

توسعه مدلی مبتنی بر تشدید گرادیان در شبکه‌های کانولوشنی عمیق به منظور شناسایی اهداف در تصاویر سنجش از دوری

نیما فرهادی^{۱*}، عباس کیانی^۲، حمید عبادی^۳

^۱ دانشجوی کارشناسی ارشد فتوگرامتری و سنجش از دور - دانشکده مهندسی نقشه برداری - دانشگاه صنعتی خواجه

نصیرالدین طوسی

farhadinima75@email.kntu.ac.ir

^۲ استادیار دانشکده عمران - دانشگاه صنعتی نوشیروانی بابل

a.kiani@nit.ac.ir

^۳ استاد دانشکده مهندسی نقشه برداری - دانشگاه صنعتی خواجه نصیرالدین طوسی

ebadi@kntu.ac.ir

(تاریخ دریافت مرداد ۱۳۹۸، تاریخ تصویب اردیبهشت ۱۴۰۰)

چکیده

پیشرفت‌های صورت گرفته در فناوری تصویربرداری ماهواره‌ای امکان تهیه اطلاعات متنوع برای شناسایی اهداف را فراهم می‌کند. چنین اطلاعاتی فرآیند تفسیر تصاویر سنجش از دوری نوری را تسهیل می‌بخشد. نوع خاصی از این تفاسیر به فعالیت‌های مربوط به شناسایی اهداف ختم می‌شود که امروزه اکثر تحقیقات انجام شده در این حوزه با استفاده از شبکه‌های عصبی و تکنیک‌های یادگیری عمیق صورت می‌گیرد. نحوه طراحی شبکه عصبی کانولوشن مورد استفاده، در دقت شناسایی نقش بسزایی دارد. تحقیقات اخیر در زمینه یادگیری عمیق و شبکه‌های کانولوشن نشان می‌دهد که عمیق تر کردن این شبکه‌ها باعث افزایش دقت آن‌ها می‌شود؛ اما گاهی بیش از حد عمیق تر شدن باعث به وجود آمدن مشکلاتی از جمله بالا رفتن پارامترهای آموزشی، محو شدن گرادیان آموزشی، بلااستفاده ماندن بسیاری از ویژگی‌های تولید شده و... می‌شود که در پی آن کاهش دقت در شناسایی اهداف مورد نظر را خواهد داشت. به این منظور در این تحقیق روشی توسعه داده شده است که در آن سعی گردید با حفظ ویژگی‌های تولید شده توسط لایه‌های کانولوشن و انتقال آن‌ها به لایه‌های بعدی، بر این مشکل غلبه گردد. این نوع ارتباط بین لایه‌ها، اجازه عمیق تر کردن شبکه‌های کانولوشنی با افت گرادیان کمتر را می‌دهد. معماری ارائه شده علاوه بر کم رنگ کردن مشکل ناپدید شدن گرادیان، باعث می‌شود تعداد پارامترها و همچنین مدت زمان مورد نیاز برای آموزش یک مدل یادگیری عمیق کاهش یابد. بدین منظور در ابتدا با استفاده از تصاویر سنجش از دوری، مجموعه‌ای از داده‌های آموزشی آماده و پس از پردازش‌های اولیه، عوارض هدف برچسب گذاری شده است. سپس روش پیشنهاد شده را به عنوان استخراج گر ویژگی مدل Faster R-CNN تعریف کرده و بر روی داده‌های آموزشی، آموزش داده می‌شود. جهت ارزیابی روش پیشنهادی نیز، بخشی از فرودگاه بین‌المللی پکن چین به عنوان مطالعه موردی اول و بخشی از فرودگاه بین‌المللی امام خمینی (ره) به عنوان منطقه مورد مطالعه دوم انتخاب شده است و مقادیر معیار F1-Measure برای هر دو منطقه به ترتیب برابر ۹۷/۹ و ۹۳/۷ می‌باشد. در نهایت نتایج حاصله از اعمال مدل پیشنهادی، با مدل‌های مختلف شبکه مطرح موجود، مقایسه شده است. نتایج به دست آمده، دلالت بر قابل اعتماد بودن و مؤثر بودن روش ارائه شده دارند.

واژگان کلیدی: یادگیری عمیق، شبکه‌های کانولوشن، شناسایی اهداف سنجش از دوری، استخراج گر ویژگی

* نویسنده رابط

۱- مقدمه

شناسایی اهداف در تصاویر ماهواره‌ای یکی از موضوعات تحقیقاتی است که در سال‌های اخیر مورد توجه بسیاری از محققان سنجش از دوری قرار گرفته است؛ از این رو، در این حوزه، روش‌های متعددی توسعه داده شده است که روش‌های مربوط به یادگیری ماشین جزء قابل اعتمادترین آن‌ها هستند. در زمینه یادگیری ماشین، شبکه‌های عصبی کانولوشن یک روش پرکاربرد طی سالیان گذشته می‌باشد که در قرن بیستم میلادی برای اولین بار معرفی شده است [۱]. این شبکه‌ها به عنوان استخراج گر ویژگی از تصاویر رقومی در مدل‌هایی همچون Faster R-CNN [۲]، SSD [۳] و YOLO [۴] برای شناسایی اهداف مورد استفاده قرار می‌گیرند که در میان آن‌ها مدل Faster R-CNN با دقت و ضریب اطمینان بالاتری عوارض موجود در تصاویر را شناسایی می‌کند. در سال‌های اخیر تحقیقات فراوانی در زمینه بهبود Faster R-CNN در تصاویر طبیعی صورت گرفته است. به عنوان مثال Chen و Gupta [۵] از NMS^۱ در انتخاب مرزهای مورد نظر از کلاسه‌بندی مرزها استفاده کردند. با این حال در این روش امکان انتخاب عارضه اشتباه و یا عوارضی که در پس زمینه قرار دارند، به عنوان عارضه مورد نظر وجود دارد [۶]. Li و همکاران [۷] با اعمال عملگرهای کانولوشن کوچک تر میزان محاسبات روش Faster R-CNN تا حد قابل توجهی کاهش دادند. دسته‌ای دیگر از محققان شناسایی اهداف را در زمینه‌های تخصصی‌تری، مانند شناسایی هواپیما در تصاویر سنجش از دوری، انجام داده‌اند. به طور مثال Farhadi و همکاران [۸] در سال ۲۰۱۹ از شبکه‌های پیش آموخته و ترکیب سلسله مراتبی از چندین شبکه Faster R-CNN برای شناسایی هواپیما استفاده کرده‌اند. Wu و همکاران [۹] برای شناسایی هواپیما از یک روش مبتنی بر گرادیان باینری و نرمال شده^۵ و شبکه‌های عصبی کانولوشن استفاده کردند. Zhang و همکاران [۱۰] به منظور شناسایی هواپیما در تصاویر سنجش از دوری، یک چهارچوب قابل آموزش نیمه نظارتی را با استفاده از شبکه‌های کانولوشن دوگانه، ارائه کردند. Zeng و همکاران [۱۱] در سال ۲۰۱۹ با

تلفیق مفاهیم یادگیری عمیق و تحلیل‌های مکانی یک روش سلسله مراتبی را به منظور شناسایی هواپیما ارائه کردند. Xu و همکاران [۱۲] از ترکیب ویژگی‌های چندلایه در شبکه‌های کانولوشنی تمام متصل باهدف شناسایی هواپیما بهره گرفتند؛ اما با این حال بیشتر این روش‌ها برای عوارض بزرگ مقیاس و با وضوح بالا مناسب هستند، اما در تصاویر با پس زمینه پیچیده و عوارض کوچک، دقت پایین تری را حاصل می‌کنند [۱۳]. از طرفی دیگر، شبکه‌های عصبی کانولوشن بکار رفته در این روش‌ها، ویژگی‌های سطح پایین استخراج شده از داده‌های ورودی را فراگرفته و ویژگی‌های سطح بالا را تولید می‌کنند و با تلفیق این ویژگی‌ها، الگوهای مناسبی را پدید می‌آورند که از این جهت گزینه مناسبی برای شناسایی اهداف بصری به حساب می‌آیند [۱۴]. عمق این شبکه‌ها تأثیر بسیار زیادی بر روی نحوه ترکیب ویژگی‌های سطح پایین دارد؛ شبکه با عمق بیشتر دارای لایه‌ها و سطوح ویژگی بیشتر است. پیشرفت‌های اخیر در زمینه قدرت پردازشی سخت‌افزارها، امکان آموزش شبکه‌های عمیق مانند VGG ۱۹ لایه [۱۵] و ResNet^۶ ۱۵۲ لایه [۱۶] را به محققان می‌دهند.

شبکه‌های عمیق تر از دقت طبقه‌بندی بالاتری نسبت به شبکه‌های کم عمق تر برخوردارند؛ اما چنانچه عمق شبکه‌ها به صورت قابل توجهی زیاد شود، مشکل جدیدی به نام محو شدن گرادیان به وجود می‌آید [۱۷]. محو شدن گرادیان زمانی رخ می‌دهد که تعداد لایه‌های کانولوشن درون شبکه افزایش یابد؛ به تعبیری دیگر با افزایش لایه‌های کانولوشن درون شبکه اطلاعات ورودی به لایه‌های بالاتر، کاهش می‌یابد و گاهی در انتهای شبکه به صفر میل می‌کند [۱۸]. مقالات زیادی که اخیراً چاپ شده‌اند به این موضوع اشاره دارند [۱۴]، [۱۹]، [۲۰]. همگی این مقالات با ارائه روش‌های مختلف سعی در حل مشکل گرادیان دارند. روش‌های موجود توپولوژی و پروسه آموزشی مختلفی دارند اما یک مفهوم اصلی را با خود به یدک می‌کشند و آن هم ایجاد یک مسیر میانبر برای اتصال لایه‌های قبل به لایه‌های بعد است که باهدف جلوگیری از محو شدن گرادیان دنبال می‌شود.

تحقیق در زمینه معماری‌های شبکه، قسمت جداناپذیر از تاریخچه شبکه‌های عصبی از بدو تولد تا به حال بوده است. تفاوت اصلی بین شبکه‌های مدرن و سنتی، افزایش تعداد لایه‌های مخفی بکار رفته در معماری و نحوه

^۱ Faster Region-based Convolutional Neural Networks^۲ Single Shot multibox Detector^۳ You Look Only Once^۴ Non-Maximal Suppression^۵ Binarized Normed Gradients (BING)^۶ RESidual NETWORKs

سال ۲۰۱۷ Huang و همکاران [۱۸] ایده‌ی جالبی را در طراحی شبکه‌های عصبی کانولوشن به کار بردند که باعث کاهش تعداد پارامترهای موردنیاز در آموزش شبکه می‌شد. آن‌ها تمام لایه‌های کانولوشن را در بلوک‌هایی به نام بلوک چگال^۳ به هم متصل کردند. با استفاده از این ایده مشکل شارش گرادیان تا حدودی حل شده و علاوه بر این، استفاده دوباره از ویژگی‌های استخراج‌شده‌ی لایه‌های قبل در لایه‌های بعدی، باعث تفکیک‌پذیری بهتر ویژگی‌ها توسط شناساگر می‌شود.

با توجه به تحقیقات متعددی که در زمینه‌ی بهبود شبکه‌های کانولوشنی صورت گرفته است، اما همچنان این شبکه‌ها با مشکلاتی همچون هزینه پردازشی بالا و محو شدن گرادیان روبه‌رو هستند. از طرفی، به دلیل این‌که شبکه‌های CNN به‌عنوان پایه و اساس روش‌های شناسایی عارضه ایفای نقش می‌کنند، بهبود عملکرد و هزینه پردازشی موردنیاز این دسته از شبکه‌ها، باعث بهبود کارایی روش‌های شناسایی عارضه می‌گردد. در همین راستا، در تحقیق پیش رو، هدف، ارائه یک معماری جدید با نحوه اتصال متفاوت لایه‌های کانولوشن است که محو شدن گرادیان در آن به حداقل مقدار خود می‌رسد. در این معماری اطلاعات هر لایه مهم و در طول آموزش شبکه مؤثر هستند. خروجی لایه آخر هر سلول در کنار دیگر خروجی‌ها قرار گرفته و تا انتهای شبکه حفظ می‌شود.

۲- مبانی نظری تحقیق

یادگیری عمیق زیرمجموعه‌ای از الگوریتم‌های یادگیری ماشین می‌باشد که هدف آن کشف چندین سطح از بازنمودهای توزیع‌شده از داده‌ی ورودی است؛ درواقع یادگیری عمیق مفاهیم انتزاعی سطح بالای استخراج‌شده از داده‌های ورودی را در لایه‌های سطح پایین پشت سر هم مدل می‌کند [۲۷]. مکانیسم استخراج مفاهیم استفاده‌شده در مجموعه‌ای به نام شبکه عصبی انجام می‌شود. نوعی از شبکه‌های عصبی، شبکه‌های عصبی کانولوشن است که مجموعه‌ای از فیلترها را در خود جای داده است. بسته به ابعاد این فیلترها، پارامترهایی تعریف می‌شود که نیازمند بهینه‌سازی هستند. این پارامترها در طول پروسه آموزش تغییر می‌کنند تا بهترین مفاهیم را از تصاویر آموزشی

اتصال آن‌ها به هم است. اولین مفاهیم به وجود آمده از شبکه‌های عصبی را می‌توان به دهه ۹۰ میلادی ارجاع داد. در سال ۱۹۹۸ LeCun و همکاران [۱] شبکه‌های عصبی کانولوشن را برای اولین بار به‌منظور دسته‌بندی اعداد مندرج در کد پستی معرفی کردند و اسم شبکه خود را LeNet گذاشتند. از شبکه عصبی آن‌ها در سرویس‌های پستی استفاده فراوانی شد اما اعداد عارضه‌های بسیار ساده‌ای هستند و شبکه LeNet قابلیت دسته‌بندی داده‌های پیچیده‌تر مانند تصاویر بصری دوبعدی را نداشت [۲۱]. به‌منظور پوشش این نقص Krizhevsky و همکاران [۲۲] در سال ۲۰۱۲ با تغییراتی اساسی در شبکه عصبی کانولوشن توانستند خطای روش‌های قبل در دسته‌بندی تصاویر بر روی مجموعه داده ImageNet [۲۳] را به میزان ۱۰ درصد کاهش دهند. Simonyan و Zisserman [۱۵] در سال ۲۰۱۴ شبکه‌ای عمیق‌تر با نام VGG را معرفی کردند. استراتژی به‌کاررفته در این معماری استفاده از فیلتر با ابعاد کوچک و تعداد لایه‌های بیشتر بوده است. در همان سال، Szegedy و همکاران [۲۴، ۲۵] بلوک‌های مفهومی^۱ را در معماری شبکه‌ها به کار بردند و شبکه GoogLeNet را با استفاده از چندین مدل مفهومی پشت سرهم تشکیل دادند. اساس طراحی این بلوک‌ها بر پایه راندمان بالا در محاسبات می‌باشد. آن‌ها پارامترهای به‌کاررفته در کل شبکه را به ۵ میلیون کاهش دادند. مادامی‌که این شبکه‌ها رفته‌رفته عمیق و عمیق‌تر می‌شدند، اندازه گرادیان نیز روبه‌زوال و صفر شدن میل می‌کرد. سرانجام در سال ۲۰۱۵، He و همکاران [۱۶] روشی را در طراحی شبکه‌ها ابداع کردند که اجازه استفاده از ۱۵۲ لایه کانولوشنی را با انتشار گرادیان مناسب، می‌دهد. آن‌ها بلوک‌های باقی‌مانده^۲ را که شامل دولایه کانولوشن بود به وجود آوردند. در این بلوک‌ها مقدار ورودی با یک عمل جمع ساده به خروجی بلوک اضافه می‌گردد؛ که همین امر باعث بهتر جریان یافتن گرادیان در کل شبکه می‌شود. Larsson و همکاران [۲۶] در سال ۲۰۱۶ با ترکیب خطی از چندین لایه موازی بلوک‌های Fractal را به وجود آوردند که با تکرار این بلوک‌ها می‌توان عمق شبکه را بدون استفاده از روش باقی‌مانده‌ها زیاد کرد. مفهوم اصلی بلوک‌های Fractal، انتقال مؤثر اطلاعات از لایه‌های اولیه به لایه‌های پایین‌تر از طریق میانبرهای کوتاه است. سرانجام در

^۱ Inception Blocks

^۲ Residual Blocks

^۳ Dense Block

استخراج کنند. سپس مفاهیم از لایه‌های اول به لایه‌های بعدی منتقل می‌شود تا مفاهیم کلی تشکیل شوند.

شناسایی اهداف یکی از مسائل مهم و پرمخاطب در زمینه‌ی یادگیری عمیق است [۲۸, ۲۹] و کاربرد فراوانی در علوم مختلف نقشه‌برداری دارد. روش‌های متنوعی برای این موضوع شده است که از جمله مهم‌ترین آن‌ها می‌توان به Faster R-CNN، SDD و YOLO اشاره کرد که همه‌ی این روش‌ها از شبکه‌های عصبی کانولوشنی به‌عنوان استخراج ویژگی استفاده می‌کنند. در این مقاله از روش Faster R-CNN برای شناسایی اهداف استفاده شده است که نخستین بار توسط Girshick در سال ۲۰۱۵ توسعه یافته است.

۲-۱- شناساگر مبتنی بر مرز

در زمینه بینایی ماشین یکی از مسائل پایه و بحث‌برانگیز شناسایی اهداف است. در سال ۲۰۱۶، Girshick و همکاران [۳۰] شبکه کانولوشنی را ارائه کردند که با استفاده از مرزهای مشخص شده در داده‌های آموزشی، اهداف موردنظر را شناسایی می‌کند. شناساگر ارائه شده توسط آن‌ها میانگین دقت حاصله از روش‌های قبل را تا ۵۰ درصد بهبود دادند [۳۱]. این روش را می‌توان ترکیبی از چندین مکانیسم مختلف شامل یک شبکه عصبی کانولوشن برای استخراج ویژگی، یک ماشین بردار پشتیبان برای دسته‌بندی و یک مکانیسم ارائه اهداف دانست [۳۲]. R-CNN با نتایج خوب خود توانست بسیاری از رقیبان خود را پشت سر بگذارد اما به دلیل افزونگی محاسبات ناشی از همپوشانی مرزهای ارائه شده، عملکردی کند را از خود به‌جای می‌گذارد. به‌منظور افزایش سرعت و دقت در شناسایی، Girshick و همکاران نسخه بروز شده‌ی R-CNN، به نام Fast R-CNN را ارائه کردند [۳۳]. این مدل در گام اول، عملیات برداشت ویژگی را قبل از انتخاب مرز عوارض انجام می‌دهد، به این معنی که یک شبکه کانولوشنی را به‌جای اعمال دنباله‌ای از شبکه کانولوشنی، بر روی تصویر ورودی اعمال می‌کند. در این مدل، گام دوم، جایگزین کردن یک لایه Softmax به‌جای مدل SVM مورد استفاده در R-CNN است؛ با این کار از تولید یک مدل جدید SVM نیز جلوگیری می‌شود. یکی از وجوه اشتراک بین R-CNN و Fast R-CNN وجود الگوریتم جستجوی انتخابی است این الگوریتم ذاتاً نیاز به عملیات محاسباتی بالایی دارد و باعث کندی شبکه می

شود [۳۲]. Ren و همکاران [۲] در سال ۲۰۱۵ مدل Faster R-CNN را ارائه کردند. آن‌ها برای بهبود سرعت یادگیری، یک شبکه RPN طراحی و به‌جای الگوریتم جستجوی انتخابی استفاده کردند. نحوه عملکرد RPN به‌گونه‌ای است که پس از اعمال شبکه عصبی به تصویر ورودی، در لایه آخر همان شبکه یک ماسک با ابعاد n در n (معمولاً n برابر ۳) به نقشه ویژگی تولید شده توسط شبکه کانولوشنی، اعمال می‌شود و ابعاد آن‌ها را کاهش می‌دهد. سپس ماسک اعمال شده، چندین مرز امکان‌پذیر را بر اساس محدوده‌های پیش‌فرض تولید می‌کند و درنهایت نیز مرزهایی که شامل عارضه موردنظر هستند، انتخاب و مختصاتی که توسط محدوده‌های اطراف عارضه مشخص شده است، ارائه می‌گردد. مدل Faster R-CNN، مقادیر به‌دست‌آمده از RPN را جایگزین مقادیر جستجوی انتخابی کرده و همین موضوع باعث افزایش کارایی و سرعت شناسایی عارضه مدل مذکور می‌گردد [۲]. دقت و کارایی قابل قبول این مدل، دلیل اصلی انتخاب آن به‌عنوان روش شناسایی عارضه می‌باشد که در این مقاله استفاده شده است. یکی از وجوه اشتراک بین هر سه مدل، شبکه‌های عصبی کانولوشن به‌عنوان استخراج‌گر ویژگی هستند که نحوه طراحی آن در سال‌های اخیر موردتوجه بسیاری از محققان قرار گرفته است.

۲-۲- طراحی شبکه کانولوشن

فرض کنید X یک تصویر آموزشی گذرنده از شبکه کانولوشن با L لایه بوده که هر کدام از این لایه‌ها یک انتقال غیرخطی $H_L(X)$ را اعمال می‌کنند. انتقال غیرخطی $H_L(X)$ می‌تواند ترکیبی از لایه‌های کانولوشن یا عملگرهایی مانند Batch Normalization [۳۴]، Rectified Linear Unit (ReLU) [۳۵] یا Pooling [۱] باشد. خروجی هر لایه $L - 1$ ام نیز، X_L در نظر گرفته می‌شود.

$$X_L = H_L(X_{L-1}) \quad (1)$$

رابطه (۱) نحوه جریان اطلاعات در لایه‌های شبکه‌های سنتی را نمایش می‌دهد که هر لایه ورودی خود را از لایه قبل دریافت می‌کند. شبکه ResNet با ایجاد یک میانبر بین ورودی هر لایه و خروجی آن، انتقال غیرخطی را به نحوی دور می‌زند که منجر به انتقال بهتر اطلاعات می‌شود [۱۸].

۲ استفاده شده است تا اندازه مکانی نقش‌های ویژگی^۲ را نصف کند. جدول (۱) تفاوت بین دو معماری ResNet و معماری پیشنهادی را نمایش می‌دهد. اعداد داخل براکت نشان‌دهنده ابعاد و تعداد لایه‌های مورد استفاده در معماری می‌باشد. همچنین متغیر N معرف تعداد فیلترهای ابتدایی لایه‌های کانولوشنی است. تعداد پارامترهای استفاده شده در شبکه ResNet تقریباً $2/5$ برابر شبکه پیشنهادی است. عمق شبکه پیشنهادی نیز با در نظر گرفتن تعداد سلول‌های استفاده شده تعیین می‌گردد. هر سلول دارای دو لایه کانولوشن است. با فرض داشتن ۴ بلوک، n سلول در شبکه، یک لایه کانولوشن در ابتدا و یک لایه FC^۳، شبکه شامل $2n + 2$ لایه می‌گردد. جزییات بیشتر در ادامه مورد بررسی قرار خواهند گرفت.

جدول ۱- ترتیب قرارگیری لایه‌ها در معماری پیشنهادی

مرحله	خروجی	ResNet-50	معماری پیشنهادی ۵۰ لایه $G = 16$
کانولوشن اول	256×256	Stride $2, 64, 7 \times 7$	
لایه Pooling	128×128	Max pool, 3×3 , Stride ۲	
بلوک اول	128×128	$\begin{bmatrix} 1 \times 1, 64 \\ 3 \times 3, 64 \\ 1 \times 1, 256 \end{bmatrix} \times 3$	$\begin{bmatrix} 1 \times 1, N+nG \\ 3 \times 3, N+nG \end{bmatrix} \times 4$
لایه Pooling	64×64	----	Average pooling 2×2 , Stride ۲
بلوک دوم	64×64	$\begin{bmatrix} 1 \times 1, 128 \\ 3 \times 3, 128 \\ 1 \times 1, 512 \end{bmatrix} \times 4$	$\begin{bmatrix} 1 \times 1, N+nG \\ 3 \times 3, N+nG \end{bmatrix} \times 6$
لایه Pooling	32×32	----	Average pooling 2×2 , Stride ۲
بلوک سوم	32×32	$\begin{bmatrix} 1 \times 1, 256 \\ 3 \times 3, 256 \\ 1 \times 1, 1024 \end{bmatrix} \times 6$	$\begin{bmatrix} 1 \times 1, N+nG \\ 3 \times 3, N+nG \end{bmatrix} \times 10$
لایه Pooling	16×16	----	Average pooling 2×2 , Stride ۲
بلوک چهارم	16×16	$\begin{bmatrix} 1 \times 1, 512 \\ 3 \times 3, 512 \\ 1 \times 1, 2048 \end{bmatrix} \times 3$	$\begin{bmatrix} 1 \times 1, N+nG \\ 3 \times 3, N+nG \end{bmatrix} \times 4$
----	1×1	Global Average Pool, ۲d-FC, Softmax	
تعداد پارامتر		$25/5 \times 10^6$	$10/6 \times 10^6$

^۲ Feature Maps

^۳ Fully Connected layer

$$X_L = H_L(X_{L-1}) + X_{L-1} \quad (2)$$

رابطه (۲)، معادله به کاررفته در شبکه ResNet می باشد. مزیت اصلی این تابع، انتقال مناسب گرادیان از لایه‌های ابتدایی به لایه‌های بعدی است؛ اما با این حال ارتباط بین ورودی و خروجی تابع با استفاده از یک عمل جمع ساده انجام می‌شود که ممکن است جریان اطلاعات در شبکه را مختل کند [۱۸].

۳- روش پیشنهادی

معماری پیشنهادی ترکیبی از چند بلوک است که پشت سر هم قرار گرفته‌اند. بلوک‌ها ورودی خود را از بلوک قبل دریافت می‌کنند و محاسبات مربوطه را بر روی آن‌ها انجام می‌دهند. هر بلوک شامل چند سلول می‌شود که این سلول‌ها دارای دو لایه کانولوشن هستند. به منظور استفاده بهینه از تمام ویژگی‌های تصاویر آموزشی، فیلترهای بکار رفته در لایه‌های کانولوشن دارای ماسک‌هایی با اندازه‌های 1×1 و 3×3 هستند که خروجی لایه 3×3 در لایه ترکیب‌کننده^۱ با اطلاعات لایه‌های قبلی ادغام می‌شود. معماری شبکه‌های سنتی به گونه‌ای است که هر لایه کانولوشن ورودی خود را از خروجی لایه قبل دریافت می‌کند. این فرآیند باعث ایجاد ویژگی جدید می‌شود اما در کنار آن اطلاعات مورد نیاز برای آموزش شبکه در فرآیند آموزش از بین می‌رود. معماری هر سلول در شبکه پیشنهادی به گونه‌ای است که تمام ویژگی‌های استخراج شده را در خود حفظ و نگهداری می‌کند تا در سلول‌های بعدی مورد استفاده قرار گیرد. رابطه (۳)، نحوه ورود اطلاعات به هر لایه کانولوشن در معماری پیشنهادی را مشخص می‌کند.

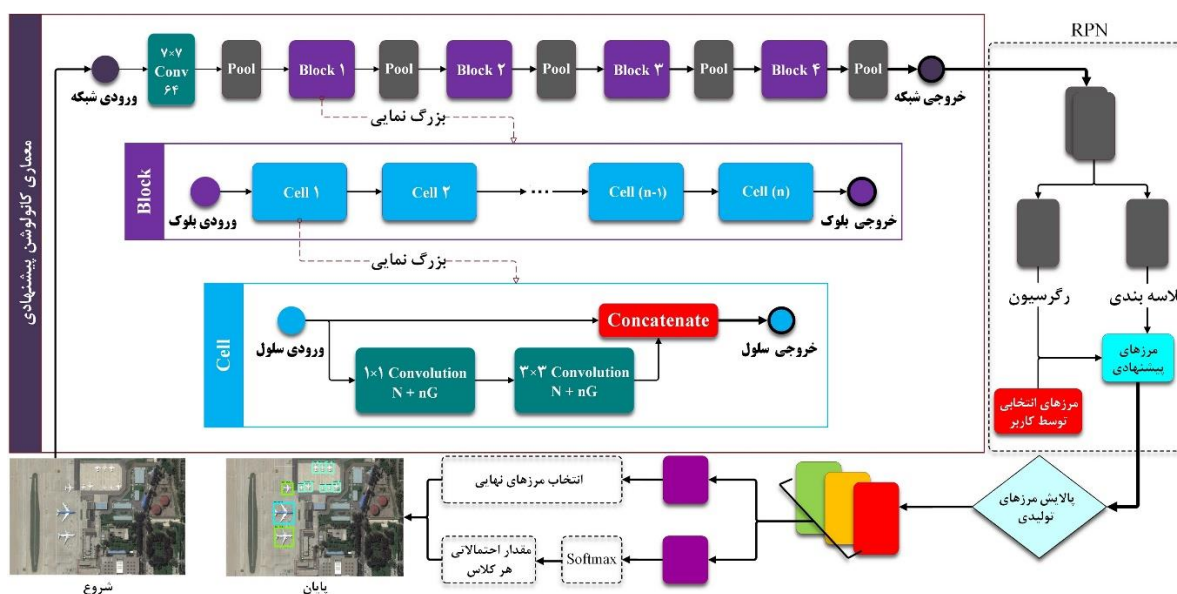
$$X_L = H_L(X, X_1, \dots, X_{L-1}) \quad (3)$$

طبق رابطه فوق، ورودی هر لایه ترکیبی از تمام خروجی لایه‌های قبلی است. همانند معماری ResNet، بلوک‌های در نظر گرفته شده برای شبکه پیشنهادی، ۴ بلوک بوده است که به منظور کاهش محاسبات و بهبود فرایند یادگیری، میان هر بلوک، از یک لایه Average Pooling با ماسک دو پیکسل در دو پیکسل و Stride برابر

^۱ Concatenation Layer

تعداد فیلترهای هر لایه کانولوشن استفاده شده در شبکه از یک عمل جمع ساده پیروی می کند که به عنوان یک پارامتر آزاد قابل تغییر و بهینه شدن است. تعداد فیلترهای هر بلوک در معماری ResNet ثابت هستند ولی نسبت به بلوک قبل خود دو برابرند درحالی که در معماری پیشنهادی روند دیگری برای افزایش فیلترها نسبت به عمیق شدن شبکه در پیش گرفته شده است. در این شبکه

تعداد فیلترهای هر لایه، حاصل جمع متغیر N با پارامتر نرخ رشد (G) است که در پایان هر سلول به آن اضافه می گردد. شکل (۱)، فلوچارت روش Faster R-CNN را نمایش می دهد که معماری کانولوشن ارائه شده در آن، وظیفه آموزش و استخراج ویژگی های سطح بالا را در یک فرآیند سلسله مراتبی، بر عهده دارد.



شکل ۱- فلوچارت روش Faster R-CNN به همراه معماری کانولوشن ارائه شده به عنوان استخراج گر ویژگی

۴- پیاده سازی و ارزیابی نتایج

۴-۱- آماده سازی داده ها

داده هایی که برای آموزش شبکه Faster R-CNN استفاده شده است، از فرودگاه های مختلف در سرتاسر دنیا اخذ شده است. فرودگاه هایی مانند فرودگاه بین المللی امام خمینی (ره)، دبی، مهرآباد و فرودگاه بین المللی پکن از جمله مهم ترین آن ها هستند. مجموعه داده جمع آوری شده شامل ۳۲۰ تصویر آموزشی و ۷۵ تصویر جهت ارزیابی معماری ارائه شده است که سه کلاس هواپیمای دو موتور، چهار موتور و هواپیما با موتور در عقب را شامل می شوند.

شبکه های مبتنی بر مرز جهت یادگیری عوارض موردنظر نیازمند مکان دقیق عارضه همراه با کلاس مرتبط با آن عارضه هستند. لازم است در داده های آموزشی، محل قرارگیری عوارض در تصویر با استفاده از یک چارچوب خاص مشخص شود تا شبکه عصبی کانولوشن به کاررفته در روش های شناسایی اهداف، بتواند عارضه هدف را از دیگر

عوارض موجود در تصویر آزمایشی تشخیص دهد. به طور معمول مشخص کردن عوارض در تصاویر آموزشی با استفاده از ایجاد یک چهارضلعی دور آن انجام می شود. هر یک از رئوس بیانگر یک x و y به عنوان مختصات پیکسلی موقعیت عارضه هستند. به هر یک از آن ها یک برچسب که نمایانگر کلاس عارضه ی موردنظر است داده می شود تا در پروسه دسته بندی عوارض مورد استفاده قرار گیرد.

در پاره ای از تحقیقات [۳۶-۳۸] برای کاهش هزینه محاسباتی، تصویر مورد ارزیابی را به چندین بخش تقسیم کرده و پردازش های موردنظر را بر روی این قطعات انجام می دهند. همچنین این روش باعث کاهش خطا در شناسایی اهداف کوچک می شود [۳۶]. تعداد تقسیمات صورت گرفته در این مقاله با توجه به اندازه تصویر سنجش از دوری، میزان همپوشانی^۱ و اندازه هر بخش محاسبه شده است.

^۱ Overlap

همپوشانی، شناسایی هواپیماهایی است که بین مرز دو بلوک قرار دارند [۳۶]. در این مقاله برای منطقه مطالعاتی اول، تصویر ماهواره‌ای رزولوشن بالا با ابعاد 11000×4800 پیکسل مورد ارزیابی قرار گرفته است. برای هر بخش از این تصویر، ابعاد و همپوشانی به ترتیب برابر 1100×1600 و ۲۰ پیکسل به صورت دلخواه منظور شده است که با توجه به روابط (۴) و (۵)، مقادیر ۱۰ و ۳ به ترتیب برای m و n حاصل می‌گردد. شکل (۲) منطقه مطالعاتی اول را به همراه بخش‌های آن، نمایش می‌دهد.

$$m = \left\lceil \frac{\text{همپوشانی} - \text{طول تصویر منطقه}}{\text{همپوشانی} - \text{طول بخش}} \right\rceil \quad (4)$$

$$n = \left\lceil \frac{\text{همپوشانی} - \text{عرض تصویر منطقه}}{\text{همپوشانی} - \text{عرض بخش}} \right\rceil \quad (5)$$

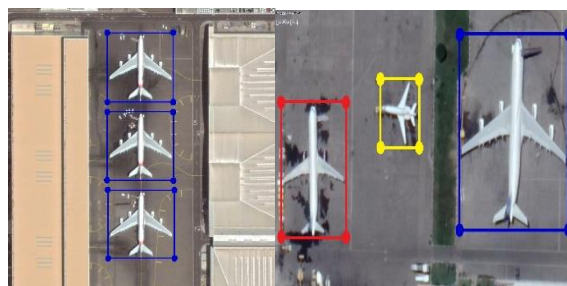
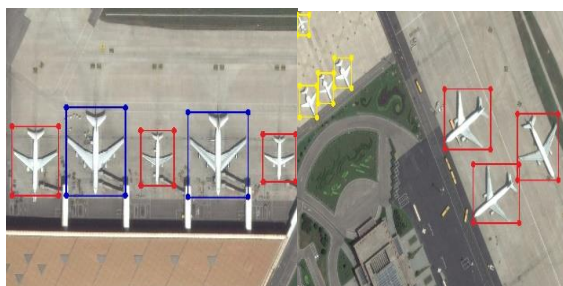
در رابطه (۴)، m تعداد بخش‌های موجود در طول و در رابطه (۵)، n تعداد بخش‌های موجود در عرض تصویر منطقه مطالعاتی می‌باشد. در روابط بالا، عبارت همپوشانی، فاصله پوششی بین بخش‌ها است. منطق و دلیل تعیین میزان



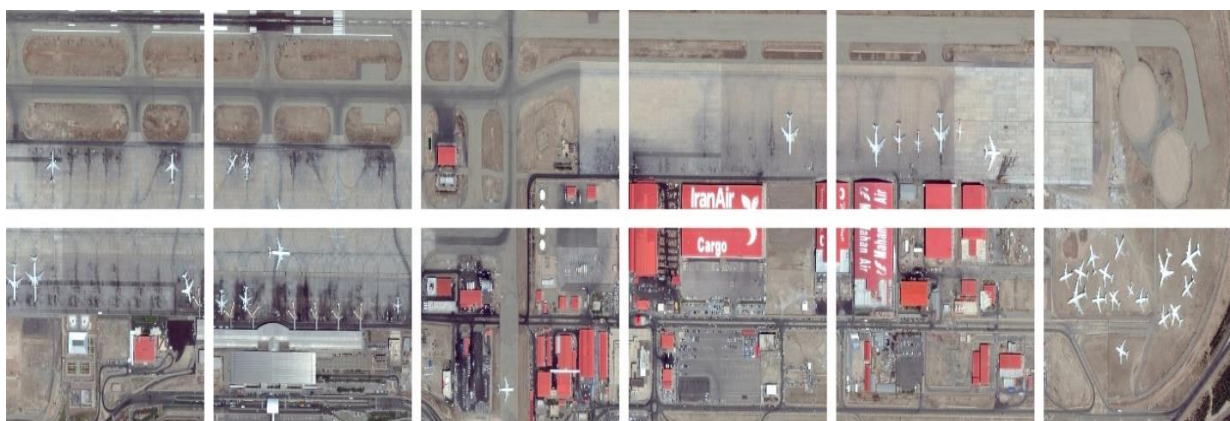
شکل ۲- تصویر منطقه مورد مطالعه اول پس از عملیات بخش‌بندی

شکل (۳) نمونه‌ای از چند تصویر آموزشی و آزمایشی را به تصویر می‌کشد که هرکدام از هواپیماها در دسته‌های خود مشخص گردیده‌اند. قدرت تفکیک مکانی تصاویر ماهواره‌ای موجود در حدود ۱۵ سانتیمتر است. عوامل مؤثر بر کیفیت عکس مانند روشنایی، قدرت تفکیک و... به‌گونه‌ای انتخاب و کنترل شده‌اند که اکثر حالت‌های رؤیت شدن هواپیما را دربر می‌گیرد. منطقه مطالعاتی اول انتخاب‌شده، فرودگاه بین‌المللی پکن است که در این

فرودگاه از سه نوع کلاس استفاده‌شده، وجود دارد. شکل (۴) تصویر ماهواره‌ای فرودگاه امام خمینی (ره) را نمایش می‌دهد که این تصویر نیز به‌عنوان تصویر مطالعاتی دوم انتخاب و پردازش شده است. ابعاد کلی این تصویر شامل 4000×3000 پیکسل می‌باشد. با توجه به روابط (۴) و (۵)، مقادیر ۲، ۶ و صفر به ترتیب برای m ، n و پارامتر همپوشانی حاصل می‌گردد که در نتیجه هرکدام از قسمت‌های ایجاد شده دارای ابعاد 2000×5000 هستند.



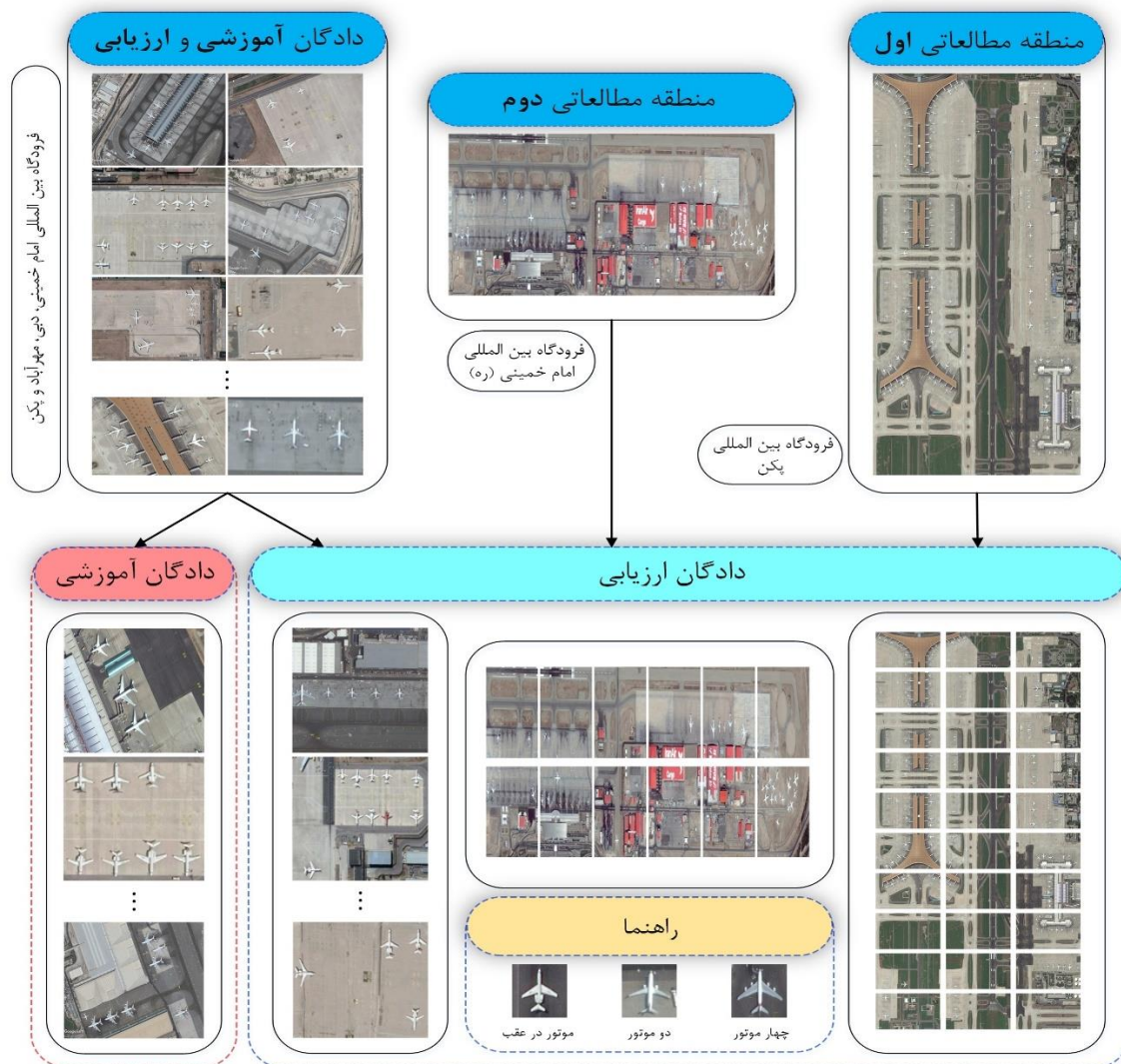
شکل ۳- نمونه ای از داده های آموزشی؛ تصاویر سمت راست، داده‌های آموزشی و تصاویر سمت چپ، داده‌های ارزیابی را نمایش می‌دهد. کادرهای آبی، قرمز و زرد به ترتیب معرف هواپیماهای چهار موتوره، دو موتوره و موتور در عقب هستند



شکل ۴- تصویر سنجش از دوری منطقه مطالعاتی دوم پس از انجام فرآیند بخش‌بندی

در این پژوهش، به‌منظور ارزیابی هرچه بهتر معماری پیشنهادی و معماری‌های دیگر، از سه دسته داده اعتبار-سنجی استفاده‌شده است. دسته اول مربوط به داده‌هایی از جنس داده‌های آموزشی است ولی با این تفاوت که در طی فرآیند آموزش شبکه، از آن‌ها استفاده‌نشده است. در این مجموعه تعداد ۷۵ نمونه، در دسترس است. دسته دوم و سوم جزء داده‌های با ابعاد بالا هستند که از فرودگاه‌های بین‌المللی پکن و امام خمینی (ره) اخذ گردیده و پس از اعمال بخش‌بندی جهت ارزیابی مورد استفاده قرار می‌گیرند.

شکل (۵) شمای کلی از نحوه بودجه‌بندی تصاویر ماهواره‌ای با قدرت تفکیک بالا به مجموعه داده‌های آموزشی و ارزیابی است. شایان‌ذکر است تصحیحات اتمسفری بر روی تمامی داده‌های آموزشی، ارزیابی، تصاویر منطقه مطالعاتی اول و دوم اعمال شده است.



شکل ۵- دسته‌بندی تصاویر سنجش‌ازدوری به داده‌های آموزشی و ارزیابی

۴-۲- آموزش و پیاده‌سازی معماری ارائه‌شده

معماری ارائه‌شده در این تحقیق با استفاده از الگوریتم بهینه‌سازی SGD^۱ [۳۹] آموزش دیده است. نرخ یادگیری برای آن در ابتدای آموزش ۰/۰۰۵ قرار داده شد و در گام‌های ۴۰۰۰ و ۴۶۰۰۰، مقدار آن تقسیم‌بر ۱۰ شد. مقدار پارامترهای Weight Decay، Momentum و Batch Size الگوریتم بهینه‌ساز به ترتیب برابر ۰/۰۰۰۱، ۰/۹ و ۱۶ قرار داده‌شده‌اند. همچنین در معماری ارائه‌شده، از Dropout [۴۰] استفاده‌نشده است. پارامتر آزاد نرخ رشد برابر ۲ و پارامتر N برابر ۲۴ مقداردهی شده‌اند تا شبکه بهترین پوشش را روی داده‌ها داشته باشد. چارچوب

^۱ Stochastic Gradient Descent

مورداستفاده جهت آموزش نیز Keras با پس‌زمینه‌ی Tensorflow، CUDA نسخه ۹ و CuDNN نسخه ۷ می‌باشد و حافظه RAM و حافظه گرافیکی (VRAM) قابل‌دسترس به ترتیب برابر ۱۲ و ۱۶ گیگابایت هستند. پس از مشخص کردن پارامترهای الگوریتم بهینه‌ساز، نوبت به انجام فرآیند آموزشی می‌رسد. در این فرآیند داده‌های آموزشی از لایه‌های کانولوشنی طراحی‌شده در شبکه عبور می‌کنند و در متغیرهای ترکیب‌کننده ذخیره می‌شوند. بعد از آن که فرایند پردازشی هر سلول به اتمام رسید، خروجی ترکیب‌کننده‌ها در درون سلول از یک لایه Average pooling عبور می‌کنند. داده‌های آموزشی، این فرآیند را دنبال می‌کنند تا به لایه چگال برسند. در این لایه به تعداد کلاس‌های موجود، نورون وجود دارد که وظیفه دسته‌بندی ویژگی‌های برداشت‌شده از دسته داده-

های آموزشی را دارند. دسته‌بندی کلاس‌ها با استفاده از مقداردهی به هر کدام از نورون‌ها انجام می‌شود. نورونی که مقدار بیشتری داشته باشد، به عنوان کلاس پیش‌بینی شده برای تصویر انتخاب می‌شود. میزان دقت کلاس پیش‌بینی شده با تابع هزینه اندازه‌گیری می‌شود.

$$L = -[y \log(z) + (1 - y) \log(1 - z)] \quad (6)$$

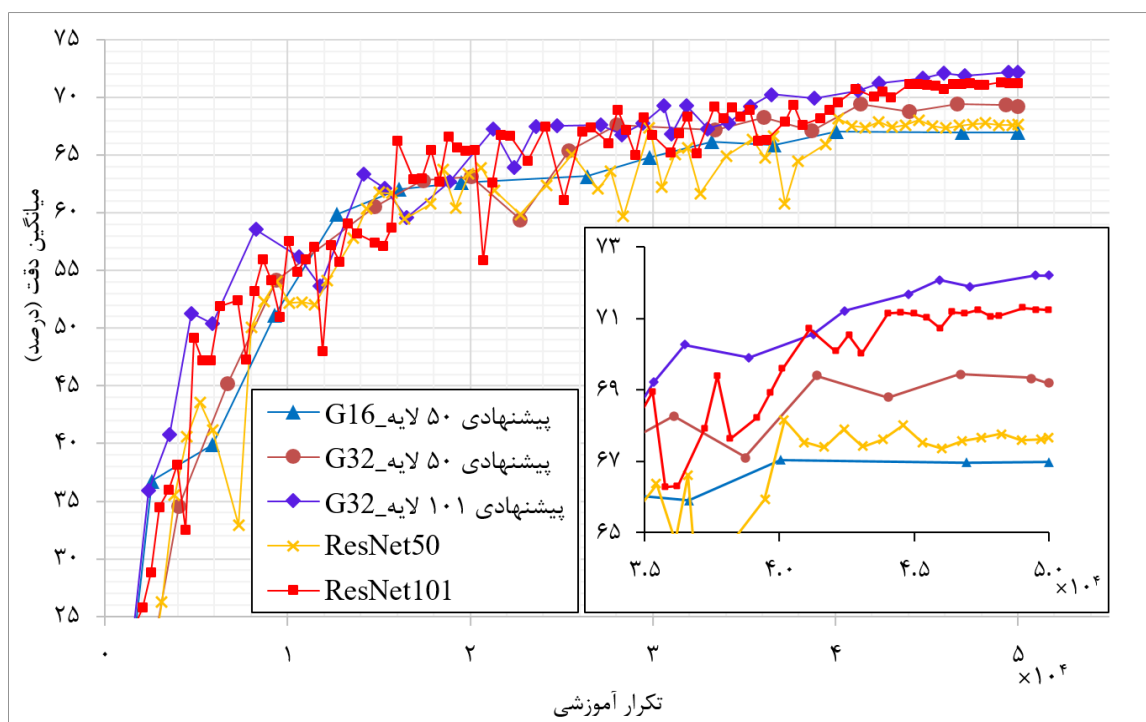
تابع هزینه استفاده شده در این تحقیق Cross-entropy است که رابطه (۶) بیانگر آن است. y و z به ترتیب مقدار صحیح برای کلاس داده و مقدار پیش‌بینی شده هستند.

۳-۴- ارزیابی و تحلیل نتایج

ارزیابی یک شبکه کانولوشن نیازمند بررسی دقت طبقه‌بندی آن بر روی کلاس‌های موردنظر است. به منظور ارزیابی دقت روش ارائه شده، قسمتی از مجموعه داده به داده‌های ارزیابی اختصاص یافته است. این داده‌های پس از آموزش، مورد ارزیابی قرار می‌گیرد تا میزان اختلاف بین معماری‌های قبلی و معماری ارائه شده با استفاده از معیارهایی مشخص شود. رقابت Microsoft COCO [۴۱] از جمله رقابت‌هایی است که به منظور ارزیابی دقت شبکه‌های کانولوشن برگزار می‌شوند. این رقابت شامل چندین معیار برای ارزیابی است

که از میانگین دقت کلی (Average Precision) با IOU (Intersection Over Union) در بازه ۰/۵ تا ۰/۹۵ به عنوان معیار اصلی یاد می‌شود. طبق [۴۲]، مقادیر AP بر روی داده‌گان ارزیابی محاسبه می‌شود.

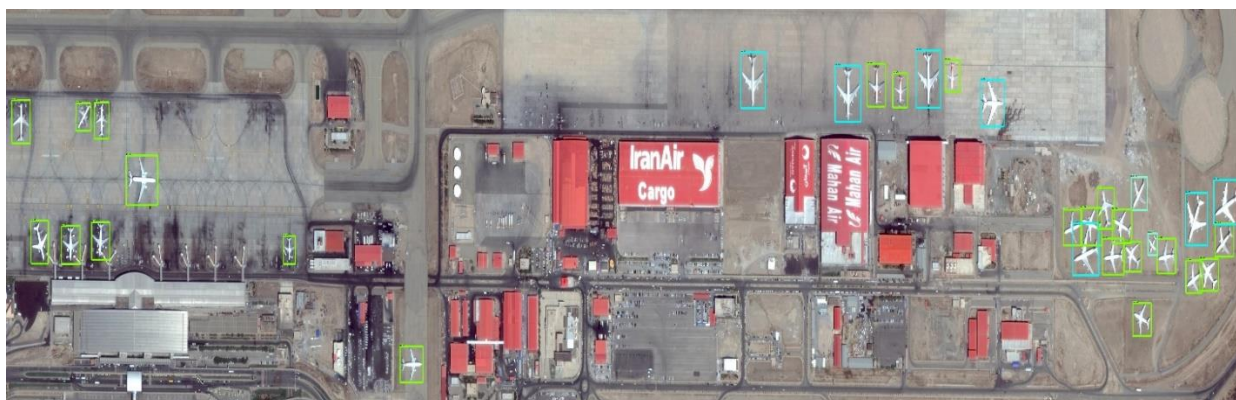
شکل (۶) میانگین دقت به دست آمده از اعمال معماری ارائه شده و معماری ResNet با تعداد لایه‌ها و نرخ رشد متفاوت در IOU بین ۰/۵ تا ۰/۹۵ را نمایش می‌دهد. با توجه به این که نرخ آموزشی در تکرار ۴۰۰۰۰ و ۴۶۰۰۰ به ده تقسیم می‌شود، الگوی رفتاری شبکه‌ها روان‌تر دنبال می‌شود. برای پایش وضعیت آموزشی، به طور منظم و هر ده دقیقه یک‌بار گزارشی از دقت متوسط هر یک از شبکه‌ها ذخیره می‌شود. به عبارتی دیگر نقاط موجود در شکل (۶) برای هر شبکه، گویای تعداد دفعات ارائه گزارش است. شبکه‌ای که در طول ۵۰۰۰۰ تکرار نقاط بیشتری دارد، حاکی از زمان آموزشی بالاتر آن برای تکمیل پروسه آموزش می‌باشد. در میان شبکه‌های مورد ارزیابی، شبکه پیشنهادی با ۵۰ لایه و نرخ رشد ۱۶، دارای کمترین و ResNet101 دارای بیشترین مدت زمان آموزشی می‌باشند. بر این اساس معماری ارائه شده در این مقاله توانایی دستیابی به دقت شبکه ResNet را در عین کمتر بودن تعداد پارامتر و زمان لازم برای آموزش، دارا است.



شکل ۶- مقدار میانگین دقت بدست آمده بر روی داده های ارزیابی

ارائه شده علاوه بر این که دقت بالاتری در تشخیص اهداف دارد، سریع تر از سایرین داده های آموزشی را پوشش می دهد. حفظ ویژگی های حاصله در لایه های ابتدایی و استفاده از آن ها در لایه های بعدی موجب می شود میزان هدر رفت در استفاده از ویژگی ها کاهش یابد. همین مورد معماری پیشنهادی را از سایر معماری ها متمایز می کند.

شکل (۷) تصویر سنجش از دوری از منطقه مطالعاتی دوم را نشان می دهد. تصویر مورد نظر پس از اعمال استخراج گر معرفی شده در این تحقیق، به دست آمده است. شناسایی ۳۰ هواپیما از ۳۲ فروند هواپیما، حاکی از دقت مناسب شبکه در مواجهه شدن با هر نوع تصویر رقومی برای شناسایی اهداف هدف دارد.



شکل ۷- نتیجه بدست آمده از اعمال شبکه پیشنهادی با ۱۰۱ لایه و $G=32$ بر روی تصویر سنجش از دوری فرودگاه امام خمینی (ره)

می شود و شبکه برای هر هدف برچسبی با عنوان های TP^۱، FP^۲، NP^۳ اختصاص می دهد. زمانی به یک هدف تشخیص داده شده، برچسب TP داده می شود که محدوده ی شناسایی شده و منسوب به آن، حداقل ۵۰٪ همپوشانی باهدفی که به صورت دستی مشخص شده است، داشته [۲۸]. در غیر این صورت به هدف مورد نظر برچسب FP داده خواهد شد. معیارهای Precision و Recall درصد آشکارسازی اهداف را نشان می دهند و به صورت زیر تعریف می شوند.

$$\text{Recall} = \frac{TP}{NP} \quad (7)$$

در رابطه (۷)، NP تعداد کل اهداف مکانی است که در داده های واقعی تولید شده اند. TP موقعیت های صحیح شناسایی شده است. این پارامتر همیشه عددی نامنفی است.

^۱ True Positives sample

^۲ False Positive samples

^۳ Negative Positive samples

جدول ۲- میانگین دقت (AP) حاصله از داده های ارزیابی

میانگین دقت (درصد)	مدت زمان آموزش (دقیقه)	تعداد پارامتر (میلیون)	نرخ رشد	تعداد لایه	معماری
۶۷/۶	۶۱۶	۲۵/۵	---	۵۰	ResNet
۷۱/۲	۹۰۲	۴۴/۶	---	۱۰۱	
۶۶/۹	۱۵۳	۱۰/۶	۱۶	۵۰	معماری پیشنهادی
۶۹/۲	۲۰۹	۱۸/۲	۳۲	۵۰	
۷۲/۱	۳۷۰	۳۰/۲	۳۲	۱۰۱	

جدول (۲) نتایج به دست آمده از اعمال Faster RCNN همراه با معماری معرفی شده در این تحقیق و معماری ResNet را نمایش می دهد. بهترین دقت به دست آمده از داده های ارزیابی با قلم درشت مشخص شده است. روش

شکل (۸) شامل تصویر خروجی از اعمال معماری پیشنهادی و معماری ResNet با ۱۰۱ لایه بر روی تصویر منطقه مطالعاتی اول است. همان طور که مشاهده می شود، در سمت راست تصویر، معماری پیشنهادی با شناسایی صحیح ۱۶۵ فروند از ۱۶۹ فروند هواپیما با توجه به کلاس های تعیین شده، دقت قابل قبولی را در شناسایی عوارض هدف از خود به جای گذاشته است. البته شایان ذکر است که معماری ارائه شده در این مقاله نسبت به معماری ResNet با ۱۰۱ لایه، ۱۴/۴ میلیون پارامتر آموزشی کاهش داشته است. این موضوع، خود باعث کاهش توان پردازشی مورد نیاز هم در فرآیند آموزش و هم در فرآیند شناسایی اهداف می گردد.

همچنین برای مقایسه شبکه پیشنهادی و دیگر شبکه ها، از معیار F1-Measure استفاده شده است. این معیار ترکیبی غیرخطی از دو معیار Precision و Recall می باشد [۴۳، ۴۴]. به منظور محاسبه معیار F1-Measure، مکان و کلاس تمام اهداف مورد نظر در داده های ارزیابی مشخص



(ب): معماری ResNet با ۱۰۱ لایه (۴۴/۶ میلیون پارامتر)



(الف): روش پیشنهادی با ۱۰۱ لایه و $G=32$ (۳۰/۲ میلیون پارامتر)

شکل ۸- نتیجه حاصل شده از اعمال معماری پیشنهادی و معماری ResNet بر روی تصویر منطقه مطالعاتی اول؛ رنگ‌های آبی فیروزه‌ای، زرد و سبز به ترتیب معرف دسته‌های هواپیماهای چهار موتوره، دو موتوره و موتور در عقب هستند.

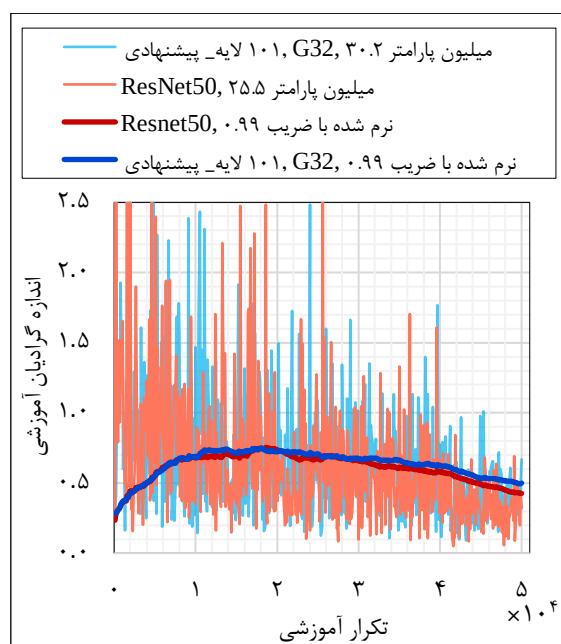
$$F_{\beta} = (1 + \beta^2) \frac{\text{Precision} \times \text{Recall}}{\beta^2 \times \text{Precision} + \text{Recall}} \quad (9)$$

رابطه‌ی (۹)، معیارهای Precision و Recall را با وزن β^2 ترکیب نموده و یک مقدار واحد و جامع حاصل می‌کند. مطابق با اکثر تحقیقاتی که در زمینه‌ی شناسایی

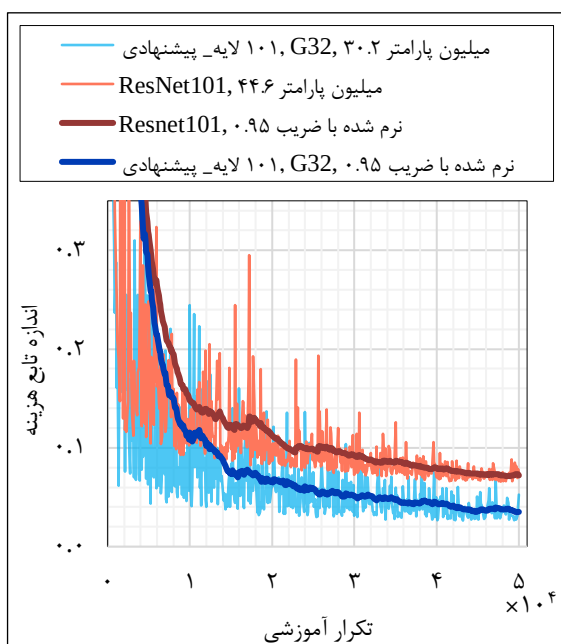
$$\text{Precision} = \frac{TP}{TP + FP} \quad (8)$$

در رابطه (۸)، FP تعداد موقعیت‌های غلط شناسایی شده در تمام تصاویر مورد آزمایش است. با استفاده از Precision و Recall، می‌توان معیار F1-Measure را محاسبه کرد.

تعداد لایه‌های برابر انجام شده است. نمودارهای موجود در شکل (۱۰) نشان‌دهنده خطای کمتر معماری طراحی شده در این تحقیق، نسبت به معماری ResNet است. این مقایسه در شرایطی کاملاً برابر انجام شده است. در تمامی آزمایش‌ها، شبکه ارائه شده با تعداد پارامترهای کمتر و مدت زمان آموزش و ارزیابی پایین‌تر توانست نتایج بهتری نسبت به معماری ResNet از خود به جای بگذارد.



شکل ۹- مقدار نرم مربعی گرادیان حاصل از دو شبکه



شکل ۱۰- مقایسه تابع هزینه دو معماری در فرآیند آموزش

اهداف انجام شده است، مقدار β^2 برابر با ۱ در نظر گرفته شده است. در نظر گرفتن این مقدار برای پارامتر ذکر شده باعث می‌شود میزان تأثیرگذاری هر دو معیار یکسان باشد [۴۵].

جداول (۳) و (۴) نتایج خروجی از اعمال شبکه‌های موردبررسی بر روی منطقه مطالعاتی اول و دوم است. بهترین مقادیر برای معیار F1-Measure با قلم درشت و آبی مشخص شده است.

جدول ۳- مقادیر F1-Measure برای منطقه مورد مطالعه اول

معماری	تعداد لایه	نرخ رشد	Precision	Recall	F1-Measure
ResNet	۵۰	----	۹۵/۶	۹۴/۸	۹۵/۲
	۱۰۱	----	۹۸/۳	۹۵/۲	۹۶/۷
معماری پیشنهادی	۵۰	۱۶	۹۲/۹	۹۳/۷	۹۳/۳
	۱۰۱	۳۲	۹۶/۷	۹۵/۹	۹۶/۳
		۳۲	۹۸/۲	۹۷/۶	۹۷/۹

جدول ۴- مقادیر F1-Measure برای منطقه مورد مطالعه دوم

معماری	تعداد لایه	نرخ رشد	Precision	Recall	F1-Measure
ResNet	۵۰	----	۸۹/۶	۹۰/۳	۸۹/۹
	۱۰۱	----	۹۲/۹	۹۳/۲	۹۳
معماری پیشنهادی	۵۰	۱۶	۸۸/۱	۸۸/۹	۸۸/۵
	۱۰۱	۳۲	۹۲/۵	۹۲/۹	۹۲/۷
		۳۲	۹۳/۷	۹۳/۷	۹۳/۷

تعداد لایه بکار رفته در هر شبکه تأثیر بسزایی در میزان محو شدن گرادیان دارد. لایه‌های بیشتر باعث کم شدن اطلاعات منتقل شده به لایه‌های انتهایی می‌شود. شکل (۹) اندازه گرادیان را در ۵۰۰۰۰ تکرار آموزشی شبکه ResNet50 و معماری پیشنهادی را نمایش می‌دهد. معماری ارائه شده در این تحقیق با وجود تعداد لایه‌های دو برابر نسبت به شبکه مورد مقایسه، از شیب کاهشی کمتری برخوردار است و به‌طور مناسبی مشکل کاهش گرادیان را کم‌رنگ کرده است. این دلیلی بر مفید بودن استفاده چندین باره از ویژگی‌های هر لایه در طول فرایند آموزش شبکه است. با توجه به نوسان زیادی که مقدار گرادیان در هر تکرار از خود نشان می‌دهد، با استفاده از توابع نمایی همراه با ضریب ۰/۹۹ نمودار موردنظر نرم شده است تا دید بهتری از اختلاف موجود بین دو معماری به وجود آید. شکل (۱۰) آزمایشی است که برای مقایسه دقت طبقه‌بندی معماری پیشنهادی و ResNet با

۵- نتیجه گیری

در این تحقیق با توجه به اهمیت شناسایی و طبقه بندی اهداف در تصاویر سنجش از دوری، فرآیندی مبتنی بر شبکه های کانولوشنی ارائه گردید که در عین کاهش تعداد پارامترهای آموزشی، با استخراج بهتر ویژگی های مناسب طبقه بندی، دقت روش های شناسایی اهداف را افزایش می دهد. افزایش تعداد لایه های کانولوشنی یکی از مؤلفه های اصلی در افزایش دقت طبقه بندی شبکه های کانولوشنی عمیق است؛ اما از طرفی افزایش بیش از حد این لایه ها موجب محو شدن گرادیان و اطلاعات وارد شده در لایه های انتهایی می گردد. این موضوع در بسیاری از تحقیقات صورت گرفته در زمینه ی شبکه های کانولوشن مطرح شده و با بررسی توپولوژی های مختلف و یا ارائه روش های آموزش نوین سعی در مقابله با چالش بوجود آمده کرده اند. پیرو دیگر محققان، در این مقاله، معماری کانولوشنی ارائه شده است که با حفظ ویژگی های تولید شده در تمام طول شبکه، از محو شدن گرادیان و از بین رفتن ویژگی های مستخرج جلوگیری می کند. نتیجه ی این فرآیند، افزایش توانایی مدل های موجود در تفکیک اهداف چند کلاسه است.

مراجع

فرآیند آموزش معماری کانولوشن پیشنهادی توسط ۳۲۰ قطعه تصویر انجام شد و پس از آن، این معماری بر روی دو منطقه مطالعاتی فرودگاه بین المللی پکن و فرودگاه بین المللی امام خمینی (ره) مورد ارزیابی قرار گرفت که به ترتیب مقادیر ۹۷/۹ و ۹۳/۷ درصد برای معیار F1-Measure به دست آمده است. این در حالی است که معماری ResNet با ۱۰۱ لایه کانولوشن و ۱۴/۴ میلیون پارامتر آموزشی بیشتر نسبت به معماری پیشنهادی، مقادیر ۹۶/۷ و ۹۳ درصد را برای معیار مذکور حاصل کرده است که این موضوع حاکی از برتری معماری پیشنهادی هم از نقطه نظر دقت شناسایی و هم از نظر میزان پارامتر آموزشی و توان پردازشی مورد نیاز دارد.

در راستای بهبود معماری ارائه شده در این مقاله می توان از اپراتورهای کانولوشن منبسط شده (Dilated Convolution) باهدف استخراج ویژگی های بارزتر استفاده کرد. از سوی دیگر باهدف توسعه و عمومی سازی هرچه بهتر روش ارائه شده، می توان آن را در دو مرحله بر روی تصاویر سنجش از دوری با مقیاس کوچک در سطح کشوری و یا حتی بیشتر اعمال کرد؛ به طوری که در مرحله ی اول هدف شناسایی مکان فرودگاه باشد و در گام بعد هواپیماهای داخل هر فرودگاه توسط روش ارائه شده شناسایی شود.

- [1] Y. LeCun et al., "Backpropagation applied to handwritten zip code recognition," Neural computation, vol. 1, no. 4, pp. 541-551, 1989.
- [2] S. Ren, K. He, R. Girshick, and J. Sun, "Faster r-cnn: Towards real-time object detection with region proposal networks," in Advances in neural information processing systems, 2015, pp. 91-99.
- [3] W. Liu et al., "Ssd: Single shot multibox detector," in European conference on computer vision, 2016: Springer, pp. 21-37.
- [4] J. Redmon, S. Divvala, R. Girshick, and A. Farhadi, "You only look once: Unified, real-time object detection," in Proceedings of the IEEE conference on computer vision and pattern recognition, 2016, pp. 779-788.
- [5] X. Chen and A. Gupta, "An implementation of faster rcnn with study for region sampling," arXiv preprint arXiv:1702.02138, 2017.
- [6] J. Lu, H. Sibai, and E. Fabry, "Adversarial examples that fool detectors," arXiv preprint arXiv:1712.02494, 2017.
- [7] Z. Li, C. Peng, G. Yu, X. Zhang, Y. Deng, and J. Sun, "Light-head r-cnn: In defense of two-stage object detector," arXiv preprint arXiv:1711.07264, 2017.
- [8] N. Farhadi, A. Kiani, and H. Ebadi, "Target detection from high-resolution remote sensing images using deep learning methods," Iranian Journal of Remote Sensing & GIS, vol. 11, no. 1, pp. 48-64, 2019.
- [9] H. Wu, H. Zhang, J. Zhang, and F. Xu, "Fast aircraft detection in satellite images based on convolutional neural networks," in 2015 IEEE International Conference on Image Processing (ICIP), 2015: IEEE, pp. 4210-4214.
- [10] F. Zhang, B. Du, L. Zhang, and M. Xu, "Weakly supervised learning based on coupled convolutional neural networks for aircraft detection," IEEE Transactions on Geoscience and Remote Sensing, vol. 54, no. 9, pp. 5553-5563, 2016.

- [11] F. Zeng et al., "A hierarchical airport detection method using spatial analysis and deep learning," *Remote Sensing*, vol. 11, no. 19, p. 2204, 2019.
- [12] Y. Xu, M. Zhu, P. Xin, S. Li, M. Qi, and S. Ma, "Rapid airplane detection in remote sensing images based on multilayer feature fusion in fully convolutional neural networks," *Sensors*, vol. 18, no. 7, p. 2335, 2018.
- [13] C. Cao et al., "Research on Airplane and Ship Detection of Aerial Remote Sensing Images Based on Convolutional Neural Network," *Sensors*, vol. 20, no. 17, p. 4696, 2020.
- [14] Y. LeCun, Y. Bengio, and G. Hinton, "Deep learning," *nature*, vol. 521, no. 7553, p. 436, 2015.
- [15] K. Simonyan and A. Zisserman, "Very deep convolutional networks for large-scale image recognition," *arXiv preprint arXiv:1409.1556*, 2014.
- [16] K. He, X. Zhang, S. Ren, and J. Sun, "Deep residual learning for image recognition," in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2016, pp. 770-778.
- [17] Y. Bengio, P. Simard, and P. Frasconi, "Learning long-term dependencies with gradient descent is difficult," *IEEE transactions on neural networks*, vol. 5, no. 2, pp. 157-166, 1994.
- [18] G. Huang, Z. Liu, L. Van Der Maaten, and K. Q. Weinberger, "Densely connected convolutional networks," in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2017, pp. 4700-4708.
- [19] R. K. Srivastava, K. Greff, and J. Schmidhuber, "Training very deep networks," in *Advances in neural information processing systems*, 2015, pp. 2377-2385.
- [20] G. Huang, Y. Sun, Z. Liu, D. Sedra, and K. Q. Weinberger, "Deep networks with stochastic depth," in *European conference on computer vision*, 2016: Springer, pp. 646-661.
- [21] M. Chen, W. Wang, S. Dong, and X. Zhou, "Video Vehicle Detection and Recognition Based on MapReduce and Convolutional Neural Network," in *International Conference on Sensing and Imaging*, 2018: Springer, pp. 552-562.
- [22] A. Krizhevsky, I. Sutskever, and G. E. Hinton, "Imagenet classification with deep convolutional neural networks," in *Advances in neural information processing systems*, 2012, pp. 1097-1105.
- [23] J. Deng, W. Dong, R. Socher, L.-J. Li, K. Li, and L. Fei-Fei, "Imagenet: A large-scale hierarchical image database," in *2009 IEEE conference on computer vision and pattern recognition*, 2009: Ieee, pp. 248-255.
- [24] C. Szegedy et al., "D., Vanhoucke, V., and Rabinovich, A. Going deeper with convolutions," *IEEE Xplore*, 2015.
- [25] C. Szegedy, V. Vanhoucke, S. Ioffe, J. Shlens, and Z. Wojna, "Rethinking the inception architecture for computer vision," in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2016, pp. 2818-2826.
- [26] G. Larsson, M. Maire, and G. Shakhnarovich, "Fractalnet: Ultra-deep neural networks without residuals," *arXiv preprint arXiv:1605.07648*, 2016.
- [27] Y. Bengio, "Learning deep architectures for AI," *Foundations and trends® in Machine Learning*, vol. 2, no. 1, pp. 1-127, 2009.
- [28] P. F. Felzenszwalb, R. B. Girshick, D. McAllester, and D. Ramanan, "Object detection with discriminatively trained part-based models," *IEEE transactions on pattern analysis and machine intelligence*, vol. 32, no. 9, pp. 1627-1645, 2009.
- [29] P. Viola and M. Jones, "Rapid object detection using a boosted cascade of simple features," *CVPR (1)*, vol. 1, pp. 511-518, 2001.
- [30] N. O. Attah-Okine, "Analysis of learning rate and momentum term in backpropagation neural network algorithm trained to predict pavement performance," *Advances in Engineering Software*, vol. 30, no. 4, pp. 291-302, 1999.
- [31] R. Girshick, J. Donahue, T. Darrell, and J. Malik, "Rich feature hierarchies for accurate object detection and semantic segmentation," in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2014, pp. 580-587.
- [32] Z. Deng, H. Sun, S. Zhou, J. Zhao, L. Lei, and H. Zou, "Multi-scale object detection in remote sensing imagery with convolutional neural networks," *ISPRS Journal of Photogrammetry and Remote Sensing*, vol. 145, pp. 3-22, 2018.
- [33] R. Girshick, "Fast R-CNN. In processing of IEEE International Conference on Computer Vision," *Santiago*, pp. 13-16, 2015.
- [34] S. Ioffe and C. Szegedy, "Batch normalization: Accelerating deep network training by reducing internal covariate shift," *arXiv preprint arXiv:1502.03167*, 2015.

- [35] X. Glorot, A. Bordes, and Y. Bengio, "Deep sparse rectifier neural networks," in Proceedings of the fourteenth international conference on artificial intelligence and statistics, 2011, pp. 315-323.
- [36] Z. Chen, T. Zhang, and C. Ouyang, "End-to-end airplane detection using transfer learning in remote sensing images," Remote Sensing, vol. 10, no. 1, p. 139, 2018.
- [37] Z. Han, H. Zhang, J. Zhang, and X. Hu, "Fast aircraft detection based on region locating network in large-scale remote sensing images," in 2017 IEEE International Conference on Image Processing (ICIP), 2017: IEEE, pp. 2294-2298.
- [38] L. Sommer, N. Schmidt, A. Schumann, and J. Beyerer, "Search Area Reduction Fast-RCNN for Fast Vehicle Detection in Large Aerial Imagery," in 2018 25th IEEE International Conference on Image Processing (ICIP), 2018: IEEE, pp. 3054-3058.
- [39] I. Sutskever, J. Martens, G. Dahl, and G. Hinton, "On the importance of initialization and momentum in deep learning," in International conference on machine learning, 2013, pp. 1139-1147.
- [40] N. Srivastava, G. Hinton, A. Krizhevsky, I. Sutskever, and R. Salakhutdinov, "Dropout: a simple way to prevent neural networks from overfitting," The Journal of Machine Learning Research, vol. 15, no. 1, pp. 1929-1958, 2014.
- [41] T.-Y. Lin et al., "Microsoft coco: Common objects in context," in European conference on computer vision, 2014: Springer, pp. 740-755.
- [42] J. Hosang, R. Benenson, P. Dollár, and B. Schiele, "What makes for effective detection proposals?," IEEE transactions on pattern analysis and machine intelligence, vol. 38, no. 4, pp. 814-830, 2015.
- [43] H. G. Akçay and S. Aksoy, "Automatic detection of geospatial objects using multiple hierarchical segmentations," IEEE Transactions on Geoscience and Remote Sensing, vol. 46, no. 7, pp. 2097-2111, 2008.
- [44] H. Sun, X. Sun, H. Wang, Y. Li, and X. Li, "Automatic target detection in high-resolution remote sensing images using spatial sparse coding bag-of-words model," IEEE Geoscience and Remote Sensing Letters, vol. 9, no. 1, pp. 109-113, 2011.
- [45] A. O. Ok, C. Senaras, and B. Yuksel, "Automated detection of arbitrarily shaped buildings in complex environments from monocular VHR optical satellite imagery," IEEE Transactions on Geoscience and Remote Sensing, vol. 51, no. 3, pp. 1701-1717, 2012.