



طبقه‌بندی گلبول‌های سفید با استفاده از شبکه عصبی کانولوشنی

امین ادراکی (BS)^{۱*}، ابوالحسن رزمی نیا (PhD)^{۲**}

^۱ گروه مهندسی برق، دانشکده مهندسی برق و کامپیوتر، دانشگاه تهران، تهران، ایران

^۲ گروه مهندسی برق، دانشکده مهندسی، دانشگاه خلیج فارس بوشهر، بوشهر، ایران

(دریافت مقاله: ۹۶/۶/۱ - پذیرش مقاله: ۹۶/۷/۲۹)

چکیده

زمینه: مشاهده، دسته‌بندی و شمارش انواع مختلف گلبول‌های سفید در نمونه خون، یکی از گام‌های اساسی در درمان بیماری‌های مختلف است. هدف از انجام این پژوهش طراحی و پیاده‌سازی سیستمی سریع، قابل اعتماد و مبتنی بر پردازش تصاویر میکروسکوپی نمونه خون برای طبقه‌بندی چهار نوع از گلبول‌های سفید است.

مواد و روش‌ها: در این مقاله، از روش خوشه‌بندی k-means اصلاح‌شده برای انجام عمل بخش‌بندی تصویر استفاده شده است. علاوه بر این، عمل طبقه‌بندی گلبول‌های سفید با استفاده از یک شبکه عصبی کانولوشنی عمیق و با کمک داده‌های موجود در پایگاه داده MISP - پایگاه داده رایگان و متشکل از تصاویر میکروسکوپی نمونه خون - انجام شده است. همچنین، روش‌های مختلف رگولاریزاسیون مثل حذف تصادفی و افزایش تعداد تصاویر پایگاه داده، برای جلوگیری از بیش‌برازش (Overfitting) مدل پیشنهادی مورد استفاده قرار گرفته‌اند. **یافته‌ها:** در بخش طبقه‌بندی، دقت شبکه عصبی برابر ۹۹ درصد اندازه‌گیری شده است که نسبت به بسیاری از پژوهش‌های پیشین موفق‌تر بوده است. همچنین در بخش بخش‌بندی، شاخص اطلاعات متقابل برابر ۰/۷۳ حاصل شد.

نتیجه‌گیری: نتایج حاصل از این پژوهش نشان می‌دهد طراحی و پیاده‌سازی سیستمی سریع و قابل اعتماد با کمک پردازش تصاویر میکروسکوپی نمونه خون با استفاده از روش‌های مختلف پردازش تصویر و یادگیری ماشین امکان‌پذیر است.

واژگان کلیدی: بخش‌بندی تصویر، طبقه‌بندی تصویر، شبکه‌های عصبی عمیق، تصاویر میکروسکوپی نمونه خون، گلبول سفید، شبکه عصبی کانولوشنی

* گروه مهندسی برق، دانشکده مهندسی، دانشگاه خلیج فارس، بوشهر، ایران

Email: razminia@pgu.ac.ir

* ORCID: 0000-0002-0843-5522

مقدمه

هم‌گام با توسعه گسترده الگوریتم‌های رایانه‌ای و یادگیری ماشین و به موازات آن، افزایش توان پردازنده‌های رایانه‌ای، استفاده از این الگوریتم‌ها در کاربردهای مختلف پزشکی و زیست‌شناسی به سرعت رو به گسترش است. استفاده از این الگوریتم‌ها حوزه‌های مختلفی را در بر می‌گیرد: از سیستم‌های خودکار و نیمه‌خودکار جهت پردازش تصاویر پزشکی تا پردازش اطلاعات ژنوم انسان با استفاده از الگوریتم‌های مربوط به پردازش داده‌های بزرگ^۱ (۳-۱).

همان‌گونه که اشاره شد، پردازش خودکار و نیمه‌خودکار تصاویر مختلف پزشکی، همانند تصاویر سی‌تی، ام‌آر‌آی^۲، سونوگرافی و تصاویر میکروسکوپی نمونه خون به کمک سیستم‌های بینایی ماشین و پردازش تصویر یکی از کاربردهای بسیار متداول در این حوزه است (۲). هدف از اعمال الگوریتم‌های پردازش تصویر، معمولاً ایجاد تصاویر جدید با ویژگی‌های متفاوت نسبت به تصویر اصلی و یا استخراج ویژگی‌هایی بر اساس تصویر اصلی است.

یکی از کاربردهای سیستم‌های پردازش تصویر در حوزه پزشکی، آنالیز کمی و کیفی تصاویر میکروسکوپی نمونه خون است (۴). هدف از بررسی میکروسکوپی نمونه خون می‌تواند شامل شمارش سلول‌های مختلف موجود در نمونه خون، مانند گلبول‌های سفید، گلبول‌های قرمز و پلاکت‌ها و یا تحلیل کیفی این سلول‌ها، همانند تشخیص انواع مختلف سرطان خون همانند لوسمی حاد لفاوی باشد (۵). در این بین، تحلیل کمی و کیفی انواع مختلف گلبول‌های سفید، بسیار حائز اهمیت است چرا که می‌تواند به تشخیص و بررسی روند درمان بیماری‌های مختلفی مانند ایدز و سرطان خون کمک کند. به همین جهت شمارش انواع مختلف گلبول سفید در نمونه خون یکی از گام‌های مهم در بررسی و آزمایش نمونه خون است.

شمارش انواع مختلف گلبول سفید در نمونه خون به دو روش کلی خودکار و غیرخودکار قابل انجام است. در روش‌های غیرخودکار، نمونه خون توسط یک متخصص بررسی می‌شود که در آن، شمارش و تحلیل گلبول‌های نمونه خون، فرایندی کند، خسته‌کننده و خطاپذیر است. در مقابل این سبک، سیستم‌های خودکار متعددی برای بررسی کمی گلبول‌های سفید خون وجود دارند که اساس کار آنها عمدتاً بر دستگاه‌های جریان‌سنج^۳ و ویژگی‌های شیمیایی گلبول‌ها استوار است. این سیستم‌ها اغلب گران‌قیمت و تا حدودی کند هستند و صرفاً اطلاعات کمی مربوط به گلبول‌ها را در اختیار قرار می‌دهند. به همین جهت طراحی و پیاده‌سازی سیستم‌هایی ارزان، سریع و قابل اعتماد، جهت شمارش انواع مختلف گلبول سفید، حائز اهمیت است. یکی از روش‌های متداول در طراحی و پیاده‌سازی این سیستم‌ها، پردازش تصاویر میکروسکوپی نمونه خون است.

به عنوان نمونه در (۶)، روشی مبتنی بر اطلاعات رنگ پیکسل‌ها برای بخش‌بندی گلبول‌ها معرفی شده است. این روش‌ها از نظر محاسباتی و پیاده‌سازی، بسیار ساده هستند اما عملکرد دقیقی در جداسازی هسته و سیتوپلاسم ندارند. در برخی پژوهش‌های پیشین روش‌های گسترش ناحیه^۴ و آب‌پخشان^۵ برای تشخیص گلبول‌های سفید مورد استفاده قرار گرفته است (۷ و ۸). به دلیل مرز مبهم سیتوپلاسم، کاربرد این روش‌ها در این مسایل به نتایج نسبتاً مطلوبی منجر شده است و امروزه جزء شیوه‌های استاندارد تشخیص محسوب می‌شوند. همچنین در (۹) روشی مبتنی بر عمودسازی گرام-اشمیت^۶ برای تشخیص گلبول سفید و مدل‌های پارامتری تغییرپذیر^۷ برای جداسازی هسته از سیتوپلاسم مورد بررسی قرار گرفته است.

برای انجام دسته‌بندی گلبول‌ها اولین گام، استخراج ویژگی‌های تصویر است. با توجه به ماهیت مسأله در بسیاری از مقالات پیشین ویژگی‌های مورفولوژیک و بافتی،

¹ Big data

² MRI

³ Flowcytometry

⁴ Region growing

⁵ Watershed

⁶ Gram-Schmidt orthogonalization method

⁷ Parametric deformable models

مشخص کننده گلبول سفید از پس‌زمینه جدا می‌گردد و در ادامه برای آموزش شبکه عصبی، تعداد تصاویر موجود در پایگاه داده به کمک روش‌های مختلف افزایش می‌یابد. در گام بعد چند ساختار متفاوت برای شبکه عصبی، در نظر گرفته شده عملکرد آنها در طبقه‌بندی نوع سلول‌ها مقایسه می‌شود و در نهایت مدل نهایی بر اساس این مقایسه انتخاب می‌شود. و به کمک داده‌های موجود آموزش داده می‌شود. همچنین الگوریتم k-means اصلاح‌شده با کمک تبدیلات مورفولوژیک برای انجام بخش‌بندی و مقایسه نتایج با ماسک‌های موجود، مورد استفاده قرار خواهد گرفت.

پس از این مقدمه کوتاه، تعریف دقیق مسأله و برخی پیش‌نیازها در بخش پیش‌نیازها ارائه می‌شوند. همچنین مدل پیشنهادی در بخش توصیف مدل مورد بررسی قرار می‌گیرد و در بخش شبیه‌سازی‌ها نتایج شبیه‌سازی به تفصیل ارائه می‌شوند. در بخش نتیجه‌گیری نیز نتیجه‌گیری‌های اصلی حاصل از این پژوهش ارائه خواهد شد.

مواد و روش‌ها

در شکل ۱، بلوک دیاگرام فرایند پردازش تصویر به منظور شمارش گلبول‌های سفید نمونه خون نشان داده شده است. همان‌گونه که در شکل مشخص است، این فرایند از چهار بخش عمده تشکیل یافته است: پیش‌پردازش، بخش‌بندی، استخراج ویژگی و طبقه‌بندی تصویر. هدف از انجام پیش‌پردازش، کاهش نویزها و تغییرات تصادفی شدت روشنایی در تصویر و ایجاد تصویر مناسب برای گام‌های بعدی است. برای مدل کردن نویز موجود در تصویر و کاهش آن بر اساس این مدل‌ها، روش‌های مختلفی وجود دارد. به عنوان نمونه، در بسیاری موارد، فیلترهای حوزه مکان، همانند فیلترهای میانگین، میانه و گاوسی و یا فیلترهای پیچیده‌تر حوزه فرکانس و فیلترهای تطبیقی^{۱۴} برای کاهش نویز مورد استفاده قرار می‌گیرند

برای مقایسه گلبول‌ها مورد استفاده قرار می‌گیرند (۱۰) و (۱۱). برای دسته‌بندی این اطلاعات استخراج شده می‌توان از ابزارها و تکنیک‌های مختلف دسته‌بندی استفاده کرد. به طور نمونه در برخی تحقیقات پیشین از شبکه عصبی پیش‌خور و روش فازی برای دسته‌بندی گلبول‌ها استفاده شده است (۱۲-۱۴). همچنین در مقالات پیشین به عنوان یکی از نمونه‌های موفق طبقه‌بندی گلبول‌های سفید، در گام نخست از دو مجموعه از ویژگی‌ها برای عمل طبقه‌بندی استفاده شده است: ویژگی‌های بافتی و ویژگی‌های مورفولوژی. برای محاسبه ویژگی‌های بافتی، از ماتریس همایش^۸ در کنار الگوهای دودویی محلی^۹ استفاده شده و پس از انتخاب ویژگی‌ها با در نظر گرفتن نسبت تفکیک‌کننده فیشر^{۱۰}، عمل طبقه‌بندی به کمک دو روش بردار ماشین پشتیبان و پرسپترون چندلایه^{۱۱} انجام شده است (۱۵). دقت نهایی ۹۵ درصد در طبقه‌بندی حاصل شده، که اگرچه بسیار موفقیت‌آمیز است، اما با توجه به اهمیت شمارش دقیق گلبول‌ها، با اعمال تغییراتی قابلیت بهبود دارد.

هدف از انجام این پژوهش، ارائه راه‌یافتی مبتنی بر پردازش تصاویر میکروسکوپی گلبول‌های سفید خون با هدف تشخیص نوع این گلبول‌هاست. در این تحقیق، تمرکز ما بر تشخیص نوع گلبول سفید در تصاویر برش‌داده شده گلبول‌های سفید، به عنوان بخشی از یک سیستم جامع برای شمارش گلبول‌ها خواهد بود. برای انجام این کار از پایگاه داده رایگان پردازش سیگنال و تصویر پزشکی (MISP)^{۱۲} و یک شبکه عصبی کانولوشنی^{۱۳} استفاده شده است. این پایگاه داده شامل ۱۴۹ تصویر میکروسکوپی از گلبول‌های سفید اتوزینوفیل، لنفوسیت، مونوسیت و نوتروفیل است (۱۵) و (۱۶). از آنجا که این پایگاه داده شامل هیچ تصویری از سلول‌های بازوفیل نیست، تشخیص این نوع گلبول در این تحقیق مورد بررسی قرار نگرفته است. در گام ابتدایی با استفاده از ماسک‌های مربوط به بخش‌بندی سلول‌ها، قسمت

⁸ Co-occurrence matrix

⁹ Local binary patterns (LBP)

¹⁰ Fisher's discriminant ratio

¹¹ Multi-layer perceptron

¹² Medical Image and Signal Processing

¹³ Convolutional neural network (ConvNet)

¹⁴ Adaptive filters

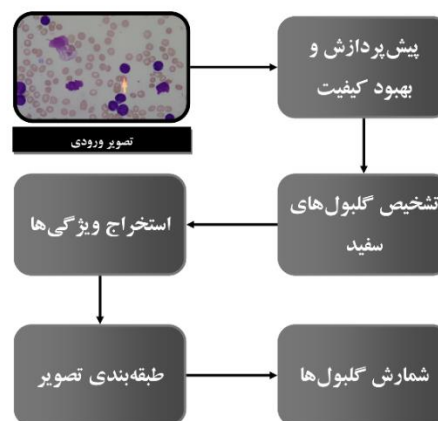
مرزها و ویژگی‌های سطحی در کنار طبقه‌بندی‌های مختلفی مانند بردار ماشین پشتیبان (۱۹)، سال‌ها برای طبقه‌بندی تصاویر مورد استفاده قرار گرفته‌اند. انتخاب و استخراج ویژگی‌ها در استفاده از این سبک برای طبقه‌بندی تصاویر با چالش‌های زیادی روبروست. یکی از روش‌های جایگزین پیشرفته برای طبقه‌بندی تصاویر که امروزه بسیار مورد توجه است، استفاده از شبکه‌های عصبی کانولوشنی است. این شبکه‌های عصبی که تغییر یافته مدل کلاسیک شبکه‌های عصبی پرسپترون چند لایه هستند، با توجه به کاهش قابل ملاحظه تعداد وزن‌ها و ساختار خود که مبتنی بر استخراج ویژگی‌های تصویر به شکل محلی است، برای استفاده در طبقه‌بندی تصاویر و مسائل مربوط به بینایی ماشین بسیار مناسب هستند و در ادبیات پزشکی مدرن، عملکرد قابل قبولی از آن گزارش شده است (۲۰). در ادامه به اختصار الگوریتم‌ها و ساختارهایی که در قسمت‌های آتی این نوشتار مورد استفاده قرار می‌گیرند مرور می‌شوند.

شبکه‌های عصبی کانولوشنی

شبکه‌های عصبی کانولوشنی نوع خاصی از شبکه‌های عصبی عمیق پیش‌خور هستند که با اعمال تغییراتی بر شبکه‌های کلاسیک پرسپترون چند لایه طراحی شده‌اند. این شبکه‌ها با الهام از سیگنال‌های طبیعی چهار ویژگی اصلی دارند که آنها را به انتخابی مناسب برای پردازش سیگنال‌های طبیعی و ماتریس‌های چندبُعدی تبدیل می‌کند. این ویژگی‌ها عبارت‌اند از: ارتباطات محلی بین نورون‌ها، وزن‌های مشترک در هر لایه، استفاده از تعداد لایه‌های زیاد و استفاده از ویژگی جمع‌آوری^{۱۸} (۲۱).

شبکه‌های عصبی کانولوشنی در بخش‌های ابتدایی، از لایه‌های کانولوشنی و لایه‌های جمع‌آوری تشکیل شده‌اند. نورون‌های لایه‌های کانولوشنی توسط مجموعه‌ای از ضرایب (وزن‌های شبکه عصبی) به بخشی از نورون‌های

(۱۷). در ادامه، هدف از بلوک بخش‌بندی، مشخص کردن و جدا کردن قسمت‌هایی از تصویر است که مایل به پردازش آن‌ها هستیم. در این تحقیق هدف از انجام این بخش، جدا کردن گلبول‌های سفید خون از پس زمینه و سایر المان‌های موجود در تصویر برای انجام پردازش‌های بیشتر است. با توجه به اهمیت این مرحله، محققین، روش‌های متعددی برای انجام آن پیشنهاد داده‌اند که از جمله آنها می‌توان به روش‌های مبتنی بر اطلاعات لبه‌ها، روش‌های خوشه‌بندی همانند k-means، الگوریتم‌های آب‌پخشان، الگوریتم‌های آستانه‌گذاری^{۱۵} و روش‌های مبتنی بر یادگیری عمیق^{۱۶} و شبکه‌های عصبی مصنوعی اشاره نمود (۱۸).



شکل (۱) بلوک دیاگرام فرایند پردازش تصویر جهت شمارش انواع گلبول‌های سفید در تصویر نمونه خون.

در مسائل طبقه‌بندی تصویر، پس از مشخص کردن قسمت‌هایی از تصویر که هدف ما پردازش و مطالعه آن است، استخراج ویژگی‌های بخش مورد نظر بایستی انجام گیرد. در اغلب روش‌های طبقه‌بندی، بخشی از داده‌ها جهت آموزش طبقه‌بند^{۱۷} استفاده می‌گردد و دقت مدل با کمک سایر داده‌ها مورد ارزیابی قرار می‌گیرد. استخراج ویژگی‌ها نیز همانند سایر قسمت‌ها به روش‌های مختلفی انجام می‌پذیرد. به عنوان نمونه، ویژگی‌های مربوط به

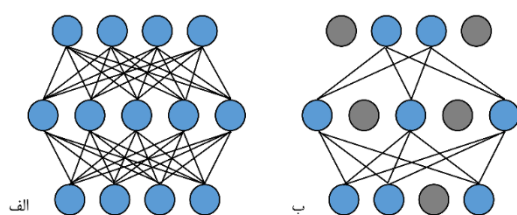
¹⁵ Thresholding

¹⁶ Deep learning

¹⁷ Classifier

¹⁸ Pooling

مقداردهی اولیه پارامترها را کمتر می‌کند که بدین ترتیب به صورتی طبیعی، به تسریع فرایند یادگیری کمک می‌کند. یکی دیگر از چالش‌هایی که در آموزش شبکه‌های عصبی عمیق با تعداد وزن‌ها و پارامترهای زیاد موضوعیت می‌یابد، بیش‌برازش^{۲۳} است. در این حالت، آزادی مدل پیشنهادی بسیار بیشتر از مدل واقعی داده‌هاست و اگرچه نتایج حاصل از استفاده از مدل بر داده‌های آموزش بسیار مطلوب است، اما عملکرد مدل روی داده‌های آزمون ضعیف خواهد بود. در واقع در این وضعیت، مدل به جای یادگیری الگوهای موجود، داده‌ها را به حافظه می‌سپارد. روش‌های مختلفی، همچون اعتبارسنجی متقابل^{۲۴} برای جلوگیری از بیش‌برازش وجود دارد. در (۲۳) روشی با عنوان حذف تصادفی^{۲۵} برای جلوگیری از این پدیده در شبکه‌های عصبی عمیق پیش‌نهاد شده است که در آن، طی آموزش شبکه عصبی، خروجی برخی نوروها به طور تصادفی از فرایند آموزش حذف می‌شوند. این موضوع از وابستگی بیش از حد مدل به داده‌های آموزش جلوگیری می‌کند. عملکرد این روش در (۲۳) در بسیاری از کاربردها نظیر بینایی ماشین و پردازش اسناد مورد بررسی قرار گرفته و پیشرفت چشمگیر عملکرد مدل‌های پیشنهادی در این موارد نشان داده شده است.



شکل ۲) مدل حذف تصادفی در آموزش شبکه‌های عصبی عمیق: در این روش خروجی برخی نوروهای شبکه عصبی در هر گام از فرایند آموزش به طور تصادفی صفر می‌شود و تأثیر آنها در آن مرحله در نظر گرفته نمی‌شود؛ الف: یک شبکه عصبی ساده پیش از اعمال حذف تصادفی. ب: پس از اعمال حذف تصادفی.

الگوریتم‌های بهینه‌سازی شبکه‌های عصبی

لایه پیشین متصل می‌شوند. جمع ورودی‌های وزن‌دار در هر قسمت، از تابعی غیرخطی، موسوم به تابع فعال‌سازی، همانند تابع سیگموئید و یا تابع رلو^{۱۹} عبور می‌کند. این لایه‌ها وظیفه استخراج ویژگی‌ها از ورودی شبکه عصبی را به عهده دارند و ویژگی اصلی آن استخراج ویژگی‌های محلی و پیوستگی بین این ویژگی‌هاست. در مقابل، لایه‌های جمع‌آوری وظیفه تجمیع ویژگی‌های مشترک را بر عهده دارند. این لایه‌ها معمولاً با انتخاب بیشینه مقادیر در یک همسایگی، علاوه بر کاهش ابعاد مسأله، کار تجمیع ویژگی‌ها را نیز انجام می‌دهند. در هر شبکه عصبی کانولوشنی تعدادی از لایه‌های جمع‌آوری و کانولوشنی به همراه توابع غیرخطی نظیر رلو در کنار تعدادی لایه چگال^{۲۰} با اتصالات کامل (و نه محلی) بین نوروها کار استخراج ویژگی و طبقه‌بندی را انجام می‌دهند. وزن‌های بانک‌های فیلتر این شبکه‌ها نیز همانند سایر شبکه‌های عصبی با استفاده از قاعده بازگشت به عقب^{۲۱} محاسبه می‌شود (۲۱).

چالش‌های آموزش شبکه‌های عصبی مصنوعی

یکی از چالش‌هایی که در آموزش شبکه‌های عصبی عمیق بروز پیدا می‌کند، کُندی روند آموزش شبکه به دلیل تغییر توزیع ورودی‌های هر لایه در طول آموزش است. این تغییر به دلیل تغییر وزن‌ها و پارامترهای لایه پیشین رخ می‌دهد که منجر به کُند شدن روند یادگیری شبکه عصبی می‌شود. سه دلیل عمده کاهش نرخ یادگیری در این شیوه عبارتند از: لزوم استفاده از نرخ‌های یادگیری پایین‌تر، مقداردهی اولیه دشوار پارامترها و برخورد با مدل‌هایی با خواص غیرخطی اشباع‌شونده. در پژوهش‌های پیشین روشی با عنوان نرمالیزاسیون دسته‌ای^{۲۲} مبتنی بر نرمال کردن ورودی هر لایه شبکه عصبی به عنوان جزئی از مدل و در طول هر دسته از داده‌های آموزش پیشنهاد شده است (۲۲). این روش امکان استفاده از نرخ‌های یادگیری بالاتر را فراهم کرده و اهمیت

¹⁹ ReLu

²⁰ Dense layers

²¹ Back propagation

²² Batch normalization

²³ Overfitting

²⁴ Cross validation

²⁵ Dropout

که در آن α نرخ هر گام و ϵ یک رقم متعادل‌کننده^{۳۱} است و $f'(\theta_t)$ مشتق تابع خطا نسبت به پارامترها در گام زمانی t است. علاوه بر این تابع g, s نیز به صورت زیر قابل تحصیل است:

$$g_t = (1 - \gamma)f'(\theta_t)^2 + \gamma g_{t-1} \quad (۳)$$

$$s_t = (1 - \gamma)\Delta\theta_t^2 + \gamma s_{t-1} \quad (۴)$$

جایی که γ نرخ نزول، $s_0 = 0$ و $g_0 = 0$ است. همان‌طور که مشاهده می‌شود این روش فقط از اطلاعات مشتق مرتبه اول تابع خطا استفاده می‌کند و به تنظیم دستی نرخ یادگیری نیاز ندارد. همچنین در کاربردهای مختلف مشاهده شده است که این روش در شرایطی که گرادیان تابع خطا نویزی باشد، عملکرد قابل قبولی از خود نشان می‌دهد (۲۴).

شاخص اطلاعات متقابل

در بررسی آماری داده‌ها، شاخص اطلاعات متقابل^{۳۲} معیاری برای سنجش میزان وابستگی دو متغیر تصادفی است. به بیان ساده، این شاخص، میزان اطلاعات به دست آمده از یک متغیر توسط متغیر دیگر را نشان می‌دهد. تعریف این شاخص برای دو متغیر تصادفی گسسته X و Y از این قرار است:

$$I(X, Y) = \sum_{y \in Y} \sum_{x \in X} p(x, y) \log \left(\frac{p(x, y)}{p(x)p(y)} \right) \quad (۵)$$

تابع خطا در شبکه‌های عصبی بسته به کاربرد آنها تعریف‌های مختلفی دارد. به طور نمونه در حالت طبقه‌بندی با دو دسته، در بسیاری موارد تابع آنتروپی متقابل دودویی^{۲۶} و در حالت دسته‌بندی با چند دسته معمولاً تابع آنتروپی متقابل قیاسی^{۲۷} به عنوان تابع خطای شبکه عصبی در نظر گرفته می‌شود. در همه این حالت‌ها هدف از آموزش شبکه عصبی، یافتن وزن‌های شبکه به طریقی است که تابع خطا کمینه گردد. در اغلب روش‌های موجود برای تغییر وزن‌ها این کار با استفاده از اطلاعات مشتق مرتبه اول و دوم تابع خطا نسبت به وزن‌ها و حرکت در جهت کاهش خطا انجام می‌گیرد. روش‌های مختلفی همانند نزول گرادیان تصادفی^{۲۸}، آدام^{۲۹} و آدالتا^{۳۰} برای تغییر وزن‌ها در هر گام از فرایند آموزش، پیشنهاد شده‌اند. در ادامه روابط ریاضی روش بهینه‌سازی آدالتا که در این تحقیق مورد استفاده قرار گرفته است به اختصار مورد بررسی قرار می‌گیرد.

اگر ماتریس وزن‌ها در گام t را با θ_t نشان دهیم آنگاه قاعده تغییر وزن‌ها در روش آدالتا از طریق رابطه (۱) و (۲) محاسبه می‌شود:

$$\Delta\theta_t = \alpha \frac{\sqrt{s_{t-1} + \epsilon}}{\sqrt{g_t + \epsilon}} f'(\theta_t), \quad (۱)$$

$$\theta_{t+1} = \theta_t - \Delta\theta_t \quad (۲)$$

^{۲۶} Binary cross entropy

^{۲۷} Categorical cross entropy

^{۲۸} Stochastic gradient descent

^{۲۹} Adam

^{۳۰} AdaDelta

^{۳۱} Offset

^{۳۲} Mutual information

دلیل در این مقاله از روش مقداردهی اولیه k -means++ استفاده شده است. در این روش نقاط اولیه (مراکز هر دسته) به ترتیب و به شکل احتمالی بر اساس فاصله از مراکز دسته‌های پیشین انتخاب می‌شوند. الگوریتم‌های مختلف خوشه‌بندی همانند k -means بسیاری از کاربردهای بخش‌بندی تصویر مورد استفاده قرار می‌گیرند. در این روش‌ها عموماً هر پیکسل به عنوان یک داده در نظر گرفته شده اطلاعات مربوط به رنگ (و در برخی موارد مکان آن) به عنوان ویژگی‌های هر داده اعلام می‌شود. در این پژوهش از معیار اقلیدسی با وزن‌های متفاوت برای سنجش فاصله دو نقطه در عمل خوشه‌بندی به روش k -means استفاده شده است. تابع فاصله در این روش به این شکل تعریف می‌شود:

$$D(x_i, C) = \sqrt{\sum_j a_j (x_i^j - C^j)^2} \quad (۷)$$

که در آن x_i^j و C^j نشان‌دهنده ویژگی j ام داده x_i و C هستند. همچنین ضرایب a_j تأثیر هر یک از ویژگی‌ها در محاسبه تابع فاصله را مشخص می‌کنند. چنانکه مشاهده می‌شود تأثیر ویژگی‌های مختلف در محاسبه تابع فاصله یکسان نیست؛ بنابراین می‌توان با انتخاب مناسب این ضرایب، تأثیر متفاوت ویژگی‌ها بر تابع فاصله را توصیف کرد. روش‌های گوناگونی برای محاسبه این ضرایب وجود دارد که در این تحقیق از الگوریتم ژنتیک بدین منظور استفاده شده است.

توصیف مدل

در شکل ۱ بلوک دیاگرام سیستم پیشنهادی برای طبقه‌بندی گلوبول‌های سفید نشان داده شده است. با توجه به اینکه فرایند پیش پردازش، تشخیص و جداسازی گلوبول‌ها روی مجموعه داده در اختیار ما از قبل انجام شده است، در این پژوهش تمرکز عمده بر طراحی و پیاده‌سازی دو بخش پایانی، یعنی

که در آن $p(x, y)$ تابع توزیع مشترک دو متغیر تصادفی و $p(x)$ و $p(y)$ به ترتیب تابع توزیع مرزی متغیرهای تصادفی x و y هستند. یکی از کاربردهای نوعی این شاخص مقایسه عملکرد بخش‌بندی تصویر است که در قسمت‌های بعدی مقاله بدان پرداخته می‌شود.

الگوریتم خوشه‌بندی k -means

هدف از ارائه الگوریتم k -means در حالت کلی طبقه‌بندی N داده با M ویژگی به K دسته است به طوری که مجموع فاصله اعضای هر دسته تا مرکز آن برای کل داده‌ها کمینه شود. به عبارت دیگر هدف نهایی از ارائه الگوریتم، کمینه کردن تابع خطای زیر با یافتن نقاط بهینه برای مرکز دسته‌هاست:

$$J = \sum_{i=1}^k \sum_{x \in C_i} D^2(x, c_i) \quad (۸)$$

که در آن $D(x, c_i)$ نشان‌دهنده فاصله مرکز دسته i ام تا داده x است. در پیاده‌سازی الگوریتم k -means می‌توان از معیارهای مختلفی همانند فاصله اقلیدسی و فاصله منتهن^{۳۳} استفاده کرد. مستقل از تعریف فاصله، روش اجرای الگوریتم مذکور به صورت زیر است (۲۵):

داده‌ها به طور تصادفی به k دسته ناتهی^{۳۴} تقسیم می‌شوند.

۱. با استفاده از تقسیم‌بندی موجود داده‌ها، مرکز هر دسته به شکل میانگین داده‌ها محاسبه می‌شود.

۲. فاصله هر داده از مرکز هر دسته محاسبه شده و تقسیم‌بندی

بر اساس نزدیک‌ترین مرکز دسته انجام می‌پذیرد.

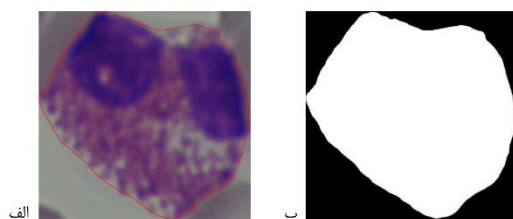
۳. اگر تغییری در دسته‌بندی ایجاد شد، الگوریتم از گام ۲

تکرار می‌شود. در غیر این صورت نتایج دسته‌بندی به عنوان خروجی الگوریتم در نظر گرفته می‌شوند.

لازم به ذکر است انتخاب دسته‌بندی اولیه و مرکز دسته‌ها به شکل کاملاً تصادفی، موجب همگرایی کند روش و افزایش احتمال قرارگیری در نقطه کمینه محلی می‌شود. به همین

³³ Manhattan

³⁴ Non-empty set



شکل ۳) نمونه تصاویر ورودی و خروجی در قسمت بخش‌بندی: هدف از انجام این بخش جداسازی کامل تصویر گلبول سفید از پس‌زمینه است - الف: تصویر ورودی نمونه. ب: تصویر ماسک ایده‌آل تولید شده در قسمت بخش‌بندی برای جداکردن گلبول سفید.

در این نوشتار ابتدا تصویر از فضای رنگی RGB به فضای رنگی HSV منتقل شده است. در فضای رنگی HSV همان‌های تشکیل‌دهنده هر پیکسل به جای میزان رنگ‌های قرمز، سبز و آبی، رنگ، میزان اشباع و شدت روشنایی هر پیکسل در نظر گرفته می‌شوند. به این ترتیب هر پیکسل تصویر، یک داده با سه ویژگی خواهد بود، اما تأثیر هر یک از این المان‌ها در خوشه‌بندی و تعیین فاصله نقطه‌ها یکسان نیست. به همین دلیل در این پژوهش از فاصله اقلیدسی وزن‌دار استفاده شده است. همچنین اثر لحاظ کردن مختصات مکانی هر پیکسل در سنجش فاصله نیز بررسی شده است. به این ترتیب تابع فاصله بین نقطه‌ای مانند $P = (h, s, v, x, y)$ که در آن h , s , v , x و y به ترتیب مقادیر hue, saturation, value, مکان عرضی و مکان طولی پیکسل هستند با مرکز دسته $C = (h_c, s_c, v_c, x_c, y_c)$ برابر است با:

$$D(P, C) = \sqrt{a_h(h - h_c)^2 + a_s(s - s_c)^2 + a_v(v - v_c)^2 + a_x\left(\frac{x - x_c}{W}\right)^2 + a_y\left(\frac{y - y_c}{H}\right)^2}$$

که در آن a_i ها ضرایب ثابت و تعیین‌کننده تأثیر هر ویژگی در سنجش فاصله و W و H به ترتیب عرض و ارتفاع تصویر هستند.

برای یافتن مقادیر مناسب ضرایب ویژگی‌های هر پیکسل، از الگوریتم ژنتیک استفاده می‌نماییم. در این گام ابتدا ۱۶

بخش‌بندی و طبقه‌بندی تصویر ورودی است. در واقع، در این تحقیق، یک سیستم کامل برای شمارش گلبول‌های سفید از تصاویر میکروسکوپی نمونه خون معرفی نخواهد شد، بلکه تنها بخش‌های مربوط به بخش‌بندی نهایی و طبقه‌بندی مورد بررسی قرار می‌گیرند که این محدودیت به واسطه‌ی تصاویر موجود در پایگاه داده اعمال شده است. در ادامه پس از معرفی پایگاه داده، به مسأله بخش‌بندی تصویر می‌پردازیم. طبقه‌بندی تصاویر نیز به عنوان گام پایانی مورد تجزیه و تحلیل قرار خواهد گرفت.

پایگاه داده

همان‌طور که پیش از این نیز اشاره شد، در این پژوهش از تصاویر موجود در پایگاه داده MISP استفاده شده است. این پایگاه داده شامل مجموعاً ۱۴۹ تصویر از چهار نوع گلبول سفید است که از آزمایش نمونه خون ۱۰ نفر به دست آمده است. این تصاویر رنگی در فضای RGB توسط یک دوربین گنون V1^{۳۵} و یک میکروسکوپ نوری گنون در ابعاد ۲۵۹۲×۳۸۷۲ ذخیره شده‌اند. همچنین برای هر تصویر، ماسک مشخص‌کننده گلبول از پس‌زمینه توسط افراد متخصص تعیین شده است.

بخش‌بندی تصویر

هدف از انجام این بخش همان‌طور که پیشتر اشاره شد، جداسازی کامل گلبول سفید از تصویر پس‌زمینه است. در واقع هدف از انجام این قسمت تولید تصویر دودویی هم‌اندازه با تصویر ورودی است که پیکسل‌های متعلق به گلبول سفید را از پس‌زمینه جدا کند. شکل ۳ نمونه ورودی و خروجی مطلوب برای این قسمت را نشان می‌دهد.

تصویر از میان داده‌های موجود انتخاب می‌شوند و سپس تابع خطای مربوط به انتخاب ضرایب a_i به شکل میانگین نرمالیزه شده شاخص اطلاعات متقابل، برای خوشه‌بندی معیار و خوشه‌بندی حاصل از اعمال الگوریتم در نظر گرفته شده و عمل بهینه‌سازی این تابع خطا با انتخاب ضرایب a_i با کمک الگوریتم ژنتیک انجام می‌گیرد. در واقع ضرایب a_i به کمک الگوریتم ژنتیک به گونه‌ای محاسبه می‌شوند که حاصل خطای خوشه‌بندی روی ۱۶ تصویر انتخاب شده کمینه شود. جزئیات مربوط به شبیه‌سازی‌های این قسمت در بخش ۴ ارائه شده است.

پس از اعمال خوشه‌بندی توسط الگوریتم k -means، دسته‌ای که میانگین اشباع پیکسل‌ها در آن بیشتر است به عنوان دسته مشخص‌کننده گلول سفید در نظر گرفته می‌شود. سپس با استفاده از تبدیلات مورفولوژیک باز کردن، بستن و فرسایش^{۳۶}، سوراخ‌های کوچک موجود در ماسک حاصل از گام قبل، پر شده لبه‌های ماسک کمی نرم می‌شوند. همچنین در انتها برای پاک‌سازی تصویر، قسمت‌های هم‌بندی^{۳۷} که مساحت آنها از یک آستانه معین کمتر یا بیشتر باشد از تصویر حذف خواهند شد.

طبقه‌بندی تصویر

در این پژوهش برای طبقه‌بندی تصویر حاصل از ترکیب ماسک بخش‌بندی و تصویر ورودی، از یک شبکه عصبی کانولوشنی استفاده می‌شود. در ادامه ساختارهای مختلفی برای شبکه عصبی در نظر گرفته شده و عملکرد آنها در بخش شبیه‌سازی مقایسه می‌گردد. شکل ۴ دو مورد از ساختارهایی که در این پژوهش مورد بررسی قرار گرفته‌اند را نشان می‌دهد. همان‌طور که مشاهده می‌شود این ساختارها از دو یا سه لایه کانولوشنی، دو لایه چگال و یک لایه جمع‌آوری‌کننده تشکیل شده‌اند. در هر دو مدل پیشنهادی لایه نرمالیزاسیون دسته‌ای، پس از لایه‌های کانولوشنی به کار گرفته می‌شود. همچنین اثر لایه حذف

تصادفی و الگوریتم‌های بهینه‌سازی مختلف در بخش شبیه‌سازی مورد بررسی قرار خواهند گرفت. در تمام مدل‌های بررسی شده در این مقاله، ابعاد هسته‌های لایه‌های کانولوشنی برابر 3×3 ، ابعاد فضای خروجی برابر ۳۲، تعداد نورون‌های لایه‌های چگال یک و دو به ترتیب برابر ۱۲۸ و چهار، و همچنین تابع فعال‌سازی برای تمامی لایه‌ها به جز لایه آخر، تابع رلو و در لایه آخر تابع سافت‌مکس^{۳۸} در نظر گرفته شده است. تابع رلو در بسیاری از پژوهش‌های پیشین حوزه بینایی ماشین، عمل‌کرد بسیار مطلوبی از خود نشان داده است. به طور مثال در (۲۶) استفاده از تابع رلو باعث تسریع همگرایی یادگیری شبکه‌ی عصبی تا شش برابر در مقایسه با سایر توابع فعال‌سازی در آموزش شبکه عصبی با پایگاه‌های داده مختلف شده است. این تابع در کنار تسریع فرایند همگرایی، به دلیل محاسبات ساده‌تر در مقایسه با توابعی مثل سیگنویید، حجم محاسبات لازم را نیز کاهش می‌دهد. به این ترتیب این تابع تقریباً در تمامی مسائل بینایی ماشین به عنوان تابع فعال‌سازی لایه‌های میانی به کار گرفته می‌شود.

برای تعیین تعداد لایه‌های چگال و همچنین تعداد نورون‌های این لایه‌ها، باید به این نکته توجه کرد که در یک شبکه عصبی کانولوشنی، بایاس‌ها و وزن‌های لایه‌های چگال، به دلیل اتصال کامل بین نورون‌ها، بخش زیادی از پارامترهای شبکه عصبی را تشکیل می‌دهند. به همین جهت و با در نظر گرفتن این نکته که جلوگیری از بیش‌برازش شبکه عصبی یکی از چالش‌های اصلی در حل این مسئله است، می‌بایست تعداد لایه‌ها و همچنین تعداد نورون‌های هر لایه برابر با کمترین میزان ممکن در نظر گرفته شود. با توجه به نوع مسئله (طبقه‌بندی با چهار کلاس) تعداد نورون‌های لایه آخر شبکه عصبی مشخص است. همچنین برای جلوگیری از بیش‌برازش داده‌ها، به جز لایه پایانی تنها از یک لایه چگال استفاده شده است. به این ترتیب با آزمایش چند مقدار مختلف (با توجه به ابعاد فضاهای ایجاد

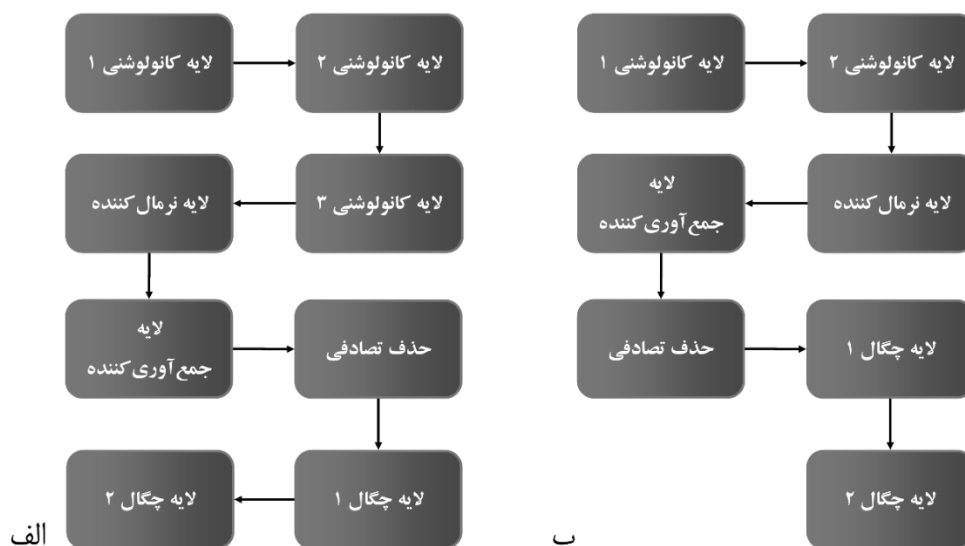
³⁶ Erosion

³⁷ Connected components

³⁸ Softmax

شده در لایه‌های کانولوشنی) به عنوان تعداد نورون‌های لایه چگال، می‌توان مقدار مناسب این هاپرپارامتر را تعیین کرد. با این حال می‌بایست عنوان کرد که همانند بسیاری از

مسائل یادگیری ماشین، بهینه‌سازی کامل این هاپرپارامتر با روش‌های تحلیلی عملاً امکان‌پذیر نیست و می‌بایست از روش‌های مبتنی بر آزمون و خطا استفاده کرد



شکل ۴) ساختارهای مختلف بررسی شده برای شبکه عصبی کانولوشنی: در این پژوهش ساختارهای متفاوتی از شبکه‌های عصبی کانولوشنی با تعداد لایه‌های متفاوت برای مقایسه کارایی آنها در فرایند طبقه‌بندی گلوبول‌ها مورد بررسی قرار گرفته‌اند که دو نمونه از این ساختارها در این شکل نشان داده شده است؛ الف: ساختار با سه لایه کانولوشنی. ب: ساختار با دو لایه کانولوشنی.

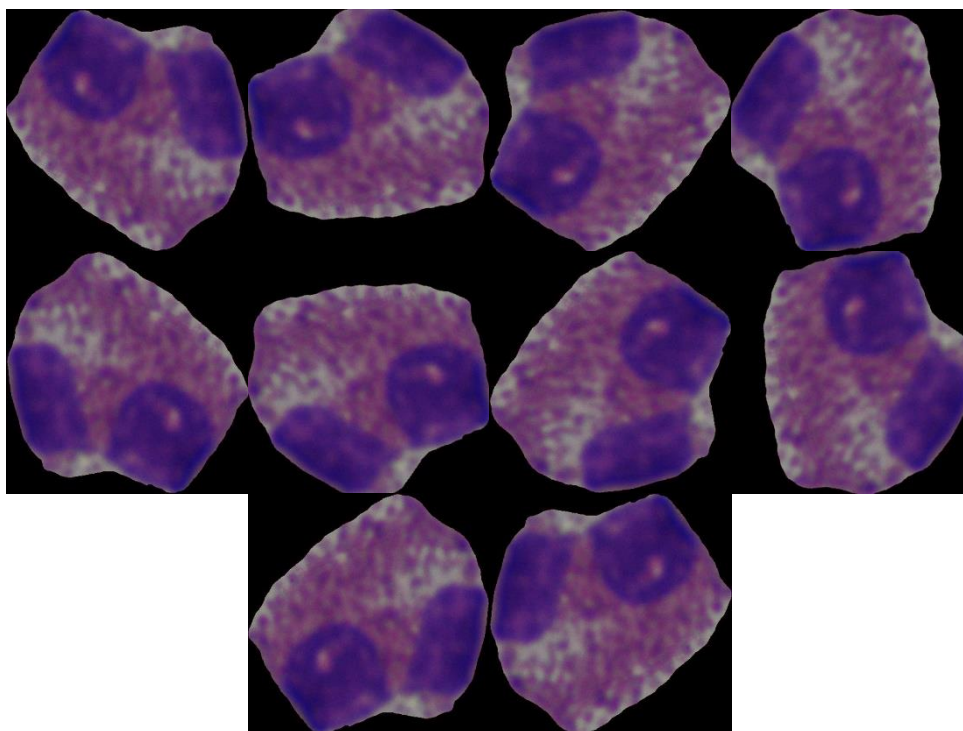
یافته‌ها

آماده‌سازی پایگاه داده

پیش از این اشاره شد که پایگاه داده MISP شامل مجموعاً ۱۴۹ تصویر از گلوبول‌های سفید نمونه خون است که این تعداد برای آموزش شبکه عصبی مورد استفاده در فرایند طبقه‌بندی کافی نخواهد بود. به همین دلیل لازم است پیش از آموزش شبکه عصبی، تعداد تصاویر موجود در پایگاه داده را به کمک روش‌های مختلف کارآمدی افزایش دهیم. با توجه به این‌که گلوبول‌های سفید خون، جهت‌گیری مشخص فضایی ندارند، یکی از روش‌هایی که می‌توان به کمک آن، تصاویر جدیدی به پایگاه داده اضافه کرد چرخش و قرینه‌سازی تصاویر اصلی است. در این پژوهش ما برای افزایش تعداد تصاویر موجود در پایگاه داده از سه

روش چرخش، قرینه‌سازی و اضافه کردن نویز گاوسی برای ایجاد تصاویر جدید استفاده کرده‌ایم. همچنین برای کاهش اندازه ورودی شبکه عصبی، سادگی مدل و کاهش تعداد وزن‌ها و بایاس‌های موجود در شبکه، ابعاد تصویر اصلی را پس از ترکیب با ماسک مشخص‌کننده گلوبول سفید به 64×64 پیکسل کاهش داده‌ایم. چنانکه نتایج شبیه‌سازی نشان می‌دهد این اندازه، مقداری بهینه است که کمتر از آن، نتیجه رضایت‌بخشی حاصل نخواهد شد.

شکل ۵ تصاویر مختلف تولید شده از یک تصویر موجود در پایگاه داده را نشان می‌دهد. با اعمال این تغییرات تعداد تصاویر موجود در پایگاه داده از ۱۴۹ تصویر به ۱۴۹۰ تصویر افزایش یافته است که مقدار قابل قبولی برای آموزش شبکه عصبی مورد استفاده می‌باشد.



شکل ۵) تغییرات اعمال شده روی یک نمونه تصویر موجود در پایگاه داده، برای افزایش تعداد تصاویر موجود در طول فرایند آموزش شبکه عصبی.

بخش‌بندی

همان‌طور که در بخش بخش‌بندی تصویر توضیح داده شد، بهینه‌سازی ضرایب معیار فاصله در الگوریتم k -means توسط الگوریتم ژنتیک انجام می‌گیرد. شکل ۶ نمودار تابع خطای تعریف شده در تکرارهای مختلف الگوریتم ژنتیک را نشان می‌دهد. پس از اتمام اجرای الگوریتم ژنتیک، بهترین ضرایب یافته شده در قسمت بعد برای سنجش عملکرد الگوریتم خوشه‌بندی مورد استفاده واقع شده و میانگین شاخص اطلاعات متقابل برای همه تصاویر موجود در پایگاه داده محاسبه شده است. این نتایج در جدول ۱ قابل مشاهده است.

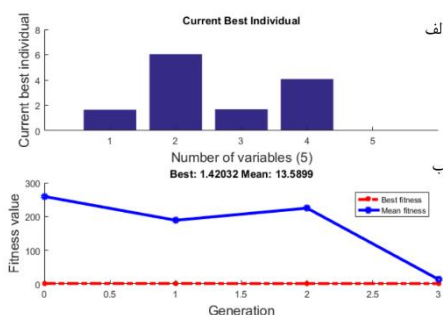
جدول ۱) مقایسه عملکرد خوشه‌بندی (بخش‌بندی تصویر) با کمک معیار شاخص اطلاعات متقابل قبل و بعد از وزن‌دهی پارامترها به کمک الگوریتم ژنتیک		
پارامترهای در نظر گرفته شده	بدون در نظر گرفتن وزن‌ها	با در نظر گرفتن وزن‌ها
میانگین شاخص اطلاعات متقابل	۰/۵۹	۰/۷۳
واریانس شاخص اطلاعات متقابل	۰/۰۷	۰/۰۶

در اجرای الگوریتم ژنتیک تعداد اعضای جمعیت برابر ۲۵، حداکثر تکرارها برابر ۳ و همچنین نرخ جهش برابر ۰/۲ در نظر گرفته شده است. با توجه به این که روشی برای تعیین پارامترهای الگوریتم ژنتیک برای عملکرد بهینه به شکل نظری و در حالت کلی وجود ندارد، در حل اغلب مسائل از مجموعه‌ای از پارامترهای مرسوم استفاده می‌شود. در این مسئله نیز از پارامترهای پیش‌فرض نرم‌افزار متلب استفاده شده است، اما با توجه به حجم زیاد محاسبات، تعداد اعضای جمعیت کاهش یافته است. همچنین با توجه به همگرایی روش پس از تکرارهای محدود (با آزمون و خطا) تعداد تکرارهای الگوریتم برابر سه در نظر گرفته شده است. همان‌طور که از جدول مزبور قابل مشاهده است، اجرای الگوریتم ژنتیک برای اصلاح وزن‌های پارامترهای مختلف، موجب بهبود چشمگیر شاخص اطلاعات متقابل در بخش‌بندی تصویر شده است. لازم به ذکر است که روش‌های ساده‌تری، مانند روش آستانه‌گذاری، نیز برای بخش‌بندی تصویر در این حالت وجود دارد. به عنوان نمونه، حاصل شاخص اطلاعات متقابل در اعمال روش ساده آستانه‌گذاری

متوالی، کاهش نیابد، آموزش متوقف شده و مدل موجود به عنوان مدل نهایی اعلام می‌شود. سپس دقت نهایی مدل با کمک داده‌های سنجش نهایی ارزیابی می‌گردد. این دقت، معیار خوبی از دقت واقعی مدل خواهد بود زیرا این داده‌ها در هیچ یکی از بخش‌های فرآیند آموزش مورد استفاده قرار نگرفته‌اند و برای مدل ناشناخته هستند.

شکل ۷ و شکل ۸ نمودار تابع خطا و دقت مدل‌های پیشنهادی در شکل ۴ را نشان می‌دهند. همان‌طور که مشاهده می‌شود، روش بهینه‌سازی آدالتا عملکرد مطلوبی در فرایند یادگیری شبکه عصبی داشته است. جدول ۲ حاصل دقت نهایی مدل‌های پیشنهادی روی داده‌های سنجش پایانی را نشان می‌دهد. همان‌طور که انتظار می‌رود، عملکرد نهایی این مدل نسبت به مدل‌های کلاسیک، مثل پرسپترون چندلایه و بردار ماشین پشتیبان، با دقت بالاتری همراه است. در پژوهش‌های پیشین مثل (۸) دقت طبقه‌بندی با کمک این روش‌ها در بهترین حالت نزدیک به ۹۵ درصد است که عملکرد بسیار مطلوب شبکه‌های عصبی کانولوشنی را در مقایسه با این روش‌ها نشان می‌دهد. همچنین در جدول ۳ ماتریس اغتشاش^{۳۲} مربوط به طبقه‌بندی داده‌های آزمون نشان داده شده است. همان‌گونه که مشاهده می‌شود و بررسی تصاویری که به اشتباه طبقه‌بندی شده‌اند شکل ۹ نیز نشان می‌دهد، طبقه‌بندی اشتباه در میان گلوبول‌های سفید با ساختار هسته نزدیک رخ داده است. در انتها خاطرنشان می‌شود شبیه‌سازی‌های این بخش با استفاده از زبان برنامه‌نویسی پایتون^{۳۳} و کتاب‌خانه‌های یادگیری عمیق تَنسورفلو^{۳۴} (۲۷) و برخی منابع دیگر انجام شده است.

اُتسو^{۳۹} بر تصاویر سیاه و سفید در این پایگاه داده برابر ۰/۷۶ می‌باشد که با وجود سادگی، عملکرد بسیار قابل قبولی دارد، در حالی که استفاده از روش خوشه‌بندی k -means با وزن‌های متغیر می‌تواند در کاربردهای پیچیده‌تر همانند جداسازی قسمت‌های نشان‌دهنده هسته و سیتوپلاسم نیز مورد استفاده قرار بگیرند.



شکل ۶) الف: ضرایب بهترین نمونه یاغته‌شده در انتخاب ضرایب

الگوریتم k -means - محور افقی: شماره‌ی ضریب. محور عمودی:

ضریب. ب: نمودار تابع خطا در تکرارهای الگوریتم ژنتیک جهت بهینه‌سازی

ضرایب تابع فاصله در الگوریتم خوشه‌بندی k -means - محور افقی:

شماره تکرار الگوریتم. محور عمودی: تابع خطا

طبقه‌بندی

به عنوان گام پایانی فرایند شمارش گلوبول‌های سفید، پس از اجرای عملیات بخش‌بندی، تصاویر موجود در پایگاه داده به طور تصادفی به سه قسمت تقسیم می‌نماییم: ۶۸ درصد داده‌ها برای آموزش شبکه عصبی، ۱۵ درصد داده‌ها برای اعتبارسنجی و ۱۷ درصد داده‌ها برای تست عملکرد مدل مورد استفاده، قرار می‌گیرند. سپس آموزش مدل‌های پیشنهادی انجام شده و حاصل دقت و تابع خطا در طول فرایند آموزش روی داده‌های آموزش و داده‌های انتخاب مدل نمایش داده می‌شود. در طول آموزش شبکه عصبی با استفاده از الگوریتم توقف سریع^{۴۰} هنگامی که حاصل تابع خطا بر داده‌های سنجش میانی در دو دوره^{۴۱}

³⁹ Otsu

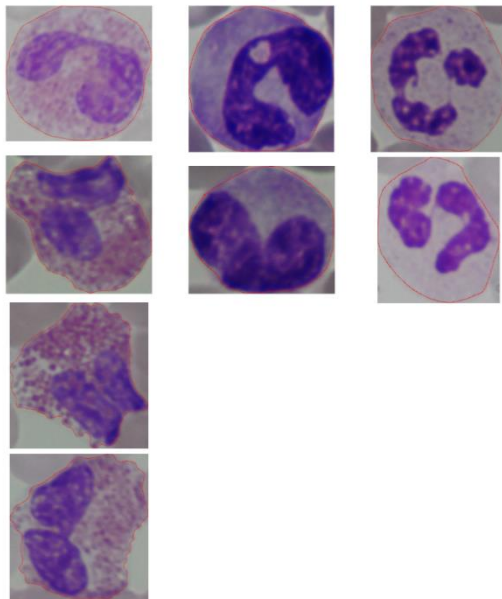
⁴⁰ Early stopping

⁴¹ Epoch

⁴² Confusion matrix

⁴³ Python

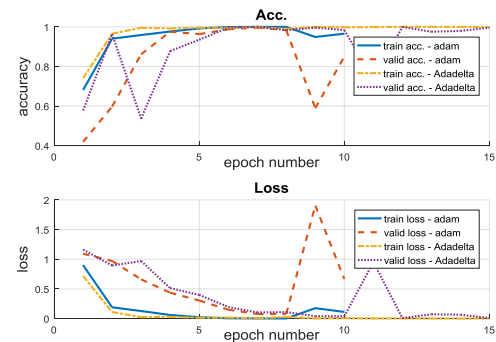
⁴⁴ Tensorflow



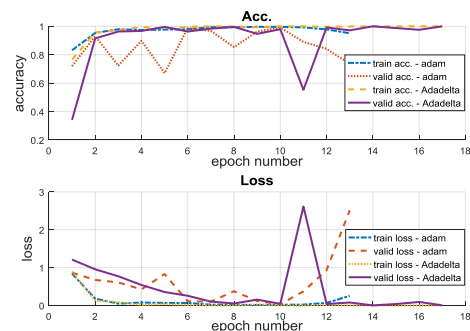
شکل ۹) تصاویر مربوط به گلبول‌هایی که توسط طبقه‌بند به اشتباه طبقه‌بندی شده‌اند.

بحث

با توجه به اهمیت شمارش سریع و قابل اعتماد انواع مختلف گلبول‌های سفید در نمونه خون، در این مقاله قسمت‌های بخش‌بندی و طبقه‌بندی یک سیستم سریع مبتنی بر پردازش تصویر برای شمارش انواع مختلف گلبول‌های سفید در نمونه خون معرفی و پیاده‌سازی شد. در قسمت بخش‌بندی، از الگوریتم خوشه‌بندی k -means با معیار اقلیدسی اصلاح شده برای سنجش فاصله در کنار تبدیلات مورفولوژیک برای انجام اصلاحات نهایی استفاده گردید. این الگوریتم با تعریف معیار مناسب‌تری از فاصله مکانی پیکسل‌ها می‌تواند در آینده در مسائل دشوار بخش‌بندی تصاویر پزشکی مورد استفاده قرار گیرد. همچنین در بخش طبقه‌بندی تصاویر، ساختار ساده‌ای از یک شبکه عصبی کانولوشنی به عنوان طبقه‌بند مورد استفاده قرار گرفت. همان‌طور که انتظار می‌رفت شبکه‌های عصبی کانولوشنی با انتخاب ساختار و داده‌های مناسب، عملکرد رضایت‌بخشی در مسائل بینایی ماشین مرتبط با تصاویر و سیگنال‌های پزشکی از خود نشان می‌دهند.



شکل ۷) نمودار دقت و تابع خطا برای مدل پیشنهادی با دو لایه کانولوشنی با استفاده از روش‌های مختلف بهینه‌سازی در فرایند طبقه‌بندی گلبول‌های سفید. محور افقی: دوره آموزش. محور عمودی: دقت طبقه‌بندی و تابع خطا



شکل ۸) نمودار دقت و تابع خطا برای مدل پیشنهادی با سه لایه کانولوشنی با استفاده از روش‌های مختلف بهینه‌سازی در فرایند طبقه‌بندی گلبول‌های سفید. محور افقی: دوره آموزش. محور عمودی: دقت طبقه‌بندی و تابع خطا

جدول ۲) دقت نهایی مدل‌های پیشنهادی روی داده‌های سنجش در طبقه‌بندی گلبول‌های سفید		
مدل پیشنهادی	دقت	تابع خطا
دو لایه کانولوشنی – adam	۰/۸۳	۰/۷۶
دو لایه کانولوشنی – Adadelta	۰/۹۹	۰/۰۳
سه لایه کانولوشنی – adam	۰/۷۱	۲/۸۷
سه لایه کانولوشنی – Adadelta	۰/۹۹	۰/۰۲

جدول ۳) ماتریس اغتشاش طبقه‌بندی روی داده‌های آزمون				
برچسب‌گذاری صحیح				
نوع گلبول سفید	نوتروفیل	مونوسیت	لنفوسیت	ائوزینوفیل
	۰/۰۰۲۶	۰/۰۰۲۷	۰	۰/۹۹۰۴
	۰	۰	۱	۰
	۰/۰۰۲۶	۰/۹۹۴	۰	۰/۰۰۲۳
نوع گلبول سفید	۰/۹۹۴۷	۰/۰۰۲۷	۰	۰/۰۰۹۳
	۰	۰	۰	۰

نتیجه گیری

در پژوهش‌های پیش رو و با هدف بهبود نتایج حاصل از طبقه‌بندی، با توجه به پیچیدگی نسبی مدل ارائه شده در این بخش، می‌بایست تعداد و گوناگونی تصاویر موجود در پایگاه داده افزایش یابد. در این راستا برای استقلال داده‌ها و افزایش توان طبقه‌بندی داده‌ها در شرایط متفاوت، بهتر است از پایگاه‌های داده فراهم‌شده در مراکز درمانی مختلف استفاده کرد. همچنین در صورت لزوم، می‌توان مدل مورد

استفاده در شبکه عصبی کانولوشنی را در صورت پیچیده‌تر بودن مرز طبقه‌بندی در فضای ویژگی‌ها، با افزایش تعداد لایه‌ها و یا افزایش ابعاد فضاهای خروجی هر لایه قدرت‌مندتر کرد.

تضاد منافع

هیچ گونه تعارض منافع توسط نویسندگان بیان نشده است.

References:

- Libbrecht MW, Noble WS. Machine Learning in Genetics and Genomics. *Nature Reviews* 2015; 16(6): 321-2.
- Hosseini MM, Safdari R, Shahmoradi L, et al. Better Diagnosis of Acute Appendicitis by Using Artificial Intelligence. *Iran South Med J* 2017; 20 (4): 339-48.
- Prince JL, Links JM. Basic Imaging Principles. In: Horton MJ. *Medical Imaging Signals and Systems*. 2nd ed. Upper Saddle River, N.J.: Pearson. 2015. 5-13.
- Gurcan MN, Boucheron LE, Can A, et al. Histopathological Image Analysis: A Review. *IEEE Reviews in Biomedical Engineering* 2009; 2: 147-71.
- Mohapatra S, Patra D, Satpathy S. An Ensemble Classifier System for Early Diagnosis of Acute Lymphoblastic Leukemia in Blood Microscopic Images. *Neural Computing and Applications* 2014; 24(7-8): 1887-1904.
- Sabino DMU, da Fontoura Costa Ldf, Rizzatti G, et al. A Texture Approach to Leukocyte Recognition. *Real-Time Imaging* 2004; 10(4): 205-16.
- Chassery JM, Garbay C. An Iterative Segmentation Method Based on Contextual Color and Shape Criterion. *IEEE Transactions on Pattern Analysis and Machine Intelligence* 1984; 6(6): 794-800.
- Jiang K, Liao QM, Dai SY. A Novel White Blood Cell Segmentation Scheme using Scale-Space Filtering and Watershed Clustering. In *Machine Learning and Cybernetics* 2002; 32(1): 48-53.
- Rezatofighi, SH, Soltanian-Zadeh H. Automatic Recognition of Five Types of White Blood Cells in Peripheral Blood. *Computerized Medical Imaging and Graphics* 2011; 35(4): 333-43
- Theera-Umpon N, Dhompangsa S. Morphological Granulometric Features of Nucleus in Automatic Bone Marrow White Blood Cell Classification. *IEEE Transactions on Information Technology in Biomedicine* 2007; 11(3): 353-9.
- Farjam R, Soltanian-Zadeh H, Jafari Khouzani K, et al. An Image Analysis Approach for Automatic Malignancy Determination of Prostate Pathological Images. *Cytometry B Clin Cytom* 2007; 72(4): 227-40.
- Long X, Cleveland WL, Yao YL. A New Preprocessing Approach for Cell Recognition. *IEEE Transactions on Information Technology in Biomedicine* 2005; 9(3): 407-12.
- Theera-Umpon N, Gader PD. System-Level Training of Neural Networks for Counting White Blood Cells. *IEEE Transactions on Systems, Man, and Cybernetics* 2002; 32(1): 48-53.
- Shitong W, Min W. A New Detection Algorithm (NDA) based on Fuzzy Cellular Neural Networks for White Blood Cell Detection. *IEEE Transactions on Information Technology in Biomedicine* 2006; 10(1): 5-10.
- Kashefpor M, Kafieh R, Jorjandi S, et al. Isfahan MISP Dataset. *Journal of Medical Signals and Sensors* 2017; 7(1): 43-8.
- Sarrafzadeh O, Rabbani H, Talebi A, et al. Selection of the Best Features for Leukocytes Classification in Blood Smear Microscopic Images. In *Proceedings of the SPIE Medical Imaging. International Society for Optics and Photonics, SPIE Medical Imaging 2014: Digital Pathology*, SPIE 9041, SPIE, San Diego, California, USA.

17. Gonzalez RC, Woods RE. Intensity Transformations and Spatial Filtering. In: Horton MJ. Digital Image Processing. 3rd ed. Upper Saddle River, N.J.: Prentice Hall 2008. 152-157.
18. Zhang YJ. A Survey on Evaluation Methods for Image Segmentation. Pattern Recognition 1996; 29(8): 1335-46.
19. Suykens JAK, Vandewalle J. Least Squares Support Vector Machine Classifiers. Neural Processing Letters 1999; 9(3): 293-300.
20. Guo Y, Liu Y, Oerlemans A, et al. Deep Learning for Visual Understanding: A Review. Neurocomputing 2016; 187: 27-48.
21. LeCun Y, Bengio Y, Hinton G. Deep Learning. Nature 2015; 521: 436-44.
22. Ioffe S, Szegedy C. Batch Normalization: Accelerating Deep Network Training by Reducing Internal Covariate Shift. International Conference on Machine Learning 2015; 37: 448-56.
23. Srivastava N, Hinton GE, Krizhevsky A, et al. Dropout: a simple way to prevent neural networks from overfitting. Journal of Machine Learning Research 2014; 15: 1929-58.
24. Schmidhuber J. Deep learning in neural networks: An overview. Neural Networks. 2015; 61: 85-117.
25. Hartigan JA, Wong MA. Algorithm AS 136: A k-means Clustering Algorithm. Journal of the Royal Statistical Society. 1979; 28(1): 100-8.
26. Krizhevsky A, Sutskever I, Hinton GE. ImageNet Classification with Deep Convolutional Neural Networks. Advances in Neural Information Processing Systems 2012; 25.
27. Abadi M, Agarwal A, Barham P, et al. TensorFlow: Large-Scale Machine Learning on Heterogeneous Distributed Systems. arXiv:1603.04467. 2016.

Original Article:

Classification of White Blood Cells Using Convolutional Neural Network

A. Edraki (BS)^{1*}, A. Razminia (PhD)^{2**}

¹ School of Electrical and Computers Engineering, University of Tehran, Tehran, Iran

² Department of Electrical Engineering, School of Engineering, Persian Gulf University, Bushehr, Iran

(Received 23 Aug, 2017

Accepted 21 Oct, 2017)

Abstract

Background: Observing, categorizing and counting various types of white blood cells in a blood sample is one of the most important steps in the treatment of various diseases. This study aimed to develop a fast and reliable system based on processing microscopic images of blood samples for classifying four types of white blood cells.

Materials and Methods: The modified k-means clustering method was used to perform image segmentation. Furthermore, white blood cells were classified using a deep convolutional neural network with the help of data in the MISP database, a free database composed of microscopic blood sample images. Moreover, several regularization techniques such as dropout and image augmentation were applied to prevent overfitting of the network.

Results: The classification accuracy of the neural network was found to be 99%, which is more successful than many earlier studies. In the segmentation section, the cross-reference index was 0.73.

Conclusion: The results of this research show that processing the microscopic images of the blood sample can help develop rapid and reliable systems using different methods of image processing and machine learning.

Key words: Image segmentation, image classification, deep neural networks, microscopic images of blood samples, white blood cell, convolutional neural network

©Iran South Med J. All rights reserved.

Cite this article as: Edraki A, Razminia A. Classification of White Blood Cells Using Convolutional Neural Network. Iran South Med J 2018; 21(1): 65-80

Copyright © 2018 Edraki, et al. This is an open-access article distributed under the terms of the Creative Commons Attribution-noncommercial 4.0 International License which permits copy and redistribute the material just in noncommercial usages, provided the original work is properly cited.

****Address for correspondence:** Department of Electrical Engineering, School of Engineering, Persian Gulf University, Bushehr, Iran, Email: razminia@pgu.ac.ir

*ORCID: 0000-0002-0843-5522

Website: <http://bpums.ac.ir>
Journal Address: <http://ismj.bpums.ac.ir>