

A Statistical Evaluation Framework with Application to YOLO-Based Models in Cross-Language License Plate Recognition: Arabic and Latin Alphabets

Israa Lewaaelhamd^{1,*}, Ahmed Elaraby^{2,3}

¹*Department of Business Administration, Faculty of Business Administration, The British University in Egypt, Cairo, Egypt*

²*Department of Computer Science, Faculty of Computers and Information, QENA University, Qena 83523, Egypt*

³*Cybersecurity Department, Engineering and Information Technology College, Buraydah Private Colleges, Buraydah, 51418, Saudi Arabia*

Abstract This study presents a statistical evaluation framework with application to YOLO-based models in cross-language License Plate (LP) recognition. Rather than relying solely on conventional accuracy metrics, the proposed framework integrates statistical performance analysis with deep learning techniques to provide a comprehensive and objective comparison of object detection models. The framework employs multiple statistical evaluation measures, including precision, recall, mean Average Precision (mAP), confusion matrices, and confidence-based performance curves, to quantify detection accuracy, classification reliability, and model robustness. Experiments were conducted using an Iraqi license plate dataset comprising 1,834 annotated images containing Arabic and Latin license plate characters. Two state-of-the-art object detection models, YOLOv5 and YOLOv8, were trained and evaluated under identical experimental settings. The results demonstrate that YOLOv8 consistently outperformed YOLOv5 across all statistical evaluation measures, achieving higher detection accuracy and improved classification performance. The proposed statistical evaluation framework provides a robust and objective approach for assessing and comparing deep learning models. By integrating statistical evaluation principles with advanced object detection techniques, this framework offers an effective methodology for model evaluation and selection in cross-language license plate recognition and can be extended to various computer vision applications.

Keywords Statistical evaluation, Applications, Statistical performance analysis, Deep learning, YOLO-based models, License plate recognition, Model comparison.

AMS 2010 subject classifications 62M45, 68T07, 62F10, 62P25, 62F10.

DOI: 10.19139/soic-2310-5070-3407

1. Introduction

In the last decade, machine learning algorithms, specifically supervised learning, have shown promising results in classification or detection problems in various tasks of computer vision. Using machine learning or, more recently, deep learning, we don't need to specify directly which features are important to differentiate between the object and its local background. Instead, we let the algorithm decide which part of the image is the best feature to represent the object. Supervised learning specifically needs annotated data as a training set to learn a mapping function from input X to output Y , or informally, to learn a decision function to map an input to its correct output. This trend also answered a research gap in the old methods, where it was difficult to design a license plate recognition system for any different plate format and character pattern due to the language dependency of the handcrafted features.

*Correspondence to: israa.lewaa@bue.edu.eg; israalewaa@feps.edu.eg

The automatic feature selection and machine learning method will adapt directly to the different characteristics of each license plate and its characters [1]. The early history of recognition methods was dominated by handcrafted algorithms based on several image-processing steps. These methods require much effort to design and produce features that potentially differentiate between an object and its local background. The handcrafted features in those methods for any OCR system were obviously language-dependent. It was proven to be a disadvantage in the sense that one algorithm cannot be straightforwardly used to recognize the plate from different countries, as it uses different patterns of characters and different plate formats. Thus, for each different license plate, it requires a significant time to redesign the system. This weakness became the main consideration to develop a modern approach in the recognition system [2]. Arabic characters are essentially cursive and have a large quantity of similarities and few diacritics. These properties make it tough to recognize. A lot of the current strategies shape the detected plate characters with a created template. Independent character recognition is employed to get the characters in form. The hypothesis is that a character of a recognized plate acquiring a major difference in writing model from the one in the template may mistakenly be rejected. A feasibility study is carried out in [20] for a partial OCR on car plate images in Jordan. In the place, individual character segmentation and recognition are carried out on a limited character set. The segmentation is done using vertical edge detection and then clustering similar regions together. Several connected element analyses are then utilized to determine the location and shape of the characters. Although it is not using Arabic characters, the technique of detecting characters in the plate using the related method may be useful [3]. Creating this system with a high level of accuracy and large variations in Arabic character text recognition will make it a possibility to later invent and implement a system that can recognize Arabic text on various types of media, such as newspapers, street signs, and even product packaging. This will enable a high level of accessibility for Arabs living in foreign countries to read Arabic text in the form of translations, without having to depend on human translators to verbally translate for them [4]. In recent years, vehicle license plate recognition has become increasingly popular, driven by its wide range of applications in security, traffic management, and law enforcement. It plays a pivotal role in automatic monitoring, ensuring compliance with regulations on public roads, and controlling access to parking facilities. Additionally, VLPR finds utility in regulating the crossing of borders and enforcing vehicle speed limits. The recognition of a license plate involves several steps, including image acquisition, LP extraction, segmentation, and recognition. Image acquisition serves as the initial stage in a VLPR system [1]. A license plate is a rectangular metal plate affixed to a vehicle, featuring numbers, characters, and words that serve as a unique identifier. Since the license plate appears as a distinct region of intensity, an effective identifying method is required to extract the captured images for license plate recognition. Consequently, various methods of deep learning approaches have been employed to identify vehicle license plates, which are explored in our paper [5, 6, 7]. The literature contains several studies focusing on license plate recognition systems, considering the diverse characteristics of LPs which vary from country to country. The Latin and Arabic alphabets used on LPs in Arab nations are especially distinctive, hence thorough research on the identification processes for these plates is required [3]. However, there have been limited studies employing deep learning methods specifically for Arab nations' license plate detection [8, 9]. Earlier works have also overlooked the inclusion of all governorates' worth of Iraqi plates. This omission can be attributed to the wide range of Iraqi plate types and variations in background color, including white, red, or yellow. To address this gap, this paper considers both Iraqi and Malaysian LPs. While many researches conducted on LPs in Malaysia [10, 11, 12, 13, 14, 15], there is a scarcity of research on Iraqi LPs utilizing conventional machine learning methods. Furthermore, no prior research has explored the use of deep learning techniques for LP recognition across all regions of Iraq. It is worth mentioning that a recent study has focused on LP recognition in the northern part of Iraq exclusively. In this study, license plate numbers based on Arabic and Latin were used to evaluate the validity of the suggested method. Deep learning was employed as a means to address this problem, providing a solution that was compared against the outcomes of two traditional machine learning techniques [37, 38, 39, 40]. The approach presented in this study is capable of handling license plates with either Arabic or Latin alphabets. Deep learning, unlike traditional methods, eliminates the need for manually engineered features as inputs, instead employing an end-to-end paradigm that automatically learns to segment numbers from the license plate. The following description outlines the general methodology employed in this study: Firstly, the recognition system is implemented using Python. This involves a series of steps. The initial step focuses on identifying the License Plate (LP). Subsequently, the license plate is segmented, meaning that the

letters, numbers, and words within the LP area are separated. Finally, a deep learning technique is employed to contrast the results obtained from the license plates with the outcomes generated by another model.

This paper is organized as follows: Section 1 presents the introduction of the study. Section 2 provides a review of related work on license plate recognition and vehicle identification. Section 3 describes the materials and methods, including the dataset, preprocessing steps, and the proposed experimental methodology. Section 4 presents the statistical evaluation framework and discusses the experimental results and performance analysis of the YOLO-based models. Finally, Section 5 concludes the paper.

2. Related Works

Since license plate recognition and vehicle identification are the foundational technologies of intelligent transportation systems, they have attracted a lot of attention in recent years. To do this, several studies in the literature have employed deep learning-based methods. Wang et al. [16] provided a thorough analysis of several DL-based vehicle re-identification methods. These methods were categorized by the authors into five groups. Wang et al. investigated the difficulties and future directions for vehicle re-identification studies while contrasting various methods. An evaluation of deep learning-based models for automatic vehicles that use vision identification (VAVR) was carried out by Boukerche and Ma [17]. They outlined the key issues and current research directions, provided an overview of the various vehicle identification datasets used in VAVR, and enumerated the features of VAVR techniques. A VAVR method was presented by Llorca et al. [18] and was based on modelling the look and shape of vehicle insignia using rearview photos. They employed a linear SVM classifier using HOG features. On a tiny sample of 1,342 images with 52 distinct car models, they achieved an accuracy of 93.75%. Conversely, Shashirangana et al. [19] investigated automatic license plate recognition (ALPR) techniques. They discussed current methods and listed the unresolved issues that the research and development community is now facing. The deep learning techniques employed in ALPR were divided into two categories by Shashirangana et al. [19], which are single-stage and multi-stage object identification DL-based techniques. Current research has been done on the identification of vehicles and the recognition of license plates. "PROVID," a deep learning method for progressive vehicle re-identification, was presented by Liu et al. [20]. Two progressive search procedures were used by the writers. A CNN model is employed to extract appearance features using a coarse-to-fine search, while a near-to-distant search is utilized for Siamese neural network-based license plate verification. Liu et al. [20] created the VeRi-776 dataset using urban surveillance footage. Three processes make up the DL-based LP detection and identification system that Selmi et al. [21] developed detection, segmentation, and character recognition. Before implementing the three earlier phases, a number of morphological operations are conducted, including geometric filtering, adaptive thresholding, and fine contours. Plate/non-plate classes are used to categorize the identified items in LP detection, which is founded on the CNN model. Then, another CNN model with 37 classes is used to recognize numerals 0–9 and characters in upper case format (A–Z). The issue of multilingual LP detection and recognition was the main focus of the research of Kessentini et al. [22]. Using YOLO proposed by Redmon and Farhadi [23], Kessentini et al. [22] suggested a technique to make the prediction based on two rounds of deep learning: (1) LP recognition on cropped photos and (2) LP detection on raw images using YOLOv2. They conducted a comparison between two recognition engines: a combined recognition strategy at the plate component level and a convolutional recurrent neural network for completing LP recognition without previous segmentation. The Tunisian LP setting was utilized to apply the Kessentini et al. technique. Using YOLO created by Redmon et al. [24], Hendry et al. [25] tackled the issue of automobile LP detection. They modified the original YOLO to create a 36-model single-class detector.

In 2003, Sarfraz et al. [26] presented the first study on the subject of Saudi Licence Plate automatic recognition. They created a dataset with 610 cars in various lighting scenarios. First, vertical edge detection, matching, and filtering are used to find the Saudi Licence Plate. Next, each character is separated by extracting the characters by counting a number of black pixels in each pixel column. Finally, template matching is used to recognize each character once it has been normalized. The Hamming distance is used to associate a character with a class. The technique is now seriously out of date. Furthermore, the format of the license plate that is being utilized does not

correspond to the present design of Saudi Licence Plate. Zidouri et al. [27] suggested an additional pipeline. The Sobel filter's intermediate uses vertical edge detection to find the SLP. A thickening mask and a set of predetermined parameters about the SLP position in the picture are then added to improve the detection. Pixel counting is the foundation for character segmentation in each column. Ultimately, multilayer perceptrons in neural networks are employed to identify the segmented characters. Using 61 plates, they evaluated their procedure and were able to identify 97% of them. Recently, several new studies have concentrated on SLP with DL-based methods. A method for automated LP identification based on DL in an unrestricted environment was described by Khan et al. [28]. The authors used a cell phone to record actual traffic footage to create their collection of still photos. The CNN model is utilized to identify the detected alphanumeric, whereas YOLOv5 is used to detect the LP. Khan et al. [28] only concentrated on English text and ignored Arabic letters in the LP to adapt their approach to other datasets. As a result, they did not utilize all of the data included in Saudi Licence Plates. Approach based on deep learning for the detection and identification of Saudi Licence Plate that was presented by Driss et al. [29]. They created a CNN for LP identification. 1150 images of Saudi cars taken in a variety of lighting and weather situations were used for the experiments. The accuracy of the suggested approach was 98% for LP identification and 92% for LP detection.

3. Materials and Methods

We proposed an approach that utilized YOLO model which introduced by Redmon et al. in [24] in 2016. In their study, they introduce a unique methodology known as the Darknet, which serves as the basis for a series of the greatest networks based on real time object detection available today, including YOLOv2 [23], YOLOv3 [30], YOLOv4 [31], YOLOv5, YOLOv7 [32], and its most recent iteration, YOLOv8 [24]. Interestingly, Region-Based Convolutional Neural Networks (R-CNN) were used for object recognition prior to the development of YOLO models. CNNs with architectures like Darknet or FPN (Feature Pyramid Network) are used to extract features from the input image. The detection head receives the output from this layer and uses it to predict the class probabilities and coordinates of the objects in the picture. The model's overall accuracy is increased. After the detection head has predicted the bounding boxes and class probabilities, identical detections are removed using a non-maximum suppression (NMS) technique. The final output of the YOLO model consists of a bounding box set, class probabilities, and objectness scores for each identified object in the image. The detecting head, backbone network, and additional innovations that increase object detection speed and accuracy—like anchor boxes—have all contributed to the evolution of the YOLO models.

3.1. Car and License Plate Detection

Identifying vehicles and license plates in every frame is the first step. In order to achieve this goal, a dataset of 1834 images of cars with license plates in Iraqi, which is available in the study proposed by Mohammed and Abdulsada [33], then manually labeling them. Often, when a car is photographed from the side, its license plate is obscured by other vehicles or other objects. 10% is set aside for validation and 90% is used for training in this dataset. The number of images in the training, validation, and testing datasets is displayed in Table 1.

Using this dataset, a YOLOv8 model was trained with an input size of 640. YOLOv8 was selected because, in contrast to alternative object detectors like YOLOv3 by Redmon and Farhadi [23], Faster R-CNN by Ren [34], or EfficientDet by Ten et al. [35], it offers a decent trade-off between precision and inference time. We employed 20 pictures for the batch size, a learning rate of 0.000333, and transfer learning from the COCO dataset based on Lin et al. [36]. A sample of car plates is shown in Figure 1.

Table 1. Number of images in the training and validation datasets

Item	Training Set	Validation Set	Testing Set	Total
Numbers	1605	153	76	1834

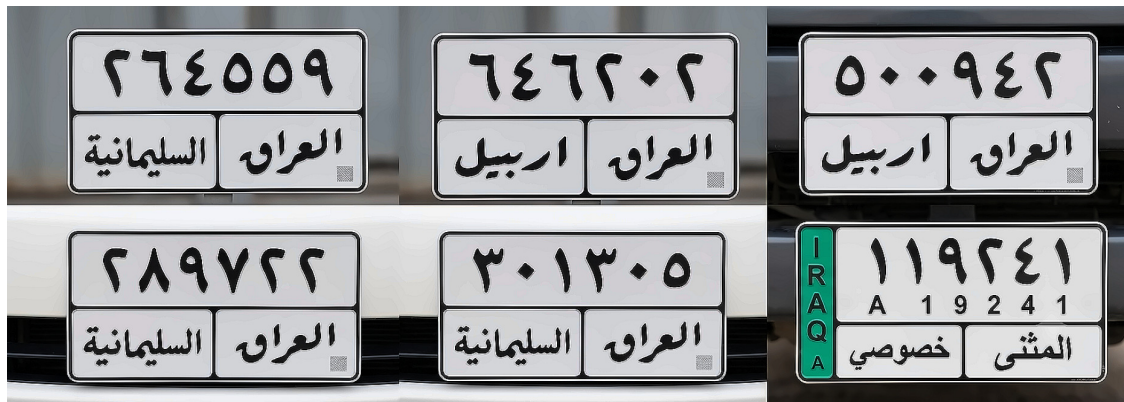


Figure 1. Sample of car plate

3.2. Car Model Recognition

A dataset of 1834 images was created to categorize car models. graphs were gathered from the Internet, and mobile phone shots were included, and the labels were manually added using the make-model-generation pattern. Three sets of this dataset were created: training, validation, and testing. Since we are using somewhat big datasets, we decided to employ a single random train-validation-test split for the other item detection and classification issue in addition to this car model classification problem's challenges that are explained below. However, for small datasets, k-fold cross-validation techniques are better. Furthermore, the image datasets used for testing and validation are just being used to provide a rough estimate of the models' accuracy.

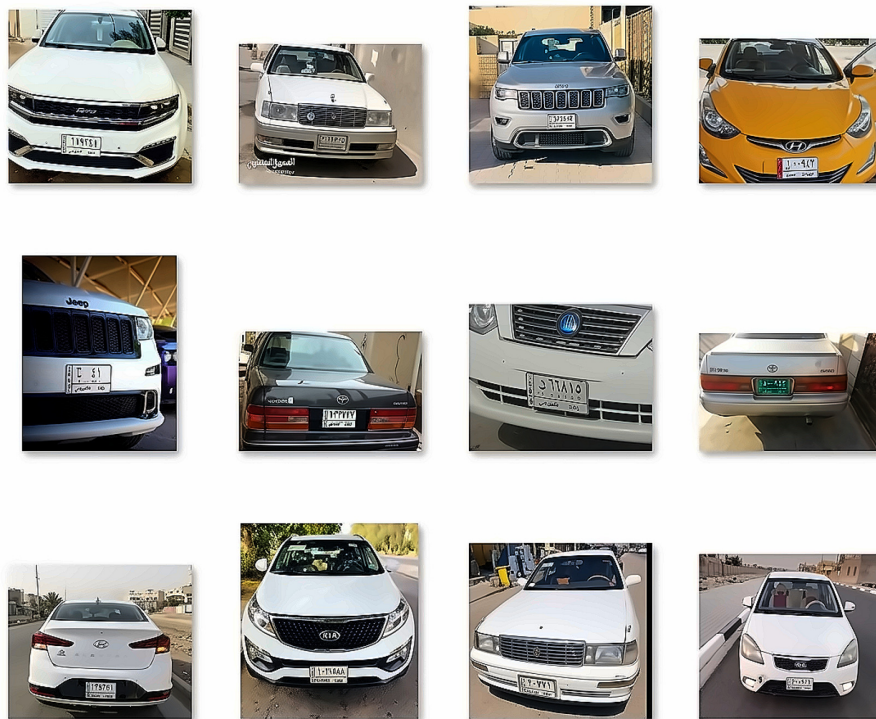


Figure 2. Sample images of the dataset used for car model classification

Figure 2 shows a selection of photos from the dataset used to categorize car models. A collection of 1834 images was created to categorize car models. Images were gathered from the Internet, mobile phone images were added, and the labels were manually added using the make-model-generation pattern. Three sets of this dataset were created: training, validation, and testing. Since we are using relatively large datasets, we decided to use a single random train-validation-test split for the other object detection and classification issues, including the car model classification problem that is detailed below. However, for small datasets, k-fold cross-validation techniques are better. Furthermore, the image datasets used for testing and validation are just being used to provide a rough estimate of the model's accuracy. Table 2 displays hyperparameters that were utilized to train the model. Following 75 training epochs, the model achieved 94.4% accuracy and 95.9% precision.

Table 2. Hyperparameters used for the training of the car model classifier for YOLOv8

Training Hyperparameter	Value
Image input size	640
Batch size	10
Learning rate	0.000333
Optimizer	Adam
Epochs	75

4. Statistical Evaluation Framework with Application to YOLO-Based Models

This section presents the statistical evaluation framework adopted to assess the performance of YOLO-based models for cross-language license plate recognition. The framework integrates deep learning architectures with statistical assessment techniques to provide a systematic and quantitative comparison of model performance. First, the architectural characteristics of the YOLOv5 and YOLOv8 models are introduced. Next, the optimization strategies employed during model training are described. Finally, the statistical performance measures used to quantify and compare the predictive capabilities of the two models are presented.

4.1. YOLO-Based Model Architecture

In the proposed statistical modeling framework, the predictive performance of a deep learning model depends not only on the optimization strategy adopted during training but also on its underlying network architecture. The architectural design directly influences feature extraction, object localization, and classification performance, which are subsequently reflected in the statistical measures used to evaluate model accuracy and reliability. Therefore, understanding the structural characteristics of the adopted YOLO models is an essential step before conducting their statistical performance evaluation.

The object recognition system known as YOLOv5 (You Only Look Once version 5) uses a single deep neural network to forecast bounding boxes and class probabilities for items that are present in a picture. The Ultralytics method is an advancement over earlier iterations of the YOLO algorithm, which were also created to aid with real-time object recognition. After being trained on a sizable collection of photos, the algorithm analyses characteristics in the pictures at different sizes to determine which items they are. In order to do this, YOLOv5 extracts attributes via a series of convolutional layers from the image that was entered, and then uses these features to forecast class probabilities and bounding boxes. The input image is normalized and resized to a predetermined size. The images' features are extracted using a backbone network. The backbone of YOLOv5 is the Feature Pyramid Network (FPN). Additionally, the YOLOv5 module consists of a collection of convolutional layers that aid in the fusion of features from various backbone network layers. In YOLOv5, the bounding boxes, objectness scores, and class probabilities are all predicted by the head module. After a series of convolutional layers, it has a set of completely

connected layers using objectness ratings, class probabilities, and bounding box predictions by the head module; fully duplicate detections are eliminated, and the overall accuracy of the model is increased by the application of the non-maximum suppression (NMS) approach. As seen in Figure 3, YOLOv5 generates a collection of bounding boxes, objectness scores, and class probabilities for each object it detects in the image.

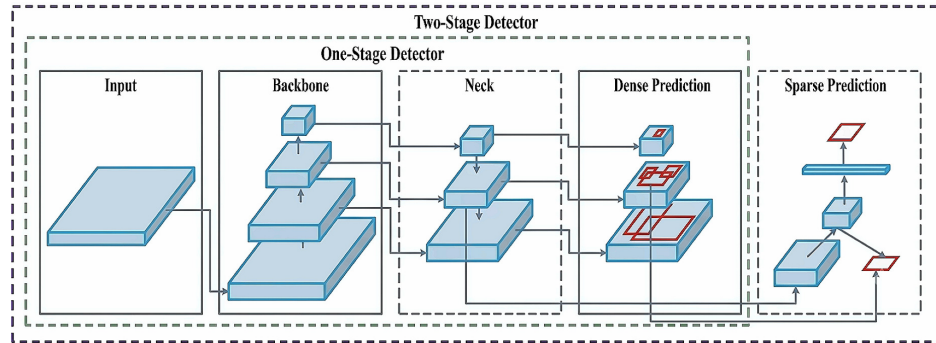


Figure 3. An illustration of the YOLO model with regard to the head pieces, neck, and backbone

The most recent model in the YOLO family, YOLOv8, is useful for instance segmentation, object recognition, and image classification. The YOLOv5 model's creator, Glenn Jocher, is also behind its development. Since the YOLOv8 model likewise has a FPN made up of head, neck, and backbone components, its operation is essentially the same as that of the YOLOv5. The modules are changed, though.

Still, the head module represented the biggest distinction between YOLOv5 and YOLOv8. In YOLOv8, the coupling structure that was initially present in YOLOv5 is replaced by a decoupling one. In addition, the YOLOv8 model does not rely on anchors, but the YOLOv5 model does. A predetermined collection of anchor boxes with different dimensions and aspect ratios is used by an anchor-based model, such as YOLOv5. The bounding boxes' size and placement in relation to these anchor boxes are predicted by the model. Next, the offset between the ground-truth boxes and the anchor boxes is used to modify the predicted bounding boxes. This technique aids the model in correctly identifying objects with varying sizes and aspect ratios. Conversely, an anchor-free model (such as YOLOv8) does not make use of anchor boxes. Rather, it forecasts the bounding box's size and centre point with precision. This approach gets rid of the requirement to define anchor boxes by hand and simplifies the model. Nevertheless, YOLOv5 and YOLOv8 carry out comparable tasks in terms of training methodologies. Since YOLOv8 is an anchor-free model, it forecasts an object's centre as opposed to its offset from a fully understood anchor box. This method reduces the quantity of box predictions, and as previously shown, this decrease results in a quicker NMS. YOLOv8's bottleneck architecture is the same as YOLOv5, but the kernel size of the first convolution is now 3×3 instead of 1×1 . Features are joined directly in the neck, without requiring the same channel dimensions, which results in a reduction of both the size of the relevant tensors and the overall number of parameters.

From a statistical perspective, the architectural differences between YOLOv5 and YOLOv8 influence the distribution of prediction errors, detection accuracy, and computational efficiency. These differences motivate the use of distinct optimization strategies during model training, which are presented in the following subsection.

4.2. Model Training and Optimization

The predictive performance of the adopted YOLO models depends on the optimization strategy employed during the training process. From a statistical learning perspective, optimization aims to estimate the network parameters by minimizing the loss function through iterative updates based on the observed training data. In this study, YOLOv5 was trained using the Stochastic Gradient Descent (SGD) optimizer, whereas YOLOv8 employed the Adam optimizer.

Optimization algorithms are employed during model training to minimize the loss function and improve the predictive accuracy of deep learning models. Depending on how the optimizer is formulated, it may be thought of as a mathematical function that modifies the network's weights given the gradients and further information.

Gradient descent, the greedy method of repeatedly lowering the loss function by following the gradient, is the foundation upon which optimizers are based. These functions might be quite complicated or as basic as deducting the gradients from the weights. In addition to being quicker and more effective than others, better optimizers are also frequently recognized for their ability to generalize well (i.e., with less overfitting). The model's performance can be significantly impacted by the optimizer selection. Stochastic gradient descent (SGD) specifically refers to the instance of vanilla GD with a batch size of 1. The most fundamental type of GD is SGD. SGD deducts from the weights the gradient multiplied by the learning rate. Despite its simplicity, SGD has strong theoretical foundations and provides an effective optimization strategy for estimating the parameters of deep neural networks through iterative minimization of the loss function. The parameter update rule for the SGD optimizer is given by Equation (1).

$$\theta_i = \theta_i - \alpha \frac{\partial L}{\partial \theta_i} \quad (1)$$

where L denotes the loss function used to measure the difference between the model predictions and the actual target values.

Since momentum is theoretically equivalent to gaining velocity, it is sometimes referred to as rolling down a ball. A momentum term, which is computed as the moving average of gradients, modifies the weights. The momentum term γ can be understood as friction or air resistance that causes the momentum to decrease proportionally. Momentum adds an extra hyperparameter but speeds up the training process. To improve convergence and reduce oscillations during optimization, a momentum term is incorporated into the update process according to Equations (2) and (3).

$$v_i = \gamma v_i + \alpha \frac{\partial L}{\partial \theta_i} \quad (2)$$

where

$$\theta_i = \theta_i - v_i \quad (3)$$

By storing both the weighted average of momentum and the individual learning rate of RMSProp, Adam effectively blends the two concepts. In essence, RMSProp is comparable to momentum. The values of v_i will rise and the learning rate will fall if the gradients are continuously big. This allows the use of higher learning rates and adaptively modifies the learning rate for every parameter. The RMSProp optimizer adaptively adjusts the learning rate for each parameter based on the moving average of squared gradients. The corresponding update equations are given in Equations (4) and (5).

$$v_i = \beta v_i + (1 - \beta) \left(\frac{\partial L}{\partial \theta_i} \right)^2 \quad (4)$$

where

$$\theta_i = \theta_i - \alpha \frac{\frac{\partial L}{\partial \theta_i}}{\sqrt{v_i} + \epsilon} \quad (5)$$

The first- and second-order moment estimates used in the Adam optimizer are computed according to Equations (6) and (7).

$$m_i = \beta_1 m_i + (1 - \beta_1) \frac{\partial L}{\partial \theta_i} \quad (6)$$

$$v_i = \beta_2 v_i + (1 - \beta_2) \left(\frac{\partial L}{\partial \theta_i} \right)^2 \quad (7)$$

To eliminate the initialization bias of the moving averages, bias-corrected estimates are computed using Equations (8) and (9).

$$\hat{m}_i = \frac{m_i}{1 - \beta_1} \quad (8)$$

$$\hat{v}_i = \frac{v_i}{1 - \beta_2} \quad (9)$$

Finally, the network parameters are updated according to Equation (10).

$$\theta_i = \theta_i - \frac{\alpha}{\sqrt{\hat{v}_i} + \epsilon} \hat{m}_i \quad (10)$$

Adam is based on RMSProp, much as in the previous equation, but in order to increase training speed, it estimates the gradient as the momentum parameter.

Adam combines adaptive learning rates with momentum-based optimization, resulting in faster convergence and improved optimization stability. These characteristics have made Adam one of the most widely adopted optimization algorithms in deep learning applications.

After defining the optimization algorithms, the proposed models were trained and evaluated using a labeled dataset of vehicle license plate images. The following experimental setup summarizes the dataset preparation, training configuration, and optimization settings adopted for both YOLOv5 and YOLOv8 models.

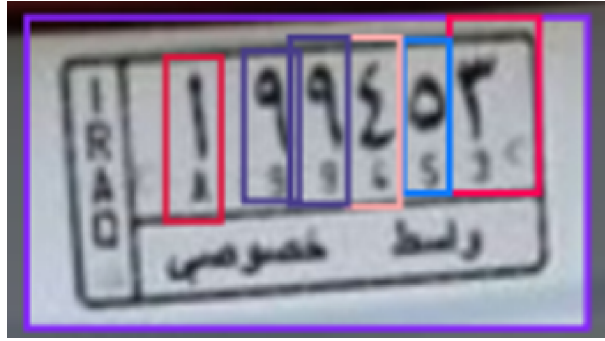


Figure 4. A license plate example using single-character detection



Figure 5. A detection for the license plate for the test image of YOLOv8

Sample images illustrating the license plate detection process are presented in Figures 4 and 5. These examples demonstrate the annotation and detection outputs obtained during the training and evaluation stages.

We first evaluated our solution on a set of 1834 images collected from Iraqi car images. We manually labeled images with Roboflow bounding box for each car license. Roboflow [26] provides a network called "Object

Detection" that offers a simple solution to the problem of training and deploying models. The actual training is carried out using Google Cloud AutoML. The length of training can change based on the number of images and annotations. It is important to note that Roboflow automatically selects the batch size and epoch number and that training ends when the model reaches an overfitting point, which in some circumstances may make the process less error-prone. The public cannot access the algorithmic details of the network, though, as it is a proprietary product. The original dataset consisted of 765 images. Data augmentation techniques were applied using Roboflow, resulting in an expanded dataset of 1,834 images to increase the number of images by performing processing (auto-oriented) with a resize of 640x640 and Augmentation of 90-rotate. The brightness is between -15% +15% and finally results in a dataset generation of 1834. License Plate Recognition is done using license plate recognition approaches by evaluating them on the set of images dataset as shown in Figures 4 and 5. The two approaches, YOLOv5 and YOLOv8, use 26 classes.

To measure the influence of the optimizer of each approach on the accuracy and inference speed, the authors also used the first model approach, yolov8 with auto Adam optimizer and a learning rate of 0.01 and momentum=0.937. The best optimizer, learning rate, and momentum are determined automatically. Then, Adam optimizer makes early stopping after 75 epochs with a learning rate of 0.000333 and momentum=0.9.

4.3. Statistical Model Performance Evaluation

From a statistical evaluation perspective, model comparison was conducted based on quantitative performance indicators derived from the detection outcomes. The trained YOLO models were evaluated using several statistical performance measures, including precision, recall, mean Average Precision (mAP), confusion matrices, and confidence-based evaluation curves. These measures provide complementary insights into detection accuracy, classification reliability, and the trade-off between precision and recall, thereby offering an objective basis for comparative analysis.

The comparative results of YOLOv5 and YOLOv8 are summarized in Table 3. The comparison considers detection performance measures, computational requirements, and model complexity. Overall, YOLOv8 achieved higher precision, recall, and mean Average Precision (mAP) values while requiring fewer model parameters and shorter training time, indicating improved detection performance and reduced computational requirements.

Table 3. Comparison of the Statistical Performance of YOLOv5 and YOLOv8

	YOLOv5	YOLOv8
Performance Metrics	Precision: 0.95 Recall: 0.905 mAP@50: 0.941 mAP@50–95: 0.441	Precision: 0.959 Recall: 0.944 mAP@50: 0.961 mAP@50–95: 0.461
Speed	Training time: 100 epochs completed in 0.811 hours.	Training time: 74 epochs completed in 0.754 hours Speed: 0.3ms preprocess, 2.5ms inference, 0.0ms loss, 5.0ms postprocess per image
Model summary	157 layers, 7080247 parameters, 0 gradients, 16.0 GFLOPs	168 layers, 3010718 parameters, 0 gradients, 8.1 GFLOPs

As shown in Table 3, YOLOv8 consistently outperformed YOLOv5 across the principal evaluation metrics while requiring fewer computational resources. These results indicate that YOLOv8 provides superior predictive performance together with improved computational efficiency.

To further investigate the predictive performance of the proposed models, confusion matrices and confidence-based evaluation curves were analyzed. When evaluating the effectiveness of a classification model, a confusion

matrix is commonly employed. It uses the test pictures to compare the predicted class label with the actual class label. Many performance indicators, including F1 score, accuracy, precision, and recall, may be calculated using confusion matrices. The confusion matrix for the YOLOv5 and YOLOv8 models is displayed in Figures 6 and 7, respectively.

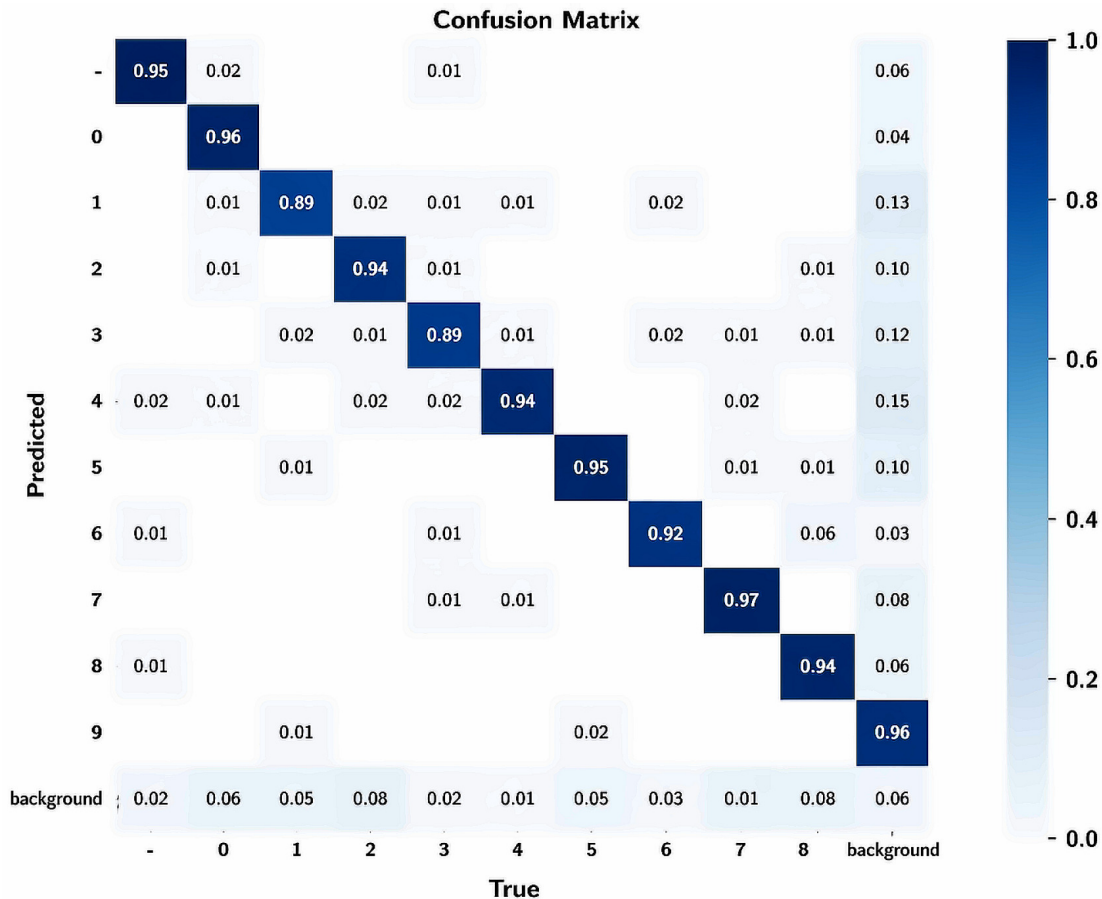


Figure 6. YOLOv5 model confusion matrix

The link between a classifier's confidence level and prediction precision is represented graphically by a precision-confidence curve. Visualizing the trade-off between a model's number of accurate predictions and number of inaccurate predictions at varying confidence levels is a common use in object recognition and classification tasks. The model's prediction quality at various confidence levels is represented by the curve. Although the classifier may be missing some real positive predictions, a high precision at a given threshold suggests that it is producing comparatively few false positive predictions at that confidence level. On the other hand, a poor accuracy at a given threshold suggests that although the classifier may be identifying more real positive predictions, it is simultaneously producing more false positive predictions. Figures 8–10 further confirm the superior predictive performance of YOLOv8, as evidenced by consistently higher precision and recall across different confidence thresholds.

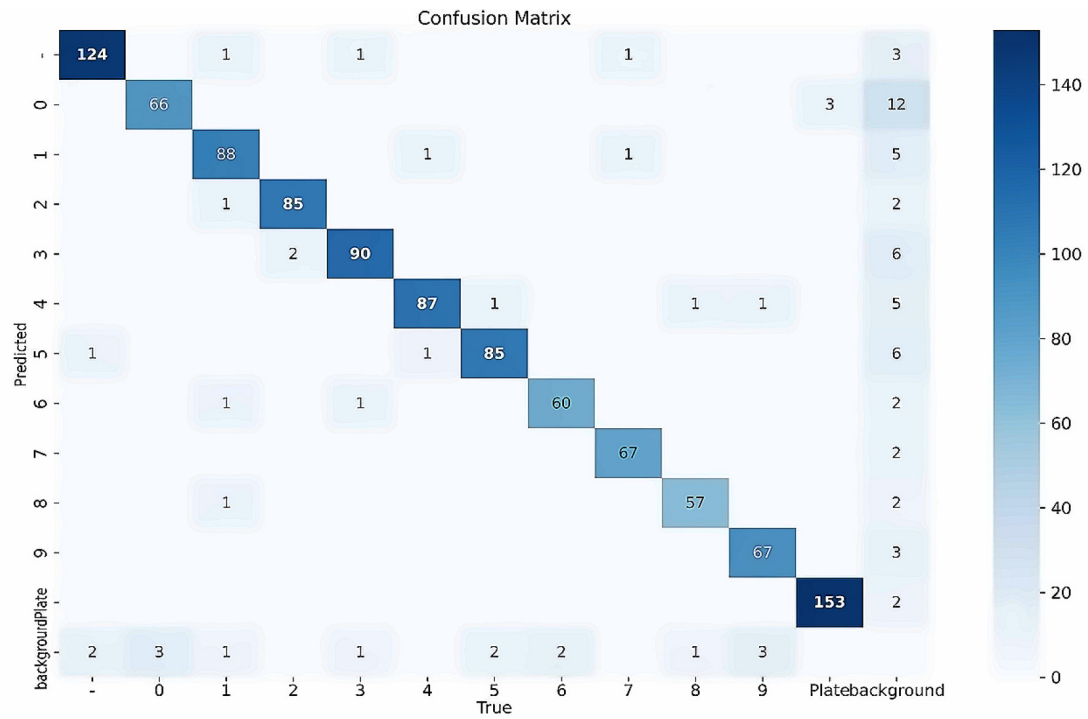


Figure 7. YOLOv8 model confusion matrix

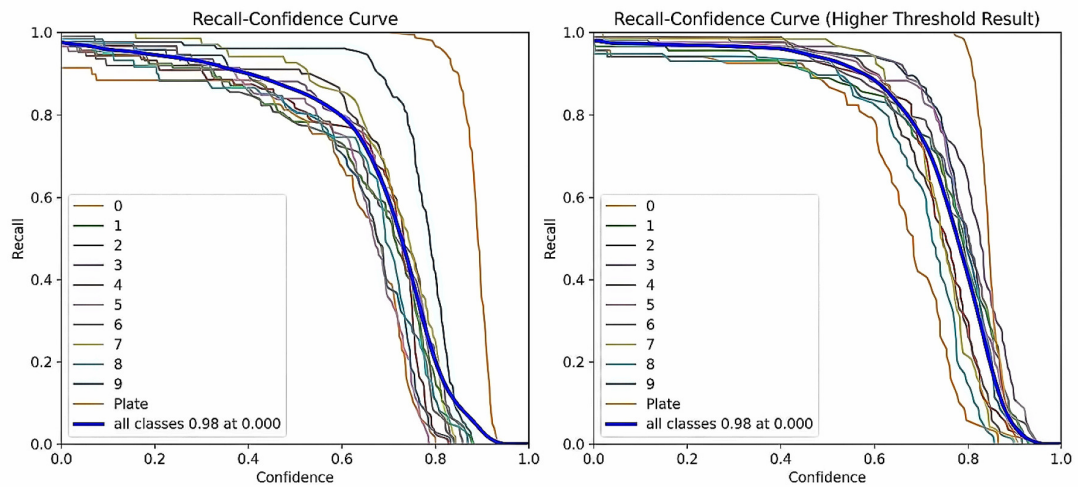


Figure 8. YOLOv5 and YOLOv8 Recall-Confidence curve

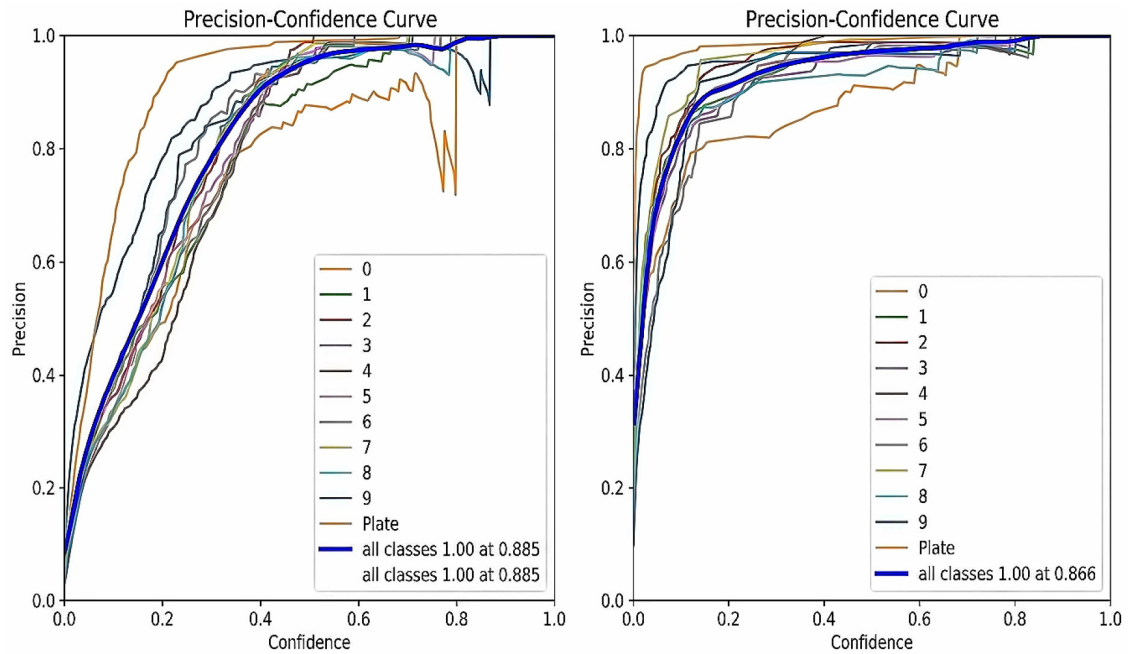


Figure 9. YOLOv5 and YOLOv8 Precision-Confidence curve

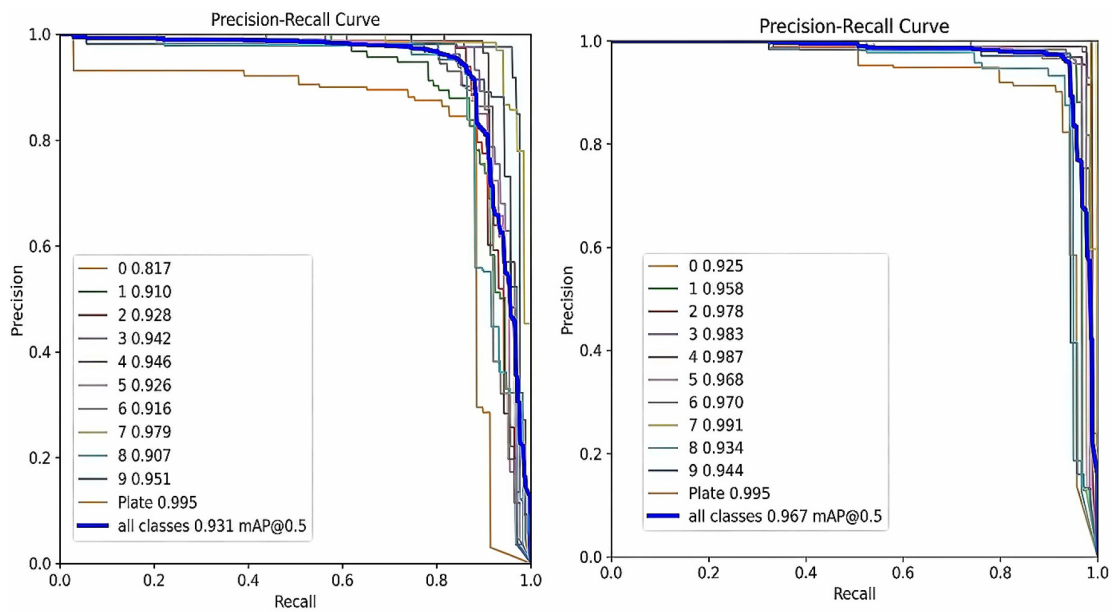


Figure 10. YOLOv5 and YOLOv8 Precision-Recall curve

5. Conclusion

This study introduced a statistical evaluation framework that was applied to YOLO-based models for cross-language license plate recognition, enabling a comprehensive and objective assessment of model performance. The proposed framework combines deep learning with statistical performance analysis to provide an objective comparison between alternative detection models.

Extensive experiments conducted on the Iraqi License Plate dataset demonstrated the effectiveness of the proposed framework, with YOLOv8 consistently outperforming YOLOv5 according to multiple statistical performance measures, including precision, recall, mean Average Precision (mAP), confusion matrices, and confidence-based evaluation curves. These findings demonstrate that employing a comprehensive statistical evaluation framework provides a more objective and reliable assessment of predictive performance than relying on a single accuracy measure. By integrating statistical modeling with deep learning, the proposed framework supports robust model comparison and facilitates informed model selection for cross-language license plate recognition tasks.

Although the proposed framework achieved promising results, further improvements can be obtained by training and validating the models on larger and more diverse datasets collected under unconstrained real-world conditions. Future research may also extend the proposed statistical evaluation framework to other object detection architectures and multilingual license plate recognition systems, enabling more comprehensive comparative analyses and enhancing the generalizability of the proposed methodology.

Data Availability Statement:

The data presented in this study are available on request from the corresponding author. The data are not publicly available due to privacy issues.

Conflicts of Interest:

The authors declare no conflict of interest.

REFERENCES

1. Sarfraz, M., Ahmed, M. J., and Ghazi, S. A. (2003). Saudi Arabian license plate recognition system. In *Proceedings of International Conference on Geometric Modeling and Graphics*, 36–41.
2. Omran, S. S. and Jarallah, J. A. (2017). Iraqi car license plate recognition using OCR. In *Proceedings of the 2nd IEEE Annual Conference on New Trends in Information & Communications Technology Applications*, 298–303.
3. Barnouti, N. H., Naser, M. A. S., and Al-Dabbagh, S. S. M. (2017). Automatic Iraqi license plate recognition system using BPNN. In *Proceedings of NTICT*, 105–110.
4. Zhang, Q. et al. (2018). Recent advances in convolutional neural network acceleration. *Neurocomputing*, **77**, 354–377.
5. Jagtap, J. and Holambe, S. (2018). Multi-style license plate recognition using ANN for Indian vehicles. In *ICICET*, 1–4.
6. Borah, P. and Gupta, D. (2017). Support vector machines in pattern recognition. *International Journal of Engineering & Technology*, **9**(3S), 43–48.
7. How, D. N. T. and Sahari, K. S. M. (2017). Character recognition of Malaysian license plates using CNN. In *International Symposium on Robotics and Intelligent Sensors*, 1–5.
8. Kessentini, Y. et al. (2019). A two-stage deep neural network for license plate detection and recognition. *Expert Systems with Applications*, **136**, 159–170.
9. Omar, N. et al. (2020). Cascaded deep learning-based approach for license plate detection. *Expert Systems with Applications*, **149**.
10. Connie, L. et al. (2018). Review of license plate recognition in mobile platforms. *Journal of Telecommunication, Electronic and Computer Engineering*, **10**(3-2), 77–82.
11. Yaacob, N. L. et al. (2021). License plate recognition for campus auto-gate system. *Indonesian Journal of Electrical Engineering*, **21**(1), 128–136.
12. Saleem, N. et al. (2016). Automatic license plate recognition using extracted features. In *IEEE ISCBI*, 221–225.
13. Babu, K. M. and Raghunadh, M. V. (2016). Vehicle number plate detection using bounding box method. In *IEEE ICACCCT*, 106–110.
14. Choong, Y. J. et al. (2020). License plate detection using simplified linear model. *Journal of Critical Reviews*, **7**(3), 55–60.
15. Omar, N. et al. (2020). Cascaded deep learning-based approach. *Expert Systems with Applications*, **149**, 20–25.
16. Wang, H., Hou, J., and Chen, N. (2019). Survey of vehicle re-identification. *IEEE Access*, **7**, 172443–172469.
17. Boukerche, A. and Ma, X. (2021). Vision-based autonomous vehicle recognition. *ACM Computing Surveys*, **54**, 1–37.
18. Llorca, D. F. et al. (2014). Vehicle model recognition using geometry. In *IEEE ITSC*, 3094–3099.
19. Shashirangana, J. et al. (2020). Automated license plate recognition survey. *IEEE Access*, **9**, 11203–11225.
20. Liu, X. et al. (2016). Deep learning-based vehicle re-identification. In *ECCV*, 869–884.
21. Selmi, Z. et al. (2017). Deep learning system for license plate recognition. In *ICDAR*, 1132–1138.
22. Kessentini, Y. et al. (2019). Two-stage deep neural network. *Expert Systems with Applications*, **136**, 159–170.

23. Redmon, J. and Farhadi, A. (2017). YOLO9000. In *CVPR*, 7263–7271.
24. Redmon, J. et al. (2016). YOLO: Unified real-time object detection. In *CVPR*, 779–788.
25. Hendry and Chen, R.-C. (2019). License plate recognition via YOLO. *Image and Vision Computing*, **87**, 47–56.
26. Sarfraz, M. et al. (2003). Saudi Arabian license plate recognition system. In *ICGMG*, 36–41.
27. Zidouri, A. and Deriche, M. (2008). Recognition of Arabic license plates. In *Image Processing Workshops*, 1–4.
28. Khan, I. R. et al. (2022). License plate recognition in real-world traffic videos. *Electronics*, **11**, 1408.
29. Driss, M. et al. (2022). Saudi license plate detection using CNN. In *International Conference on RICT*, 3–15.
30. Redmon, J. and Farhadi, A. (2018). YOLOv3: An incremental improvement. *arXiv preprint arXiv:1804.02767*.
31. Bochkovskiy, A. et al. (2020). YOLOv4. *arXiv preprint arXiv:2004.10934*.
32. Wang, C.-Y. et al. (2022). YOLOv7. *arXiv preprint arXiv:2207.02696*.
33. Mohammed, H. M. and Abdulsada, A. I. (2021). Secure multi-keyword similarity search. *Iraqi Journal for Electrical & Electronic Engineering*, **17**(2).
34. Ren, S. et al. (2015). Faster R-CNN. In *NeurIPS*, Volume 28.
35. Tan, M. et al. (2020). EfficientDet. In *CVPR*, 10781–10790.
36. Lin, T. Y. et al. (2014). Microsoft COCO dataset. In *ECCV*, 740–755.
37. Lewaaelhamd, I., & George Iskander, M. (2025). Poverty prediction using machine learning models: Insights from HICES survey in Egypt. *Statistics, Optimization & Information Computing*, **13**(4), 1623–1635.
38. Lewaaelhamd, I. (2024). Customer segmentation using machine learning model: An application of RFM analysis. *Journal of Data Science and Intelligent Systems*, **2**(1), 29–36.
39. George Iskander, M., & Lewaaelhamd, I. (2025). A computational study for probabilistic fuzzy linear programming using machine learning in the case of Poisson distribution. *Statistics, Optimization & Information Computing*, **14**(6), 3804–3815.
40. Lewaaelhamd, I., Hafez, M. S., & Ismail, M. A. (2021). Data integration using statistical matching techniques: A review. *Statistical Journal of the IAOS*, **37**(4), 1391–1410.