

تشخیص نوع خودرو به منظور کنترل هوشمند عبور و مرور خودروها توسط یادگیری عمیق

مرتضی احساندوست^۱، عاطفه سلیمی شهرکی^{۲*}، محمد روح ا... یزدانی^۳، محمد لعلی دستجردی^۱

^۱دانشجوی کارشناسی ارشد رشته مهندسی برق، واحد اصفهان (خوراسگان)، دانشگاه آزاد اسلامی، اصفهان، ایران

^۲استادیار رشته مهندسی برق، واحد اصفهان (خوراسگان)، دانشگاه آزاد اسلامی، اصفهان، ایران

^۳دانشیار رشته مهندسی برق، واحد اصفهان (خوراسگان)، دانشگاه آزاد اسلامی، اصفهان، ایران

چکیده

با افزایش خودروها در سراسر جهان و توسعه فناوری، سیستم‌های حمل و نقل هوشمند به عنوان یکی از راه‌حل‌های کلیدی برای کنترل ترافیک و افزایش ایمنی راه‌ها مطرح شده‌اند. یکی از بخش‌های مهم این سیستم‌ها، تشخیص نوع خودرو است. در این مقاله از ترکیب مدل‌های EfficientNet-B0 و YOLO-V11 جهت تشخیص نوع خودرو استفاده شده است. در این مدل، EfficientNet-B0 به عنوان ستون فقرات مدل YOLO-V11 به کار رفته است. این مدل، ویژگی‌های تصاویر را استخراج کرده و آن‌ها را به مدل YOLO اعمال می‌کند تا مکان و نوع خودرو را به طور دقیق تشخیص دهد. برای ارزیابی عملکرد مدل پیشنهادی از مجموعه داده تصویری BVMMR که شامل بیش از ۵۰۰۰ تصویر از انواع خودروهای ایرانی است، استفاده شده است. مدل مورد استفاده با زبان برنامه‌نویسی پایتون نوشته شده و در محیط کولی (Colab) اجرا شده است. نتایج اجرای کدها نشان می‌دهد که مدل پیشنهادی دارای مقدار میانگین دقت در هم پوشانی پنجاه درصد (mAP50) برابر با ۹۹/۳٪ و میانگین دقت در هم پوشانی بیشتر از پنجاه درصد (mAP50-95) برابر با ۹۸/۳٪ است. این مدل، از نظر سرعت پردازش نیز توانسته است با میانگین زمان ۰/۲ میلی ثانیه برای پیش پردازش، ۲/۸ میلی ثانیه برای استنتاج و ۲/۵ میلی ثانیه برای پس پردازش در هر تصویر، عملکرد مطلوبی را ارائه دهد.

مقاله پژوهشی

تاریخ دریافت:

۱۴۰۳/۱۲/۲۱

تاریخ پذیرش:

۱۴۰۴/۲/۶

کلیدواژه‌ها:

بینایی رایانه، سیستم حمل و نقل

هوشمند، یادگیری

عمیق، YOLO، Efficient

Net

نویسنده مسئول:

a.salimi@khuif.ac.ir

doi : 10.22034/ABMIR.2025.22898.1112

۱- مقدمه

شناسایی، شمارش و ردیابی می‌شوند. در نهایت، مسیر حرکت خودروها با مقایسه مکان آن‌ها در فریم‌های مختلف پیش‌بینی می‌شود. از مجموعه داده خودروهای هوایی^۸ [۸] و مجموعه داده وسایل نقلیه هوایی بدون سرنشین چندوجهی برای نظارت بر ترافیک در ارتفاع پایین (AU-AIR)^۹ [۹] برای ارزیابی رویکرد پیشنهادی استفاده شده و به ترتیب دقت‌های ۸۷٪، ۹۲٪، ۸۴٪ و ۸۸٪ به دست آمده است.

زوهیب و همکاران در [۱۰] رویکردی را برای تشخیص وسایل نقلیه اضطراری با ترکیب داده‌های صوتی و تصویری ارائه دادند. این روش با استفاده از شبکه طیف موقتی مبتنی بر توجه (ATSN)^{۱۰} برای تشخیص صدای آژیر و معماری فیوژن فضایی چندسطحی (MLSF-YOLO)^{۱۱} برای تشخیص بصری وسایل نقلیه، دقت و زمان پاسخ را بهبود می‌بخشد. نتایج نشان داد که این سیستم با دقت ۹۶/۱۹٪ و میزان خطای ۳/۸۱٪ عملکرد بهتری نسبت به روش‌های سنتی دارد. نتایج نشان می‌دهد که مدل پیشنهادی توانسته است به‌خوبی وسایل نقلیه امدادی را شناسایی و به بهتر شدن سرعت واکنش و کمتر شدن حوادث در شهر کمک کند.

سونگ و همکاران در [۱۱] مدلی جدید برای شناسایی بهتر خودروها و کاهش هزینه‌های محاسباتی ارائه دادند. آن‌ها با بهره‌گیری از معماری Mamba_ViT^{۱۲}، ویژگی‌های تصویر را در چند سطح مختلف و به‌صورت جداگانه بررسی کردند تا اطلاعات دقیق‌تری از صحنه به دست آورند و عملکرد مدل را بهبود بخشیدند. جهت ارزیابی مدل پیشنهادی از مجموعه داده شناسایی و ردیابی دانشگاه آلبانی (UA-DETRAC)^{۱۳} [۱۳]، استفاده شده است. نتایج نشان داد که دقت این مدل ۳/۲٪ بیشتر از مدل

تشخیص نوع وسایل نقلیه بخشی از سیستم‌های حمل‌ونقل هوشمند (ITS)^۱ است [۱]. شناسایی دقیق وسایل نقلیه و ردیابی آن‌ها توسط دوربین‌های نصب‌شده در اتوبان‌ها، باعث افزایش نظارت بر ایمنی شهری و مدیریت ترافیک می‌شود [۲]. تشخیص وسایل نقلیه می‌تواند زمان‌بندی چراغ‌های راهنمایی و رانندگی را براساس حجم ترافیک تنظیم کند و باعث کاهش تأخیر ناشی از ازدحام و تخلفات ترافیکی شود [۳]. تشخیص نوع وسیله نقلیه می‌تواند خدماتی را برای نظارت بر وضعیت جاده و سیستم‌های جمع‌آوری عوارض بزرگراه‌ها ارائه دهد [۴]. سیستم‌های رانندگی خودمختار^۲ نیز شامل ادراک محیطی، تصمیم‌گیری رفتاری و برنامه‌ریزی و کنترل حرکت هستند. تشخیص دقیق وسایل نقلیه برای درک محیط آن‌ها ضروری است [۵]. استفاده از روش‌های مبتنی بر یادگیری عمیق مانند مدل YOLO^۳، باعث بهبود دقت و سرعت تشخیص خودروها می‌شود.

باکرسی در [۶] از مدل YOLO-V8 برای تشخیص وسایل نقلیه در تصاویر هوایی پهپادهای خودمختار استفاده کرده و با جمع‌آوری داده‌های متنوع از زوایا و ارتفاعات مختلف، تعمیم‌پذیری مدل را بهبود داده است. در آزمایش‌ها، دو نسخه YOLO-V8N و YOLO-V8X در شرایط نوری و محیطی متنوع مقایسه شدند. نتایج نشان داد که YOLO-V8N با دقت^۴ ۸۳٪ و فراخوانی^۵ ۷۹٪ برای کاربردهای بلندرنج مناسب‌تر است. درحالی‌که YOLO-V8X با دقت ۹۶٪ و فراخوانی ۸۹٪ عملکرد بهتری دارد.

یوسف و همکاران در [۷] رویکرد جدیدی را برای تشخیص و ردیابی وسایل نقلیه ارائه دادند. در این روش، ابتدا موقعیت مکانی تصاویر تعیین و نویز آن‌ها حذف می‌شود. سپس با تفکیک تصویر و استفاده از الگوریتم‌هایی مانند جنگل تصادفی^۶، تطبیق الگو و شمارش با هیستوگرام گرادین‌های جهت‌دار (HOG)^۷ خودروها

^۸ Aerial Cars

^۹ A Multi-modal Unmanned Aerial Vehicle Dataset for Low Altitude Traffic Surveillance

^{۱۰} Attention-Based Temporal Spectrum Network

^{۱۱} Multi-Level Spatial Fusion YOLO

^{۱۲} University at Albany Detection and Tracking

^۱ Intelligent Transportation Systems

^۲ Autonomous driving systems

^۳ You Only Look Once

^۴ Precision

^۵ Recall

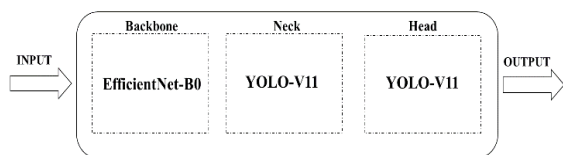
^۶ Random Forest Algorithm

^۷ Histogram Of Oriented Gradients

پیشنهادی معرفی شده است. در این بخش، ابتدا مدل‌های YOLO-V11 و EfficientNet-B0 و نحوه ادغام آن‌ها ذکر شده است. در بخش ۳ مجموعه داده مورد استفاده در این مقاله معرفی شده و در بخش ۴ نیز نتایج ارزیابی مدل نشان داده شده است. در نهایت، در بخش ۵ یک نتیجه‌گیری کلی از این مقاله ارائه شده و پیشنهادهایی برای کارهای آینده بیان شده است.

۲- روش پیشنهادی

در این مقاله، جهت تشخیص و دسته‌بندی نوع خودرو از ترکیب مدل‌های EfficientNet-B0 و YOLO-V11 استفاده شده است. مدل EfficientNet-B0 جایگزین ستون فقرات^۸ مدل YOLO-V11 شده تا دقت تشخیص را افزایش دهد و عملکرد مدل در شرایط مختلف را بهبود ببخشد. شکل (۱) بلوک دیاگرام روش پیشنهادی را نشان می‌دهد. این مدل ترکیبی [۲۱] بر مبنای مدل YOLO طراحی شده و به صورت یکپارچه عمل می‌کند. یعنی هر ۳ بخش ستون فقرات، گردن^۹ و سر^{۱۰} مدل در یک بلوک قرار گرفته‌اند و با کمک یکدیگر عملیات مورد نظر را انجام می‌دهند. قسمت ستون فقرات وظیفه استخراج ویژگی‌های داده‌ها را برعهده دارد. قسمت گردن نیز، ویژگی‌های استخراج شده جهت تشخیص نوع خودرو را پردازش کرده و آن‌ها را بهبود می‌بخشد. در نهایت، قسمت سر خروجی نهایی که شامل مکان و نوع خودرو است را تولید کرده و به خروجی شبکه منتقل می‌کند.



شکل (۱): بلوک دیاگرام روش پیشنهادی

۲-۱ مدل YOLO-V11

مدل YOLO-V11 در سال ۲۰۲۴ معرفی شده و جهت کلاس‌بندی، بخش‌بندی و تشخیص اشیا به کار می‌رود [۲۲]. در شکل (۲) بلوک

YOLO-V8-tiny است و تنها از ۶۰٪ پارامترهای خود استفاده می‌کند.

ژانگ و همکاران در [۱۴] الگوریتمی به نام YOLO با آگاهی از بافت محلی و نمای کلی تصویر (LGA-YOLO^۱) را برای تشخیص وسایل نقلیه معرفی کردند. این مدل شامل دو بلوک افزایش میدان دید و تقویت ویژگی‌های محلی (MLKM)^۲ و استخراج اطلاعات زمینه‌ای کلی (DGAM)^۳ است که باعث بهبود تمایز وسایل نقلیه از پس‌زمینه‌های پیچیده می‌شود. این مدل با استفاده از شبکه ترکیب ویژگی‌های سطح بالا و پایین، اشیا را چندمقیاسی را بهتر شناسایی می‌کند. مجموعه داده‌های یونیکورن ویژه شناسایی اشیا کوچک (USOD)^۴ [۱۵]، تشخیص وسیله نقلیه در تصاویر هوایی (VEDAI)^۵ [۱۶] و تشخیص اشیا در تصاویر هوایی (DOTA)^۶ [۱۷] برای ارزیابی این مدل به کاررفته‌اند. نتایج ارزیابی نشان دادند که میانگین دقت متوسط در هم‌پوشانی پنجاه درصد (mAP50)^۷ به ترتیب عبارت از ۹۳٪، ۸۰٪ و ۷۸٪ هستند.

پارک و همکاران در [۱۸] از ترکیب مدل‌های YOLO-V8 و MobileNet جهت تشخیص و طبقه‌بندی وسایل نقلیه و پلاک آن‌ها استفاده کرده‌اند. این دو الگوریتم به ترتیب برای تشخیص اشیا در زمان واقعی و پیش‌پردازش کارآمد به کار می‌روند و می‌توانند در دستگاه‌هایی که منابع محدودی دارند (مانند تلفن همراه) کار کنند. نویسندگان از این کار برای مدیریت یک پارکینگ خصوصی استفاده نموده و به دقت‌های ۹۷٪/۹۵٪ و ۹۷٪/۹۵٪ رسیده‌اند.

در این مقاله، از یک مدل ترکیبی کارآمد جهت شناسایی و دسته‌بندی انواع خودروها استفاده شده است. در این مدل ترکیبی، مدل EfficientNet-B0 [۱۹] با مدل YOLO-V11 [۲۰] ادغام شده است و عملکرد قابل قبولی را از خود نشان داده است. این مدل به صورت یکپارچه عمل می‌کند و هدف اصلی آن تشخیص انواع خودروها با کمترین خطا است. در بخش ۲ این مقاله روش

⁷ Mean Average Precision at 50% Intersection over Union (mAP50)

⁸ Backbone

⁹ Neck

¹⁰ Head

¹ Local And Global Aware YOLO

² Multi scale Large Kernel Local Aware Module

³ Directional Global Context Aware Module

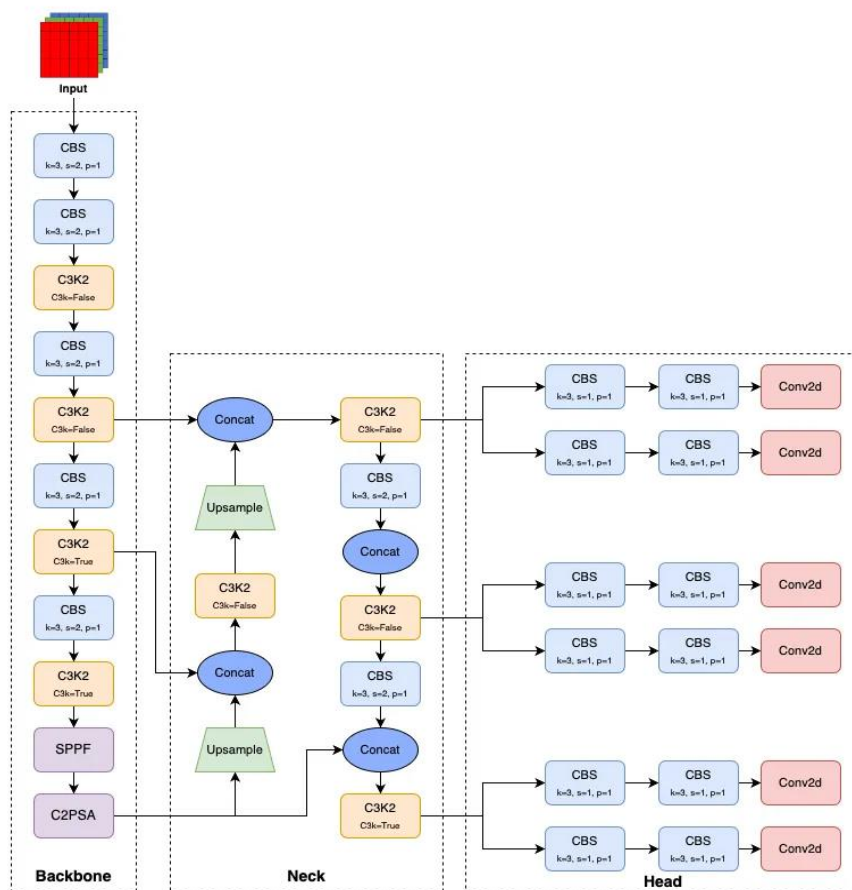
⁴ Unicorn Small Object Dataset (USOD)

⁵ Vehicle detection in aerial imagery (VEDAI)

⁶ Dataset for object detection in aerial images (DOTA)

بهتر مدیریت کند و ویژگی‌ها را از مقیاس‌های گوناگون استخراج کند. بلوک مکانیزم توجه فضایی نقطه‌ای با اشتراک‌گذاری جزئی میان مرحله‌ای (C2PSA)^۵ با استفاده از مکانیزم توجه^۶ ویژگی‌های مکانی تصویر را بهتر درک کرده و دقت تشخیص اشیا را افزایش می‌دهد. پس از استخراج ویژگی‌ها، قسمت گردن این ویژگی‌ها را ترکیب و تقویت می‌کند تا دقت تشخیص بهبود یابد. درنهایت، قسمت سر مختصات جعبه‌های محدودکننده^۷، امتیاز حضور اشیا^۸ و کلاس‌های احتمالی آن‌ها را پیش‌بینی می‌کند.

دیاگرام YOLO-V11 نشان داده شده است. ستون فقرات مدل ویژگی‌های تصاویر را از طریق بلوک‌های مختلفی استخراج می‌کند. بلوک کانولوشن، نرمال‌سازی دسته‌ای و تابع فعال‌سازی غیرخطی سیگموئید (CBS)^۱ به شبکه کمک می‌کند تا الگوهای پیچیده موجود در تصاویر را یاد بگیرد. بلوک میان مرحله‌ای جزئی با اندازه هسته ۲×۲ (C3k2)^۲ با استفاده از لایه‌های گلوگاه^۳، ابعاد ویژگی‌ها را کاهش داده و آن‌ها را ترکیب می‌کند. این کار باعث می‌شود تا ویژگی‌های بیشتری استخراج شوند. بلوک ادغام هرم فضایی سریع (SPPF)^۴ به شبکه کمک می‌کند تا تصاویر با اندازه‌های مختلف را



شکل (۲): بلوک دیاگرام مدل YOLO-V11 [۲۲]

⁵ Cross-Stage Partial Point-wise Spatial Attention (C2PSA)

⁶ Attention Mechanism

⁷ Bounding Boxes

⁸ Object Presence Score

¹ Convolution, Batch Normalization, Sigmoid Linear Unit (CBS)

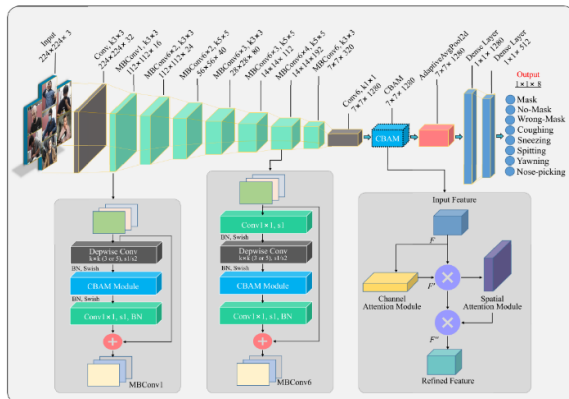
² Cross Stage Partial with kernel size 2 (c3k2)

³ Bottleneck Layers

⁴ Spatial Pyramid Pooling – Fast (SPPF)

۲-۲ مدل EfficientNet-B0

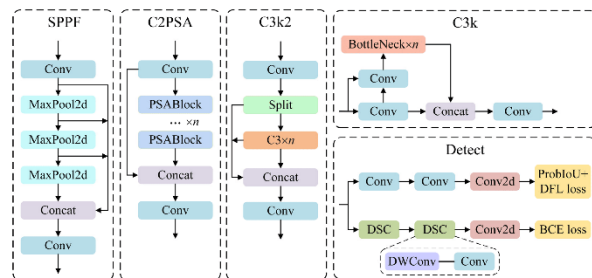
مدل EfficientNet یک شبکه عصبی کانولوشنی است که توسط تان و لی [۱۹] طراحی شده و به صورت هم‌زمان، عرض و عمق شبکه و وضوح تصاویر ورودی به آن را مقیاس‌بندی می‌کند. پایه‌ای‌ترین نسخه این خانواده مدل EfficientNet-B0 است. شکل (۴) بلوک دیاگرام این مدل رسم شده است. مشاهده می‌شود که در این مدل از تعداد زیادی بلوک MBConv جهت استخراج ویژگی‌ها استفاده شده است. ابتدا، لایه $Conv1 \times 1$ ، تعداد کانال‌های ورودی را افزایش می‌دهد تا لایه‌های بعدی بتوانند ویژگی‌های بیشتری از داده‌ها را استخراج کنند. سپس، بلوک Depthwise $Conv3 \times 3$ یک فیلتر کانولوشن 3×3 را به طور جداگانه روی هر کانال اعمال می‌کند که از نظر محاسباتی بسیار کارآمد است. بعد از آن، بلوک فشرده‌سازی و تحریک (SE)^۱ اهمیت هر کانال را تنظیم می‌کند تا شبکه بتواند روی ویژگی‌های مهم‌تر تمرکز کند. سپس، لایه $Conv1 \times 1$ آخر، تعداد کانال‌ها را کاهش می‌دهد تا حجم محاسبات کم‌تر شود. در نهایت، اتصال باقی‌مانده^۲ نیز با حفظ اطلاعات ورودی و انتقال مستقیم آن به لایه‌های بعدی، از بروز مشکلاتی مانند محو شدن گرادینت^۳ (عدم به‌روزرسانی مناسب وزن‌ها) جلوگیری کرده و موجب پایداری بیشتر در فرآیند یادگیری شبکه می‌شود.



شکل (۴): بلوک دیاگرام مدل EfficientNet-B0 [۲۴]

مدل EfficientNet-B0 یک مدل سبک‌وزن و از پیش آموزش‌دیده است و می‌تواند ویژگی‌های مهم تصاویر را با دقت

شکل (۳) ساختار داخلی بلوک‌های فوق را نشان می‌دهد. بلوک SPPF از لایه‌های کانولوشن و $MaxPool2d$ برای پردازش ویژگی‌ها در چندین سطح استفاده می‌کند تا مدل بتواند اشیاء با اندازه‌های مختلف را تشخیص دهد. بلوک C2PSA با استفاده از مکانیزم توجه، بر بخش‌های مهم تصویر تمرکز کرده و اطلاعات اشیاء را پردازش می‌کند. بلوک‌های C3k2 و C3k از لایه‌های کانولوشن و گلوگاه برای پردازش عمیق‌تر ویژگی‌ها استفاده نموده و نتایج حاصل‌شده را با یکدیگر ترکیب می‌کنند. این کار به تقویت ویژگی‌های استخراج‌شده کمک می‌کند. در نهایت، بلوک Detect مسئول شناسایی اشیاء و نام‌کلاس آن‌ها است و این عمل را طریق دو مسیر پردازش موازی انجام می‌دهد. مسیر اول، دقت تشخیص و مکان‌یابی جعبه‌های محدودکننده اشیاء را بهبود بخشیده و مسیر دوم، پیش‌بینی کلاس‌های اشیاء را بهینه می‌کند. با استفاده از این بلوک‌ها، دقت نهایی YOLO-V11 افزایش می‌یابد.



شکل (۳): معماری بخش‌های مختلف مدل YOLOV-11 [۲۳]

لازم به ذکر است که مدل YOLO-V11 نیز همانند مدل‌های قبلی خانواده YOLO، تشخیص اشیاء را به صورت هم‌زمان و با سرعت بالایی انجام می‌دهد. این امر، باعث کاهش دقت مدل در تصاویر خیلی پیچیده می‌شود. جهت بهبود دقت مدل در مواجهه با تصاویر بسیار پیچیده می‌توان ستون فقرات مدل را برداشته و آن را با یک مدل از پیش آموزش‌دیده قوی‌تر جایگزین کرد. این مدل‌ها جهت استخراج ویژگی‌های تصاویر و دسته‌بندی کردن آن‌ها استفاده می‌شوند. برای انجام این کار، از مدل از پیش آموزش‌دیده EfficientNet-B0 استفاده شده است.

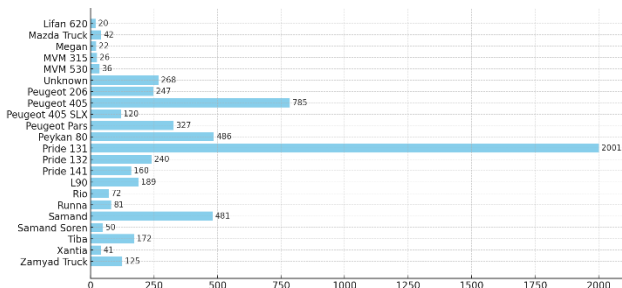
³ Vanishing Gradient

¹ Squeeze-and-Excitation

² Residual Connection



شکل (۵): تصاویر نمونه موجود در مجموعه داده BVMMR [۲۵]



شکل (۶): تعداد تصویر هر مدل در مجموعه داده BVMMR [۲۵]

۴- ارزیابی روش پیشنهادی

کد نویسی مدل ترکیبی مورد استفاده در این مقاله با زبان برنامه‌نویسی پایتون و در محیط کولی^۱ انجام گرفته است. این محیط توسط کمپانی گوگل پشتیبانی می‌شود و پژوهشگران را به صورت آنلاین به یک سیستم گرافیکی قدرتمند متصل می‌کند.

۴-۱ مراحل اجرای مدل

شکل (۷) نحوه پیاده‌سازی مدل ترکیبی YOLO-V11 و EfficientNetB0 را نشان می‌دهد. ابتدا مدل‌ها معرفی شده و پارامترهای آن تنظیم می‌گردند. مهم‌ترین پارامترهای مدل شامل نرخ یادگیری اولیه^۲، نرخ یادگیری نهایی^۳، تعداد دوره‌های آموزش^۴ و اندازه بسته‌های داده^۵ هستند. تنظیم دقیق این پارامترها باعث افزایش دقت مدل می‌شود. در این مقاله مقادیر پارامترهای فوق به ترتیب عبارت از ۰/۰۰۰۱، ۰/۰۰۱، ۱۰۰ و ۱۶ هستند.

بسیار بالایی استخراج کرده و آن‌ها را دسته‌بندی کند. وقتی که این مدل در بخش ستون فقرات مدل YOLO-V11 قرار گیرد، سرعت و دقتش متعادل‌تر می‌شود. این ادغام مزایا و معایبی نیز دارد که در ادامه به آن اشاره می‌شود.

۲-۳ مزایا و معایب مدل ترکیبی

یکی از مهم‌ترین مزایای ادغام مدل EfficientNet-B0 در ستون فقرات مدل YOLO-V11 افزایش دقت تشخیص اشیا بدون افزایش حجم مدل است. EfficientNet-B0 با استفاده از تکنیک مقیاس‌بندی مرکب ویژگی‌های غنی‌تری از تصاویر را استخراج می‌کند و منجر به بهبود توانایی مدل در تشخیص اشیا با جزئیات بیشتر می‌شود. این ترکیب می‌تواند باعث بهبود عملکرد سیستم‌های نظارتی و امنیتی شود. همچنین، استفاده از این مدل ترکیبی باعث می‌شود که مدل در مواجهه با داده‌های مختلف انعطاف‌پذیرتر باشد. این ادغام معایبی نیز دارد که باید مورد توجه قرار گیرد. یکی از چالش‌های اصلی، سازگاری معماری این دو مدل با یکدیگر است. ساختار مدل EfficientNet جهت طبقه‌بندی تصاویر طراحی شده و ممکن است مستقیماً با ساختار مدل YOLO که برای تشخیص اشیا بهینه‌سازی شده همخوانی کامل نداشته باشد و باعث افزایش زمان آموزش مدل شود. چون استخراج ویژگی‌های پیچیده‌تر نیاز به پردازش بیشتری دارند. بنابراین برای بهره‌برداری کامل از این ترکیب باید هزینه‌های محاسباتی، میزان افزایش دقت و نیازهای سخت‌افزاری را به دقت بررسی کرد تا بهترین تعادل بین دقت و سرعت حاصل شود.

۳- مجموعه داده

جهت ارزیابی روش پیشنهادی، از مجموعه داده تصویری BVMMR استفاده شده است. این مجموعه داده توسط آقای بیگلری [۲۵] تهیه شده و شامل بیش از ۵۰۰۰ تصویر از مدل‌های مختلف خودروهای ایرانی است. شکل (۵) چند نمونه از تصاویر موجود در مجموعه داده و شکل (۶) تعداد تصاویر مربوط به هر مدل را نشان می‌دهد.

³ Final Learning Rate

⁴ Number of Epochs

⁵ Batch Size

¹ Google Collaboratory (Colab)

² Initial Learning Rate



شکل (۷): فلوچارت مربوط به مدل ترکیبی مورد استفاده

- (۱) نمودار خطای مکان‌یابی جعبه (box_loss) نشان می‌دهد که مدل چقدر در پیدا کردن مکان دقیق اشیا در تصاویر بهتر می‌شود.
- (۲) نمودار خطای شناسایی کلاس (cls_loss)^۲ نشان‌دهنده میزان خطای مدل در تشخیص نام کلاس‌های مختلف اشیا است.
- (۳) نمودار خطای تمرکزی توزیعی در پیش‌بینی مکان اشیا (dfl_loss)^۳ نشان‌دهنده میزان خطای مدل در پیش‌بینی لبه‌های جعبه‌های محدودکننده است.
- لازم به ذکر است که میزان کم بودن خطا در هر ۳ نمودار فوق نشان‌دهنده عملکرد خوب مدل در مرحله آموزش آن است. به همین ترتیب، مقدار پایین خطای آن‌ها در نمودارهای اعتبارسنجی نیز نشان‌دهنده موفق بودن مدل در تست داده‌های جدید است.
- (۴) نمودار دقت نشان می‌دهد که چند درصد از اشیایی که مدل شناسایی کرده واقعاً درست هستند. این معیار برای ارزیابی کیفیت مدل در تشخیص اشیا اهمیت زیادی دارد. مقدار دقت از رابطه (۱) به دست می‌آید، که در آن (TP)^۴ تعداد اشیای درست شناسایی شده توسط مدل و (FP)^۵ تعداد اشیایی است که مدل به اشتباه شناسایی کرده است.

$$P = \frac{TP}{TP + FP} \quad (1)$$

پس از تعریف مدل و تنظیم پارامترهای آن، مجموعه داده تصویری موردنظر بارگذاری شده و ابعاد تصاویر آن به ۶۴۰×۶۴۰ (ابعاد موردنظر مدل YOLO) تغییر می‌یابد. در مرحله بعد نوبت به آموزش مدل پیشنهادی می‌رسد. این مدل برپایه مدل YOLO طراحی شده است و به صورت یکپارچه عمل می‌کند. یعنی نیازی به آموزش جداگانه مدل‌ها نیست زیرا انتهای یک مدل به ابتدای مدل بعدی متصل شده است. با استفاده از این روش ترکیبی، تصاویر ورودی ابتدا به مدل EfficientNet-B0 وارد شده تا ویژگی‌های آن‌ها استخراج شوند. سپس، این ویژگی‌ها به YOLO-V11 ارسال می‌شوند تا اشیای موجود در تصاویر و جعبه‌های محدودکننده اطراف آن‌ها را شناسایی و رسم کند. برای حذف جعبه‌های اضافی و بهبود دقت، از بلوک سرکوب غیر حداکثری^۱ استفاده می‌شود. درنهایت، تصویر نهایی به همراه اشیای شناسایی شده به کاربر نمایش داده می‌شود و فرآیند اجرای مدل پایان می‌یابد.

۴-۲ نتایج ارزیابی مدل

شکل (۸) نتایج ارزیابی مدل ترکیبی درازای مجموعه داده BVMMR را نشان می‌دهد. نمودارها به شرح زیر هستند:

^۱ Non-Maximum Suppression

^۲ Classification Loss

^۳ Distribution Focal Loss

^۴ True Positives

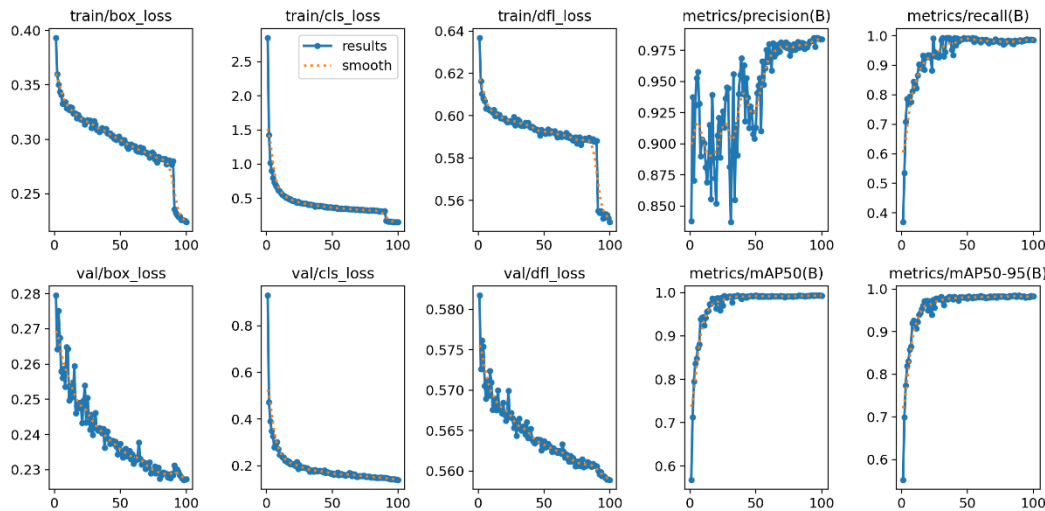
^۵ False Positives

دقت پیش‌بینی را برای تمام کلاس‌ها در نظر می‌گیرد و نشان می‌دهد که مدل چقدر در تشخیص اشیا موفق بوده است. شکل (۹) نحوه محاسبه IoU را نشان می‌دهد. درواقع IoU از تقسیم مساحت مشترک بین دوجعبه بر کل مساحت آن دوجعبه به‌دست می‌آید. مقدار میانگین دقت متوسط mAP از رابطه‌های (۳) و (۴) محاسبه می‌شود که در آن‌ها $AP^{(2)}$ میانگین دقت برای هرکلاس، $P(R)$ میزان دقت در یک مقدار مشخص از فراخوانی، n تعداد کل کلاس‌های یک مجموعه داده و i شماره هرکلاس هستند.

(۵) نمودار فراخوانی نشان می‌دهد که مدل چه درصدی از اشیا واقعی موجود در داده‌ها را شناسایی کرده است. مقدار فراخوانی از رابطه (۲) به‌دست می‌آید، که در آن TP تعداد اشیا درست شناسایی‌شده توسط مدل و FN^۱ تعداد اشیا اشتباهی است که مدل آن‌ها را تشخیص نداده است.

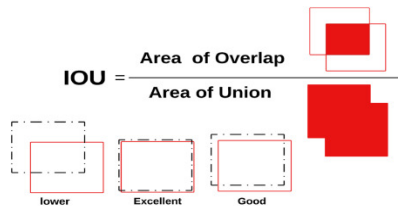
$$R = \frac{TP}{TP + FN} \quad (۲)$$

(۶) mAP50 میانگین دقت مدل در شناسایی اشیا را محاسبه می‌کند. به شرطی که میزان هم‌پوشانی بین جعبه پیش‌بینی‌شده و جعبه واقعی (IoU)^۲ حداقل ۰.۵۰ باشد. معیار mAP50



شکل (۸): نتایج ارزیابی مدل ترکیبی مورد استفاده

هرچه بالاتر باشد، یعنی مدل هم در شناسایی اشیا و هم در تعیین دقیق موقعیت و اندازه آن‌ها عملکرد بهتری دارد.



شکل (۹): نحوه محاسبه IoU [۲۶]

شکل (۱۰) رابطه بین دقت مدل و میزان اطمینان آن به پیش‌بینی‌های خود را نمایش می‌دهد. خط آبی ضخیم که میانگین دقت همه

$$AP = \int_0^1 P(R) dR \quad (۳)$$

$$mAP = \frac{1}{n} \sum_{i=1}^n AP_i \quad (۴)$$

(۷) $mAP50-95$ ^۴ میانگین دقت مدل را در یک بازه گسترده از IoU (از ۰.۵۰ تا ۰.۹۵ با گام‌های ۰.۰۵) محاسبه می‌کند. این معیار سخت‌گیرانه‌تر از mAP50 است، زیرا مدل نه تنها اشیا را باید درست تشخیص دهد، بلکه موقعیت و ابعاد آن‌ها را نیز باید با دقت بیشتری پیش‌بینی کند. مقدار mAP50-95

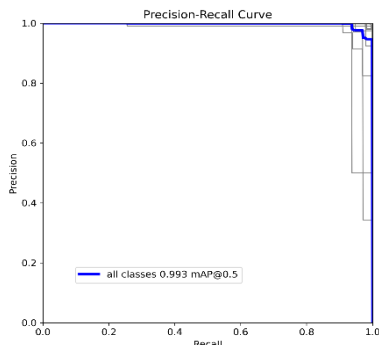
^۴ Mean Average Precision at IoU thresholds from 0.5 to 0.95

^۱ False Negatives

^۲ Intersection over Union (IoU)

^۳ Average Precision

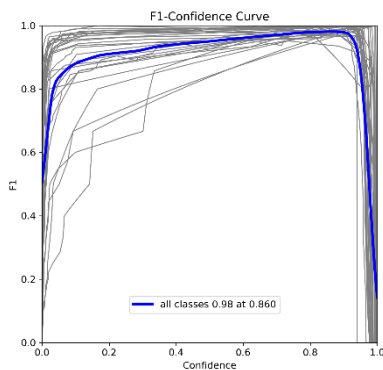
و محور عمودی، نسبت نمونه‌های مثبت شناسایی شده را نشان می‌دهند. این نمودار برای ارزیابی تعادل بین مقادیر دقت و فراخوانی استفاده می‌شود.



شکل (۱۲): نمودار دقت - فراخوانی

شکل (۱۳) رابطه بین مقدار F1 و اطمینان مدل به پیش‌بینی‌های خود را نمایش می‌دهد. محور افقی، سطح اطمینان مدل در پیش‌بینی‌ها و محور عمودی، میانگین هماهنگی بین دقت و فراخوانی را نشان می‌دهد. مقدار $F1 = 0.98$ در مقدار اطمینان 0.86 نشان می‌دهد که در این مقدار اطمینان، تعادل بین دقت و فراخوانی برقرار است. مقدار F1 از رابطه (۵) به دست می‌آید و در آن P دقت و R فراخوانی مدل است.

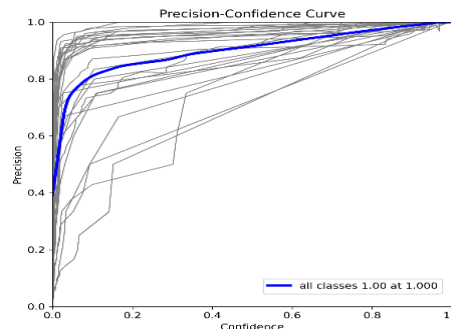
$$F1 = 2 \times \left(\frac{P \times R}{P + R} \right) \quad (5)$$



شکل (۱۳): نمودار F1 - اطمینان

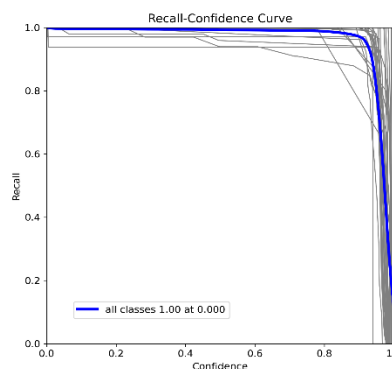
شکل (۱۴) ماتریس درهم‌ریختگی^۱ مدل را نشان می‌دهد. مقادیر روی قطر اصلی نشان‌دهنده دقت مدل در تشخیص صحیح هر کلاس هستند. هرچه مقدار آن‌ها به یک نزدیک‌تر باشند، مدل

کلاس‌ها را نشان می‌دهد، بیانگر این است که با افزایش اطمینان مدل، دقت آن نیز افزایش می‌یابد و در بازه‌ای بالاتر از 0.76 ، دقت به بیش از 95% می‌رسد. منحنی‌های خاکستری نیز عملکرد هر کلاس را نشان می‌دهند که با وجود نوسان در سطوح پایین میزان اطمینان مدل، روندی صعودی داشته و با افزایش سطح آن به مقدار قابل قبولی دست‌یافته‌اند.



شکل (۱۰): نمودار دقت - اطمینان

شکل (۱۱) رابطه بین فراخوانی و میزان اطمینان مدل به شناسایی نمونه‌های درست را نشان می‌دهد. درواقع این شکل نشان می‌دهد که مدل در شناسایی درست نمونه‌ها چقدر مطمئن است. منحنی آبی که میانگین عملکرد همه دسته‌ها را نشان می‌دهد، بیان می‌کند که مدل در بیشتر مواقع نمونه‌های درست را به خوبی شناسایی کرده و فقط در اطمینان‌های خیلی بالا کمی افت دارد. خطوط خاکستری هم مربوط به هر دسته به صورت جداگانه هستند که بیشتر آن‌ها روندی مشابه دارند و عملکرد خوبی را نشان می‌دهند.



شکل (۱۱): نمودار فراخوانی - اطمینان

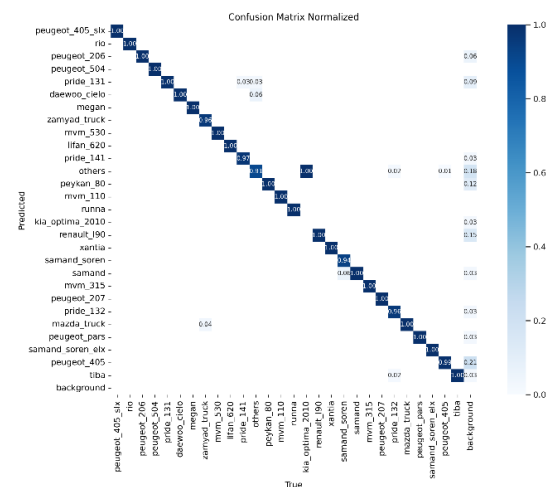
شکل (۱۲) رابطه بین دقت و فراخوانی را نمایش می‌دهد. محور افقی، درصد نمونه‌های مثبت واقعی که به درستی شناسایی شده‌اند

^۱ Confusion Matrix

و ... به همراه میزان اطمینان مدل (۱/۰، ۰/۹ و ...) را مشخص نموده است.

جدول (۱) نتایج روش پیشنهادی این مقاله را با جدیدترین روش‌های موجود مقایسه می‌کند. روش پیشنهادی، ترکیبی از مدل‌های EfficientNet-B0 و YOLO-V11 است. این مدل در معیار ارزیابی mAP50 دارای مقدار ۹۹/۳٪ و در معیار ارزیابی mAP50-95 دارای مقدار ۹۸/۳٪ است. این مقادیر نشان‌دهنده بهبود ۴/۱٪ در mAP50 و ۲۴/۱٪ در mAP50-95 نسبت به نسخه‌های استاندارد YOLO است. همچنین، دقت و فراخوانی این مدل به ترتیب ۹۸/۴٪ و ۹۸/۵٪ است، که بیانگر خطای کمتر و دقت بالاتر آن در مقایسه با سایر روش‌ها است. روش پیشنهادی با استفاده از مدل EfficientNet-B0 به عنوان ستون فقرات، ویژگی‌های تصویری را با جزئیات بیشتری استخراج کرده و دقت کلی تشخیص خودروها را بین ۳٪ تا ۷٪ و دقت تشخیص جزئیات آن‌را بیش از ۲۵٪ افزایش داده است. علاوه بر این، زمان پردازش این روش حدود ۵/۵ میلی‌ثانیه برای هر تصویر است و نشان می‌دهد که این مدل حداقل ۳۰٪ سریع‌تر از مدل‌های قبلی مانند YOLO-V9 و YOLO-V8 عمل می‌کند.

دقیق‌تر عمل می‌کند. مقادیر خارج از قطر اصلی نشان می‌دهند که ممکن است برخی از کلاس‌ها با کلاس‌های دیگر اشتباه گرفته شوند.



شکل (۱۴): ماتریس درهم‌ریختگی برای ارزیابی مدل ترکیبی

شکل (۱۵) تصویر خروجی مدل ترکیبی مورد استفاده در این مقاله را نشان می‌دهد. مشاهده می‌شود که مدل برای هر خودرو یک جعبه محدودکننده رسم کرده و بالای آن، نوع خودرو (پراید، پژو



شکل (۱۵): تصویر خروجی مدل ترکیبی مورد استفاده

باعث شد تا مدل ترکیبی در شناسایی خودروها دقیق‌تر عمل کند. نتایج آزمایش‌ها نشان داد که مدل پیشنهادی در معیار ارزیابی mAP50 به مقدار ۹۹/۳٪ و در معیار ارزیابی mAP50-95 به مقدار ۹۸/۳٪ رسیده که این مقادیر از مقادیر مشابه مدل‌های مورد مقایسه، بالاتر بوده است. همچنین، مقادیر دقت و فراخوانی مدل به ترتیب برابر با ۹۸/۴٪ و ۹۸/۵٪ بودند. این مقادیر نشان می‌دهند که مدل پیشنهادی نه تنها خودروها را به‌خوبی تشخیص می‌دهد، بلکه خطای کمتری نیز دارد. یکی از مزیت‌های مهم این مدل تعادل بین دقت و سرعت پردازش است. زمان پردازش هر تصویر در محیط کولی، تنها ۵/۵ میلی‌ثانیه است که شامل ۰/۲ میلی‌ثانیه برای پیش‌پردازش، ۲/۸ میلی‌ثانیه برای استنتاج و ۲/۵ میلی‌ثانیه برای پس‌پردازش می‌شود. این سرعت حداقل ۳۰٪ سریع‌تر از مدل‌های قبلی مانند YOLO-V8 و YOLO-V9 است. زمان پردازش مدل‌های قبلی حدود ۷ تا ۸ میلی‌ثانیه محاسبه شده است. علاوه بر این، استفاده از EfficientNet-B0 باعث شد که مدل پیچیدگی کمتری نسبت به مدل‌های مبتنی بر ترانسفورمرها و شبکه‌های عصبی گراف (GNN) داشته باشد و در شرایط مختلف محیطی عملکرد بهتری از خود نشان دهد. مدل پیشنهادی در مقایسه با مدل‌های دیگر مانند YOLO-V9، YOLO-V11، DGNN-YOLO و YOLO-V9 + CBAM، نه تنها در شناسایی کلی خودروها بهتر عمل کرده بلکه در تشخیص جزئیات و کلاس‌بندی آن‌ها نیز برتری قابل‌توجهی داشته است. به‌طورکلی، این مدل با دقت بالا، سرعت پردازش سریع و توانایی کار در شرایط متنوع، گزینه ایده‌آلی برای استفاده در سیستم‌های حمل‌ونقل هوشمند و نظارت شهری است.

سپاس‌گزاری

نویسندگان بر خود لازم می‌دانند تا مراتب تشکر و قدردانی خود را از عوامل محترم این نشریه اعلام نمایند

References

- [1] M. A. R. Alif, "Yolov11 for vehicle detection: Advancements, performance, and applications in

جدول (۱) نشان می‌دهد که روش پیشنهادی در مقایسه با مدل‌های پیچیده‌تر، در معیار mAP50-95 که سخت‌گیرانه‌ترین معیار برای ارزیابی میانگین دقت متوسط شبکه YOLO است، حداقل ۹٪ پیشرفت داشته است. این معیار نشان می‌دهد که روش پیشنهادی نسبت به سایر روش‌های ذکرشده در جدول، در شناسایی جزئیات تصاویر عملکرد بهتری داشته است. به‌طورکلی، این مدل با دقت بالا، سرعت پردازش بهینه و تعادل مناسب بین دقت و کارایی، گزینه‌ای ایده‌آل برای سیستم‌های حمل‌ونقل هوشمند و نظارت شهری است.

جدول (۱): مقایسه مدل پیشنهادی با برخی از مدل‌های ارائه‌شده

منبع	مدل پیشنهادی	mAP50 (%)	mAP50-95 (%)	Precision (%)	Recall (%)
*	مدل پیشنهادی این مقاله	۹۹/۳	۹۸/۳	۹۸/۴	۹۸/۵
[۲۷]	YOLOv11	۹۵/۲	۶۵/۷	۹۶/۹	۸۹/۸
[۲۸]	YOLOv9	۹۳/۴	۷۳/۲	۹۲/۱	۸۶/۲
[۲۹]	DGNN-YOLO	۷۸/۳۰	۶۴/۷۶	۸۳/۸۲	۶۸/۷۵
[۳]	YOLOv11	۹۵/۱	۸۵/۴	۹۵/۱	۸۹/۴
[۳۰]	YOLOv9+CBAM	۹۸/۴	۸۹/۲	۹۸/۷	۹۷/۳
	Co-DETR	۹۸/۸	۹۰/۱	۹۹/۳	۹۸/۵
[۳۱]	YOLOv8s	۹۷	۸۹/۹	۹۷/۳	۹۲/۱
[۳۲]	YOLOv7	۸۲/۲	۵۹/۶	۷۸/۲	۷۶/۹
[۳۳]	YOLOv7-RAR	۹۸/۳	۸۶/۱	-	-
[۳۴]	YOLOv8-FDD	۹۷/۸	۸۴/۹	-	-

۵- نتیجه‌گیری

در این مقاله، برای تشخیص نوع خودرو از یک مدل ترکیبی جدید استفاده شده است که عملکرد بهتری نسبت به مدل‌های دیگر دارد. این مدل ترکیبی شامل مدل‌های EfficientNet-B0 و YOLO-V11 است. از مدل EfficientNet-B0 به‌عنوان بخش اصلی مدل (ستون فقرات) استفاده شد و به مدل ترکیبی کمک کرد تا ویژگی‌های تصویری را با دقت بیشتری استخراج کند. این کار

intelligent transportation systems," *arXiv preprint arXiv:2410.22898*, 2024, doi: 10.48550/arXiv.2410.22898.

- [2] X. Yi, Q. Wang, Q. Liu, Y. Rui, and B. Ran, "Advances in vehicle re-identification techniques: A survey," *Neurocomputing*, vol. 614, p. 128745, 2025, doi: 10.1016/j.neucom.2024.128745.
- [3] M. Ashkanani, A. AlAjmi, A. Alhayyan, Z. Esmael, M. AlBedaiwi, and M. Nadeem, "A Self-Adaptive Traffic Signal System Integrating Real-Time Vehicle Detection and License Plate Recognition for Enhanced Traffic Management," *Inventions*, vol. 10, no. 1, p. 14, 2025, doi: 10.3390/inventions10010014.
- [4] J. Zheng and J. Ren, "Multi Self-Supervised Pre-Finetuned Transformer Fusion for Better Vehicle Detection," *IEEE Transactions on Automation Science and Engineering*, 2024, doi: 10.1109/TASE.2024.3374759.
- [5] J. Karangwa, J. Liu, and Z. Zeng, "Vehicle detection for autonomous driving: A review of algorithms and datasets," *IEEE Transactions on Intelligent Transportation Systems*, vol. 24, no. 11, pp. 11568-11594, 2023, doi: 10.1109/TITS.2023.3292278.
- [6] M. Bakirci, "Enhancing vehicle detection in intelligent transportation systems via autonomous UAV platform and YOLOv8 integration," *Applied Soft Computing*, vol. 164, p. 112015, 2024, doi: 10.1016/j.asoc.2024.112015.
- [7] M. O. Yusuf *et al.*, "Enhancing vehicle detection and tracking in UAV imagery: a pixel Labeling and particle filter approach," *IEEE Access*, 2024, doi: 10.1109/ACCESS.2024.3401253.
- [8] "Dataset for car detection on aerial photos applications." <https://github.com/jekhor/aerial-cars-dataset> (accessed 2018).
- [9] I. Bozcan and E. Kayacan, "Au-air: A multi-modal unmanned aerial vehicle dataset for low altitude traffic surveillance," in *2020 IEEE International Conference on Robotics and Automation (ICRA)*, 2020: IEEE, pp. 8504-8510, doi: 10.1109/ICRA40945.2020.9196845.
- [10] M. Zohaib, M. Asim, and M. ELAffendi, "Enhancing emergency vehicle detection: a deep learning approach with multimodal fusion," *Mathematics*, vol. 12, no. 10, p. 1514, 2024, doi: 10.3390/math12101514.
- [11] Z. Song, Y. Wang, S. Xu, P. Wang, and L. Liu, "Lightweight Vehicle Detection Based on Mamba_ViT," *Sensors*, vol. 24, no. 22, p. 7138, 2024, doi: 10.3390/s24227138.
- [12] Z. Wang and C. Ma, "Weak-mamba-unet: Visual mamba makes cnn and vit work better for scribble-based medical image segmentation," *arXiv preprint arXiv:2402.10887*, 2024, doi: 10.48550/arXiv.2402.10887.
- [13] L. Wen *et al.*, "UA-DETRAC: A new benchmark and protocol for multi-object detection and tracking," *Computer Vision and Image Understanding*, vol. 193, p. 102907, 2020, doi: 10.1016/j.cviu.2020.102907.
- [14] Y. Zhang, W. Wang, M. Ye, J. Yan, and R. Yang, "LGA-YOLO for Vehicle Detection in Remote Sensing Images," *IEEE Journal of Selected Topics in Applied Earth Observations and Remote Sensing*, 2025, doi: 10.1109/JSTARS.2025.3535090.
- [15] "data-unicorn-2008." <https://github.com/AFRL-RY/data-unicorn-2008> (accessed 2019).
- [16] S. Razakarivony and F. Jurie, "Vehicle detection in aerial imagery: A small target detection benchmark," *Journal of Visual Communication and Image Representation*, vol. 34, pp. 187-203, 2016, doi: 10.1016/j.jvcir.2015.11.002.
- [17] G.-S. Xia *et al.*, "DOTA: A large-scale dataset for object detection in aerial images," in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2018, pp. 3974-3983, doi: 10.1109/cvpr.2018.00418.
- [18] H. Park, K. Kim, I. Jeong, J. Jung, and J. Cho, "Special Vehicle Classification Algorithm-Based System for Dedicated Parking Zone Violation Detection in South Korea," *IEEE Access*, 2025, doi: 10.1109/ACCESS.2025.3526862.
- [19] M. Tan and Q. Le, "Efficientnet: Rethinking model scaling for convolutional neural networks," in *International conference on machine learning*, 2019: PMLR, pp. 6105-6114, doi: 10.48550/arXiv.1905.11946.
- [20] R. Khanam and M. Hussain, "Yolov11: An overview of the key architectural enhancements," *arXiv preprint arXiv:2410.17725*, 2024, doi: 10.48550/arXiv.2410.17725.
- [21] "YOLO11_EfficientNet." https://github.com/JYe9/YOLO11_EfficientNet/tree/main (accessed 2025).
- [22] N. Jegham, C. Y. Koh, M. Abdelatti, and A. Hendawi, "Evaluating the evolution of yolo (you only look once) models: A comprehensive benchmark study of yolo11 and its predecessors," *arXiv preprint arXiv:2411.00201*, 2024, doi: 10.48550/arXiv.2411.00201.
- [23] J. Huang, K. Wang, Y. Hou, and J. Wang, "LW-YOLO11: A Lightweight Arbitrary-Oriented Ship Detection Method Based on Improved YOLO11," *Sensors*, vol. 25, no. 1, p. 65, 2024, doi: 10.3390/s25010065.
- [24] S. U. Amin, M. S. Abbas, B. Kim, Y. Jung, and S. Seo, "Enhanced anomaly detection in pandemic surveillance videos: An attention approach with



- EfficientNet-B0 and CBAM integration," *IEEE Access*, 2024, doi: 10.1109/ACCESS.2024.3488797.
- [25] M. Biglari, A. Soleimani, and H. Hassanpour, "Part-based recognition of vehicle make and model," *IET image processing*, vol. 11, no. 7, pp. 483-491, 2017, doi: 10.1049/iet-ipr.2016.0969.
- [26] M. Y. Shams *et al.*, "Automated On-site Broiler Live Weight Estimation Through YOLO-Based Segmentation," *Smart Agricultural Technology*, p. 100828, 2025, doi: 10.1016/j.atech.2025.100828.
- [27] Y. Gao, Y. Xin, H. Yang, and Y. Wang, "A Lightweight Anti-Unmanned Aerial Vehicle Detection Method Based on Improved YOLOv11," *Drones*, vol. 9, no. 1, p. 11, 2024, doi: 10.3390/drones9010011.
- [28] S. A. Fahim, "Finetuning YOLOv9 for Vehicle Detection: Deep Learning for Intelligent Transportation Systems in Dhaka, Bangladesh," *arXiv preprint arXiv:2410.08230*, 2024, doi: 10.48550/arXiv.2410.08230.
- [29] S. Soudeep, M. Mridha, M. A. Jahin, and N. Dey, "DGNN-YOLO: Dynamic Graph Neural Networks with YOLO11 for Small Object Detection and Tracking in Traffic Surveillance," *arXiv preprint arXiv:2411.17251*, 2024, doi: 10.48550/arXiv.2411.17251.
- [30] E. Arthur, A. Aboah, and Y. Huang, "A Novel FHWA-Compliant Dataset for Granular Vehicle Detection and Classification," *IEEE Access*, 2024, doi: 10.1109/ACCESS.2024.3486603.
- [31] Q.-A. N. Duc *et al.*, "Optimizing Traffic Light Control using YOLOv8 for Real-Time Vehicle Detection and Traffic Density," in *2024 9th International Conference on Integrated Circuits, Design, and Verification (ICDV)*, 2024: IEEE, pp. 119-124, doi: 10.1109/ICDV61346.2024.10616901.
- [32] J. Su, F. Wang, and W. Zhuang, "An Improved YOLOv7 Tiny Algorithm for Vehicle and Pedestrian Detection with Occlusion in Autonomous Driving," *Chinese Journal of Electronics*, vol. 34, no. 1, pp. 282-294, 2025, doi: 10.23919/cje.2023.00.256.
- [33] Y. Zhang, Y. Sun, Z. Wang, and Y. Jiang, "YOLOv7-RAR for urban vehicle detection," *Sensors*, vol. 23, no. 4, p. 1801, 2023, doi: 10.3390/s23041801.
- [34] X. Liu, Y. Wang, D. Yu, and Z. Yuan, "YOLOv8-FDD: A real-time vehicle detection method based on improved YOLOv8," *IEEE Access*, 2024, doi: 10.1109/ACCESS.2024.3453298

Vehicle Models Recognition for Intelligent Control of Vehicle Traffic by Deep Learning

Morteza Ehsandoust¹, Atefeh Salimi Shahraki^{2*}, Mohammad Rohollah Yazdani³,
Mohammad Lali Dastjerdi¹

¹MSc. student Department of Electrical Engineering, Isfahan (Khorasgan) Branch, Islamic Azad University, Isfahan, Iran

²Assistant Professor Department of Electrical Engineering, Isfahan (Khorasgan) Branch, Islamic Azad University, Isfahan, Iran

³Associate Professor Department of Electrical Engineering, Isfahan (Khorasgan) Branch, Islamic Azad University, Isfahan, Iran

Article Information

Original Research Paper

Received:
2025 February 9

Accepted:
2025 April 26.

Keywords:
Computer Vision, Intelligent Transportation System, Deep Learning, YOLO, Efficient Net

Corresponding Author*:
a.salimi@khuisf.ac.ir

Abstract

With the increasing number of vehicles worldwide and the advancement of technology, intelligent transportation systems have emerged as a key solution for traffic control and road safety enhancement. One of the crucial components of these systems is vehicle type recognition. In this paper, a combination of EfficientNet-B0 and YOLO-V11 models is used for vehicle type classification. In this model, EfficientNet-B0 serves as the backbone of YOLO-V11, extracting image features and feeding them into the YOLO model to accurately determine the location and type of vehicles. To evaluate the performance of the proposed model, the BVMMR image dataset—containing over 5,000 of various Iranian vehicles—was utilized. The model was implemented using Python and executed in the Colab online environment. The results demonstrate that the proposed model achieves a mean Average Precision at 50% IoU (mAP50) of 99.3% and a mean Average Precision across IoU thresholds from 50% to 95% (mAP50-95) of 98.3%. Additionally, in terms of processing speed, the model delivers an efficient performance with an images average preprocessing time of 0.2 milliseconds, inference time of 2.8 milliseconds, and post-processing time of 2.5 milliseconds per image.

 : 10.22034/ABMIR.2025.22898.1112

E-ISSN: [2821-2037](#) /© 2024. Published by Yazd University This is an open access article under the CC BY 4.0 License (<https://creativecommons.org/licenses/by/4.0/>).

