

International Journal of Applied Mathematics in Control Engineering

Journal homepage: <http://www.ijamce.com>

Detection and Identification of Substation Pointer Meter Based on Lightweight Yolov4 Model

Yuheng Song^{a,*}, Liangliang Sun^b, Mantong Zhao^c, Yue Wang^d^a School of Electrical and Control Engineering, Shenyang Jianzhu University, 110168, China^b School of Electrical and Control Engineering, Shenyang Jianzhu University, 110168, China^c School of Economics and Management, Shenyang Institute of Technology, 113000 China^d School of Sports, Shenyang Jianzhu University, 110168, China

ARTICLE INFO

Article history:

Received 10 September 2022

Accepted 24 December 2022

Available online 25 December 2022

Keywords:

Pointer instrument detection

YOLOv4

attention mechanism

mAP

ABSTRACT

Due to its simple structure, high reliability, small error and high efficiency, the pointer meter is widely used in the detection of the operating state of power equipment in substations in the process of industrial testing. In this paper, a target detection algorithm model based on deep learning is established based on the identification of pointer-type meters in the actual substation scene to achieve automatic detection and identification of pointer-type meters. Its specific design is that tensorflow is a large framework. On the basis of the YOLOv4 model, the deep separable convolution and multi-scale target detection network are used as the backbone feature extraction network, and the attention mechanism module is added. The leakyrelu activation function is used instead of the Mish activation function to achieve higher feature extraction speed with fewer modules and network layers, while retaining the spatial pooling pyramid (spp-net) feature layer fusion mechanism. Due to the particularity of instrument detection, this paper uses k-means clustering to optimize the original a priori frame, which has dealt with the shortcomings of the existing data set. Experiments show that the improved model greatly reduces the amount of system parameters, and its average accuracy The mean value (mAP) and detection speed (FPS) are both improved, and meet the actual detection requirements of the substation.

Published by Y.X.Union. All rights reserved.

1. Introduction

The substation is a key place in the power system to transform, concentrate and distribute the voltage and current of electric energy. There are a large number of instruments in it to complete the status monitoring and data measurement of industrial instruments. The pointer meter is widely used in the detection of the operating state of the power equipment in the industrial detection process due to its simple structure, high reliability, small error and high efficiency. At present, due to the high labor intensity, low work efficiency, scattered detection quality, and single means of manual inspection, there is a certain risk of missed inspection. Inspection robots are an important part of industrial robots. With the development and progress of industrialization, it has become the mainstream of the current era to replace humans with robots to complete various high-risk and difficult tasks. In the process of automatic instrument identification, due to the disturbance factors such as uneven illumination and noise interference in the substation environment, as well as the changing

environment of the instrument, most of the image information captured by the inspection robot is faced with unclear and multi-occluded pointer tilt characteristics. Inconspicuous and other problems, which caused the difficulty of shooting the image itself. At the same time, due to the variability of the instrument status of the substation, the detection algorithm used must meet the requirements of high-precision and high-speed detection at the same time, which cannot be achieved at the same time with the current detection capability. Therefore, it is of great practical significance to study an intelligent identification algorithm so that inspection robots can replace manpower to complete the inspection work required by industrial systems^[1].

At present, there are two schools of pointer instrument detection. One is the traditional machine vision detection method represented by the template matching method. First, the template matching is performed on the instrument panel, and then the outline of the dial area is extracted. The common area is the cumulative gradient method. The minimum the second method is to fit the straight-line method, SIFT and SUFT transformation method; the second is to

* Corresponding author.

E-mail addresses: sunliangliang@sjzu.edu.cn (L. Sun)

doi:

locate the dial outline and pointer through the target detection algorithm based on deep learning, and complete the pointer feature information extraction. The common target detection algorithm has two stages represented by Faster RCNN. Target detection algorithm and one stage target detection algorithm represented by YOLO.

In 2020, the YOLOv4 target detection algorithm will be launched, which is a masterpiece of target detection algorithms based on deep learning. In this paper, the target detection algorithm based on deep learning is mainly applied to the substation inspection robot instrument identification, which has important theoretical and practical significance, mainly in the following aspects^[2]:

1. Realize inspection robots instead of manual operations.

The disturbance factors caused by the substation environment are analyzed, taking into account the many disturbance factors of the image itself caused by illumination and noise. At the same time, considering the detection requirements required by the substation, the application of the inspection robot to the instrument detection operation can enable the substation to have safe, efficient and accurate instrument detection capabilities.

2. Provide driving force for the reform of robot enterprises

The target detection algorithm based on deep learning has been developed for only a few years, and many enterprises are still waiting to see its practical application. For inspection robots, most of them still use traditional identification and detection algorithms, and their actual efficiency is difficult to meet most enterprises. Therefore, robots urgently need a widely applicable and efficient detection algorithm. This paper can promote the development of the high-tech industry of robots. Even more complex object recognition provides a theoretical basis.

3. Provide a new innovative improvement scheme for the target detection algorithm based on deep learning.

The target detection algorithm is a key research problem in the field of computer vision. Its purpose is to identify the target object in the picture, and to determine its type and orientation, which is closely related to the development trend of the current social computer and artificial intelligence research. Although many researchers have improved and innovated, there is still a lot of room for improvement. In this paper, a new algorithm improvement scheme is proposed, which improves the shortcomings of the detection algorithm and achieves the coexistence of real-time detection and accuracy.

2. Convolutional Neural Network

Convolutional neural network is composed of input layer, convolution layer, activation function, pooling layer, and fully connected layer^[2]

2.1 Convolution

Convolution is an operation that uses convolution to check each pixel of the image. The elements of convolutional neural network are generally composed of input image, convolution kernel, stride, padding, channel change, and output image (feature map). The convolution kernel is the image processing, given the input image, the weighted average of the pixels in a small area in the input image becomes each corresponding pixel in the output image, where the weight is defined by a function, this function is called the convolution kernel.

The step size is the pixel unit that the convolution kernel moves each time. After the convolution operation, the output image size is reduced. The more pixels on the edge, the smaller the impact on the

output, because the convolution operation ends when it moves to the edge. The middle pixels may participate in multiple calculations, but the edge pixels may only participate once. So, our results may lose edge information, so we introduce padding to fill the pixels around the image pixels with 0.

From this, we get the operation method of convolution:

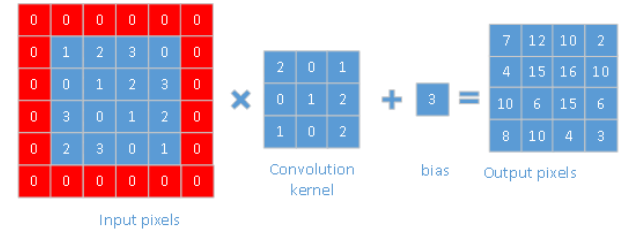


Fig.1 Convolution

2.2 pooling

Pooling is essentially down sampling. This is done in many different forms of nonlinear pooling functions, such as (Max pooling) being the most common. It divides the input image into several rectangular areas, and outputs the maximum value for each sub-area. The problem faced by pooling mechanisms is that after a feature is discovered, its precise location is far less important than its relative location to other features. The pooling layer will continuously reduce the size of the data space, so the number of parameters and the amount of computation will also decrease. Therefore, the pooling layer is periodically inserted between the convolutional layers of the neural network, which also prevents overfitting to a certain extent.

2.3 Activation Function

In the neural network layer, the input of each layer node is a linear function of the output of the upper layer. In order to increase the expressive ability and nonlinear ability of the system, we introduce an activation function. Early research neural networks mainly use sigmoid function or tanh function, the output is bounded, and it is easy to serve as the input of the next layer. The formula is as follows:

$$f(x) = \frac{1}{1 + e^{-x}} \quad (1)$$

$$\tanh = \frac{e^x - e^{-x}}{e^x + e^{-x}} \quad (2)$$

However, these two activation functions are easy to cause the problem of gradient explosion or gradient disappearance.

In recent years, the Relu function and its improvements (such as Leaky-ReLU, P-ReLU, R-ReLU, etc.) have been widely used in major artificial intelligence algorithms due to their fast convergence speed. Its formula is as follows:

$$\text{Relu} = \max(0, x) \quad (3)$$

Deep learning often requires a lot of time to process a large amount of data, and the convergence speed of the model is particularly important. Therefore, in general, training deep learning networks should try to use zero-centered data (which can be achieved through data preprocessing) and zero-centered output. Therefore, try to choose an activation function with zero-centered characteristics to speed up the convergence of the model.

2.4 Fully connected layers

The fully connected layer acts as a classifier in a neural network.

In actual use, the fully connected layer can be implemented by convolution operations: the fully connected layer that is fully connected to the front layer can be converted into a convolution kernel of 1*1 convolution; the fully connected layer whose front layer is a convolutional layer can be converted to a volume. The product kernel is the global convolution of hw, and hw is the height and width of the convolution result of the previous layer, respectively. The full connection layer is actually a matrix-vector product. Its structure is as follows:

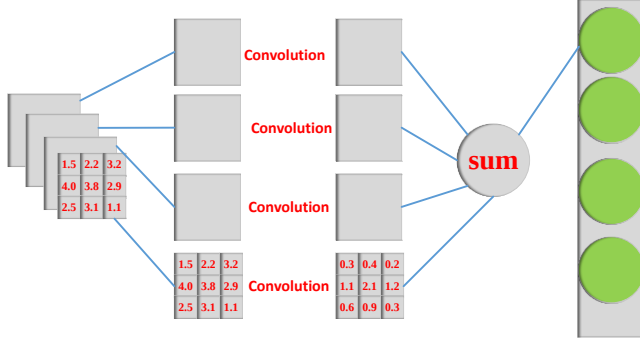


Fig.2 Fully connected layers

3. YOLOv4

2.1 Basic principles of YOLOv4

The YOLOv4 algorithm proposed by ochkovskiy et al. came out in 2020 and is a regression-based target detection algorithm [4]. Its network model proposes many innovative improvements based on the YOLOv3 network. Its network model consists of four parts, namely the input end, the Backbone benchmark network: the backbone feature network CSPDarkNet53, the Neck end SPP and PANet detection networks, and the three output end YOLO Head effective feature layers. It achieves the current optimum in balancing detection accuracy and speed, and the Pareto optimal curve in target detection is at the upper right [8].

2.1 Backbone

YOLOv4 uses CSPDarkNet53 as the feature extraction network, which consists of 5 resblock-body structures and a darknet module. The resblock-body is the CSPX residual module, and the CBM is composed of convolution + BN normalization + Mish activation function. Resunit - Drawing on the residual structure in the ResNet network, the double residual structure is shown in Figure 1. Compared with other activation functions, the Mish activation function is smoother, which makes the system have better nonlinear expression ability [9].

$$\text{Mish} = x \tanh \quad (4)$$

2.2 Neck

The Neck network of YOLOv4 consists of SPP and PAN network. SPP SPP-uses 1×1, 5×5, 9×9 and 13×13 maximum pooling methods to perform multi-scale feature fusion to achieve robust feature representation, which increases the perception field of the system. The PAN network is an innovation of the PANet algorithm in image segmentation and is a feature pyramid structure. By building a pyramid on the feature map, the target detection mesoscale problem

can be better solved [10]

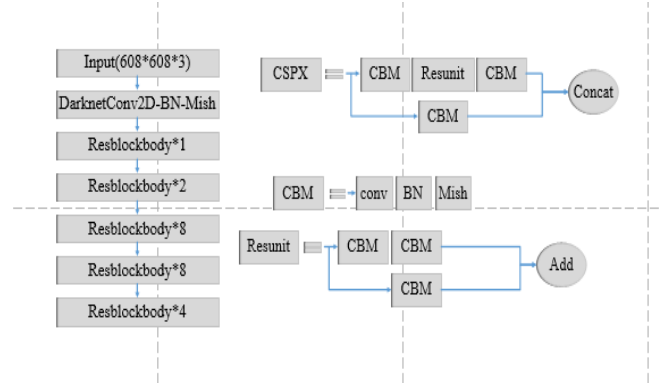


Fig.3 CSPDarkNet53 structure

2.3 loss function

The loss function of YOLOv4 is divided into three parts: bounding box regression (border regression) loss Loss_{iou} , confidence loss L_{conf} and classification loss L_{cla} [5]

The confidence loss L_{conf} adopts the cross entropy loss function:

$$\begin{aligned} L_{\text{conf}} = & \sum_{i=n}^{s^2} \sum_{j=0}^B I_{ij}^{\text{obj}} \left[\hat{c}_i \log_2(c_i^j) + (1 - \hat{c}_i^j) \log(1 - \hat{c}_i^j) \right] \\ & + \lambda_{\text{nobi}} \sum_{i=0}^{s^2} \sum_{j=0}^B I_{ij}^{\text{obj}} \left[\hat{c}_i \log_2(c_i^j) + (1 - \hat{c}_i^j) \log(1 - \hat{c}_i^j) \right] \end{aligned} \quad (4)$$

In the formula, I_{ij}^{obj} represents the jth prediction box of the ith grid is responsible for predicting the target. I_{ij}^{nobi} represents the box is not responsible for predicting this target. c_i^j represents the fitted value of the probability score of the object contained in the prediction box, $\hat{c}_i^j$ represents actual value. λ_{nobi} represents weight factor, which means the weight of the confidence error in the loss function.

The class loss also uses the cross entropy loss function:

$$L = \sum_{i=0}^{s^2} I_{ij}^{\text{obj}} \sum_{c \in \text{class}} [\hat{p}_i^j \log_2(\hat{p}_i^j) + (1 - \hat{p}_i^j)] \quad (5)$$

Classes is the target category. $\hat{p}_i^j$ represents The predicted probability of the jth forecaster in the ith grid. $\hat{p}_i^j$ represents the true value to which the box belongs, 1 otherwise 0.

The regression loss uses the mean square error function, which is expressed as follows:

$$\begin{aligned} L_{\text{loc}} = & \lambda_{\text{corr}} \sum_{i=0}^{s^2} \sum_j^B I_{ij}^{\text{obj}} [(x_i - \hat{x}_i)^2 \\ & + (y_i - \hat{y}_i)^2] + \sum_{i=0}^{s^2} \sum_j^B I_{ij}^{\text{obj}} [(\sqrt{w_i^j} - \sqrt{\hat{w}_i^j})^2 \\ & + (\sqrt{h_i^j} - \sqrt{\hat{h}_i^j})^2] \end{aligned} \quad (6)$$

In the formula I_{ij}^{obj} represents if the jth prediction box of the ith grid predict this target. λ_{corr} represents weight parameter. x, y, w, h represent the coordinates, width and height of the box which Considers the aspect ratio, overlap area, and center point distance between the predicted frame and the real frame, it better reflects the gap between the real value and the predicted value.

4. The improvements to YOLOv4

The pointer meter of the substation has the following two characteristics: the position of the target pointer in the image is uncertain, and the different sizes lead to different positions in the image; the information structure of the target pointer is relatively simple and there are few types.

The number of parameters of the YOLOv4 backbone network is roughly 60 million. For instrument detection, the system is a bit large. Therefore, the pursuit of light weight is the direction of improvement of the backbone network. Therefore, we need to establish a network depth based on the reduction of the detection accuracy. A detection system that meets the requirements.

3.1 The change of Backbone

In this paper, aiming at the identification of substation meters, aiming at light weight simplification without reducing the detection accuracy, the system residual structure, convolution block and attention mechanism are improved, and the improvement effect is proved by experiments. The backbone network of this paper is composed of Resnet residual structure, bneck convolution block and SE attention module. A total of 5 downsampling operations are performed, three of which are stride 3 and two max pooling operations with stride 2.

The residual structure Resunit of YOLOv4 has two residual edges. Although the feature extraction ability is improved, this structure causes a lot of parameter calculations. Therefore, we changed the Resunit structure to the ordinary residual of the original Resnet. Its specific structure is shown in Figure 3. .

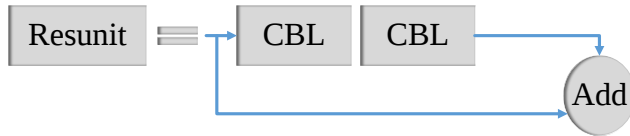


Fig.4 Resunit

The CBL consists of convolution + BN normalization + Leaky-RELU activation function. The Leaky-RELU activation function expression is as follows:

$$y = \begin{cases} x_i & x \geq 0 \\ \frac{x_i}{a_i} & x < 0 \end{cases} \quad (7)$$

Although the expressive power is weaker than the Mish activation function, the particularity of the instrument detection, it has a faster convergence speed and a more stable gradient, which is easier to train.

The convolution operation is intended to perform feature extraction on the feature layer through the convolution layer. The standard convolution is to use N 3×3 convolution kernels and each channel of the input feature map to convolve to obtain a feature map with N channels. The amount of parameters used: $P1 = C \times H \times 3 \times N$. The depthwise separable convolution block is to first use three 3×3 convolution kernels to convolve with each channel of the input feature map, and then use N 1×1 convolution kernels to perform channel expansion on the feature map to obtain the N channel. Convolutional feature map. The amount of parameters used in the depthwise convolution: $P2 = C \times H \times 3 + C \times H \times N$. Its process is as shown

in Figure 4:

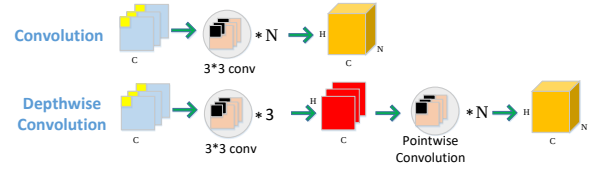


Fig.5. Depthwise convolution

It can be seen that the use of depthwise separable convolution blocks greatly reduces a number of parameters without changing the effect.

The original intention of the attention mechanism is that humans pay different attention to the items in the scene. In the neural network, it is intended to assign higher weight information to the important features of the feature map, so as to enhance the backbone network's higher attention to the important features^[6]. In the substation-oriented instrument identification system, the pointer position and deflection angle are given higher weight information. Therefore, this paper introduces the attention mechanism module SE. The SE structure implements the attention mechanism by training the weights of each channel of the feature map.

The channel with a larger contribution to the extracted features will have a higher weight. The SE structure is as follows: Figure 5

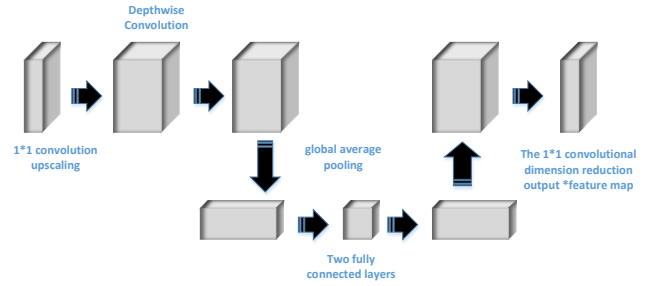


Fig. 6. Attention mechanism

The channel with a larger contribution to the extracted features will have a higher weight. The SE structure is as follows: Figure 4:

$$f(x) = \frac{1}{1 + e^{-x}} \quad (8)$$

The detection network structure uses the Bi-FPN detection network to replace the original FPN detection network. The SPP module and the improved Bi-FPN are used to screen and extract useful information features.

Since the types of instruments are not very complicated, and the original YOLOv4 has three feature layers for feature extraction, it is too complicated. In this paper, the improved YOLOv4 uses two feature layers for classification and prediction. When the input image feature is $416 \times 416 \times 3$, two feature layers of $13 \times 13 \times 255$ and $26 \times 26 \times 255$ will be output. Therefore, 6 a priori boxes are initially set, and each feature layer is more than 3.:

The prior frame was first proposed in Faster-RCNN^[7], which is equivalent to a kind of annotation information of the target category, which is more convenient for the prediction frame to move closer to the real frame. The instrument prediction mainly outputs three kinds of information: the coordinates of the target center position, the length and width of the prior frame, and the confidence. In this paper, an improved k-means clustering method is used to select the prior box size^[11]. The original k-means clustering method uses Euclidean

distance, but it will generate unnecessary errors due to the size of the prior frame. Therefore, a new clustering method is needed to overcome this problem. We introduce the minimization objective function as follows:

$$d(box, centroid) = 1 - iou(box, centroid) \quad (9)$$

The intersection and union ratio IOU uses AIOU (average intersection and union ratio), and the larger the value, the better the clustering effect. Select $k=3\sim 12$, as shown in Figure 7 below, when $k=9$, the function increase tends to be stable. However, since the number of a priori boxes will also affect the model complexity, the number of a priori boxes selected in this paper is 6

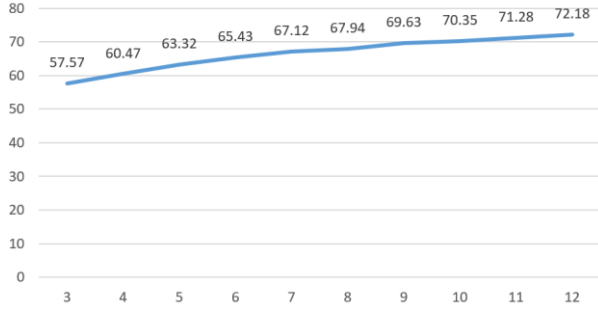


Fig.7. K-mean

The improved version of YOLOv4 of this article is obtained by merging the above improvements with the original network as shown in Figure 7

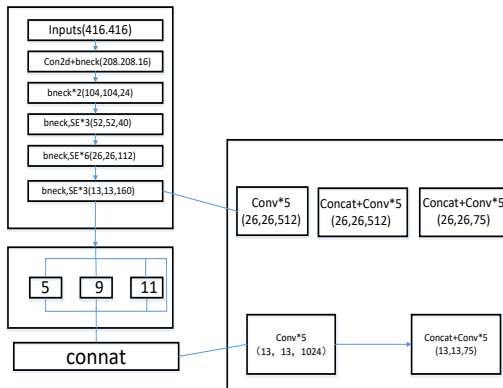


Fig.8. I-YOLOv4

5. Experimental results and analysis

5.1 Date set

The data set contains a large amount of training annotation data, which is biased [12]. In the target detection algorithm, the data set will directly affect the effect of target detection, so the data set needs to have accurate annotation information and high adaptability, and the types of pictures contained should meet most detection scenarios. The data set used in the experiment in this paper is the instrument photos taken by the actual substation. After being manually marked by Labeling, it is divided into two parts, which are used as training set and test set respectively. The data excerpt of the dataset is shown in Figure 8 below.



Fig.8. date set

Since there are only 1000 instrument pictures in the data set, 200 cat and dog classification pictures are added to the data set, so that the total number of the data set is 1200, which is divided into training set and test set by 9:1. The platform used for training is Intel(R) Xeon(R) CPU, 2080 Ti x2 GPU, cloud server with 128G memory, and the system is WIN7. The experimental platform of the test algorithm is a computer with Intel(R) Core(TM) i5-6700HQ CPU, 960M GPU and 16G memory.

4.2 Training situation

The training epoch is 100 rounds, the number of pictures in each iteration is set to 2, and the loss value (loss) represents the coordinate error between the predicted target position and the actual position of the target. In this paper, the training statistics are listed in Table 1.

Tab. 1. training situation

totaltraining rounds	training times	number of training rounds	The beginning of training loss	The end of training loss
10	1	10	1171.359	150.150
20	2	10	6.065	5.198
10	3	10	2.583	2.320
40	4	10	1.742	1.600
50	5	10	1.404	1.308
60	6	10	1.214	1.106
70	7	10	0.015	0.014
80	8	10	0.007	0.008
90	9	10	0.005	0.014
100	10	10	0.003	0.003

4.3 Evaluation indicators

The data set target detection algorithm is based on IOU (Intersection of Union), detection speed FPX (Frames Per Second), and mean average precision mAP (Mean Average Precision) as its algorithm model evaluation indicators. Summary

IOU refers to the ratio of the intersection area to the union area of the predicted frame b_f and the real frame b_r [13], and its calculation formula is as follows:

$$Iou = \frac{Area(b_f \cap b_r)}{Area(b_f \cup b_r)} \quad (10)$$

It mainly reflects the regression ability of the model, and judges whether there is a target in the box. Generally, IOU0.5 or above is a reasonable model.

The detection speed FPX is the speed of processing pictures. The higher the FPX, the stronger the model computing power.

mAP is the most intuitive performance indicator to measure the target detection model [14]. The IOU threshold should be specified when calculating mAP. AP is the average accuracy, and mAP is the average value of AP under all categories, which can correspond to the detection quality of multiple categories. The calculation formula is as follows:

$$AP = \int_0^1 precision(t)dt \quad (11)$$

$$mAP = \frac{\sum_{n=1}^N AP_N}{N} \quad (12)$$

precision(t) represents the detection accuracy under the specified IOU value. Most of the current target detection algorithm models use the coco data set to calculate mAP, showing the ability of the model for localization tasks and classification tasks.

4.4 Performance comparison

The static image detection results and video snapshot detection results in this paper are as follows:



Fig.9. Image detection



Fig.10. video detection

It can be seen from the above that no matter whether the picture is blurred or not, the target can be detected as the instrument image, and its IOU is above 0.9. At the same time, this paper trains and tests YOLOv4, YOLOv4-tiny and the improved version of YOLOv4 on the same dataset, and the results are listed in Table 2 below:

Tab. 2. training situation

algorithm	FPX	mAP	IOU
YOLOv4	4.13	80.01	0.83
YOLOv4-tiny	33.8	77.80	0.54
I-YOLOv4	35.2	89.00	0.91

The above results all prove that the improved YOLOv4 model can better detect the instrument condition, and its performance is higher than the original system, its detection speed and detection accuracy are higher

6. Experimental results and analysis

This paper proposes an improved YOLOv4 detection model for the problem of automatic detection of substation pointer meters. This algorithm meets the actual needs of substations in terms of detection rate and detection accuracy. Compared with other algorithms, its map is greatly improved. promote. In the next step, the author will conduct in-depth research on the simultaneous detection of multiple pictures and the detection of small objects where the instrument position occupies a small area of the picture.

Acknowledgements

Thanks to the tutor for the careful guidance, and the support of the substation for taking pictures of the dataset on the spot, and the General Fund Project of the National Natural Science Foundation of China (61873174).

References

- Luo, X., et al., Time Delay Estimation-based Adaptive Sliding-Mode Control for Nonholonomic Mobile Robots. International Journal of Applied Mathematics in Control Engineering, 2018. 1(1): p. 1-8.
- Li, S. and L. Chen, Research on Object Detection Algorithm Based on Deep Learning. International Journal of Applied Mathematics in Control Engineering, 2018. 1(2): p. 127-135.
- Yao, J. and X. Wu, Reconstruction of Arteriovenous Images based on CNN and Multi-electrode Electromagnetic Measurement. International Journal of Applied Mathematics in Control Engineering, 2019. 2(2): p. 202-209.
- Alexey Bochkovskiy, Chien-Yao Wang, Hong-Yuan Mark Liao. YOLOv4 optimal speed and accuracy of object detection EB/OL. (2020). <https://arxiv.org/abs/2004.10934>.
- Shu Liu et al. Path Aggregation Network for Instance Segmentation J. CoRR, 2018, abs/1803.01534.
- Diganta Misra. Mish: A Self Regularized Non-Monotonic Neural Activation Function J. CoRR, 2019, abs/1908.08681
- Kaiming He et al. Spatial Pyramid Pooling in Deep Convolutional Networks for Visual Recognition J. CoRR, 2014, abs/1406.4729.
- R. Hadsell, S. Chopra, Y. Lecun, Dimensionality Reduction by Learning an Invariant Mapping, CVPR (2006)
- Bahdanau, Dzmitry, Kyunghyun Cho, and Yoshua Bengio. "Neural machine translation by jointly learning to align and translate." arXiv preprint arXiv:1409.0473 (2014).
- Girshick, Ross. "Fast r-cnn." Proceedings of the IEEE International Conference on Computer Vision, 2015.
- REDMON J, FARHADI A. YOLO9000: better, faster, stronger C. // Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, 2017: 6517-6525.
- Torralba A, Efros A A. Unbiased look at dataset bias. In: Proceedings of the 2011 IEEE Conference on Computer Vision and Pattern Recognition (CVPR). Colorado Springs, USA: IEEE, 2011: 1521-1528.
- WANG S. Research Towards Yolo-Series Algorithms: Comparison and Analysis of Object Detection Models for Real-Time UAV Applications J. Journal of Physics: Conference Series, 2021, 1948(1): 012021(8pp).
- KANG H J. Real-Time Object Detection on 640x480 Image With VGG16+ SSD[C]//2019 International Conference on Field-Programmable Technology (ICFPT). IEEE, 2019.



Song Yuheng was born in Shenyang, Liaoning, China in 1996. He has received the B.S degree in Electrical Engineering from Shenyang Jianzhu University Shenyang in 2018. He was admitted to Shenyang Jianzhu University in 2020. He has been engaged in the research of data detection and pattern recognition, and has been involved in the background of robot research.



Sun Liangliang received his PhD degree in Control Theory and Control Engineering from Northeastern University, China, in 2015. He is currently a Professor of Information and Control Engineering at Shenyang Jianzhu University. His current research interests include Intelligent Manufacturing Systems, planning, scheduling, and coordination of design, manufacturing.



Zhao Mantong was born in Shenyang, Liaoning, China. She is being in the School of Economics and Management, Shenyang Institute of Technology in 2020. She has been engaged in the research of data detection and pattern recognition, and has been involved in the background of robot research.



Wang Yue received her master's degree in Physical education and training from Beijing Sport University in 2012. She is now a lecturer in Sports Department of Shenyang Jianzhu University. Her current research interests include vocational education, higher education etc.