

Development of object identification model with deep reinforcement learning algorithm

P. Ramesh Naidu [‡]

Department of Computer Science and Engineering

Nitte Meenakshi Institute of Technology

Bangalore

Karnataka

India

Avinash Sharma [†]

Department of Computer Science & Engineering

M. M. Engineering College

Maharishi Markandeshwar (Deemed To Be) University

Mullana

Ambala

India

Supriya P. Diwan [§]

Department of Electronics and Telecommunication Engineering

Government College of Engineering

Karad, Satara

Maharashtra

India

V. Dankan Gowda^{*}

Department of Electronics and Communication Engineering

BMS Institute of Technology and Management

Bangalore

Karnataka

India

[‡] E-mail: ramesh.naidu@nmit.ac.in

[†] E-mail: asharma@mmumullana.org

[§] E-mail: supriya.diwan8@gmail.com

^{*} E-mail: dankan.v@bmsit.in (Corresponding Author)

Parth M. Pandya [@]

*Department of Mathematics
Indus University
Ahmedabad
Gujarat
India*

Anand Kumar Gupta [#]

*Department of Information Technology
BlueCrest University
Monrovia
Liberia*

Abstract

This research work presents an object identification method based on the machine learning technique based on human vision system. The objective is to prevent processing a complete image in sort to locate objects. Presently, the-state-of-the-art techniques divide an image into sub-regions and search for an object in all the subparts. This is ineffective for applications like embedded systems where the computation power is restricted or the resolution of the images are high. To address this issue, an object identification task was formulated as a decision-making problem. Followed the concept of DRL proposed, accepted RL algorithm DQL was applied to learn a policy from input data, i.e. images, to identify objects in a scene. In this manner, with the policy learned, a set of actions that transforms a box was apply in order to make tighter a bounding box around the target object.

Subject Classification: Primary 93A30, Secondary 49K15, 93B30, 68Q87.

Keywords: Object detection, Reinforcement learning, Feature extraction, Markov decision process, Localization.

1. Introduction

Computer vision produce power to define task computationally what we see in real world, the capability of human being to view identify and classify makes person to take quick conclusion about the objects. In some occasion object identification is complex task when a single image contains many objects. Object identification is a task which is combination of classification and localization, method identifies the instance of object and compose a window over object for localization than with the

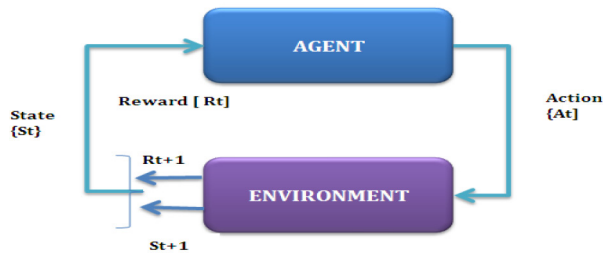
[@] E-mail: pandyaparth05@gmail.com

[#] E-mail: Ganand40@yahoo.co.in

instance score it describe the class of object [1]. Improving accuracy of object detectors are the big concern which solve by convolution neural network (CNN) .CNN plays a vital role in image visualization and deep learning. A novel model given by J. Donahue, T. Darrell, R. Girshick called R-CNN which uses image features and object proposals to detect object using Image net dataset. There are also some good results to predict the bounding box using high capable CNN architecture using regression technique [2]. Object detection is a technique in which we need to assess every instance of object because on the basis of object instance we predict class of object, very deep CNN has only power to detect object with high exactness. But need to implement supporting algorithm to complex models so increment in processing time can be saved. In literature there exist many research which work hard for object detection but the challenge is two minimize computation cost and training time of model also for a very large dataset its hard to execute that technique on resource constrained devices [3]. The method defines a geometrical box to wrap an object for localization by series of steps can be view as control problem, solving control problem we can go for different algorithm exist in literature. Identifying the accurate position of object in given image requires active study to convert point towards ground truth, identify features which can direct task to identification [4]. Proposed work presents an improved model for object identification .for any object identification task we need to extract image features and training of model should done using a feasible training methodology so we can get good results and solve the problem of identification accuracy and time efficiency [5]. Also it's necessary to store different data and model in minimum space so they can implemented in resource constrained device easily. Finally making image average evaluation time efficient we can give a better approach for identification. Proposed approach follows all standard protocols and parameter to define methodology [6]. Accuracy of model is calculated by precision and recall which are modern way of calculating accuracy of identification. One more metric which examine quality of identification given by IOU percentage, with given model we will prove the quality which can help community in modern world.

2. Related work

Object detection algorithms consist of two groups the algorithms that conduct object localization and detection jointly. primarily object detection algorithms, which perform object localization and recognition jointly, are

**Figure 1****Reinforcement Learning Concepts**

reviewed [7]. The common feature between these categories of algorithms is that they have to process all sub-regions of an image in order to detect objects. In other words, these algorithms have to examine great amounts of region proposals which have caused them to slow down. In the second part, the methods based on RL algorithms are explained. Before describing these algorithms, an introduction to RL is given. The background of RL based object localization algorithms comes from the human vision system that finds objects in a scene with a top-down search and turning visual attention to the densest part of the scene. In this section, the main topic is the recent object detection algorithms that apply deep learning methods to defeat the problem [8]. Each technique has used deep learning in a special way to enhance effectiveness and efficiency. As the techniques evolve, they move closer to the idea of end to-end learning. Because the subject of this project is object identification with RL, the focus is mostly on the identification which involves Softmax classification as well.

RL is considered as a mathematical structure for experience-driven autonomous learning. Using RL, an agent is able to learn the formation of an environment by trial and error over time to achieve its goal(s). RL challenges the idea of machine learning approaches which assume there always should be a teacher (annotated data) to learn a task [9]. In RL context, an agent observes state (S_t) at time step t . Having taken action a , the environment brings the agent to a new state S_{t+1} at time step $t + 1$ based on a transition function [10]. Also, the Environment gives response R_{t+1} to the agent by a reward function to illustrate the agent how Superior the taken action was. This process is shown in Figure 1. Every state contains of adequate information about the environment so that the agent can make a decision what action is best to take in a given state. In common, Deep Reinforcement Learning (DRL) refers to the methods that use neural networks coupled to RL algorithms. By means of a tree structured search

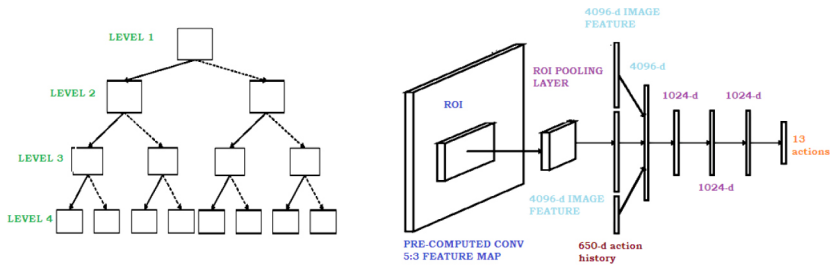


Figure 2
Tree Structured RL

strategy; a new method called Tree-RL was projected to localize various objects in an image. Tree-RL is a technique based on tree search where at each step using RL decides how the tree search should be separated. The hunt is in the form of top-down and begins from the whole image down to localizing an object in a leaf. Figure 2 shows the Tree Structured RL. The agent is allowed to take an action from two predefined set of actions, one for scaling the current window to a sub-window, and the next for translating the existing window locally. Given a state, the agent selects the top actions from both groups of actions and then the elected actions are performed in two part branches of a tree search.

In this manner, the agent takes both scaling action and local translation action concurrently at every state.

3. Active Object Localization with DRL

Active Object Localization is the most important paper by which proposed project is inspired. The key idea of the paper is to prepare the object localization task as a sequential decision making process where the agent's playground is images to localize objects [11]. The agent is going

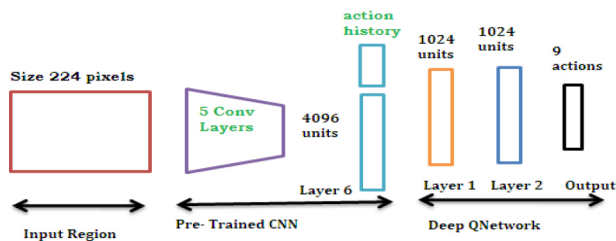


Figure 3
Active Object Localization with DRL

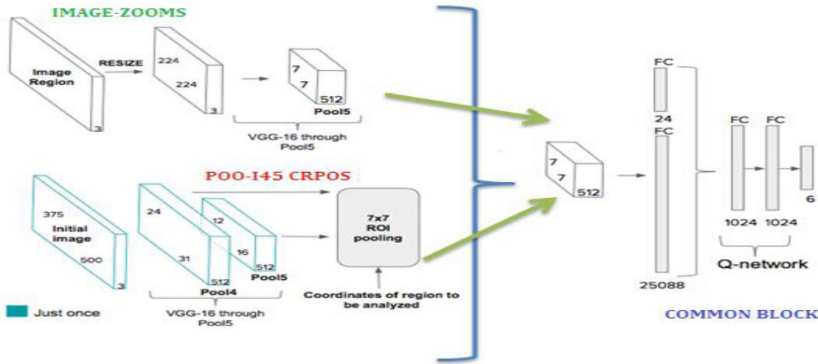


Figure 4

Hierarchical Object Detection Model

to learn focusing on thick regions in images where there might be a few objects. Figure 3 shows the Active Object Localization with DRL. The agent is estimated to determine the coordinates of bounding boxes tied to the objects in an image. More particularly, the agent initiates by analyzing the complete scene and narrowing its current observed window down to an object by taking a sequence of actions.

The states of the MDP are an illustration of the agent's window [12]. In the case of an action rising IOU the agent would get a positive reward +1 and otherwise -1. If the agent effectively localizes an object it will receive +3. Finally, to get an optimal policy for object detection, DQL is applied to solve the MDP which makes it feasible to use high dimensional inputs, i.e. images. It formulates the problem as an MDP and tries to address it using Q-learning. Unlike "Active Object Localization" and "Tree Object Detection" the agent is not free to shift its window where it decides over the input image given, but there are a set of pre-defined windows that when the agent put an action its window will be moved to the one of the predestined regions equivalent to the taken action [13]. Figure 4 shows the Hierarchical Object Detection Model. The agent starts by analyzing the whole image and take decision where to focus between the sub-region windows.

4. Proposed ODDRL-Net Model

The technique proposed is entirely different from most detection methods given in literature, sliding window approach is used by many

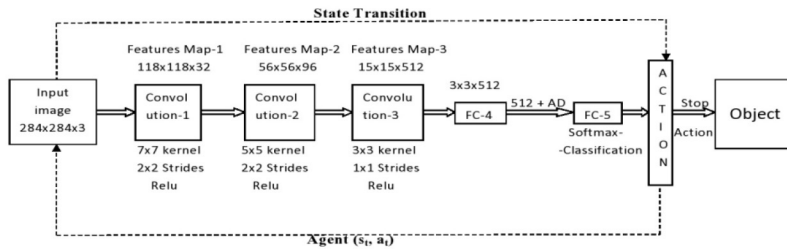


Figure 5
Proposed ODDRLnet

researcher which uses a fixed size window for scanning objects in one direction but our proposed object detection having not any fixed path the mask can select any action to move in any direction. Another technique used widely is object proposals in which candidate region predicted on some small cue but we define a systematic reasoning strategy for proposals compare to other region proposal algorithm. Proposed object detection deep reinforcement learning - network (ODDRL-Net) an improved and efficient model which predicts fast object identification as shown in Figure 5. Model contains three convolution layer {convol1, convol2, convol3} which are identical to model using features for object identification like VGG-M net [15] etc. Sang doo Yun used to initialize their model with VGG-M net as pretrained model also they take the different layer as motivated by this model for giving a tracker. The next FC4 fully connected layer combined with dropout Relu and has 512 nodes when we calculate output [16].

The final layer FC5 having Softmax classification score with action probabilities which decide what action must be applied if agent repeats action and come again to previous than oscillation occurs already we proposed the way in upper section how to deal the situation .agent updates states while selecting sequential action until it reaches final state and stop the process. Following table 1 gives the step by step process which will clear the process followed by proposed architecture, as compared to previous architecture proposed model uses less computational cost [17,18]. Proposed architecture uses the customized concept and region proposal algorithm given in literature while changing some parameter and nature for achieving dynamic implementation of model. Training and testing for proposed model have done with customized Pascal Voc Dataset.

5. Results and Discussion

Precision tell the exact quality of identification .normally we having two types of regions which is attended by agent one is all attended and well identified with threshold and other is terminal region when stop action is activated by agent that tells the final identification. To design a detector inspired by the Softmax classifier, which scores every area in a single search phase, we employed a shallower version of the RCNN. Following Figure 6. will present the actual processed image for top five classes. As different technique we have also made mAP calculation on 0.5 IoU threshold. Approximately we have tested more than 20 images per category to calculate average precision and than its mean that collectively called mAP [19]. We compare proposed method with different method in literature RCNN having same layers in architecture like our model and work without bounding box regression technique .Region lets predicts bounding box with deep CNN using regression method for input image it uses deep CNN in efficient way using few box for prediction is good strategy because it saves computing power.

Table 1
Deep Q-Learning Architecture (ODDRL-Net)

Layer	Input	Filter Size	Strides	Activation	Output
Convolution 1	284x284x3	7x7	2	Relu	118x118x32
Convolution 2	118x118x32	5x5	2	Relu	56x56x96
Convolution 3	56x56x96	3x3	1	Relu	15x15x512
FC 4	15x15x512	0	-	Relu	512
FC 5	512	-	-	-	11

Table 2
Comparison mAP Data between ODDRL-net and Different Model

Methods	Aero	Bike	Bird	Boat	Car	Cow	Horse	Person	mAP
RCNN	64.2	69.7	50	41.9	71	58.5	60.6	54.2	54.2
Region lets	44.6	55.6	24.7	23.5	51	35.9	61.9	44.1	40.2
RPNSS	77	78.1	69.3	59.4	78.6	78.8	82	69.9	70.4
RPN- 300	87.4	83.6	76.8	62.9	82	82.6	85.5	84.1	75.9
Proposed method	88.9	86	79	74	87.6	86.4	88	87.5	78.3

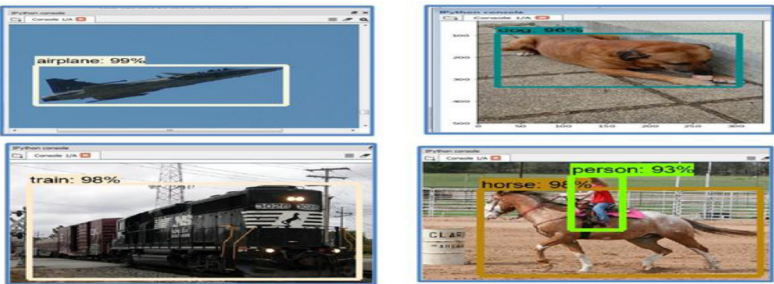


Figure 6
Processed Images for Best Result Four Classes

Table 2. gives the comparison data between ODDRL-net and different model discussed. RPN with selective search strategies also giving good results but when RPN uses 300 region proposals than it gives better results this architecture uses same no of layer as we have used in proposed method [20] . Overall our method having good results compare to other but in some image we get better results with RPN-SS, the drawback of this approach is it needs large dataset, maximum no of region and large computing power to predict object.

Figure 7. presents the Performance comparison between ODDRL-net and different model. overall if we look we found proposed method come with 76.6 MAP compare to other baseline methods it gives better results next better method is RPN having 75.9 MAP. Best results showed by underline bold fonts. Proposed method run for 200 steps per category it means we have 20 category in Pascal voc dataset so there will be 4000 candidate region reaching average 78% recall it means the quality of detection is good because with recall only we say how many relevant

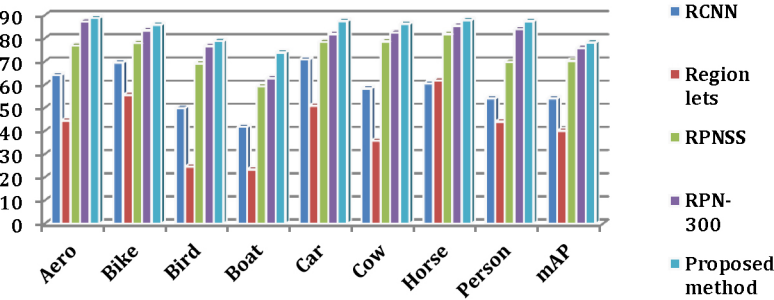


Figure 7
Performance comparison between ODDRL-net and different model

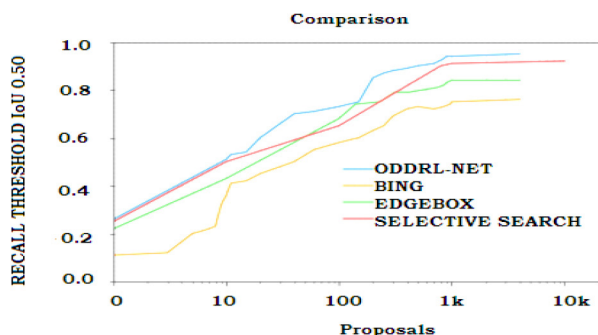


Figure 8

Plot Recall at Threshold IoU 0.50

items we have reached [21] also for top hundred candidate recall of our proposal is more than 70 percent. We also use same scoring function in place of classification score for making evaluation fare.

Figure 8. shows the comparison plot of recall vs. number of proposal of different methods in literature. We compared with Edge box, selective search and Bing result shows that proposed method take 60 percent recall on 10 proposals while other methods giving 20 to 30 percentages. Every object detection work have challenge to have better recall rate that is we have to identify every single instance of object proposed approach present. Better recall rate [22]. Like other methods we not set focus to improve recall with maximizing region, work on minimum region so fixing overall recall will improves at early interval as well. Number of action required to identify object successfully we can see median in this distribution which tells that in less than 5 actions we can get correct instance of object. Also it tells that agent has capability to detect 11 without processing more regions because they occupied big part of object.

6. Conclusion

In this paper, first the proposed architecture is explained clearly related to different layer function and process of identification. Then used dataset is discussed with respect of training and testing images. Proposed work uses model named ODDRL-net which contains three convolutions and two fully connected layers. Convolution layer used for feature extraction and identification of object is done with deep reinforcement learning. Softmax classifier is used for defining the label of class. The main contribution of this research is to defining model and

combining the localization with Softmax classification also implemented the algorithm by means of the novel deep learning framework, Tensor flow. This implementation can simply be customized by future researchers in sort to conduct additional experiments with different datasets, neural network architecture, or the agent architecture. The customization of Pascal VOC Dataset was used for training which has a considerable effect on the performance. Reinforcement learning used in proposed work is an efficient strategy to learn policy for identification because every object not found the common search path. After this agent are able to learn with its mistakes and experience to optimize the policy to find objects. ODDRL-net is comparative model with other in literature and has better ability to identify object.

References

- [1] C. C. Alves, S. Rabello, and K. M. de Carvalho, "Assistive technology applied to education of students with visual impairment," *Salud Pública*, vol. 26, no. 2, pp. 148-152 (2019).
- [2] D. McAllester, R. B. Girshick., "Object detection with discriminatively trained part based models" *IEEE Trans on Pattern Analysis and Machine Intelligence*, 1627-1645 (2018).
- [3] RaviPrakash, "Small object detection using retinanet with hybrid anchor box hyper tuning using interface of Bayesian mathematics," *Journal of Information and Optimization Sciences*, 43:8, 2099-2110 (2022).
- [4] G. Bai, Q. Dong., "Cube CNN SVM A novel hyper spectral image classification method" *Proceeding IEEE 28 ICTAI*, November 2019, page 1027-1034 (2019).
- [5] P. Pavankumar, N. K. Darwante, "Performance Monitoring and Dynamic Scaling Algorithm for Queue Based Internet of Things," *2022 International Conference on Innovative Computing, Intelligent Communication and Smart Electrical Systems (ICSSES)*, pp. 1-7 (2022).
- [6] Jyothi Varanasi & M. M. Tripathi, "K-means clustering based photo voltaic power forecasting using artificial neural network, particle swarm optimization and support vector regression," *Journal of Information and Optimization Sciences*, 40:2, 309-328 (2019).

- [7] Huiqi Zhu, "Strategies for adopting unified object identifiers in logistics resource integration environments," *Journal of Discrete Mathematical Sciences and Cryptography*, 21:4, 991-1003 (2018).
- [8] A. Singla, N. Sharma, "IoT Group Key Management using Incremental Gaussian Mixture Model," 2022 3rd International Conference on Electronics and Sustainable Communication Systems (ICESC), pp. 469-474 (2022).
- [9] Yang Jiao & Song Zhao, "Object tracking from airborne video using particle filters algorithm on dynamic feature fusion," *Journal of Discrete Mathematical Sciences and Cryptography*, 19:3, 787-799 (2016).
- [10] Jan K Chorowski, and Yoshua Bengio., "Attention-based models for speech recognition" In Advances in neural information processing systems, pages 577-585 (2020).
- [11] A. S. Naik, R. S. Meena, J. M. Kudari and S. Purushotham, "Design and Implementation of a System for Vehicle Accident Reporting and Tracking," 2022 7th International Conference on Communication and Electronics Systems (ICCES), pp. 349-353 (2022).
- [12] Juan C Caicedo and Svetlana Lazebnik., "Active object localization with deep reinforcement learning" In Proceedings of the IEEE Inter. Conference on ComputerVision, 2018, pages 2488-2496 (2018).
- [13] Liu, W., Anguelov, D., Erhan, D., Szegedy, C., Reed, S., Fu, C.-Y., and Berg, A. C. Ssd., Single shot multibox detector. In European Conference on Computer Vision (2016), Springer, pp. 21-37 (2016).
- [14] M. Oquab I. Laptev, and J. Sivic., "Learning and transferring mid-level image representations using convolutional neural networks" in Proc. IEEE Conf. Comput. Vis. Pattern Recognit., Jun. 2014, pp. 1717-1724 (2014).
- [15] N. Srivastava, A. Krizhevsky and R. Salakhutdinov., "Dropout: A simple way to prevent neural networks from over fitting," *J. Mach. Learn. Res.*, vol. 15, no. 1, pp. 1929-1958 (2020).
- [16] M. Nagabushanam, H. G. Govardhana Reddy & K. Raghavendra, "Vector space modelling-based intelligent binary image encryption for secure communication," *Journal of Discrete Mathematical Sciences and Cryptography*, 25:4, pp.1157-1171 (2022).
- [17] Ross Girshick., "Convolutional neural net" In Proceedings of the IEEE international conference on computer vision, pages 1440-1448 (2020).

- [18] S. Gupta and J. Malik, "Perceptual organization and recognition of indoor scenes from RGB-D images," in *Proc. IEEE Conf. Comput. Vis. Pattern Recognit.*, Jun. 2019, pp. 564-571 (2019).
- [19] M. Sahana, R. S. Varun, and T. Rajesh, "Implementation of swarm intelligence in obstacle avoidance," in *2nd IEEE International Conference on Recent Trends in Electronics, Information and Communication Technology, Proceedings*, vol. 2018-January, pp. 525-528 (2017).
- [20] Sangdoo Yun, Yoo, Jongwon Choi, "Action-Decision Networks for Visual Tracking with Deep Reinforcement Learning" *IEEE Conf. on Computer Vision and Pattern Recognition (CVPR)*, pp. 101-120 (2017).
- [21] Steinkraus, D., Buck, I., and Simard, P. "GPU for machine learning algorithms." In *Document Analysis and Recognition, 2005. Proceedings. Eighth International Conference on* (2015), IEEE, pp. 1115-1120 (2015).
- [22] Triggs, B. and Dalal, N., "Histograms of oriented gradients for human detection" In *Computer Vision and Pattern Recognition, CVPR 2018. IEEE Computer Society Conference on* 2018, vol. 1, IEEE, page 886-893 (2018).

