



Computational Intelligence in Electrical Engineering
Vol. 13, No. 4, 2023
Research Paper

Offline Identification of the Author using Heterogeneous Data based on Deep Learning

Seyed Nadi Mahamed Khosroshahi, Seyed Naser Razavi^{*}, Amin Babazadeh Sangar, Kambiz Majidzadeh

Ph.D. Student, Department of IT and Computer Engineering, Urmia Branch,
Islamic Azad University, Urmia, Iran.

Abstract:

Handwriting recognition has always been a challenge; therefore, it has attracted the attention of many researchers. The present study presents an offline system for the automatic detection of human handwriting under different experimental conditions. This system includes input data, image processing unit, and output unit. In this study, a right-to-left dataset is designed based on the standards of the American Society for Experiments and Materials (ASTM). An improved deep convolution neural network (DCNN) model based on a pre-trained network is designed to extract features hierarchically from raw handwritten data. A significant advantage in this study is the use of heterogeneous data. Another significant aspect of the present study is that the proposed DCNN model is independent of any particular language and can be used for different languages. The results show that the proposed DCNN model has a very good performance for identifying the author based on heterogeneous data.

Keywords: Offline Identification of the Author, Heterogeneous Data, Feature Learning, Deep Neural Network.



This is an open access article under the CC BY-NC-ND/4.0/ License (<https://creativecommons.org/licenses/by-nc-nd/4.0/>).



<https://doi.org/10.22108/isee.2021.127816.1459>

مقاله پژوهشی

شناسایی آفلاین (غیر برخط) نویسنده با استفاده از داده‌های نامتجانس دستخط بر پایه

یادگیری عمیق

سید نادى محامد خسروشاهی^۱، سید ناصر رضوی^{۲*}، امین بابازاده سنگر^۳ و کامبیز مجیدزاده^۴

۱- دانشجوی دکتری، دانشکده مهندسی کامپیوتر و فناوری اطلاعات - واحد ارومیه - دانشگاه آزاد اسلامی - ارومیه - ایران

n.khosroshahi@iaurmia.ac.ir

۲- استادیار، دانشکده مهندسی کامپیوتر و فناوری اطلاعات - واحد ارومیه - دانشگاه آزاد اسلامی - ارومیه - ایران

n.razavi@tabrizu.ac.ir

۳- استادیار، دانشکده مهندسی کامپیوتر و فناوری اطلاعات - واحد ارومیه - دانشگاه آزاد اسلامی - ارومیه - ایران

Bsamin2@liveutm.onmicrosoft.com

۴- استادیار، دانشکده مهندسی کامپیوتر و فناوری اطلاعات - واحد ارومیه - دانشگاه آزاد اسلامی - ارومیه - ایران

K.Majidzadeh@iaurmia.ac.ir

چکیده: تشخیص دستخط همواره مسئله چالش‌برانگیزی بوده است؛ از این رو، توجه محققان زیادی را به خود جلب کرده است. مطالعه حاضر یک سیستم آفلاین (غیر برخط) تشخیص خودکار دست‌نوشته‌های انسان را در شرایط آزمایشی مختلف ارائه می‌دهد. این سیستم شامل داده‌های ورودی، واحد پردازش تصویر و واحد خروجی است. در این مطالعه، یک مجموعه داده راست به چپ بر پایه استانداردهای آمریکایی (ASTM) طراحی شده است. یک مدل شبکه عصبی کانولوشن عمیق (DCNN) بهبودیافته بر پایه شبکه از پیش آموزش‌دیده، برای استخراج ویژگی‌ها به صورت سلسله‌مراتبی از داده‌های خام دستخط طراحی شده است. یک مزیت درخور توجه در این مطالعه استفاده از داده‌های نامتجانس است. یکی دیگر از جنبه‌های شایان توجه مطالعه حاضر این است که مدل پیشنهادی DCNN مستقل از هر زبان خاصی است و می‌تواند برای زبان‌های مختلف استفاده شود. نتایج نشان می‌دهند مدل پیشنهادی DCNN، عملکرد بسیار خوبی برای شناسایی نویسنده بر پایه داده‌های نامتجانس دستخط دارد.

واژه‌های کلیدی: شناسایی آفلاین نویسنده، داده نامتجانس، یادگیری ویژگی، شبکه عصبی عمیق.

۱- مقدمه

زبان‌های مختلف است [۱-۲]. دست‌نوشته‌ها ویژگی‌های قابل اندازه‌گیری دارند که می‌توانند نویسندگان را توصیف کنند [۳]. شناسایی دست‌نوشته‌ها یکی از فعال‌ترین حوزه‌های تحقیقاتی در زمینه بینایی 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. در تمامی این کاربردها، همواره مسئله دقت و سرعت سیستم اهمیت بسیار زیادی داشته است. جرم جعل در اسناد

در طول زمان‌های گذشته، از دستخط برای ایجاد ارتباط میان مردم استفاده شده که شامل نمادهایی برای نمایش

^۱ تاریخ ارسال مقاله: ۱۳۹۹/۱۲/۲۴

تاریخ پذیرش مقاله: ۱۴۰۰/۰۸/۱۷

نام نویسنده مسئول: سید ناصر رضوی

نشانی نویسنده مسئول: ایران - ارومیه - دانشگاه آزاد اسلامی - دانشکده مهندسی کامپیوتر و فناوری اطلاعات

دست‌نویس به سرعت در حال رشد است؛ بنابراین، شناسایی اسناد دست‌نویس جعلی به عنوان یک مشکل بسیار حیاتی برای سازمان‌های دولتی (مانند دادگستری و دفاتر ثبت اسناد) در نظر گرفته می‌شود. به طور 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 پایه استخراج ویژگی‌ها از سند اصلی و مقایسه آنها با سند جعلی استوار است. ویژگی‌های دستخط به سبک‌های نوشتاری زبان بستگی دارد. این به عنوان یک نشانه اساسی در تجزیه و تحلیل اسناد دست‌نویس در نظر گرفته می‌شود.

بیشتر محققان زبان انگلیسی را برای تجزیه و تحلیل اسناد دست‌نویس انتخاب کرده‌اند [۴-۵]. 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 به طور 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 درباره زبان‌های راست به چپ این است که آنها از روش‌های پیچیده‌ای برای اتصال حروف استفاده می‌کنند؛ از این رو، آنها به عنوان سبک‌های نوشتاری ناشناخته و کم مطالعه باقی مانده‌اند.

به طور کلی، مطالعات تحقیقاتی در زمینه اسناد دست‌نویس، به سه گروه عمده طبقه‌بندی می‌شوند: تشخیص دستخط^۱، شناسایی نویسنده^۲ و تأیید نویسنده^۳. 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 پایه رویکردهای سنتی یادگیری ماشین و برخی دیگر بر پایه رویکردهای جدید یادگیری ماشین، مانند یادگیری عمیق [۶-۷] استوارند. بسیاری از محققان روش‌های طبقه‌بندی را برای مسئله شناسایی نویسنده پیشنهاد داده‌اند [۸-۹]. 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 به صورت سلسله‌مراتبی ویژگی‌های مطلوب را از داده‌های خام استخراج کنند. به طور کلی، تشخیص دست‌نوشته‌ها به صورت مرسوم به دو دسته آنلاین (برخط) و آفلاین (غیر برخط) تقسیم می‌شوند. در تشخیص آنلاین ترتیب زمانی از مختصات دریافت می‌شود که بیان‌کننده حرکات نوک قلم شخص است؛ در حالی که در روش‌های آفلاین تنها تصویر متن در دسترس است [۱۰-۱۱].

آوایدا و همکاران [۱۲] 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 [۱۳] روشی را برای شناسایی آفلاین اسناد فارسی دست‌نویس با استفاده از فیلتر چندکاناله گابور^۴ 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. باغشاه و همکاران [۱۴] روش جدیدی را برای شناسایی آفلاین (برخط) دست‌نوشته‌های فارسی ارائه دادند. در روش بیان‌شده، ویژگی‌ها با توجه به معیارهای فاصله اقلیدسی

مختلف دستخط آموزش دید. در این سیستم تصاویر دست‌نویس به متون دیجیتالی تبدیل شدند. نتایج نشان دادند این سیستم برای متن‌هایی که نویز کمتری دارند بهترین دقت را ارائه می‌دهد. همچنین دقت سیستم بیان‌شده کاملاً به مجموعه داده بستگی دارد و در صورت افزایش داده‌ها می‌توان با این سیستم به دقت بیشتری دست یافت. آداک و همکاران [۸] روش شناسایی و تأیید نویسنده را از دست‌نوشته‌های آفلاین بنگالی^۹ بررسی کردند. در این روش، برخی از ویژگی‌های مهندسی از این دست‌نوشته‌ها، استخراج و با استفاده از مدل ماشین بردار پشتیبان ارزیابی شدند. ویژگی‌های خودکار با استفاده از مدل شبکه عصبی کانولوشنال عمیق^{۱۰} (DCNN) نیز از این دست‌نوشته‌ها استخراج شدند. در این مطالعه دو پایگاه داده از دو مجموعه مختلف با ۱۰۰ نویسنده برای آزمایش طراحی شدند. پس از آزمایش مشاهده شد مدل شبکه عصبی کانولوشنال عمیق در مقایسه با سایر مدل‌ها نتایج بهتری ارائه می‌دهد. ژانگ و همکاران [۱۸] یک شبکه عصبی بازگشتی^{۱۱} (RNN) را برای شناسایی آنالاین نویسنده بررسی کردند. داده‌های دست‌نویس هر نویسنده با مجموعه‌ای از RHSها^{۱۲} نشان داده شدند. از یک مدل شبکه عصبی با حافظه کوتاه‌مدت دو جهت برای رمزگذاری هر RHS (در یک بردار با طول ثابت) برای طبقه‌بندی استفاده شد. آزمایش‌های مربوطه روی مجموعه داده‌های انگلیسی (۱۳۳ نویسنده) و چینی (۱۸۶ نویسنده) انجام و مزایای روش آنها در مقایسه با سایر روش‌های پیشرفته تأیید شدند. کاربون و همکاران [۱۹] یک سیستم تشخیص آنالاین دستخط را ارائه دادند که می‌تواند ۱۰۲ زبان مختلف را با استفاده از شبکه عصبی عمیق^{۱۳} پشتیبانی کند. این سیستم روش‌های شناسایی متوالی را با رمزگذاری جدید ورودی با استفاده از منحنی‌های Bézier ترکیب کرده است. نتایج نشان دادند سیستم بیان‌شده در مقایسه با روش‌های دیگر نتایج بهتری ارائه می‌دهد. جاویدی و همکاران [۲۰] روشی را برای شناسایی نویسنده (مستقل از متن) بر پایه یادگیری عمیق ارائه دادند. در این مطالعه یک نسخه توسعه‌یافته از ResNet با ترکیب شبکه‌های residual عمیق و یک توصیف‌کننده دستخط سستی برای تجزیه و تحلیل دستخط استفاده شده است.

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 دست یافت. احمد و همکاران [۴] یک طبقه‌بند مبتنی بر فاصله را با مجموعه‌ای از ویژگی‌های استخراج‌شده با الگوریتم مور^{۱۴} به‌منظور شناسایی آفلاین نویسنده بررسی کردند. این روش به چهار مجموعه داده جامع اعمال شد و به میزان دقت چشمگیری در شناسایی اسناد دست‌نویس دست یافت. وو و همکاران [۱۵] روش شناسایی آفلاین متون را بر پایه طبقه‌بند مبتنی بر فاصله پیشنهاد کردند. در این مطالعه برای استخراج ویژگی‌های ساختاری از شش مجموعه داده مختلف، از فیلتر ایزوتروپیک استفاده شد. سه مورد از این مجموعه داده‌ها مربوط به زبان انگلیسی، یکی مربوط به زبان چینی و دو مورد دیگر مربوط به ترکیبی از این زبان‌ها بود. نتایج تجربی نشان دادند روش بیان‌شده در شناسایی متون بهتر از سایر روش‌های مقایسه‌ای عمل می‌کند. کومار و همکاران [۱۶] یک روش شناسایی آفلاین دستخط را بررسی کردند. در این مطالعه، پنج نوع ویژگی استخراج‌شده از متن‌های دست‌نویس ارزیابی شدند. برای کاهش ابعاد ویژگی‌ها، از روش تجزیه و تحلیل خطی فشر^{۱۵} و آنالیز اجزای اصلی^{۱۶} استفاده شد. با استفاده از ماشین بردار پشتیبان^{۱۷} (SVM) و شبکه عصبی کارایی روش بیان‌شده تأیید شد. روش‌های شناسایی نویسنده مبتنی بر یادگیری ماشین سنتی به‌طور گسترده استفاده شده‌اند؛ اما مشکلات متعددی دارند. برخی از این مشکلات عبارت‌اند از وابسته‌بودن به دانش تخصصی، حساسیت به تغییرات شرایط محیطی و محدودیت استخراج ویژگی‌های جدید. براساس مشکلات بیان‌شده، لازم است روش‌های خودکار شناسایی نویسنده براساس رویکردهای جدید یادگیری ماشین، مانند رویکردهای یادگیری عمیق بررسی شوند که ویژگی‌های مطلوب مربوط به هر مسئله را از طریق داده‌های خام به‌صورت سلسله‌مراتبی می‌توانند بیاموزند [۶-۷].

منچالا و همکاران [۱۷] سیستمی را برای شناسایی دستخط با استفاده از یادگیری عمیق ارائه دادند. این سیستم برای یافتن شباهت‌ها و همچنین تفاوت‌ها در میان نمونه‌های

توصیف‌کننده ضخامت دستخط را به‌عنوان یک ویژگی اولیه و ضروری دستخط تجزیه و تحلیل می‌کند. این روش می‌تواند هویت نویسنده مستقل از متن را ارائه دهد که برای یادگیری مدل خود نیازی به محتوای دست‌نویس یکسان ندارد. رویکرد پیشنهادی روی مجموعه داده‌های عمومی و مشهور ارزیابی شد. نتایج نشان دادند شبکه ترکیبی پیشنهادی نسبت به روش‌های مقایسه‌ای بهتر عمل می‌کند و می‌تواند برای برنامه‌های کاربردی در دنیای واقعی استفاده شود. یانگ و همکاران [۲۱] یک روش یادگیری عمیق را برای شناسایی نویسنده بر پایه زبان چینی با استفاده از ترکیب ویژگی‌ها ارائه دادند. در این مطالعه از ترکیب ویژگی‌های عمیق و ویژگی‌های دستی برای به دست آوردن ویژگی‌های دستخط از تصاویر دست‌نویس استفاده شد. نتایج نشان دادند این روش عملکرد بهتری در شناسایی حروف چینی نسبت به سایر روش‌های مقایسه‌ای دارد. وانگ و همکاران [۲۲] یک روش شناسایی خودکار نویسنده را بر پایه یادگیری عمیق ارائه دادند. در این مطالعه، ترکیبی از شبکه‌های u-net و resnet به‌عنوان مدل پیشنهادی در نظر گرفته شد. روش پیشنهادی آنها روی مجموعه داده ICDAR17 ارزیابی شد و نتایج بهتری نسبت به مدل‌های مقایسه‌ای ارائه داد.

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 جامعی 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. 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 فواصل زمانی مختلف در شرایط آزمایشی مختلف براساس استانداردهای آمریکایی (ASTM) جمع‌آوری شده است [۲۳]. 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^{۱۴} 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.

ادامه مقاله به‌صورت زیر تدوین شده است؛ در بخش ۲، شبکه‌های عصبی کانولوشنال و بازگشتی بررسی می‌شوند. در بخش ۳، روش پیشنهادی برای شناسایی نویسنده ارائه می‌شود. در بخش ۴ نتایج شبیه‌سازی بررسی می‌شوند و درنهایت، بخش ۵ مربوط به نتیجه‌گیری است.

۲- مواد و روش‌ها

در این بخش ابتدا شبکه‌های عصبی کانولوشنال^{۱۵} (CNN) و پس از آن، شبکه‌های حافظه طولانی کوتاه مدت^{۱۶} (LSTM) بررسی می‌شوند که زیرمجموعه‌ای از شبکه‌های عصبی بازگشتی‌اند.

۲-۱- 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 [۲۴]. این روش، بسیار کارآمد بوده و یکی از رایج‌ترین روش‌ها در کاربردهای مختلف بینایی ماشین است. همانند شبکه‌های عصبی مصنوعی^{۱۷}، تصمیم خروجی نهایی شبکه عصبی کانولوشنال براساس وزن و بایاس لایه‌های قبلی در ساختار شبکه است. در این شبکه، دو مرحله برای آموزش وجود دارد؛ مرحله انتشار پیش‌رو^{۱۸} و مرحله پس‌انتشار^{۱۹} (BP) [۲۵]. BP روشی برای محاسبه گرادیان تابع اتلاف نسبت به وزن‌ها است. BP سیگنال‌های خطا را در شبکه حین آموزش پس می‌زند و باعث به‌روزرسانی وزن‌ها می‌شود. در مرحله اول، داده‌های ورودی به شبکه اعمال می‌شوند و این عمل چیزی به جز ضرب نقطه‌ای بین ورودی و پارامترهای هر نورون و اعمال عملیات کانولوشن در هر لایه نیست و درنهایت، خروجی شبکه محاسبه می‌شود. به‌منظور تنظیم پارامترهای شبکه یا به عبارت دیگر آموزش شبکه، از نتیجه خروجی برای محاسبه میزان خطای شبکه استفاده می‌شود. برای این کار، خروجی شبکه با استفاده از یک تابع خطا^{۲۰} با پاسخ صحیح، مقایسه و به این ترتیب، میزان خطا محاسبه می‌شود. در مرحله بعد، براساس میزان خطای محاسبه‌شده، مرحله پس‌انتشار آغاز می‌شود. در این مرحله، گرادیانت هر پارامتر با توجه به قاعده زنجیره‌ای محاسبه می‌شود و تمامی پارامترها، با توجه به تأثیرشان بر خطای ایجادشده در شبکه، به‌روزرسانی می‌شوند. بعد از به‌روزرسانی پارامترها، مرحله بعدی انتشار پیش‌رو آغاز خواهد شد. بعد از تکرار تعداد مناسبی از این مراحل، آموزش شبکه به پایان می‌رسد. در این شبکه، خروجی هر لایه همان ویژگی‌ها هستند که بعد

کمتری نسبت به داده اصلی دارند.

به‌طور کلی، یک شبکه کانولوشنال از سه لایه اصلی تشکیل می‌شود که عبارت‌اند از لایه کانولوشنال، لایه ادغام^{۲۱} و لایه تمام متصل^{۲۲} (FC) [۲۴]. برای جلوگیری از فرایند بیش‌برازش^{۲۳} و بهبود عملکرد شبکه از لایه‌های حذف تصادفی^{۲۴} و نرمال‌ساز دسته‌ای^{۲۵} نیز استفاده می‌شود. همچنین در شبکه‌های عصبی نیاز است پس از هر لایه از تابع فعال‌سازی استفاده شود که در ادامه، این لایه‌ها و توابع به‌طور خلاصه معرفی می‌شوند.

لایه کانولوشنال: شامل فیلترهایی (کرنل‌هایی) است که روی داده‌های ورودی می‌لغزند. یک کرنل، یک ماتریس است که با داده ورودی کانوالو می‌شود. این لایه عمل کانولوشن را روی داده‌های ورودی با استفاده از کرنل انجام می‌دهد. خروجی کانولوشن را نگاشت ویژگی می‌نامند. عملکرد کانولوشن به شرح زیر است:

$$y_k = \sum_{n=0}^{N-1} x_n h_{k-n} \quad (1)$$

که x سیگنال، h فیلتر، N تعداد عناصر در x و y بردار خروجی است.

لایه ادغام: این لایه که به کاهش نمونه^{۲۶} نیز معروف است، ابعاد نورون‌های خروجی از لایه کانولوشنال را کاهش می‌دهد و باعث کاهش محاسبات و همچنین جلوگیری از پدیده بیش‌برازش می‌شود. در این پژوهش از لایه ادغام بیشینه^{۲۷} استفاده شده که فقط مقادیر بیشینه در هر نگاشت ویژگی را انتخاب کرده است و باعث کاهش تعداد نورون‌های خروجی می‌شود.

لایه FC: دارای اتصال کامل به تمامی فعال‌سازی‌ها در لایه قبلی است.

لایه حذف تصادفی: از این لایه به‌منظور جلوگیری از پدیده بیش‌برازش استفاده می‌شود [۲۵]. نحوه کار آن به این صورت است که در هر مرحله از آموزش، هر نورون با احتمالی از شبکه بیرون انداخته شده است؛ به طوری که درنهایت، یک شبکه کاهش داده شده باقی می‌ماند.

لایه نرمال‌سازی دسته‌ای: این لایه به‌منظور نرمال‌سازی داده‌ها در داخل شبکه انجام می‌شود [۲۶]. زمانی که

۲-۲- شبکه‌های عصبی بازگشتی

شبکه‌های عصبی بازگشتی شاخه‌ای مهم از شبکه‌های عصبی عمیق‌اند که به منظور تحلیل سیستم‌های پیچیده استفاده می‌شوند. این شبکه‌ها می‌توانند با کاهش ابعاد داده ورودی X_t ، بار محاسباتی را کاهش دهند و همچنین باعث بهبود عملکرد آموزش شوند. علاوه بر این، این شبکه‌ها امکان تلفیق اطلاعات بین ورودی‌های مختلف را به منظور دستیابی به ویژگی‌هایی فراهم می‌کنند که نمی‌توان با استفاده از روش‌های سنتی استخراج کرد [۲۸-۳۰]. شبکه حافظه طولانی کوتاه مدت از جمله شبکه‌های عصبی بازگشتی‌اند که به منظور رفع ضعف‌های شبکه‌های بازگشتی از جمله حل مشکل پراکندگی گرادیان یا مشکلات انفجاری گرادیان به کار برده می‌شوند [۲۸-۳۰]. 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. یک سلول حافظه از چهار عنصر اصلی تشکیل شده است: یک دروازه ورودی یا دروازه به‌روزرسانی Γ_u ، یک نورون با اتصال خودبازگشتی، یک دروازه فراموشی Γ_f و یک دروازه خروجی Γ_o . فعال‌سازی واحد LSTM به صورت رابطه زیر تعریف می‌شود [۲۸-۳۰]:

$$h_t = \Gamma_o \cdot \tanh(C_t) \quad (5)$$

که در آن، Γ_o دروازه خروجی و کنترل‌کننده میزان محتوایی است که از طریق حافظه ارائه می‌شود. دروازه خروجی با رابطه زیر محاسبه می‌شود [۲۸-۳۰]:

$$\Gamma_o = \sigma(W_o \cdot [h_{t-1}, X_t] + b_o) \quad (6)$$

که در آن، σ تابع فعال‌سازی سافت‌مکس است و W_o و b_o به ترتیب ماتریس وزن و بردار بایاس اولیه‌اند. سلول حافظه C_t نیز با فراموشی نسبی حافظه فعلی و اضافه کردن محتوای حافظه جدید به صورت $\tilde{C}_t$ از رابطه (۸) به‌روزرسانی می‌شود که در آن، محتوای حافظه جدید از رابطه (۷) به دست می‌آید [۲۸-۳۰]:

محاسبات مختلف روی داده ورودی اعمال شود، توزیع داده‌ها تغییر خواهد کرد. این لایه با هدف کاهش تغییر کوواریانس داخلی، سرعت آموزش شبکه را افزایش می‌دهد و باعث تسریع در همگرایی می‌شود. تبدیل لایه نرمال‌سازی دسته‌ای به شرح زیر است:

$$\begin{aligned} \mu_B &= \frac{1}{n} \sum_{i=1}^n y_i^{(l-1)} \\ \sigma_B^2 &= \frac{1}{n} \sum_{i=1}^n (y_i^{(l-1)} - \mu_B)^2 \\ \hat{y}^{(l-1)} &= \frac{y^{(l-1)} - \mu_B}{\sqrt{(\sigma_B^2 + \epsilon)}} \\ z^{(l)} &= \gamma^{(l)} \hat{y}^{(l-1)} + \beta^{(l)} \end{aligned} \quad (2)$$

که μ_B و σ_B^2 به ترتیب میانگین و واریانس دسته هستند. ϵ یک ثابت کوچک برای ثبات عددی، l شماره لایه، $y^{(l-1)}$ بردار ورودی به لایه نرمال‌ساز و $z^{(l)}$ بردار خروجی نرمال مربوط به یک نورون است و $\gamma^{(l)}$ و $\beta^{(l)}$ به ترتیب پارامترهای مربوط به مقیاس و تغییر نرخ یادگیری‌اند. تابع فعال‌سازی: پس از هر لایه کانولوشن، یک تابع فعال‌سازی اعمال می‌شود. تابع فعال‌سازی یک عملگر است که خروجی را به مجموعه‌ای از ورودی‌ها نگاشت می‌کند و برای غیرخطی کردن ساختار شبکه استفاده می‌شود [۲۷]. تابع Relu یکی از پرکاربردترین توابع فعال‌سازی است و این ویژگی را دارد که غیرخطی بودن را به ساختار شبکه اعمال کند؛ بنابراین، در برابر تغییرات جزئی در ورودی مقاوم است. رابطه ۳ تابع Relu را نشان می‌دهد.

$$f(x) = \begin{cases} x & x > 0 \\ 0 & x \leq 0 \end{cases} \quad (3)$$

تابع سافت‌مکس^{۲۸}: این تابع توزیع احتمالی کلاس‌های خروجی را محاسبه می‌کند که رابطه آن به فرم زیر است:

$$p_i = \frac{e^{x_j}}{\sum_{k=1}^k e^{x_k}} \quad \text{for } j=1, \dots, k \quad (4)$$

که در آن، x ورودی شبکه است و مقادیر خروجی p بین صفر و یک هستند که مجموع آنها برابر با یک است.

توسعه داده شده است. همچنین، در این مطالعه علاوه بر پایگاه داده طراحی شده از پایگاه داده‌های CVL، JAM، KHATT و IFN/ENIT نیز استفاده شده است تا روش پیشنهادی بهتر ارزیابی شود. پایگاه داده IAM [۳۱] شامل ۴۸۹۹ نمونه از ۱۵۰ نویسنده و پایگاه داده CVL [۳۲] شامل ۱۸۵۴ نمونه از ۳۰۹ نویسنده به زبان انگلیسی است. پایگاه داده KHATT [۳۳] شامل ۱۰۸۹۸ نمونه از ۸۲۸ نویسنده و پایگاه داده IFN/ENIT [۳۴] شامل ۲۶۴۵۹ نمونه از ۴۱۱ نویسنده به زبان عربی است. این پایگاه داده‌ها در مطالعات اخیر بسیار استفاده شده‌اند و در زمینه شناسایی نویسنده از پایگاه داده‌های قابل اعتماد و کاربرد هستند.

در پایگاه داده پیشنهادی، بر پایه استانداردهای ASTM، نمونه‌های دستخط از ۶۲ شرکت‌کننده در بازه‌های زمانی مختلف در شرایط محیطی مختلف جمع‌آوری شدند (شکل ۳ را ببینید). از ۶۲ شرکت‌کننده در این آزمایش، ۳۴ نفر مرد و ۲۸ نفر زن، با میانگین سنی ۲۲ تا ۵۴ سال‌اند. همچنین ۶ نفر از این شرکت‌کنندگان چپ‌دست و ۵۶ نفر راست‌دست بودند. درنهایت، براساس استانداردهای از پیش تعریف‌شده، ویژگی‌های بافتی و ساختاری دستخط به دست آمدند. ویژگی‌های بافتی مربوط به یک سند دست‌نویس به عواملی مانند کاغذ، جوهر، ابزارهای نوشتن و غیره بستگی دارد؛ در صورتی که ویژگی‌های ساختاری به سبک‌های نوشتاری، نحوه پیوستن حروف و غیره اشاره دارد. گفتنی است زبان فارسی به‌عنوان زبان مقصد انتخاب شده که مجموعه داده مربوط به آن به دست آمده است. انتخاب این زبان به این علت است که محققان به نمونه‌های این زبان دسترسی آسان و کافی داشتند.

هر حرف از زبان فارسی چهار فرم نوشتاری مختلف دارد؛ فرم جداگانه^{۲۹} (S)، فرم شروع کلمه^{۳۰} (BOW)، فرم میانی کلمه^{۳۱} (MOW) و فرم انتهایی کلمه^{۳۲} (EOW). انتخاب فرم به موقعیت آن در یک کلمه بستگی دارد. از شرکت‌کنندگان چهار جمله مربوط به چهار فرم بیان‌شده 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

$$\tilde{C}_t = \tanh(W_C \cdot [h_{t-1}, X_t] + b_C) \quad (7)$$

$$C_t = \Gamma_f \cdot C_{t-1} + \Gamma_u \cdot \hat{C}_t \quad (8)$$

میزان حافظه فعلی که باید فراموش شود، توسط دروازه فراموشی Γ_f کنترل می‌شود و مقدار حافظه جدید که باید به سلول حافظه اضافه شود، توسط دروازه به‌روزرسانی (دروازه ورودی) Γ_u انجام می‌گیرد. این عملیات در رابطه‌های (۹) و (۱۰) نشان داده شده است [۲۸-۳۰]:

$$\Gamma_f = \sigma(W_f \cdot [h_{t-1}, X_t] + b_f) \quad (9)$$

$$\Gamma_u = \sigma(W_u \cdot [h_{t-1}, X_t] + b_u) \quad (10)$$

شکل ۱ ساختار یک شبکه عصبی بازگشتی LSTM را نشان می‌دهد. در این شبکه که یک ورودی یا همان X_t دارد، دو خروجی تولید می‌شود: یک خروجی C_t و خروجی دیگر h_t ؛ به دو بخش تقسیم می‌شود؛ بخشی به گام زمانی بعد، منتقل و بخشی نیز در صورت نیاز به تولید خروجی در گام زمانی فعلی استفاده می‌شود. دروازه فراموشی Γ_f وظیفه کنترل جریان اطلاعات از گام زمانی قبلی را دارد. این دروازه مشخص می‌کند 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 به‌روزرسانی Γ_u وظیفه کنترل جریان اطلاعات جدید را بر عهده دارد. این دروازه مشخص می‌کند آیا در گام زمانی فعلی باید از اطلاعات جدید استفاده شود یا خیر و اگر بلی به چه میزان. دروازه خروجی Γ_o نیز مشخص می‌کند چه میزان از اطلاعات گام زمانی قبل با اطلاعات گام زمانی فعلی به گام زمانی بعد منتقل شود.

۳- روش پیشنهادی

در این بخش، روش پیشنهادی مقاله ارائه می‌شود. شکل ۲ بلوک دیاگرام الگوریتم پیشنهادی را نشان می‌دهد.

۳-۱- جمع‌آوری داده‌ها

در روش پیشنهادی یک پایگاه داده جامع طراحی و

نوشته شدند. از دو نوع کاغذ استاندارد متفاوت «PaperOne» و «Double-A» استفاده شد که مشخصات آنها در جدول ۱ نشان داده شده است [۳۵].

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 [۳۶-۳۷]. 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 نشان‌دادن میزان فشار قلم در نظر گرفته شده است. نمونه‌های جمع‌آوری‌شده توسط RICOH Aficio MP 6001 با رزولوشن ۳۰۰ dpi در حالت رنگی اسکن شدند. ترتیب نمونه‌های جمع‌آوری‌شده، اطلاعات و جزئیات موردنیاز به‌عنوان کتاب‌کد در مجموعه داده‌ها ذخیره شده است. مجموعه داده جمع‌آوری‌شده 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. از این پس، این مجموعه داده خاص «DANA_HW» نام‌گذاری می‌شود و در دسترس همه محققان در پلتفرم GitHub قرار می‌گیرد.

۳-۲- پیش‌پردازش داده‌ها

در این مطالعه، به‌منظور کاهش زمان اجرا و حجم محاسبات، پس از جداسازی ۴۲۰۳ جمله (مربوط به ۶۲ شرکت‌کننده با اندازه ۲۳۶ پیکسل و عرض متغیر)، اندازه جملات ابتدا به ۱۱۲ پیکسل و عرض متغیر تبدیل می‌شود. سپس با استفاده از روش تقسیم‌بندی^۳، ۴۲۰۳ جمله به ۸۶۳۰۴ نمونه با اندازه ۱۱۲×۱۱۲ پیکسل تقسیم می‌شوند؛ پس از آن، نمونه‌ها نرمال می‌شوند. عملیات تقسیم‌بندی یکی از جملات در شکل ۵ نشان داده شده است. مطابق

شکل ۵، هر جمله به نمونه‌های ۱۱۲×۱۱۲ به‌صورت اتوماتیک تقسیم شده است. در برخی از نمونه‌های تقسیم‌شده، تصویر حاوی شکاف‌هایی بین کلمات است که با نماد (a) نمایش داده شده‌اند. در برخی از نمونه‌های دیگر تصویر یا حاوی هیچ داده ارزشمندی برای پردازش نیست که با نماد (b) یا حاوی داده‌های غیرقابل توجه یا کمی برای پردازش‌اند که با نماد (c) مشخص شده‌اند؛ حتی گاهی شرکت‌کنندگان در نمونه‌گیری با هدف تصحیح، بخشی از دست‌نوشته را مخدوش کرده‌اند که در تقسیم‌بندی با نماد (d) نمایش داده شده‌اند. چنین بخش‌هایی به کاهش دقت در شبکه پیشنهادی منجر می‌شود؛ با وجود این، آنها از مجموعه داده حذف نمی‌شوند تا نمونه‌ها اصلی و دست‌نخورده باقی بمانند. چنین بخش‌هایی «بخش‌های فریبده» نام‌گذاری شده‌اند. برخی از این بخش‌های فریبده در شکل ۶ نشان داده شده‌اند.

۳-۳- شبکه عمیق پیشنهادی

شبکه عمیق پیشنهادی در این مطالعه از ترکیب یک شبکه از پیش آموزش‌دیده کانولوشنال Resnet-152 [۳۸] با شبکه LSTM ایجاد شده است. با ترکیب شبکه Resnet-152 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 [۲۹-۳۰]. شبکه پیشنهادی بر پایه شبکه از پیش آموزش‌دیده Resnet-152 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. 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^۴ به

درصد افزایش می‌یابند.

به‌عنوان آخرین مرحله برای افزایش و به حداکثر رساندن دقت پیش‌بینی، از تکنیک TTA^{۴۱} 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 آموزش دیده، چندین بار تصاویر تقویت شده 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 [۴۱-۴۲]. 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.

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).

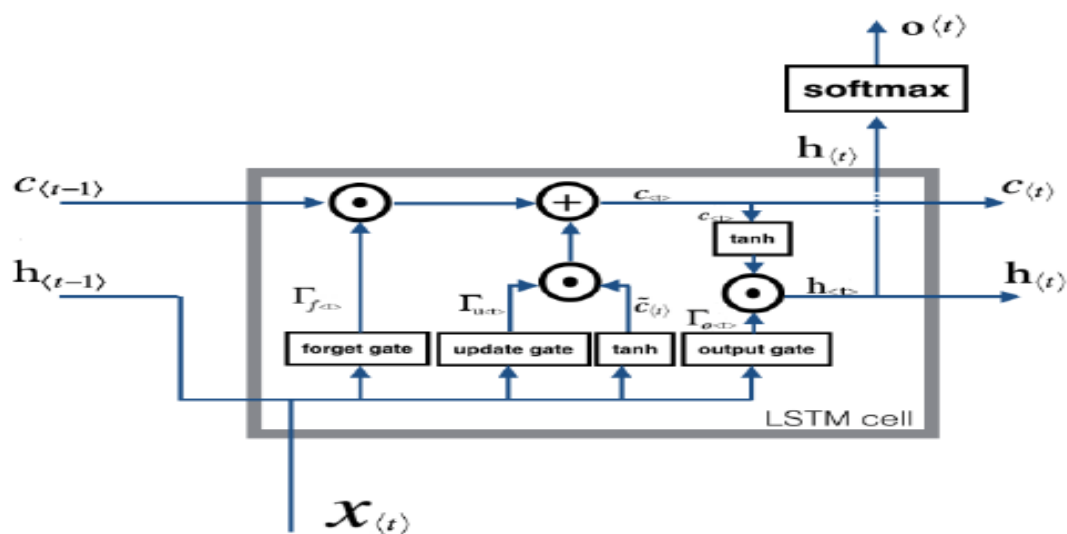
از کل نمونه‌های جمع‌آوری شده (۸۶۳۰۴ نمونه)، ۵۱۷۸۲ نمونه برای داده‌های آموزش (۶۰ درصد)، ۲۵۸۹۰ نمونه (۳۰ درصد) برای داده‌های اعتبارسنجی و ۸۶۳۲ نمونه برای داده‌های آزمون (۱۰ درصد) استفاده می‌شود. علاوه بر این، تمام نمونه‌های اختصاص داده شده به مجموعه‌های آموزش و ارزیابی به‌طور تصادفی انتخاب می‌شوند.

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 آموزش دیده یک بردار ویژگی با اندازه 512×256 است. در اولین لایه بلوک پیشنهادی، یعنی FC، تابع خطی روی وزن‌های قابل یادگیری ویژگی‌های به‌دست آمده (w) اعمال می‌شود. مقادیر پیش‌بینی شده بایاس

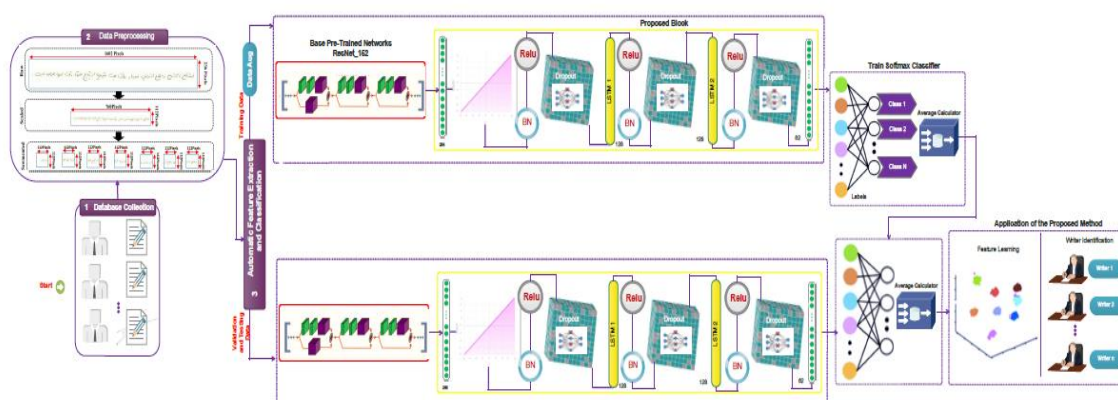
در نظر گرفته می‌شود تا ابعاد بردار ویژگی را به 1×256 تغییر دهد. همان‌طور که ملاحظه می‌شود، کاهش ابعاد در لایه‌های پنهان از 112×112 (اندازه ورودی) به 128 (بردار ویژگی انتخاب شده) ادامه یافته است که در نهایت، بردار ویژگی انتخاب شده به یک لایه FC با تابع غیرخطی سافت‌مکس متصل می‌شود (شکل ۲ را ببینید).

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 آنروپی^{۳۵} و بهینه‌ساز SGD^{۳۶} با نرخ یادگیری 0.01 انتخاب شده‌اند. روش مرسوم پس‌انتشار خطا با اندازه دسته ۱۰۰ برای آموزش شبکه استفاده شده است. ابر پارامترهای بهینه انتخاب شده برای مدل پیشنهادی در جدول ۳ نشان داده شده است.

براساس مطالعات صورت گرفته در ادبیات پیشین، از تکنیک‌های افزایش داده برای بهبود صحت، جلوگیری از پدیده بیش‌برازش و بهبود روند آموزش در شبکه‌های عصبی استفاده می‌شود [۳۹-۴۰]؛ 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^{۳۷}، تغییر رنگ^{۳۸} و چرخش تصادفی^{۳۹} به‌عنوان تکنیک‌های افزایش داده انتخاب شده‌اند. نمونه‌هایی از داده‌های افزوده شده در شکل ۷ نشان داده شده‌اند. پس از استفاده از تکنیک‌های افزایش داده، مجموعه داده‌های آموزش ۵۰



شکل (۱): ساختار یک شبکه عصبی بازگشتی LSTM [۲۸].



شکل (۲): بلوک دیاگرام الگوریتم پیشنهادی.

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

خواص	واحد	PaperOne			DoubleA		
		هدف	تولرانس	روش	هدف	تولرانس	روش
وزن پایه	g/m^2	۸۰	$\pm 4\%$	ISO ۵۳۶	۸۰	$\pm 3\%$	ISO ۵۳۶
ضخامت	μm	۱۱۰	± 3	ISO ۵۳۴	۱۰۶/۵	$\pm 3/5$	ISO ۵۳۴
زبری	ml/min	۱۴۰	± 40	ISO ۸۷۹۱-۲	۱۵۰	± 50	ISO ۸۷۹۱
روشنایی ISO	%	۹۹	± 2	ISO ۲۴۷۰	۱۰۲/۵	$\pm 1/5$	ISO ۲۴۷۰
تیرگی ISO	%	۹۵	± 2	ISO ۲۴۷۱	۹۴	-	ISO ۲۴۷۱
سفیدی CIE	#	۱۶۷	± 2	ISO ۱۱۴۷۵	۱۶۰	± 2	ISO ۱۱۴۷۵

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

(الف) مشخصات خودکار Schneider، (ب) مشخصات خودکار Faber-Castell

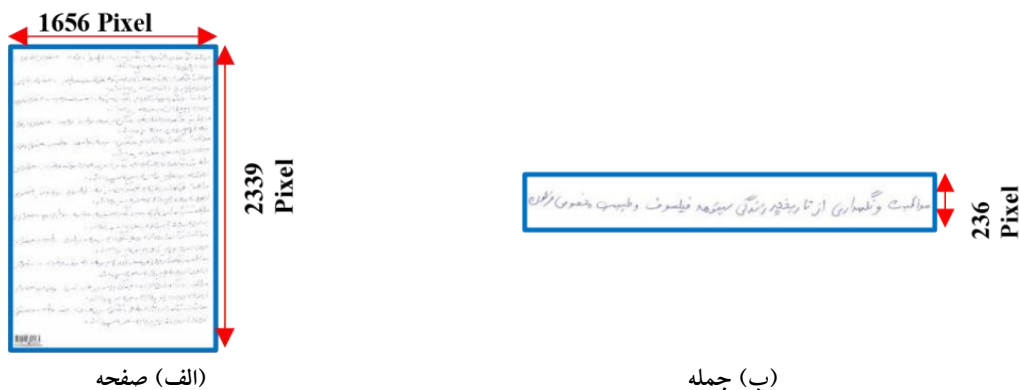
نوع	GTIN	مدل	رنگ	عرض خط	شماره مقاله
Ballpoint	400467 5004 567	Tops 505	Blue	Fine 0.4 mm	150503
Ballpoint	400467 5004 529	Tops 505	Black	Fine 0.4 mm	150501

(ب) مشخصات خودکار Faber-Castell.

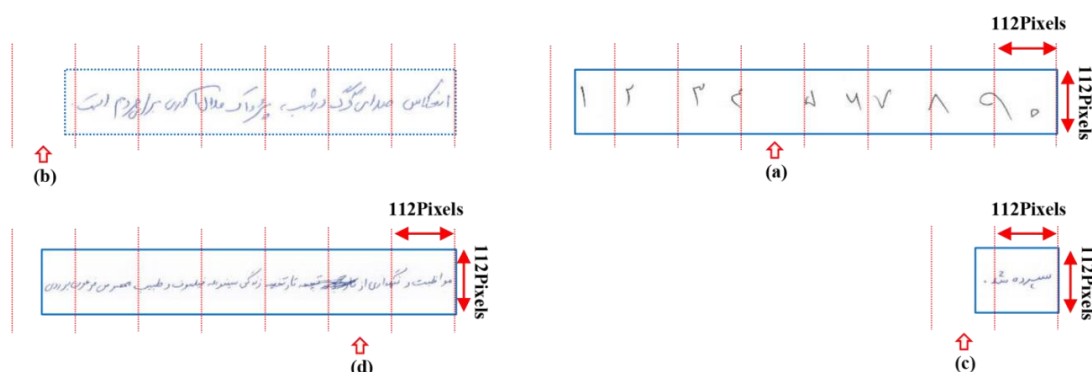
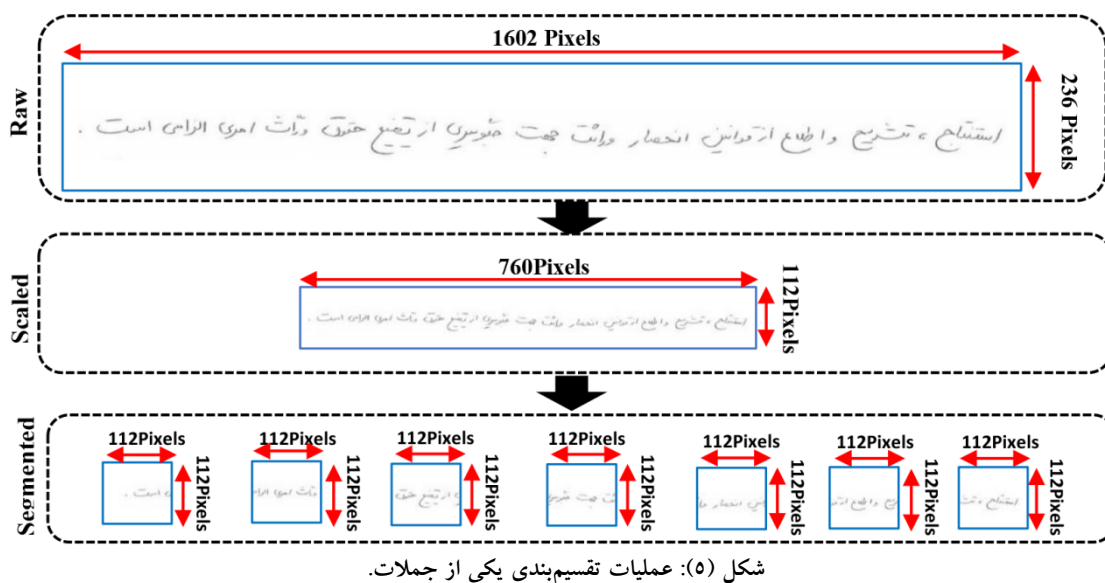
نوع	GTIN	مدل	رنگ	عرض خط	شماره مقاله
Ballpoint	8901180407516	TRI-Flow	Blue	Fine 0.7 mm	34 07 98
Ballpoint	8901180407998	TRI-Flow	Black	Fine 0.7 mm	34 07 50

[illegible]

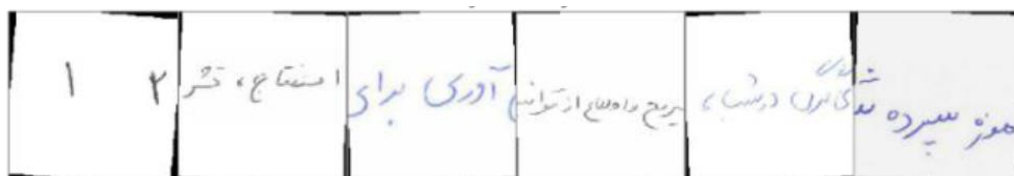
شکل (۳): نمونه‌هایی از فرم نوشتاری شرکت کنندگان

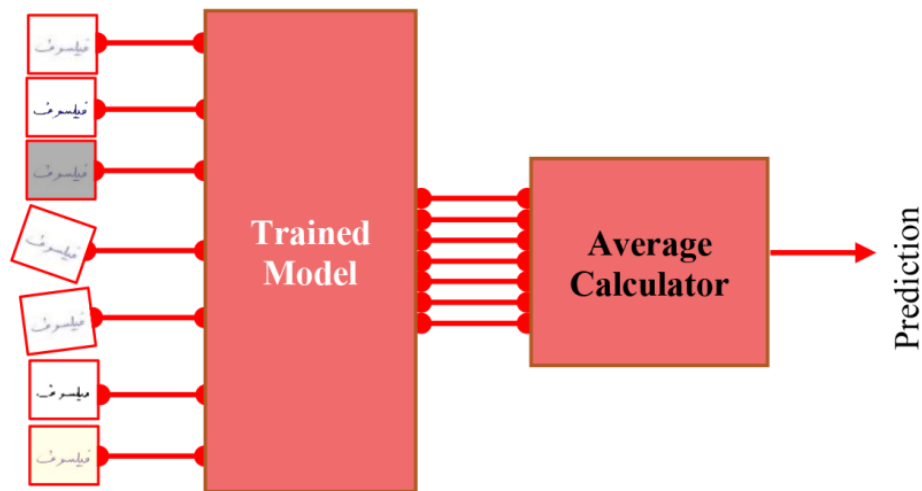


شکل (۴): اندازه نمونه‌ها؛ (الف) صفحه، (ب) جمله.

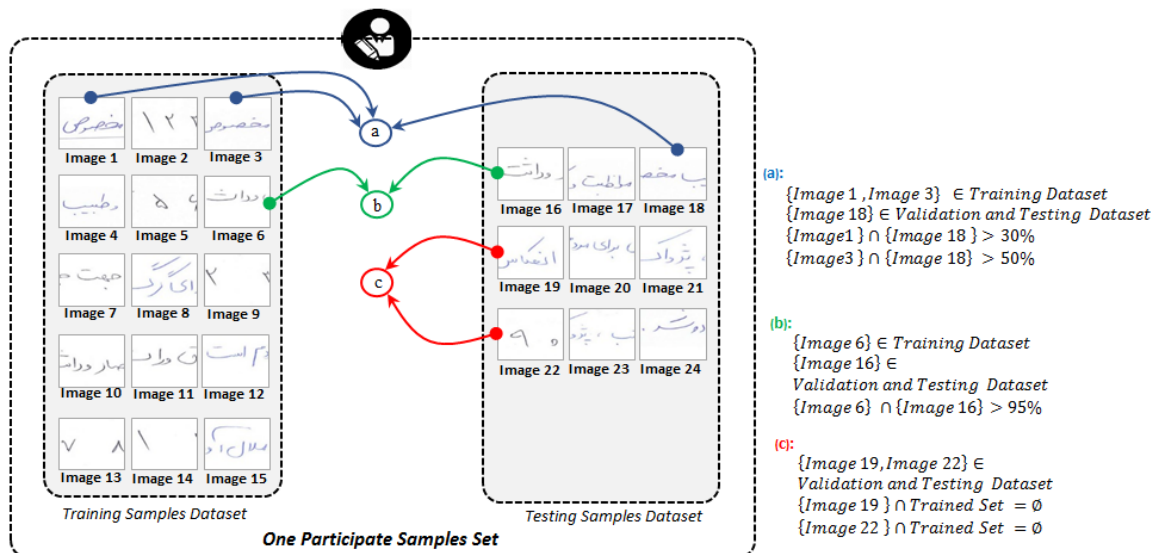


جدول (۳): ابر پارامترهای بهینه استفاده‌شده در مدل پیشنهادی.		
پارامتر	فضای جستجو	مقدار بهینه
بهینه‌ساز	RMSProp, Adam, Adamax, SGD, Adadelta	SGD
تابع خطا	MSE, Cross-entropy	Cross-entropy
نرخ لایه حذف تصادفی	۰/۵, ۰/۴, ۰/۳, ۰/۲, ۰	۰/۲
اندازه دسته	۱۰۰, ۶۴, ۳۲, ۱۶, ۱۰, ۸, ۴	۱۰۰
نرخ یادگیری	۰/۰۰۱, ۰/۰۰۱, ۰/۰۱	۰/۰۱
تابع فعال‌ساز در لایه FC اول بلوک پیشنهادی	Leaky-Relu, Sigmoid, Relu, Linear	Linear
تابع فعال‌ساز در لایه‌های مخفی بلوک پیشنهادی	Leaky-Relu, Sigmoid, Relu, Linear	Relu
تابع فعال‌ساز در لایه آخر بلوک پیشنهادی	Softmax, Sigmoid	Softmax





شکل (۸): ساختار TTA.



شکل (۹): نمونه‌های دستخط نامتجانس.

که به صورت زیر بیان می‌شود [۴۳]:

$$\text{Accuracy} = \frac{T_P + T_N}{T_P + F_P + T_N + F_N} \quad (11)$$

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-152 همراه با بلوک پیشنهادی) و شبکه از پیش

۴- نتایج و بحث

روش پیشنهادی^{۴۱} (P-M) 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 i7-6700K، پردازنده گرافیکی GeForce GTX TITAN X 12 GB، رم ۶۴ گیگابایت DDR IV و هارددیسک ۱ ترابایت SSD. به منظور ارزیابی عملکرد روش پیشنهادی، از رابطه مربوط به صحت استفاده می‌شود

همان‌طور 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^{۴۲} (BPNN) و ماشین بردار پشتیبان (SVM) به عنوان مدل‌های مقایسه‌ای انتخاب شده‌اند [۵۵-۵۳]. تابع پایه شعاعی گوسین^{۴۳} به عنوان تابع کرنل ماشین بردار پشتیبان بوده و از روش جستجوی شبکه^{۴۴} 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^{۴۵}، مختصات مرکزی^{۴۶}، گریزازمرکز^{۴۷}، کشیدگی^{۴۸} و چولگی^{۴۹} [۵۶]. 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 در بخش ۳-۳ 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

آموزش دیده Resnet-152 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-152 با تکنیک TTA، ۹۶/۵۱ درصد است؛ در حالی که صحت ارزیابی Resnet-152 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-152 بیشتر است.

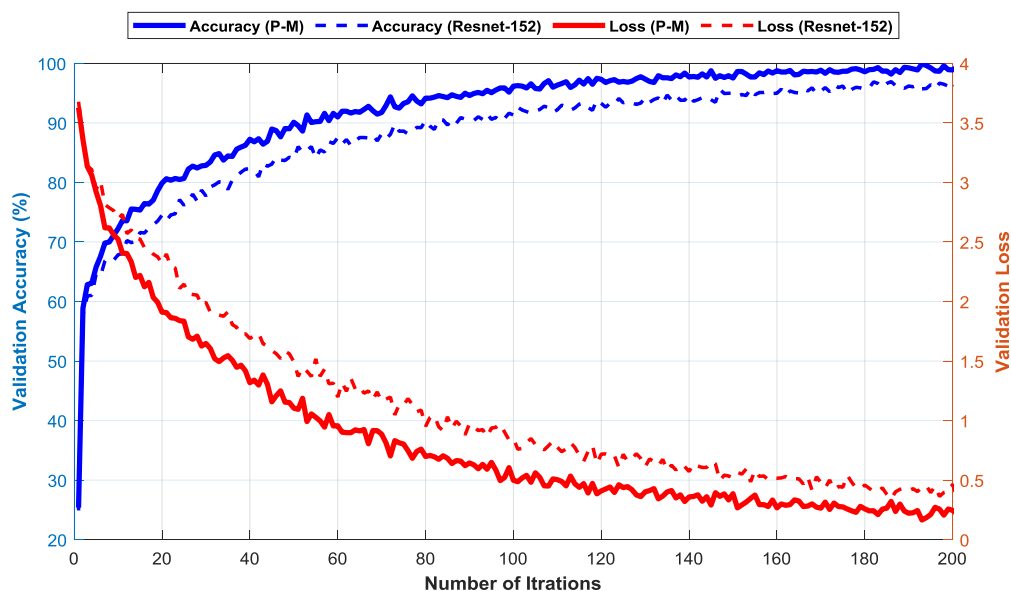
همچنین نمودار صحت و خطا برای داده‌های اعتبارسنجی مدل پیشنهادی (P-M) و مدل Resnet-152 با تکنیک TTA در شکل ۱۰ ارائه شده است. مطابق این شکل، خطای شبکه مدل پیشنهادی و مدل Resnet-152 با افزایش تکرار الگوریتم کاهش می‌یابد. همچنین مشاهده می‌شود مدل پیشنهادی و مدل Resnet-152 پس از ۲۰۰ تکرار تقریباً به صحت ۹۹ درصد و ۹۶ درصد می‌رسند. برای ارزیابی مدل پیشنهادی، از چهار مجموعه داده بیان‌شده در بخش ۳-۱ استفاده شده است. نتایج شناسایی مدل پیشنهادی و مدل Resnet-152 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 Resnet-152، از مدل Resnet-152 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 بیان‌شده، 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 بیان‌شده، به دلیل در دسترس بودن آنها است.

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 بررسی‌شده، 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 Resnet-152 همراه با بلوک پیشنهادی (مدل پیشنهادی) در مقایسه با Resnet-152 بدون بلوک پیشنهادی.

نام شبکه		ارزیابی بدون TTA		ارزیابی با TTA		زمان ارزیابی با TTA (ms)
		صحت آزمون		صحت آزمون		زمان آموزش
Resnet-152 همراه با بلوک پیشنهادی (P-M)		۹۵/۷۸		۹۹/۶۶		۱/۸۵
Resnet-152		۹۳/۴۵		۹۶/۵۱		۱/۱۵



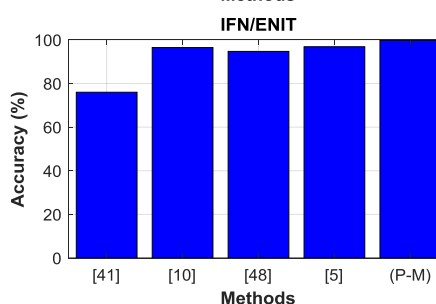
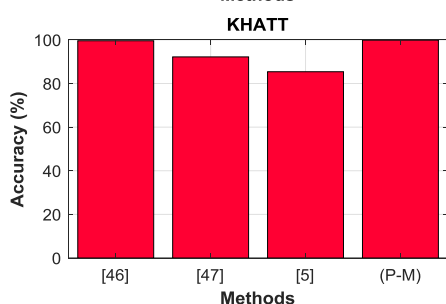
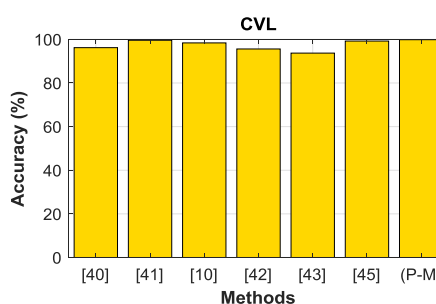
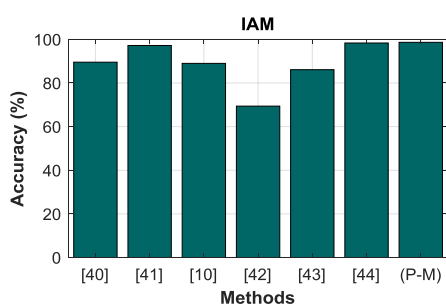
شکل (۱۰): نمودار صحت و خطا برای داده‌های اعتبارسنجی

جدول (۵): نتایج ارزیابی مدل پیشنهادی و مدل Resnet-152 برای چهار مجموعه داده جامع.

نام شبکه	پایگاه داده				ارزیابی با TTA
	نام داده	زبان	تعداد نویسندگان	تعداد نمونه‌ها	
Resnet-152 همراه با بلوک پیشنهادی (P-M)	IAM	انگلیسی	۱۵۰	۴۸۹۹	۹۸/۶۴
	CVL	انگلیسی	۳۰۹	۱۸۵۴	۹۹/۸۵
	KHATT	عربی	۸۲۸	۱۰۸۹۸	۹۹/۸۸
	IFN/ENIT	عربی	۴۱۱	۲۶۴۵۹	۹۹/۷۶
Resnet-152	IAM	انگلیسی	۱۵۰	۴۸۹۹	۹۷/۹۲
	CVL	انگلیسی	۳۰۹	۱۸۵۴	۹۹/۳۴
	KHATT	عربی	۸۲۸	۱۰۸۹۸	۹۹/۳۸
	IFN/ENIT	عربی	۴۱۱	۲۶۴۵۹	۹۹/۱۱

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

ارزیابی		پایگاه داده		
صحت (%)	روش	تعداد نویسندگان	زبان	نام
۸۹/۵۴	هاناد و همکاران [۴۴]	۶۵۰	English	IAM
۹۷/۲	خان و همکاران [۴۵]	۶۵۷		
۸۸/۹۹	چاهی و همکاران [۱۰]	۶۵۷		
۶۹/۴	والبرگ و همکاران [۴۶]	۶۵۷		
۸۶/۱	هی و همکاران [۴۷]	۶۵۷		
۹۸/۳۴	کوتزور و همکاران [۴۸]	۲۲۰		
۹۸/۶۴	مدل پیشنهادی (P-M)	۱۵۰		
۹۶/۲	هاناد و همکاران [۴۴]	۳۱۰	English	CVL
۹۹/۶	خان و همکاران [۴۵]	۳۱۰		
۹۸/۳۸	چاهی و همکاران [۱۰]	۳۰۹		
۹۵/۶	والبرگ و همکاران [۴۶]	۳۱۰		
۹۳/۷	هی و همکاران [۴۷]	۳۱۰		
۹۹/۲۳	بننور و همکاران [۴۹]	۳۱۱		
۹۹/۸۵	مدل پیشنهادی (P-M)	۳۰۹		
۹۹/۶	کریستلین و همکاران [۵۰]	۱۰۰۰	Arabic	KHATT
۹۲/۲	رحمان و همکاران [۵۱]	۱۰۰۰		
۸۵/۴	هاناد و همکاران [۵]	۱۰۰۰		
۹۹/۸۸	مدل پیشنهادی (P-M)	۸۲۸		
۷۶	خان و همکاران [۴۵]	۴۱۱	Arabic	IFN/ENIT
۹۶/۴۷	چاهی و همکاران [۱۰]	۴۱۱		
۹۴/۷۲	صبا و همکاران [۵۲]	۴۱۱		
۹۶/۸۶	هاناد و همکاران [۵]	۴۱۱		
۹۹/۷۶	مدل پیشنهادی (P-M)	۴۱۱		



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
۸۶/۴۲	۸۷/۲۶	SVM
۸۴/۲۳	۸۶/۱۸	BPNN
۸۵/۲۷	۹۹/۶۶	DCNN

۵- نتیجه‌گیری

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.

مراجع

- [1] S. N. Srihari, S.-H. Cha, H. Arora, and S. Lee, "Individuality of handwriting," *Journal of forensic science*, vol. 47, no. 4, pp. 1-17, 2002.
- [2] 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. 12, no. 1, pp. 134-143, 2016.
- [3] H. E. Said, T. N. Tan, and K. D. Baker, "Personal identification based on handwriting," *Pattern Recognition*, vol. 33, no. 1, pp. 149-160, 2000.
- [4] A. 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. 2, no. 3, pp. 178-185, 2017.
- [5] 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. 8, no. 3, pp. 221-229, 2019.
- [6] 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. 1475921720932614, 2020.
- [7] 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. 168, p. 107402, 2020.
- [8] C. Adak, B. B. Chaudhuri, and M. Blumenstein, "An empirical study on writer identification and verification from intra-variable individual handwriting," *IEEE Access*, vol. 7, pp. 24738-24758, 2019.

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 از پیش 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

- End-to-End Deep Learning-based Writer Identification," *INFORMATIK 2020*, 2021.
- [23] E2290-07a, A., Standard Guide for Examination of Handwritten Items, in *ASTM International*. 2007: West Conshohocken.
- [24] I. Goodfellow, Y. Bengio, A. Courville, and Y. Bengio, *Deep learning* (no. 2). MIT press Cambridge, 2016.
- [25] 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:1207.0580*, 2012.
- [26] S. Ioffe and C. Szegedy, "Batch normalization: Accelerating deep network training by reducing internal covariate shift," in *International conference on machine learning*, 2015: PMLR, pp. 448-456.
- [27] N. Siddique and H. Adeli, *Computational intelligence: synergies of fuzzy logic, neural networks and evolutionary computing*. John Wiley & Sons, 2013.
- [28] S. Hochreiter and J. Schmidhuber, "Long short-term memory," *Neural computation*, vol. 9, no. 8, pp. 1735-1780, 1997.
- [29] S. Sheykhiwand, Z. Mousavi, T. Y. Rezaii, and A. Farzamnia, "Recognizing Emotions Evoked by Music Using CNN-LSTM Networks on EEG Signals," *IEEE Access*, vol. 8, pp. 139332-139345, 2020.
- [30] 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. 29, no. 11, p. 113126, 2019.
- [31] 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. 5, no. 1, pp. 39-46, 2002.
- [32] F. Kleber, S. Fiel, M. Diem, and R. Sablatnig, "Cvl-database: An off-line database for writer retrieval, writer identification and word spotting," in *2013 12th international conference on document analysis and recognition*, 2013: IEEE, pp. 560-564.
- [33] 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. 12, no. 1, 2018.
- [34] M. Pechwitz, S. S. Maddouri, V. Märgner, N. Ellouze, and H. Amiri, "IFN/ENIT-database of handwritten Arabic words," in *Proc. of CIFED*, 2002, vol. 2: Citeseer, pp. 127-136.
- [35] DoubleApaper. 2017; Available from: <http://igepa-allcart.com/myuploads/WDAcjSnxTpKf2iBf1415718830.pdf>.
- [36] Schneiderpen. 2017; Available from: https://schneiderpen.com/en_us/office/tops-505-black-f-4004675004529.pdf/.
- [37] Faber-Castell. 2017; Available from:
- [9] 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. 14, no. 8, pp. 2153-2166, 2019.
- [10] A. Chahi, Y. Ruichek, and R. Touahni, "Block wise local binary count for off-line text-independent writer identification," *Expert Systems with Applications*, vol. 93, pp. 1-14, 2018.
- [11] 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. 47, no. 2, pp. 285-292, 2016.
- [12] S. M. Awaida and S. A. Mahmoud, "Writer identification of arabic text using statistical and structural features," *Cybernetics and Systems*, vol. 44, no. 1, pp. 57-76, 2013.
- [13] F. Shahabi and M. Rahmati, "A new method for writer identification of handwritten Farsi documents," in *2009 10th International Conference on Document Analysis and Recognition*, 2009: IEEE, pp. 426-430.
- [14] M. S. Baghshah, S. B. Shouraki, and S. Kasaei, "A novel fuzzy classifier using fuzzy LVQ to recognize online Persian handwriting," in *2006 2nd International Conference on Information & Communication Technologies*, 2006, vol. 1: IEEE, pp. 1878-1883.
- [15] 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. 9, no. 3, pp. 526-536, 2014.
- [16] 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)*, 2017.
- [17] 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. 9, no. 5, 2020.
- [18] 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. 47, no. 2, pp. 285-292, 2016.
- [19] V. Carbune et al., "Fast multi-language LSTM-based online handwriting recognition," *International Journal on Document Analysis and Recognition (IJDAR)*, pp. 1-14, 2020.
- [20] M. Javidi and M. Jampour, "A deep learning framework for text-independent writer identification," *Engineering Applications of Artificial Intelligence*, vol. 95, p. 103912, 2020.
- [21] 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*, 2021, vol. 1883, no. 1, p. 012142: IOP Publishing.
- [22] Z. Wang, A. Maier, and V. Christlein, "Towards

- handwritten word images," *Pattern Recognition*, vol. 88, pp. 64-74, 2019.
- [48] 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. 8, no. 4, p. 391, 2019.
- [49] A. Bennour, C. Djeddi, A. Gattal, I. Siddiqi, and T. Mekhaznia, "Handwriting based writer recognition using implicit shape codebook," *Forensic science international*, vol. 301, pp. 91-100, 2019.
- [50] V. Christlein and A. Maier, "Encoding CNN activations for writer recognition," in *2018 13th IAPR International Workshop on Document Analysis Systems (DAS)*, 2018: IEEE, pp. 169-174.
- [51] A. Rehman, S. Naz, M. I. Razzak, and I. A. Hameed, "Automatic visual features for writer identification: a deep learning approach," *IEEE access*, vol. 7, pp. 17149-17157, 2019.
- [52] T. Saba, "Fuzzy ARTMAP Approach for Arabic Writer Identification using Novel Features Fusion," *J. Comput. Sci.*, vol. 14, no. 2, pp. 210-220, 2018.
- [53] P. Santos, L. F. Villa, A. Reñones, A. Bustillo, and J. Maudes, "An SVM-based solution for fault detection in wind turbines," *Sensors*, vol. 15, no. 3, pp. 5627-5648, 2015.
- [54] M. Hagan, H. Demuth, and M. Beale, "Neural Network Design (PWS, Boston, MA)," *Google Scholar Google Scholar Digital Library Digital Library*, 1996.
- [55] S. Sheykhivand *et al.*, "Developing an efficient deep neural network for automatic detection of COVID-19 using chest X-ray images," *Alexandria Engineering Journal*, vol. 60, no. 3, pp. 2885-2903, 2021.
- [56] A. Karouni, B. Daya, and S. Bahlak, "Offline signature recognition using neural networks approach," *Procedia Computer Science*, vol. 3, pp. 155-161, 2011.
- http://www.faber-castell.in/40526/Products/Exports/default_news.aspx.
- [38] Z. Wu, C. Shen, and A. Van Den Hengel, "Wider or deeper: Revisiting the resnet model for visual recognition," *Pattern Recognition*, vol. 90, pp. 119-133, 2019.
- [39] P. Domingos, "Bayesian averaging of classifiers and the overfitting problem," in *ICML*, 2000, vol. 747, pp. 223-230.
- [40] D. M. Hawkins, "The problem of overfitting," *Journal of chemical information and computer sciences*, vol. 44, no. 1, pp. 1-12, 2004.
- [41] D. Shanmugam, D. Blalock, G. Balakrishnan, and J. Guttag, "When and why test-time augmentation works," *arXiv preprint arXiv:2011.11156*, 2020.
- [42] D. 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. 25, no. 6, pp. 2029-2040, 2021.
- [43] A. 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. 24, pp. 1-10, 2016.
- [44] Y. Hannad, I. Siddiqi, and M. E. Y. El Kettani, "Writer identification using texture descriptors of handwritten fragments," *Expert Systems with Applications*, vol. 47, pp. 14-22, 2016.
- [45] F. A. Khan, M. A. Tahir, F. Khelifi, A. Bouridane, and R. Almotary, "Robust off-line text independent writer identification using bagged discrete cosine transform features," *Expert Systems with Applications*, vol. 71, pp. 404-415, 2017.
- [46] F. Wahlberg, "Gaussian process classification as metric learning for forensic writer identification," in *2018 13th IAPR International Workshop on Document Analysis Systems (DAS)*, 2018: IEEE, pp. 175-180.
- [47] S. He and L. Schomaker, "Deep adaptive learning for writer identification based on single

¹ Handwriting Recognition² Writer Identification³ Writer Verification⁴ Multi-Channel Gabor Filter (MGF)⁵ Moor⁶ Fisher's Linear Discriminant Analysis (LDA)⁷ Principal Component Analysis (PCA)⁸ Support Vector Machine (SVM)⁹ Bengali¹⁰ Deep Convolutional Neural Network (DCNN)¹¹ Recurrent Neural Network (RNN)¹² Random Hybrid Strokes (RHSs)¹³ Deep Neural Network (DNN)¹⁴ pre-trained network¹⁵ Convolutional Neural Network (CNN)¹⁶ Long Short Term Memory (LSTM)¹⁷ Artificial neural networks (ANNs)¹⁸ Feedforward

- ¹⁹ Back Propagation (BP)
- ²⁰ Loss Function
- ²¹ Pooling Layer
- ²² Fully Connected
- ²³ Over Fitting
- ²⁴ Dropout
- ²⁵ Batch Normalization (BN)
- ²⁶ Down Sampling
- ²⁷ Max-Pooling
- ²⁸ Softmax
- ²⁹ Separated form (S)
- ³⁰ Beginning Of Word form (BOW)
- ³¹ Middle Of Word form (MOW)
- ³² End Of Word form (EOW)
- ³³ Segmentation
- ³⁴ Linear
- ³⁵ Cross-Entropy
- ³⁶ Stochastic Gradient Descent (SGD)
- ³⁷ RandomGrayscale
- ³⁸ ColorJitter
- ³⁹ RandomRotation
- ⁴⁰ Test Time Augmentation (TTA)
- ⁴¹ Proposed Method (P-M)
- ⁴² Back-Propagation Neural Network (BPNN)
- ⁴³ Gaussian Radial Basis Function (RBF)
- ⁴⁴ Grid search method
- ⁴⁵ Area
- ⁴⁶ Centroid Coordinates
- ⁴⁷ Eccentricity
- ⁴⁸ Kurtosis
- ⁴⁹ Skewness