده است. ساختار TTA در شکل ۸ نشان داده شده است. در این مطالعه از چرخش تصادفی به‌عنوان تکنیک افزایش داده در داده‌های آزمون استفاده شده است (خسروشاهی و همکاران، ۱۴۰۱).

همان‌طور که گفته شد، در این مطالعه آموزش و ارزیابی مدل پیشنهادی با استفاده از داده‌های نامتجانس انجام می‌شود. شکل ۹ نمونه‌های دست‌خط نامتجانس را برای فرایند آموزش و ارزیابی نشان می‌دهد. با توجه به این واقعیت که تمام تصاویر ۱ تا ۲۴ متعلق به یک فرد است، تصاویر ۱ تا ۱۵ در مجموعه داده‌های آموزشی و تصاویر ۱۶ تا ۲۴ به مجموعه داده‌های اعتبارسنجی و ارزیابی تعلق دارند؛ برای مثال، تصاویر ۱ و ۳ بیشترین شباهت را با تصویر ۱۸ (a) دارند. همچنین، تصاویر ۶ و ۱۶ یکسان‌اند (b)؛ در مقابل، تصاویر ۱۹ و ۲۲ هیچ‌گونه هم‌تابی در مجموعه داده‌های آموزشی ندارند (c).

از کل نمونه‌های جمع‌آوری‌شده (۸۶۳۰۴ نمونه)، ۵۱۷۸۲ نمونه برای داده‌های آموزش (۶۰ درصد)، ۲۵۸۹۰ نمونه (۳۰ درصد) برای داده‌های اعتبارسنجی و ۸۶۳۲ نمونه برای داده‌های آزمون (۱۰ درصد) استفاده می‌شود. علاوه بر این، تمام نمونه‌های اختصاص داده شده به مجموعه‌های آموزش و ارزیابی به‌طور تصادفی انتخاب می‌شوند.

¹³ Cross-Entropy

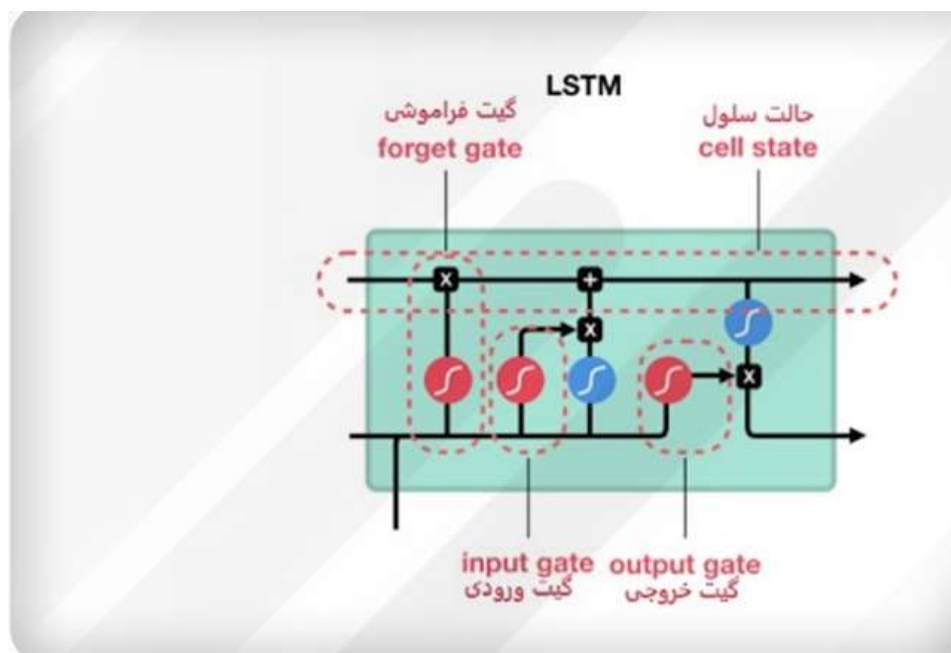
¹⁴ Stochastic Gradient Descent (SGD)

¹⁵ RandomGrayscale

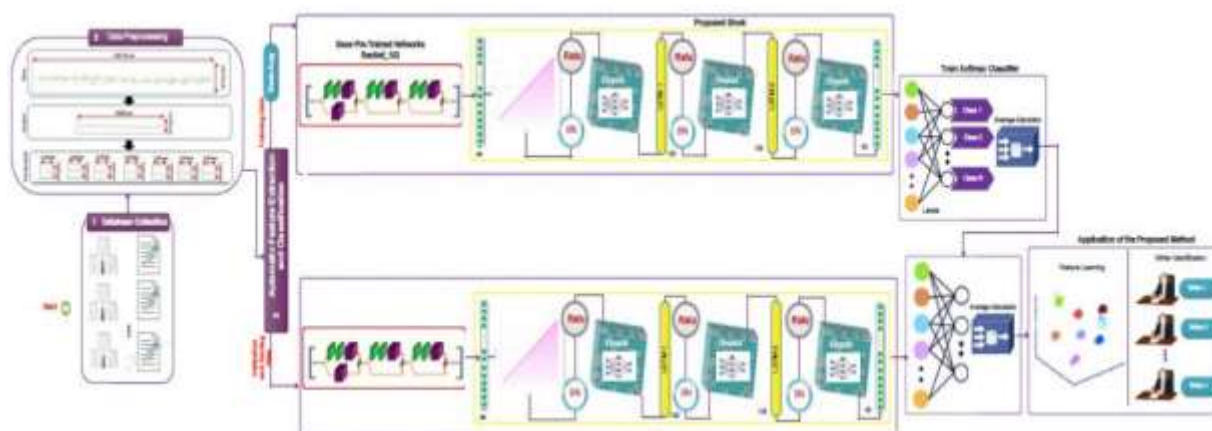
¹⁶ ColorJitter

¹⁷ RandomRotation

¹⁸ Test Time Augmentation (TTA)



شکل (۱): ساختار یک شبکه عصبی بازگشتی LSTM.



شکل (۲): بلوک دیاگرام الگوریتم پیشنهادی.

جدول (۱): جزئیات دو نوع کاغذ استاندارد استفاده شده در این مطالعه

DoubleA			PaperOne			واحد	خواص
روش	تولرانس	هدف	روش	تولرانس	هدف		
ISO 536	±3%	80	ISO 536	±4%	80	g/m ²	وزن پایه
ISO 534	±5/3	5/106	ISO 534	±3	110	μm	ضخامت
ISO8791	±50	150	ISO 2-8791	±40	140	ml/min	زبری
ISO 2470	±5/1	5/102	ISO 2470	±2	99	%	روشنایی ISO

ISO 2471	-	94	ISO 2471	± 2	95	%	تیرگی ISO
ISO 11475	± 2	160	ISO 11475	± 2	167	#	سفیدی CIE

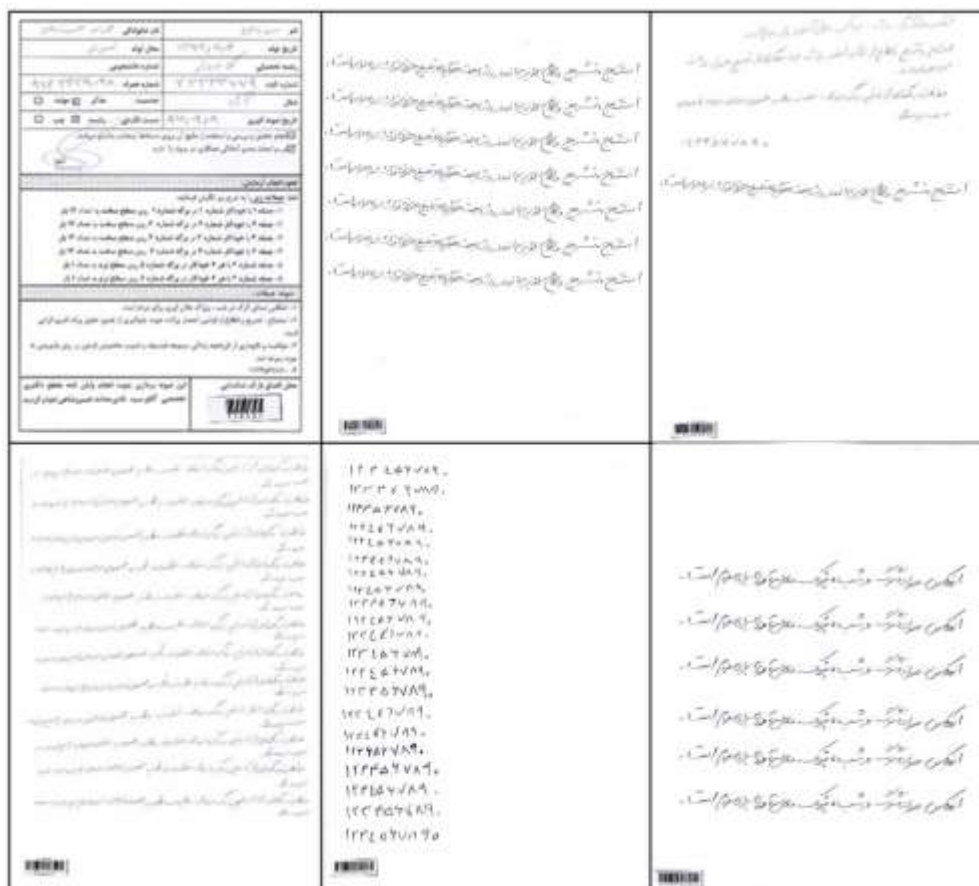
جدول (۲): مشخصات خودکارهای استاندارد استفاده شده در این مطالعه؛

(الف) مشخصات خودکار Schneider.

شماره مقاله	عرض خط	رنگ	مدل	GTIN	نوع
150503	Fine 0.4 mm	Blue	Tops 505	400467 5004 567	Ballpoint
150501	Fine 0.4 mm	Black	Tops 505	400467 5004 529	Ballpoint

(ب) مشخصات خودکار Faber-Castell.

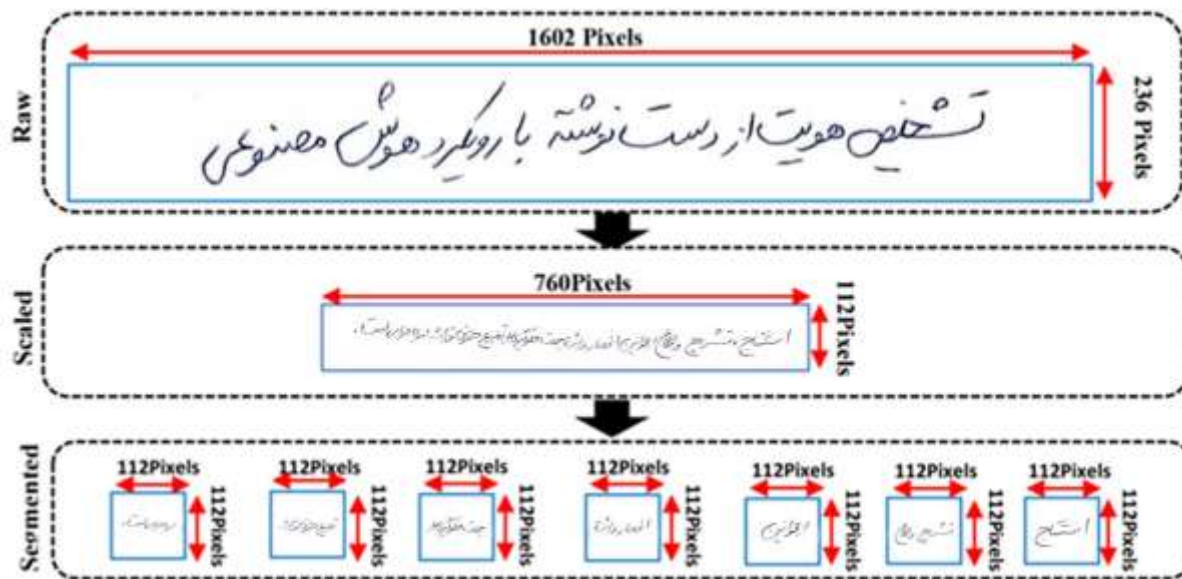
شماره مقاله	عرض خط	رنگ	مدل	GTIN	نوع
34 07 98	Fine 0.7 mm	Blue	TRI- Flow	8.90118E+12	Ballpoint
34 07 50	Fine 0.7 mm	Black	TRI- Flow	8.90118E+12	Ballpoint



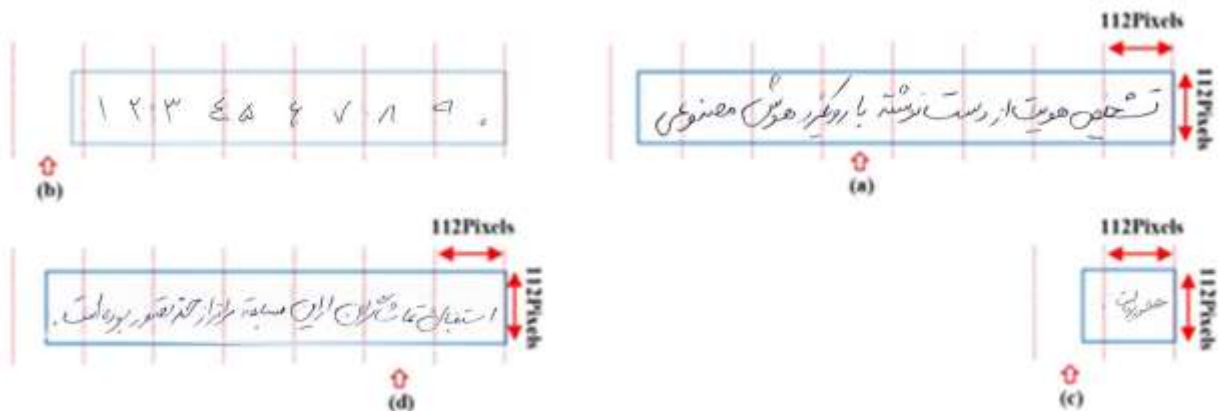
شکل (۳): نمونه‌هایی از فرم نوشتاری شرکت‌کنندگان



شکل (۴): اندازه نمونه‌ها؛ (الف) صفحه، (ب) جمله.



شکل (۵): عملیات تقسیم‌بندی یکی از جملات.

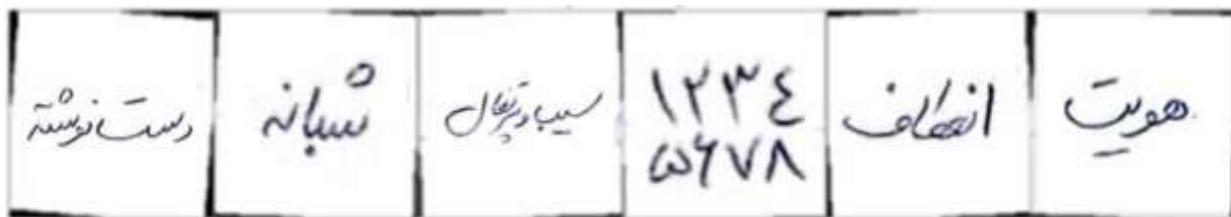


شکل (۶): بخش‌های فریبنده.

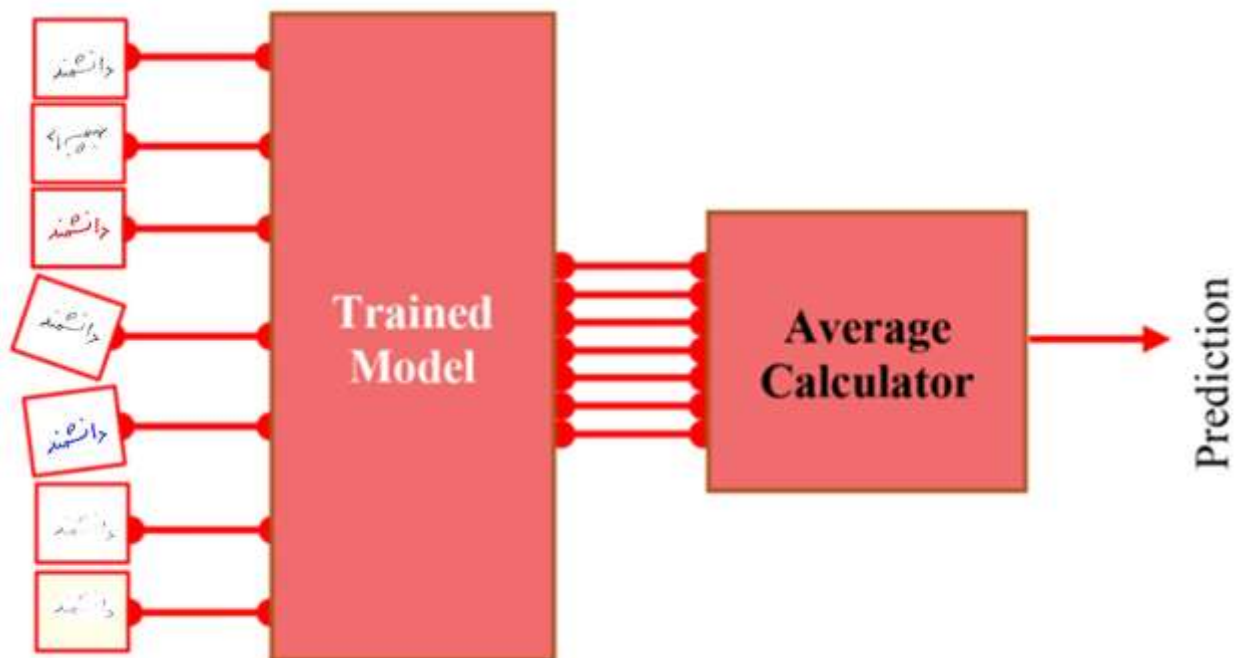
جدول (۳): ابر پارامترهای بهینه استفاده شده در مدل پیشنهادی.

پارامتر	فضای جستجو	مقدار بهینه
بهینه‌ساز	RMSProp, Adam, Adamax, SGD, AdaDelta	SGD
تابع خطا	MSE, Cross-entropy	Cross-entropy
نرخ لایه حذف تصادفی	0, 0.1, 0.3, 0.5, 0.7, 0.9	Feb-00
اندازه دسته	4, 32, 16, 10, 8, 64, 100	100
نرخ یادگیری	0.01, 0.001, 0.0001, 0.00001	Jan-00

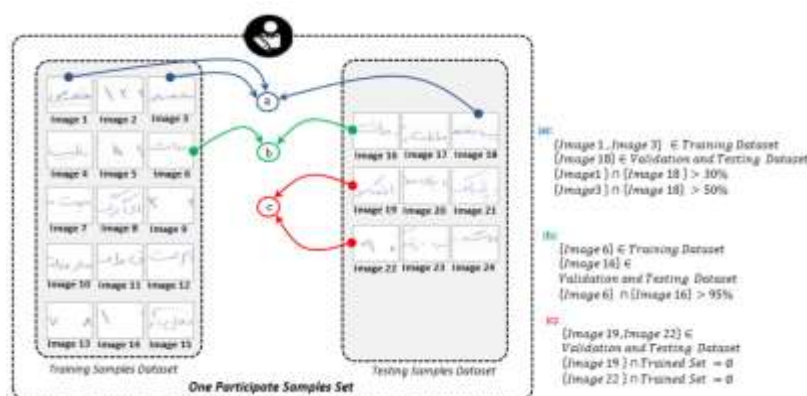
تابع فعال ساز در لایه اول بلوک پیشنهادی	Leaky-Relu, Sigmoid,Relu, Linear	Linear
تابع فعال ساز در لایه های مخفی بلوک پیشنهادی	Leaky-Relu, Sigmoid,Relu, Linear	Relu
تابع فعال ساز در لایه آخر بلوک پیشنهادی	Softmax,Sigmoid	Softmax



شکل (۷): نمونه هایی از داده های افزوده شده با استفاده از تکنیک داده افزایی.



شکل (۸): ساختار TTA.



شکل (۹): نمونه‌های دستخط نامتجانس.

نتایج و بحث

روش پیشنهادی ($P-M$)^{۱۹} شناسایی نویسنده و کلیه نتایج و بررسی‌ها در پایتون با استفاده از کتابخانه‌های متنوعی انجام شده‌اند که مهم‌ترین آنها PyTorch و NumPy هستند. این بررسی‌ها روی یک سیستم رایانه‌ای با مشخصات زیر انجام شده‌اند: پردازنده مرکزی Intel Core i-۷-۶۷۰۰K، پردازنده گرافیکی GTX TITAN X ۱۲GB، رم ۶۴ گیگابایت DDR IV و هارد دیسک ۱ ترابایت SSD. به منظور ارزیابی عملکرد روش پیشنهادی، از رابطه‌ی مربوط به صحت استفاده می‌شود که به صورت زیر بیان می‌شود:

$$\text{Accuracy} = \frac{T_P + T_N}{T_P + F_P + T_N + F_N}$$

که در آن، TP موارد مثبتی است که به درستی مثبت تشخیص داده شده است. FP موارد منفی است که به اشتباه مثبت تشخیص داده شده است. TN موارد منفی است که به درستی منفی تشخیص داده شده است. FN موارد مثبتی است که به اشتباه منفی تشخیص داده شده است.

نتایج تجربی مدل پیشنهادی (شبکه از پیش آموزش دیده Resnet-۱۵۲ همراه با بلوک پیشنهادی) و شبکه از پیش آموزش دیده Resnet-۱۵۲ بدون بلوک پیشنهادی در جدول ۴ آمده‌اند. مطابق این جدول، هر دو مدل هنگام استفاده از تکنیک TTA، عملکرد بهتری نسبت به استفاده از تکنیک TTA دارند. صحت ارزیابی مدل پیشنهادی ($P-M$) با تکنیک TTA، ۹۹/۶۶ درصد است؛ در حالی که صحت ارزیابی مدل پیشنهادی بدون استفاده از تکنیک TTA، ۹۵/۷۸ درصد است. همچنین صحت ارزیابی Resnet-۱۵۲ با تکنیک TTA، ۹۶/۵۱ درصد است؛ در حالی که صحت ارزیابی Resnet-۱۵۲ بدون استفاده از تکنیک TTA، ۹۳/۴۵ درصد است. همان‌طور که در جدول ۴ ملاحظه می‌شود، صحت و زمان اجرای مدل پیشنهادی از مدل Resnet-۱۵۲ بیشتر است.

همچنین نمودار صحت و خطا برای داده‌های اعتبارسنجی مدل پیشنهادی ($P-M$) و مدل Resnet-۱۵۲ با تکنیک TTA در شکل ۱۰ ارائه شده است. مطابق این شکل، خطای شبکه مدل پیشنهادی و مدل Resnet-۱۵۲ با افزایش تکرار الگوریتم کاهش می‌یابد. همچنین مشاهده می‌شود مدل پیشنهادی و مدل Resnet-۱۵۲ پس از ۲۰۰ تکرار تقریباً به صحت ۹۹ درصد و ۹۶ درصد می‌رسند.

¹⁹ Proposed Method (P-M)

نتایج شناسایی مدل پیشنهادی و مدل -Resnet۱۵۲ با تکنیک TTA برای شناسایی نویسندگان با استفاده از هریک از چهار مجموعه داده در جدول ۵ آورده شده‌اند. مطابق جدول ۵، مدل پیشنهادی مبتنی بر شبکه بهبودیافته -Resnet۱۵۲، از مدل -Resnet۱۵۲ برای شناسایی نویسندگان با استفاده از هریک از چهار مجموعه داده بهتر عمل می‌کند.

صحت ارزیابی روش‌های مختلف برای شناسایی نویسنده در جدول ۶ ارائه و در شکل ۱۱ مقایسه شده‌اند که نتایج مدل پیشنهادی (P-M) در جدول ۶ به صورت برجسته مشخص شده‌اند. مقایسه‌ها بر پایه چهار مجموعه داده بیان شده، یعنی IAM، CVL، KHATT و IFN/ENIT انجام می‌شوند. تعداد نویسندگان در هر بررسی نیز در جدول ۶ آمده است. گفتنی است تفاوت‌های ارائه شده در جدول ۶، درباره تعداد نمونه‌های بررسی شده در مجموعه داده‌های بیان شده، به دلیل در دسترس بودن آنها است. همان طور که در جدول ۶ و شکل ۱۱ نشان داده شده است، درباره تمام مجموعه داده‌های ارزیابی شده، مقادیر صحت طبقه‌بندی نشان‌دهنده عملکرد بهتر مدل پیشنهادی در مقایسه با سایر روش‌هاست. به منظور نشان دادن عملکرد مدل شبکه عصبی کانولوشن عمیق (DCNN) با مجموعه داده‌های DANA_HW به عنوان ورودی، صحت ارزیابی با استفاده از مدل‌های دیگر نیز به دست آمده است؛ براساس این، داده‌های خام DANA_HW و چندین ویژگی مهندسی از مجموعه داده‌های DANA_HW همراه با شبکه پس‌انتشار خطا (BPNN)^{۲۰} و ماشین بردار پشتیبان (SVM) به عنوان مدل‌های مقایسه‌ای انتخاب شده‌اند. تابع پایه شعاعی گوسین^{۲۱} به عنوان تابع کرنل ماشین بردار پشتیبان بوده و از روش جستجوی شبکه^{۲۲} برای بهینه‌سازی پارامترهای کرنل استفاده شده است. معماری شبکه BPNN از یک لایه مخفی تشکیل شده که در آن از تابع فعال‌ساز سیگموئید استفاده شده است. به منظور دستیابی به نتایج بهتری از مدل‌های BPNN و DCNN، ابر پارامترهای آنها با توجه به داده‌های مختلف تنظیم می‌شوند. پنج ویژگی استاندارد که نمی‌توانند از تغییرات زمانی تأثیر بگیرند به عنوان ویژگی‌های مهندسی انتخاب شده‌اند: مساحت^{۲۳}، مختصات مرکزی^{۲۴}، گریزازمرکز^{۲۵}، کشیدگی^{۲۶} و چولگی^{۲۷}. صحت ارزیابی روش‌های مختلف بر پایه یادگیری ویژگی از داده‌های خام و ویژگی‌های مهندسی در جدول ۷ ارائه شده است که نتایج مدل DCNN پیشنهادی با داده‌های خام به عنوان ورودی، یعنی روش پیشنهادی، در جدول ۷ برجسته شده‌اند. در روش پیشنهادی از معماری ارائه شده استفاده شده است. مقایسه عملکرد یادگیری ویژگی‌ها و ویژگی‌های مهندسی ارائه شده در جدول ۷ نشان می‌دهد یادگیری ویژگی از داده‌های خام با مدل DCNN پیشنهادی، نتایج بهتری نسبت به ویژگی‌های مهندسی ارائه می‌دهد (با افزایش صحت در حدود ۱۳ درصد). این نتیجه کاملاً به معماری منحصربه‌فرد DCNN پیشنهادی مربوط می‌شود که می‌تواند به صورت خودکار ویژگی‌های مفید را از داده‌های خام استخراج کند. علاوه بر این، استخراج ویژگی‌های مهندسی به دانش و تخصص قبلی نیاز دارد؛ در حالی که یادگیری ویژگی از داده‌های خام وابستگی کمتری به دانش تخصصی دارد. مدل DCNN پیشنهادی با یادگیری ویژگی از داده‌های خام نتایج بهتری ارائه می‌دهد؛ در حالی که همه مدل‌های بررسی شده، یعنی DCNN، BPNN و SVM نتایج مشابهی را برای ویژگی‌های مهندسی ارائه می‌دهند. این نشان می‌دهد DCNN بدون توانایی یادگیری ویژگی نمی‌تواند نتایج بهتری در شناسایی نویسنده نسبت به روش‌های سنتی ارائه دهد.

جدول (۴): نتایج ارزیابی مدل -Resnet۱۵۲ همراه با بلوک پیشنهادی (مدل پیشنهادی) در مقایسه با -Resnet۱۵۲ بدون بلوک پیشنهادی.

²⁰ Back-Propagation Neural Network (BPNN)

²¹ Gaussian Radial Basis Function (RBF)

²² Grid search method

²³ Area

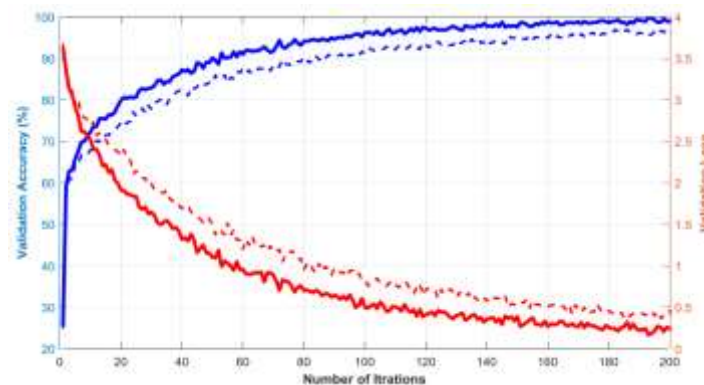
²⁴ Centroid Coordinates

²⁵ Eccentricity

²⁶ Kurtosis

²⁷ Skewness

نام شبکه	ارزیابی بدون TTA	ارزیابی با TTA	زمان ارزیابی با TTA (ms)	
	صحت آزمون	صحت آزمون	زمان آموزش	زمان آزمون
همراه با بلوک Resnet-152 (P-M) پیشنهادی	79.210535	67.211067	54.348761	85.38083
Resnet-152	46.081569	52.131468	32.246062	15.14983



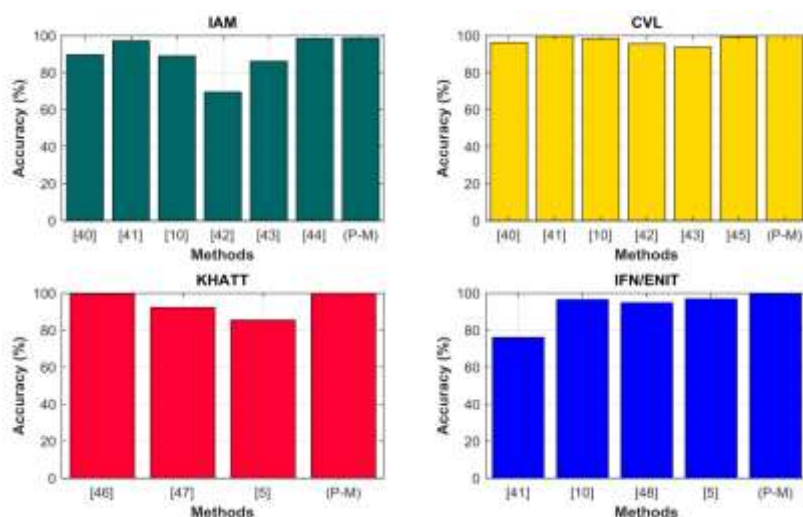
شکل (۱۰): نمودار صحت و خطا برای داده‌های اعتبارسنجی

جدول (۵): نتایج ارزیابی مدل پیشنهادی و مدل Resnet-۱۵۲ برای چهار مجموعه داده جامع.

نام شبکه	پایگاه داده				ارزیابی با TTA
	نام داده	زبان	تعداد نویسندگان	تعداد نمونه‌ها	صحت آزمون (%)
Resnet-152 همراه با بلوک پیشنهادی (P-M)	IAM	انگلیسی	150	4899	64/98
	CVL	انگلیسی	309	1854	85/99
	KHATT	عربی	828	10898	88/99
	IFN/ENIT	عربی	411	26459	76/99
Resnet-152	IAM	انگلیسی	150	4899	92/97
	CVL	انگلیسی	309	1854	34/99
	KHATT	عربی	828	10898	38/99
	IFN/ENIT	عربی	411	26459	Nov-99

جدول (۶): صحت روش پیشنهادی در مقایسه با سایر روش‌ها.

پایگاه داده		ارزیابی		
نام	زبان	تعداد نویسندگان	روش	صحت (%)
IAM	English	650	هاناد و همکاران [۴۴]	54/89
		657	خان و همکاران [۴۵]	Feb-97
		657	چاهی و همکاران [۱۰]	99/88
		657	والبرگ و همکاران [۴۶]	Apr-69
		657	هی و همکاران [۴۷]	Jan-86
		220	کوتزئر و همکاران [۴۸]	34/98
		150	مدل پیشنهادی (P-M)	64/98
CVL	English	310	هاناد و همکاران [۴۴]	Feb-96
		310	خان و همکاران [۴۵]	Jun-99
		309	چاهی و همکاران [۱۰]	38/98
		310	والبرگ و همکاران [۴۶]	Jun-95
		310	هی و همکاران [۴۷]	Jul-93
		311	بننور و همکاران [۴۹]	23/99
		309	مدل پیشنهادی (P-M)	85/99
KHAT	Arabic	1000	کریستلین و همکاران [۵۰]	Jun-99
		1000	رحمان و همکاران [۵۱]	Feb-92
		1000	هاناد و همکاران [۵]	Apr-85
		828	مدل پیشنهادی (P-M)	88/99
IFN/ENIT	Arabic	411	خان و همکاران [۴۵]	76
		411	چاهی و همکاران [۱۰]	47/96
		411	صبا و همکاران [۵۲]	72/94
		411	هاناد و همکاران [۵]	86/96



شکل (۱۱): مقایسه صحت آزمون روش‌های مختلف.

جدول (۷): صحت آزمون روش پیشنهادی در مقایسه با سایر مدل‌ها

ویژگی‌های مهندسی	یادگیری ویژگی از داده خام	مدل‌ها
43.142876	26.958671	SVM
23.918672	18.922238	BPNN
27.941905	67.211067	DCNN

نتیجه‌گیری

با توجه به پیچیدگی‌های سبک‌های نوشتاری و نیاز سازمان‌های دولتی (مانند دادگستری و دفاتر ثبت اسناد) به شناسایی دستخط نویسندگان، هدف این مطالعه ارائه یک روش جدید برای شناسایی آفلاین نویسنده با استفاده از نمونه‌های دستخط در شرایط آزمایشی مختلف است. دو ویژگی درخور توجه و مهم مطالعه حاضر استفاده از داده‌های نامتجانس و استقلال روش پیشنهادی برای هر زبان خاص است. در این مطالعه یک مجموعه داده جامع بر پایه استانداردهای ASTM طراحی شده است. یک مدل DCNN مبتنی بر شبکه از پیش آموزش‌دیده برای استخراج ویژگی‌ها به صورت سلسله‌مراتبی از دست‌نوشته‌های خام طراحی و توسعه یافته است.

مطالعه حاضر نشان داد روش پیشنهادی می‌تواند ویژگی‌ها را از روی داده‌های خام دستخط بیاموزد و به صحت قابل قبولی برای شناسایی نویسنده دست یابد. مدل پیشنهادی بر پایه شبکه از پیش آموزش‌دیده به همراه مجموعه داده طراحی شده و چهار نوع مجموعه داده جامع بررسی شد. نتایج نشان دادند مدل پیشنهادی (شبکه از پیش آموزش‌دیده همراه با بلوک پیشنهادی) از شبکه از پیش آموزش‌دیده بدون بلوک پیشنهادی در شناسایی نویسنده برای هریک از پنج مجموعه داده بیان‌شده، بهتر عمل می‌کند.

همچنین، صحت روش‌های مختلف برای چهار نوع مجموعه داده جامع با مدل پیشنهادی مقایسه شد. نتایج نشان‌دهنده صحت بالاتر مدل پیشنهادی در مقایسه با سایر روش‌ها برای همه مجموعه داده‌ها بود. علاوه بر این، مجموعه داده طراحی شده همراه با DCNN بررسی و با ویژگی‌های مهندسی و دو روش هوشمند BPNN و SVM مقایسه شد. نتایج نشان دادند روش پیشنهادی قادر به یادگیری ویژگی‌ها و به دست آوردن نتایج شناسایی قانع‌کننده‌ای است. در مقایسه با ویژگی‌های مهندسی، روش پیشنهادی صحت شناسایی را حدوداً ۱۳ درصد افزایش می‌دهد و همچنین وابستگی کمتری به دانش تخصص دارد. با توجه به نتایج ارائه شده، می‌توان بیان کرد روش پیشنهادی برای شناسایی خودکار نویسنده بسیار رضایت‌بخش و مناسب است و می‌تواند با ورود به حوزه کاربردی، دستیار خوبی برای متخصصان شناسایی دستخط باشد.

منابع

۱. جوادى مقدم، سيد محمد، غلامعلی نژاد، حسين، نوروزی، اعظم، عبدی، محمد حسن و مرتضی پور، حمید . (۱۴۰۳). تخمین بازده الکتریکی صفحه‌های فتوولتائیک با روشی مبتنی بر یادگیری عمیق با استفاده از تصویر. مجله مهندسی برق دانشگاه تبریز، ۵۴(۳)، ۲۸۱-۲۹۰. doi: ۱۰.۲۲۰۳۴/tjee.۱۰.۲۲۰۳۴.۲۹۰-۲۸۱, ۵۸۴۰۰, ۴۷۱۹
۲. حرّی، ابوالفضل . (۱۴۰۳). از هوش انسانی تا هوش مصنوعی: بررسی برگردان‌های فارسی فنون ادبی غزلواره ش. ۶۲ شکسپیر از منظر راهکارهای انسانی و ماشینی. فنون ادبی، ۱۶(۳)، ۲۹-۵۰. doi: ۱۰.۲۲۱۰۸/liar.۱۰.۲۲۱۰۸.۲۰۲۴, ۱۴۱۶۱۶, ۲۳۸۰
۳. خسروشاهی، سید نادى محامد، رضوی، سید ناصر، بابازاده سنگر، امین و مجیدزاده، کامبیز . (۱۴۰۱). شناسایی آفلاین (غیر برخط) نویسنده با استفاده از داده‌های نامتجانس دستخط بر پایه یادگیری عمیق. هوش محاسباتی در مهندسی برق، ۱۳(۴)، ۱۱۵-۱۳۴. doi: ۱۰.۲۲۱۰۸/isee.۱۰.۲۲۱۰۸.۲۰۲۱, ۱۲۷۸۱۶, ۱۴۵۹
۴. زارع چاهوکی، محمدعلی و کارگرشورکی، حمید. (۱۳۹۵). روشی مبتنی بر ماشین یادگیری سریع با هسته غیرخطی برای انتخاب نمونه‌های اولیه در یادگیری چندبرچسبه مقیاس بزرگ. مجله بینایی ماشین و پردازش تصویر، ۳(۲)، ۳۹-۵۷.
۵. زارع لاهیجان، لیدا، مشکینی، سعید و افروزیان، رضا . (۱۴۰۳). تقسیم‌بندی خودکار ضایعات سکنه مغزی ایسکمیک در نقشه‌های پرفیوژن CT با استفاده از شبکه‌های یادگیری عمیق. هوش محاسباتی در مهندسی برق، ۱۵(۳)، ۹۹-۱۱۶. doi: ۱۰.۲۲۱۰۸/isee.۱۰.۲۲۱۰۸.۲۰۲۴, ۱۴۰۹۹۶, ۱۶۸۱
۶. زرندي مریم، عزتی منصور، زرندي محمد (۱۴۰۳)، هوش مصنوعی تحولی در یادگیری و آموزش، (حانن/ تهران) فارسی.
۷. زرندي، محمد، ترجمه: تجزیه و تحلیل دیجیتال مارکتینگ (۱۴۰۰) (در تئوری و در عمل) (حانن / تهران) فارسی.
۸. شریف زاده، زهرا، میرکوشش، امیرهوشنگ و حسینی، محمد مهدی . (۱۴۰۳). بررسی آثار سیاست‌های توسعه فناوری‌های نوین و هوش مصنوعی در گسترش راهبردهای سیاسی کلان با رویکرد سیاست‌های کلی نظام. سیاست‌های راهبردی و کلان، ۱۲(۴۵)، ۲۴-۴۷. doi: ۱۰.۳۰۵۰۷/jmsp.۱۰.۳۰۵۰۷.۲۰۲۳, ۳۸۸۲۴۴, ۲۵۵۵
۹. طاهری حاجی وند، عادل، شیرینی، کیمیا و صمدی قره ورن، سینا . (۱۴۰۳). تشخیص علف های هرز در مزارع با استفاده از شبکه عصبی کانولوشنال بر اساس یادگیری عمیق. مهندسی زراعی، ۴۷(۱)، ۱۲۹-۱۴۲. doi: ۱۰.۲۲۰۵۵/agen.۱۰.۲۲۰۵۵.۲۰۲۴, ۴۵۳۲۷, ۱۶۸۸
۱۰. مسیبي اطاقسرا، امیر حسین، آراسته، عبدالله و جوادیان، نیکبخش . (۱۴۰۳). پیش‌بینی روزانه تقاضای کالای فسادپذیر با استفاده از یادگیری عمیق (مطالعه موردی: ماست سون شرکت فرآورده‌های لبنی کاله آمل). تصمیم‌گیری و تحقیق در عملیات، (۱)، - . doi: ۱۰.۲۲۱۰۵/dmor.۱۰.۲۲۱۰۵.۲۰۲۴, ۴۶۷۸۱۸, ۱۸۵۹

۱۱. مشایخی، محمد، سلیمانی، حسین و آریانپان، ایمان. (۱۴۰۳). شکل‌دهی پرتو آرایه آنتن مدوله شده زمانی با استفاده از شبکه‌های عصبی کانولوشن. علوم، فناوری و کاربردهای فضایی، ۴(۱)، ۹-۱. doi: ۲۰۲۴,۴۲۷۸۰۴,۱۱۴۷/jssta.۱۰,۲۲۰۳۴

12. A. Ahmed, H. R. Hasan, F. A. Hameed, and O. I. Al-Sanjary, "Writer identification on multi-script handwritten using optimum features," Kurdistan Journal of Applied Research, vol. ۲, no. ۳, pp. ۱۸۵-۱۷۸, ۲۰۱۷
13. A. Khan, M. A. Tahir, F. Khelifi, A. Bouridane, and R. Almotaeryi, "Robust off-line text independent writer identification using bagged discrete cosine transform features," Expert Systems with Applications, vol. ۷۱, pp. ۴۱۵-۴۰۴, ۲۰۱۷
14. Adak, B. B. Chaudhuri, and M. Blumenstein, "An empirical study on writer identification and verification from intra-variable individual handwriting," IEEE Access, vol. ۷, pp. ۲۴۷۵۸-۲۴۷۳۸, ۲۰۱۹
15. Bennour, C. Djeddi, A. Gattal, I. Siddiqi, and T. Mekhaznia, "Handwriting based writer recognition using implicit shape codebook," Forensic science international, vol. ۳۰۱, pp. ۱۰۰-۹۱, ۲۰۱۹
16. Chahi, Y. Ruichek, and R. Touahni, "Block wise local binary count for off-line text-independent writer identification," Expert Systems with Applications, vol. ۹۳, pp. ۱۴-۱, ۲۰۱۸
17. Ding, H., Hou, H., Wang, L., Cui, X., Yu, W., & Wilson, D. I. (2025). Application of Convolutional Neural Networks and Recurrent Neural Networks in Food Safety. Foods, 14(2), 247.
18. DoubleApaper. ۲۰۱۷; Available from: <http://igepa-allcart.com/myuploads/WDAcjSnxTpKfziBf\۴۱۵۷۱۸۸۳۰.pdf>.
19. E۰۷-۲۲۹۰a, A., Standard Guide for Examination of Handwritten Items, in ASTM International. ۲۰۰۷: West Conshohocken.
20. F. Kleber, S. Fiel, M. Diem, and R. Sablatnig, "Cvl-database: An off-line database for writer retrieval, writer identification and word spotting," in ۱۲ ۲۰۱۳th international conference on document analysis and recognition, ۲۰۱۳: IEEE, pp. ۵۶۴-۵۶۰.
21. F. Shahabi and M. Rahmati, "A new method for writer identification of handwritten Farsi documents," in ۱۰ ۲۰۰۹th International Conference on Document Analysis and Recognition, ۲۰۰۹: IEEE, pp. ۴۳۰-۴۲۶
22. Faber-Castell. ۲۰۱۷; Available from: http://www.faber-castell.in/۴۰۵۲۶/Products/Exports/default_news.aspx.
23. G. E. Hinton, N. Srivastava, A. Krizhevsky, I. Sutskever, and R. R. Salakhutdinov, "Improving neural networks by preventing co-adaptation of feature detectors," arXiv preprint arXiv:۱۲۰۷,۰۵۸۰, ۲۰۱۲
24. Goodfellow, Y. Bengio, A. Courville, and Y. Bengio, Deep learning (no. ۲). MIT press Cambridge, ۲۰۱۶

25. H. E. Said, T. N. Tan, and K. D. Baker, "Personal identification based on handwriting," *Pattern Recognition*, vol. ۳۳, no. ۱, pp. ۱۶۰-۱۴۹, ۲۰۰۰.
26. Jha et al., "A comprehensive study on colorectal polyp segmentation with ResUNet++, conditional random field and test-time augmentation," *IEEE journal of biomedical and health informatics*, vol. ۲۵, no. ۶, pp. ۲۰۴۰-۲۰۲۹, ۲۰۲۱.
27. Karouni, B. Daya, and S. Bahlak, "Offline signature recognition using neural networks approach," *Procedia Computer Science*, vol. ۳, pp. ۱۶۱-۱۵۵, ۲۰۱۱.
28. L. G. Hafemann, R. Sabourin, and L. S. Oliveira, "Characterizing and evaluating adversarial examples for Offline Handwritten Signature Verification," *IEEE Transactions on Information Forensics and Security*, vol. ۱۴, no. ۸, pp. ۲۱۶۶-۲۱۵۳, ۲۰۱۹.
29. M. Hagan, H. Demuth, and M. Beale, "Neural Network Design (PWS, Boston, MA)," *Google Scholar Google Scholar Digital Library Digital Library*, ۱۹۹۶.
30. M. Hawkins, "The problem of overfitting," *Journal of chemical information and computer sciences*, vol. ۴۴, no. ۱, pp. ۱۲-۱, ۲۰۰۴.
31. M. Javidi and M. Jampour, "A deep learning framework for text-independent writer identification," *Engineering Applications of Artificial Intelligence*, vol. ۹۵, p. ۱۰۳۹۱۲, ۲۰۲۰.
32. M. Pechwitz, S. S. Maddouri, V. Märgner, N. Ellouze, and H. Amiri, "IFN/ENIT-database of handwritten Arabic words," in *Proc. of CIFED*, ۲۰۰۲, vol. ۲: Citeseer, pp. ۱۳۶-۱۲۷.
33. M. S. Baghshah, S. B. Shouraki, and S. Kasaei, "A novel fuzzy classifier using fuzzy LVQ to recognize online Persian handwriting," in *۲۲۰۰th International Conference on Information & Communication Technologies*, ۲۰۰۶, vol. ۱: IEEE, pp. ۱۸۸۳-۱۸۷۸.
34. N. Pokhriyal, K. Tayal, I. Nwogu, and V. Govindaraju, "Cognitive-biometric recognition from language usage: A feasibility study," *IEEE Transactions on Information Forensics and Security*, vol. ۱۲, no. ۱, pp. ۱۴۳-۱۳۴, ۲۰۱۶.
35. N. Siddique and H. Adeli, *Computational intelligence: synergies of fuzzy logic, neural networks and evolutionary computing*. John Wiley & Sons, ۲۰۱۳.
36. P. Domingos, "Bayesian averaging of classifiers and the overfitting problem," in *ICML*, ۲۰۰۰, vol. ۷۴۷, pp. ۲۳۰-۲۲۳.
37. P. Santos, L. F. Villa, A. Reñones, A. Bustillo, and J. Maudes, "An SVM-based solution for fault detection in wind turbines," *Sensors*, vol. ۱۵, no. ۳, pp. ۵۶۴۸-۵۶۲۷, ۲۰۱۵.
38. R. Hassan and M. I. H. Bhuiyan, "Computer-aided sleep staging using complete ensemble empirical mode decomposition with adaptive noise and bootstrap aggregating," *Biomedical Signal Processing and Control*, vol. ۲۴, pp. ۱۰-۱, ۲۰۱۶.
39. R. Kumar and M. Kaur, "A character based handwritten identification using neural network and SVM," *International Journal of Scientific Research in Science, Engineering and Technology (IJSRSET)*, ۲۰۱۷.
40. Rehman, S. Naz, M. I. Razzak, and I. A. Hameed, "Automatic visual features for writer identification: a deep learning approach," *IEEE access*, vol. ۷, pp. ۱۷۱۵۷-۱۷۱۴۹, ۲۰۱۹.

41. S. A. Mahmoud, H. Luqman, B. M. Al-Helali, G. BinMakhashen, and M. T. Parvez, "Online-KHATT: An Open-Vocabulary Database for Arabic Online-Text Processing," The Open Cybernetics & Systemics Journal, vol. ۱۲, no. ۱, ۲۰۱۸
42. S. He and L. Schomaker, "Deep adaptive learning for writer identification based on single handwritten word images," Pattern Recognition, vol. ۸۸, pp. ۷۴-۶۴, ۲۰۱۹
43. S. Hochreiter and J. Schmidhuber, "Long short-term memory," Neural computation, vol. ۹, no. ۸, pp. ۱۷۸۰-۱۷۳۵, ۱۹۹۷
44. S. Ioffe and C. Szegedy, "Batch normalization: Accelerating deep network training by reducing internal covariate shift," in International conference on machine learning, ۲۰۱۵: PMLR, pp. ۴۵۶-۴۴۸
45. S. M. Awaida and S. A. Mahmoud, "Writer identification of arabic text using statistical and structural features," Cybernetics and Systems, vol. ۴۴, no. ۱, pp. ۷۶-۵۷, ۲۰۱۳
46. S. N. Srihari, S.-H. Cha, H. Arora, and S. Lee, "Individuality of handwriting," Journal of forensic science, vol. ۴۷, no. ۴, pp. ۱۷-۱, ۲۰۰۲
47. S. Sheykhivand et al., "Developing an efficient deep neural network for automatic detection of COVID- ۱۹ using chest X-ray images," Alexandria Engineering Journal, vol. ۶۰, no. ۳, pp. ۲۹۰۳-۲۸۸۵, ۲۰۲۱
48. S. Sheykhivand, Z. Mousavi, T. Y. Rezaii, and A. Farzamnia, "Recognizing Emotions Evoked by Music Using CNN-LSTM Networks on EEG Signals," IEEE Access, vol. ۸, pp. ۱۳۹۳۴۵-۱۳۹۳۳۲, ۲۰۲۰
49. S. Y. Manchala, J. Kinthali, K. Kotha, J. Kumar, and J. Jayalaxmi, "Handwritten text recognition using deep learning with Tensorflow," International Journal of Engineering and Technical Research, vol. ۹, no. ۵, ۲۰۲۰
50. Schneiderpen. ۲۰۱۷; Available from: https://schneiderpen.com/en_us/office/tops-۵۰۵-black-f-۴۰۰۴۶۷۵۰۰۴۵۲۹.pdf/
51. Shanmugam, D. Blalock, G. Balakrishnan, and J. Gutttag, "When and why test-time augmentation works," arXiv preprint arXiv:۲۰۱۱.۱۱۱۵۶, ۲۰۲۰
52. T. Kutzner, C. F. Pazmiño-Zapatier, M. Gebhard, I. Bönninger, W.-D. Plath, and C. M. Travieso, "Writer identification using handwritten cursive texts and single character words," Electronics, vol. ۸, no. ۴, p. ۳۹۱, ۲۰۱۹
53. T. Saba, "Fuzzy ARTMAP Approach for Arabic Writer Identification using Novel Features Fusion," J. Comput. Sci., vol. ۱۴, no. ۲, pp. ۲۲۰-۲۱۰, ۲۰۱۸
54. U.-V. Marti and H. Bunke, "The IAM-database: an English sentence database for offline handwriting recognition," International Journal on Document Analysis and Recognition, vol. ۵, no. ۱, pp. ۴۶-۳۹, ۲۰۰۲
55. V. Christlein and A. Maier, "Encoding CNN activations for writer recognition," in ۲۰۱۸ ۱۳th IAPR International Workshop on Document Analysis Systems (DAS), ۲۰۱۸: IEEE, pp. ۱۷۴-۱۶۹

56. Wahlberg, "Gaussian process classification as metric learning for forensic writer identification," in ۱۳ ۲۰۱۸th IAPR International Workshop on Document Analysis Systems (DAS), ۲۰۱۸: IEEE, pp. ۱۸۰-۱۷۵
57. X. Wu, Y. Tang, and W. Bu, "Offline text-independent writer identification based on scale invariant feature transform," IEEE Transactions on Information Forensics and Security, vol. ۹, no. ۳, pp. ۵۳۶-۵۲۶, ۲۰۱۴
58. X.-Y. Zhang, G.-S. Xie, C.-L. Liu, and Y. Bengio, "End-to-end online writer identification with recurrent neural network," IEEE Transactions on Human-Machine Systems, vol. ۴۷, no. ۲, pp. ۲۹۲-۲۸۵, ۲۰۱۶
59. Y. Hannad, I. Siddiqi, and M. E. Y. El Kettani, "Writer identification using texture descriptors of handwritten fragments," Expert Systems with Applications, vol. ۴۷, pp. -۱۴ ۲۲, ۲۰۱۶
60. Y. Hannad, I. Siddiqi, C. Djeddi, and M. E.-Y. El-Kettani, "Improving Arabic writer identification using score-level fusion of textural descriptors," IET Biometrics, vol. ۸, no. ۳, pp. ۲۲۹-۲۲۱, ۲۰۱۹
61. Y. Xu, Y. Chen, Y. Cao, and Y. Zhao, "A Deep Learning Method for Chinese writer Identification with Feature Fusion," in Journal of Physics: Conference Series, ۲۰۲۱, vol. ۱۸۸۳, no. ۱, p. ۰۱۲۱۴۲: IOP Publishing.
62. Z. Mousavi, M. M. Ettefagh, M. H. Sadeghi, and S. N. Razavi, "Developing deep neural network for damage detection of beam-like structures using dynamic response based on FE model and real healthy state," Applied Acoustics, vol. ۱۶۸, p. ۱۰۷۴۰۲, ۲۰۲۰
63. Z. Mousavi, S. Varahram, M. M. Ettefagh, M. H. Sadeghi, and S. N. Razavi, "Deep neural networks-based damage detection using vibration signals of finite element model and real intact state: An evaluation via a lab-scale offshore jacket structure," Structural Health Monitoring, p. ۱۴۷۵۹۲۱۷۲۰۹۳۲۶۱۴, ۲۰۲۰
64. Z. Wang, A. Maier, and V. Christlein, "Towards End-to-End Deep Learning-based Writer Identification," INFORMATIK ۲۰۲۰, ۲۰۲۱
65. Z. Wu, C. Shen, and A. Van Den Hengel, "Wider or deeper: Revisiting the resnet model for visual recognition," Pattern Recognition, vol. ۹۰, pp. ۱۳۳-۱۱۹, ۲۰۱۹
66. Z.-K. Gao, Y.-L. Li, Y.-X. Yang, and C. Ma, "A recurrence network-based convolutional neural network for fatigue driving detection from EEG," Chaos: An Interdisciplinary Journal of Nonlinear Science, vol. ۲۹, no. ۱۱, p. ۱۱۳۱۲۶, ۲۰۱۹
67. Zarandi, M., & Jivan, S. E. (2015). Effective Factors on Managers Decision-Making Style in Organizations (Viewpoints of Senior Executives using AHP Technique).
68. Zarandi, Mohammad and Zarandi, Maryam, (1402), The role of competitive advantage with strategic management in the education industry, Third National Conference on Organization and Business Improvement and Restructuring, Tehran, <https://civilica.com/doc/2017181>.