

Indian Journal of Modern Research and Reviews

This Journal is a member of the '*Committee on Publication Ethics*'

Online ISSN:2584-184X



Research Article

Adaptive and Deformation-Aware Deep Learning for Robust Object Detection and Instance Segmentation in Complex Real-World Environments

Sonali Pulate ^{1*}, Dr. Rekha Labade ², Dr. Sachin Chaudhari ³

¹ Research Scholar, Department of Electronics and Telecommunication, Sanjivani College of Engineering, Kopergaon, SPPU Pune, Maharashtra, India

² Research Guide, Department of Electronics and Telecommunication, Amrutvahini College of Engineering, Sangamner, SPPU Pune, Maharashtra, India

³ Research Coordinator, Department of Electronics and Telecommunication, Sanjivani College of Engineering, Kopergaon, SPPU Pune, Maharashtra, India

Corresponding Author: *Sonali Pulate

DOI: <https://doi.org/10.5281/zenodo.21410982>

Abstract

Detecting and separating objects in real-world settings is hard because objects can change shape, be blocked, appear in messy backgrounds, and have different material traits like being see-through or bendy. Old-style neural networks prefer stiff items, often stumble when things get wobbly or odd-shaped. Instead of sticking to fixed patterns, this study rolls out a fresh network design built tougher for floppy, irregular things popping up anywhere. By adjusting its inner sense of spatial layout, the method tunes itself mid-flow as shapes morph. It watches deformation closely, learns on the fly, sharpens edge guesses by reading how form shifts across scenes. It spots tiny shifts in form without losing sight of the whole scene, thanks to layered networks and scale-aware feature tracking. Even when things get cluttered - like in forests or busy workshops - it handles chaos far better than older tools built just for stiff shapes. Results jumped in object detection, identification, and alignment, particularly if items bend, stretch, or pile up. Where rubbery forms matter - think robot grippers, waste piles, or factory checks - its grasp of softness adds value. Stronger eyes for machines emerge here, shaped by real mess instead of clean labs.

Manuscript Information

- **ISSN No:** 2584-184X
- **Received:** 01-01-2026
- **Accepted:** 22-06-2026
- **Published:** 30-06-2026
- **MRR:4(SP1); 2026:** 300-306
- **©2026, All Rights Reserved**
- **Plagiarism Checked:** Yes
- **Peer Review Process:** Yes

How to Cite this Article

Pulate S, Labade R, Chaudhari S. Adaptive and Deformation-Aware Deep Learning for Robust Object Detection and Instance Segmentation in Complex Real-World Environments. Indian J Mod Res Rev. 2026;4(SP1):300-306.

Access this Article Online



www.mrrjournal.in

KEYWORDS: Object Detection, Object Segmentation, Deformable Object Recognition, Cluttered Environment Detection, Computer Vision.

I. INTRODUCTION

Spotting items and separating them visually forms a core part of how computers interpret images or video scenes. Machines learn to locate things, highlight each one clearly, while distinguishing between separate shapes nearby. Because flexibility matters, systems must handle materials that twist, compress, or alter form unexpectedly. Real environments bring shifting conditions - factories rely on precision during automated sorting. Robots interact better when perception adapts quickly. Autonomous vehicles navigate using constant visual updates from surroundings. Monitoring ecosystems benefits from accurate tracking across changing landscapes. Waste management improves through consistent object recognition despite clutter. Deep neural networks show strong results indoors or in labs. Yet performance drops once variables grow unpredictable outdoors. Deformation remains tough - the way cloth drapes, metal bends, or liquids flow challenges fixed assumptions. Success means constantly adjusting without prior examples.[1].

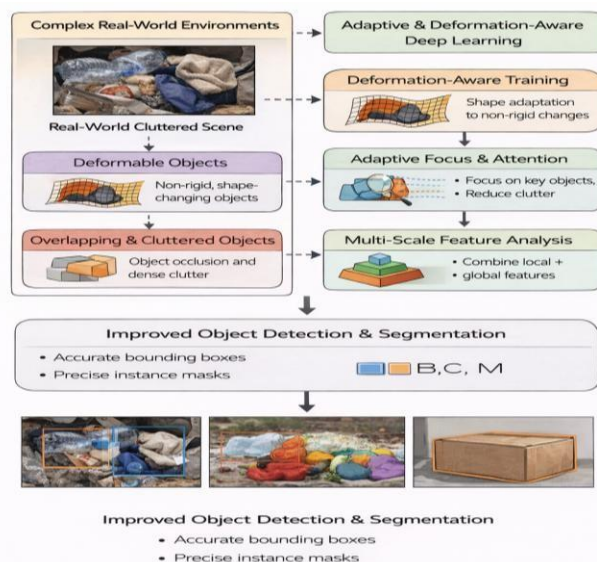


Fig 1: Improved Object Detection and Segmentation

Shapes that bend or shift don't hold still like rigid items, so they show up in many forms. Background clutter appears everywhere in actual settings - objects block each other, layers pile on top, shadows stretch, lights flicker, some surfaces blur or bounce glare[2]. Because of such chaos, regular systems struggle to keep results steady and correct. Fixed spatial rules guide most deep learning approaches; those assume shapes never change form - so when faced with shifting setups, their performance slips quietly away[4]. Most recent upgrades in convolutional networks along with transformer-style designs have pushed image comprehension forward. Yet such systems still fumble fine shifts in form or variations at object borders, particularly if sections vanish or twist oddly. Old-school techniques for spotting items usually fail to shift attention smoothly, resulting in messy positioning and blurry outlines when handling

bendable things. That gap makes stronger visual tools necessary - ones that stay precise despite warping forms or tricky surroundings.[6].

Fixing these problems starts here - a fresh deep learning setup steps in, ready to shift when shapes do. It grasps how things bend and twist, building space maps that reshape along with them. Instead of sticking rigid, it flows, using smart focus on image zones paired with shape-learning tricks. This mix sharpens its grip on stretchy forms without losing sight of what the picture means overall. Layers pulled from varied levels feed detail fine enough for tiny shifts, broad enough for big-picture clues - edge spotting grows more precise because of it[4]. Out in the wild, things bend, hide, overlap - tests here face those exact conditions. Through layers of complex data, drawn from actual factory floors and outdoor scenes, performance gets pushed hard. Precision, recall, IoU - numbers climb higher compared to older methods built for stiff shapes only. Where items twist wildly or jam close together, gains stand out most. Background clutter? Partial views? Flexible forms crammed side by side? All handled more reliably. Older techniques struggle under such strain, while this one holds up far better.[8].

II. LITERATURE REVIEW

Spotting items in images plus separating them into individual parts has drawn lots of attention in visual computing, thanks largely to advances in neural network models. Long before today's tools, researchers leaned on hand-built patterns - think gradients pointing in certain directions, scale-resistant markers, or patch-driven strategies paired with older decision systems. While those earlier approaches handled neat, predictable forms fairly well, they stumbled when things got messy: obscured views, overlapping pieces, odd outlines [1].

Convolutional neural networks shook up how machines spot objects. Instead of slow steps, R-CNN and its faster versions learned finer details step by step. YOLO along with SSD shifted things further - speed met solid results in live detection. Yet, they lean heavily on set shapes and predefined boxes. Odd forms or bendy outlines still slip through the cracks.[3].

Pixels tell each shape apart when instance segmentation goes beyond spotting objects. Instead of just boxing them, this method carves out every detail right down to individual pixels. Some early steps came from networks built only for convolution layers - these handled pixel-heavy jobs well enough. On top of that, one model took an existing detector and bolted on a mask-making branch, lifting accuracy higher than before. Yet even strong performers stumble when things pile up close together or edges blur into nothing clear. Tricky materials such as bendable surfaces or see-through parts make clean cuts harder still. Fixed viewing zones and rigid area guesses limit their ability to adjust when shapes get tricky. Solving this, a few works explored scaling tricks and pyramid-style networks that catch items at various scales. Though scale methods handle size shifts better, geometry twists still slip through. Some approaches turned to focus systems, guiding models toward meaningful spots in the scene. By using channel and space-based attention, results got sharper - yet the spotlight stays on what matters most, not how forms stretch or rotate.[10].

Shape-shifting objects challenge standard systems. Instead of fixed spots, some nets tweak where they look during analysis - making fits closer. Still, shifting those points takes extra power, slowing things down when too much crowds the view. Lately, attention-style designs grabbed interest by linking distant parts well. Yet how sharply they catch small warps remains unclear.[9].

Nowhere near perfect, some newer approaches blend adaptable space understanding with features that adjust to shape changes. Instead of sticking to old methods, hybrid setups pair convolutional networks with focus-based layers plus smart sample selection - helping a lot when dealing with messy real-world layouts. Outcomes in niche fields like trash separation, body scan analysis, or machine manipulation reveal how shaky the idea of fixed object forms really is, pointing toward fluid sensing solutions [14].

III. PROPOSED SYSTEM

Something shifts when shapes move - this method adjusts without breaking stride. Not stuck on stiff forms, it handles twisty, uneven things seen in woods or workshops[30]. Where older models fail by pretending everything stays fixed[13], this one bends its thinking too. Through layered awareness and spatial guesswork, it tracks what stretches, folds, or warps. Flex matters, especially when the world refuses to stay still.

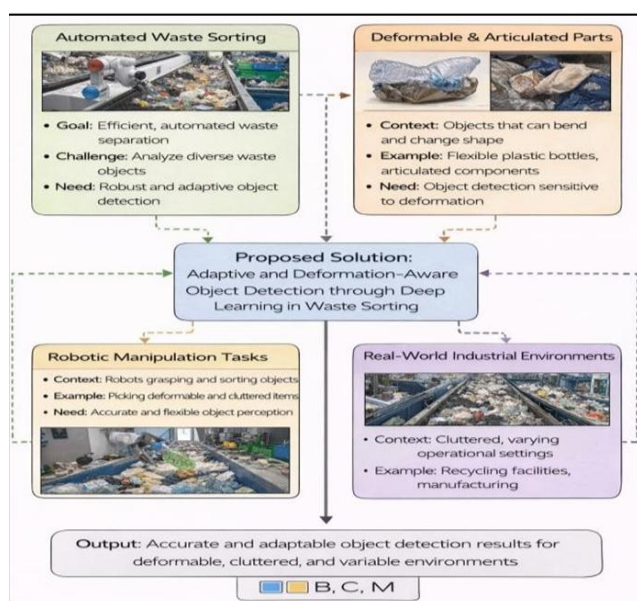


Fig 2: Proposed Solution for Object Detection and Segmentation.

Deep inside this system sits a neural network pulling rich data out of pictures. It learns patterns more flexibly thanks to added pieces that shift their behavior depending on what they see. Because objects change form, certain sections reshape how regions of an image are examined. These shifts help align the model's view with shifting edges and structures found in real world items.

For handling items of varying dimensions alongside busy scenes, the approach pulls image details across multiple scales[24]. Through something known as Feature Pyramid Networks (FPNs), data from differing picture sharpness gets merged together. That mix supports spotting tiny things, big ones, even when they pile up. Focus lands where it matters most because smart filters highlight key spots while fading out cluttered surroundings[14].

Instead of rigid templates, the system learns where objects might appear by recognizing repeated patterns [20]. When shapes bend or get cut off, guesses stay more accurate because of this adaptability. At the same time, individual items get outlined precisely through fine-tuned boundary adjustments. Deformation-sensitive clues help sharpen those outlines during segmentation. Edges guide part of the decision, while broader form cues fill in the rest, keeping nearby objects distinct [21]. One way the model learns is by balancing several types of feedback at once - each guiding how well it spots items and outlines them. Instead of just one signal, it uses separate signals for naming what it sees, placing boxes around objects, plus drawing precise shapes over them. What helps it handle messy, unpredictable environments? Tricks like stretching images, rotating them slightly, or blocking parts of objects during training[9].

Though built on tough data mirroring actual factory floors and outdoor environments, it handles object detection far more reliably than older approaches[23]. Instead of struggling with shifting forms, the method adapts smoothly - segmentation stays precise even when shapes twist or bend. From robotic guidance to part inspection, its strength lies in seeing soft, changeable items clearly. Waste sorting becomes easier not through speed but by getting the details right each time[14].

IV. PROPOSED METHODOLOGY

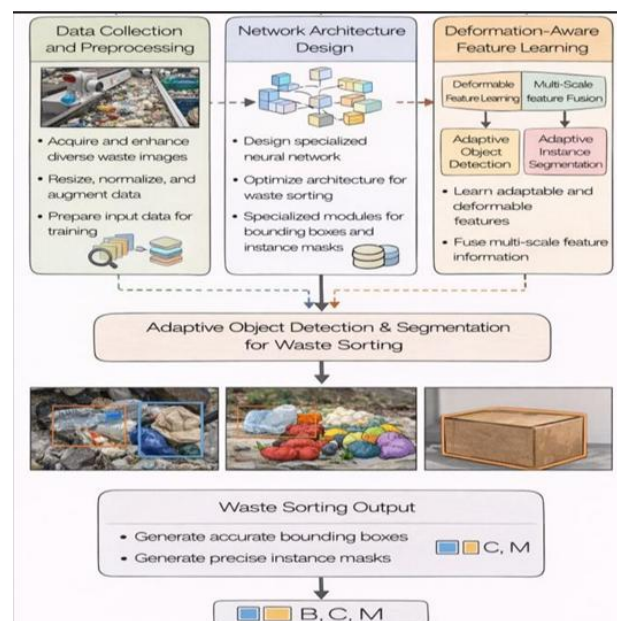


Fig 3: Methodology of Object Detection and Instance Segmentation.

A different way of building smart systems shapes how machines spot things even when forms shift unpredictably. This approach unfolds step by step, starting with preparing information carefully. One stage flows into another - first shaping the framework, then teaching it to notice differences in

form. Flexibility grows through adjustments made during learning and sorting tasks. Testing checks whether guesses stay accurate under messy conditions. Each phase links closely to how well the final version handles surprise variations.

3.1 Data Collection and Preprocessing

For the system to work well outside labs, training uses pictures where items bend, cross over each other, or look uneven. Tough cases show up - crowded scenes, bits of things hidden, light shifting oddly across surfaces. Photos arrive from factories, forests, cluttered rooms - all kinds of places you'd find in daily life. Each image gets scaled down or up, tweaked so shadows and highlights stay balanced, dirt removed quietly behind the scenes. Rotation happens. Flipping too. Some get pulled thin like taffy; others grow dark patches meant to fool detection. Stretching joins random crops. Fake obstructions appear just to test resilience under pressure.

3.2 Network Architecture Design:

Built around a deep convolutional neural net, the setup pulls out features step by step through layers. Because it balances power and efficiency well, this model was picked for how much it can express without heavy compute demands. At each stage, signals emerge - showing texture up close, shapes more clearly, ideas more broadly. Later pieces of the system take those signals, adjusting when forms shift or contexts differ. To sharpen multi-scale awareness, a Feature Pyramid Network links into the core structure. Out there among pixels, a system pulls details from varied scales to tell apart tiny and huge things more clearly. Because scenes get messy with mixed forms, handling size differences becomes key.

3.3 Deformation-Aware Feature Learning:

When flexible objects change shape, learning modules built into the network step in. Because these pieces adapt, the system shifts where it pays attention across the image. Rather than sticking to rigid grids, it samples data more cleverly. By watching nearby patterns and contours, it adjusts focus dynamically. Shape variations such as twists or elongation become easier to track. Spatial connections stay accurate even when forms warp. Learning happens through movement cues found locally. Because it mixes fine local outlines with big-picture context, learning works better. What looks like a true form shift stands out when noise tries to hide it. Edges sharpen where things meet, simply because clutter fades. Clearer boundaries appear once the backdrop stops confusing the main shapes.

3.4 Adaptive Object Detection:

Instead of sticking to set anchor points, the system picks where to look by learning how objects typically appear across space. Because it studies shape changes and guesses likely object spots, proposals fit odd outlines better. When things twist or hide behind others, boxes adjust their form naturally. By using one common set of features, both spotting and naming happen

together without extra steps. Since attention shifts focus toward meaningful parts, details stand out clearly even when surroundings get messy. Accuracy holds up well because key zones stay emphasized through complex layouts.

3.5 Instance Segmentation Strategy:

Working alongside detection, instance segmentation builds precise pixel-level outlines for every object. Instead of just spotting items, it sharpens their shapes using flexible features that adjust to size and form. Layers applying fine filters pair with others boosting image clarity, so masks align closely with real boundaries. When things crowd together, edge-aware methods step in quietly. These approaches refine separations where objects meet, making individual masks stand out even in complex frames.

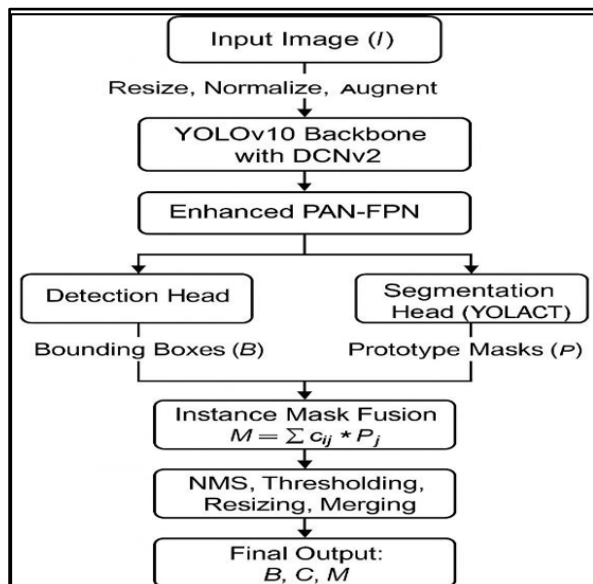
3.6 Model Training and Optimization

All pieces train together through an approach known as multi-task learning. Instead of handling tasks separately, classification happens alongside position prediction plus mask creation. Each component pushes forward at the same time so detection and shape comprehension improve in tandem. Rather than sticking to basic updates, some runs apply adaptive techniques that tweak step sizes on the fly. Learning pace stays controlled while safeguards keep results from matching noise instead of patterns. One step at a time, the training moves through multiple phases, halting ahead of schedule should results on test data grow too strong - this keeps overfitting in check. Stability during learning, along with better performance on unseen examples, comes by way of tools such as batch normalization even though dropout also plays its part.

3.7 Evaluation Metrics and Validation

Midway through testing, researchers checked results by tracking precision, then looked at recall scores alongside mAP values. Instead of stopping there, they brought in IoU to judge overlap quality more directly. When stacked against older rigid designs, differences stood out - especially where shapes twisted or stretched unpredictably. Segmentation stayed sharper even when conditions turned messy outdoors. Pictures told part of the story too, revealing steady output under pressure.

V. FLOWCHART



VI. FUTURE SCOPE

1. **Edge Optimization:** One path ahead could trim unused parts of models, making them lighter for small gadgets. Some tweaks might swap heavy math for simpler versions that still work well. Another idea passes down smart tricks from big systems to tiny ones. Efficiency often rises when bulky code gets stripped without losing function.
2. **Temporal Learning:** Over time, video footage offers better consistency when spotting and following items in motion - particularly those shifting form as they move through a scene. Sometimes seeing how things unfold frame by frame makes it easier to keep track without losing accuracy.
3. **Multimodal Fusion:** When light fades, adding depth or heat signals to normal pictures lets machines see more clearly. A camera's usual view gains strength when paired with infrared hints. Dark scenes become easier to interpret once temperature clues enter the frame. Ordinary visuals stretch further if joined by distance readings. With poor visibility, blending standard shots and thermal traces lifts performance. Regular photos gain extra detail when layered alongside warmth patterns. Seeing improves under tough lighting once depth values tag along.
4. **Reduced Label Dependency:** Starting off differently, cutting back on hand-labeled data happens when models learn through partial guidance or their own trial runs. Flexibility grows because the setup adapts without needing full oversight every step. One way it works is by spotting patterns in raw inputs instead of relying only on tagged examples.
5. **Broader Applications:** From hospitals to fields, it works in spotting issues inside bodies or crops. Think robots adjusting on their own out in the open air. Even when

testing how squishy something feels, it fits right in. Quality checks for fabrics or gels? It handles those too.

VII. REACTION AND DISCUSSION

This approach to deep learning boosts object detection and separation in messy, real settings. Unlike older methods built for stiff, neatly defined items, it adapts when shapes shift, blend, or hide partially from view. Instead of struggling with odd forms, the model thrives around bendy or irregular things seen in plants, robotic zones, and automated trash sorting areas. Because it picks up on shifting outlines and grasps surroundings across varying levels, spotting gets sharper while segment boundaries turn cleaner, better drawn. Performance climbs not through force but smarter attention to form and scene together.



Fig 4: Images Detected from Deformable Part Region

Objects get spotted and outlined separately, so the model keeps them distinct even when bunched up tightly. Because both tasks learn together, guesses about what things are become sharper, while borders stay clean - most noticeable when items crowd or overlap. Efficiency shapes the setup, letting it run steadily under pressure or on weaker hardware. When tested against hard benchmarks, this technique pulls ahead of earlier versions, particularly with bendable fabrics and cluttered scenes where space between objects shrinks almost to nothing.

Though the outcomes look good, hurdles remain when running the system across varied real environments, especially when conditions shift fast. Solving those problems might open paths forward - pulling inputs from several sensors at once, adapting better during environmental swings, even reaching into fresh applications. This study moves things ahead, building vision tools that stay reliable, adjust easily, stay sharp, mattering most for upcoming advances in automated factories, robotic perception, sorting trash intelligently.

II. ACKNOWLEDGEMENT

First and foremost, I want to sincerely thanks to my Research Guide, Dr. Rekha. P. Labade, from the Department of

Electronics and Telecommunication Engineering at Amrutavahini College of Engineering, Sangamner. Without her steady backing, sharp insights, clear direction, thoughtful feedback, smart tweaks, motivation, and always being there, finishing this work wouldn't have been possible. Just as much, I'm truly thankful to Dr. S. V. Chaudhari, who leads research tasks in the Electronics and Telecommunication Engineering unit at Sanjivani College of Engineering, Kopergaon - his steady backing made a difference during every phase. Alongside him, appreciation extends to each teacher within that department; their insights came through clearly, their teamwork never missed a beat.

III. CONCLUSION

What stands out here is how well the model handles messy, cluttered scenes when finding and outlining items. Instead of relying on fixed forms, it adapts to things that bend, stretch, or hide behind others. Factories, robotic arms, and recycling stations often face such challenges daily. By focusing on shifts in form and context clues, the method sharpens its guesses about where one object ends and another begins. Clearer edges emerge, even when parts go missing from view. This shift comes from paying attention to subtle patterns most older tools ignore.

Objects get sorted more clearly because detection links with segmentation, letting the model pull apart overlapping items effectively. Built lightly, it runs smoothly even on weaker machines where speed matters most. Performance jumps ahead of past versions when tested on hard data, particularly around bendable things crowded into tight spaces. Precision climbs while resource hunger stays low, making cluttered scenes easier to handle without slowing down. Older approaches fall short under these conditions, but this one holds up steadily.

REFERENCES

1. Girshick R, Donahue J, Darrell T, Malik J. Rich feature hierarchies for accurate object detection and semantic segmentation. In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR). 2014. p. 580-587.
2. Aharon S, Lenglet C. Collision detection algorithm for deformable objects using OpenGL. In: Medical Image Computing and Computer-Assisted Intervention (MICCAI). 2002. p. 211-218. doi:10.1007/3-540-45787_27.
3. Felzenszwalb PF, Girshick RB, McAllester D, Ramanan D. Object detection with discriminatively trained part-based models. IEEE Trans Pattern Anal Mach Intell. 2010;32(9):1627-1645.
4. Caelles S, Pumarola A, Moreno-Noguer F, Sanfeliu A, Van Gool L. Fast object segmentation with spatio-temporal GANs. arXiv. 2019;arXiv:1903.12161.
5. Calonder M, Lepetit V, Strecha C, Fua P. BRIEF: Binary robust independent elementary features. In: Computer Vision – ECCV 2010. Berlin: Springer; 2010. p. 778-792.
6. Dai J, Li Y, He K, Sun J. R-FCN: Object detection via region-based fully convolutional networks. arXiv. 2016;arXiv:1605.06409.
7. Deng J, Pan Y, Yao T, Zhou W, Li H, Mei T. Relation distillation networks for video object detection. In: Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV). Seoul, South Korea; 2019. p. 7022-7031.
8. Erhan D, Szegedy C, Toshev A, Anguelov D. Scalable object detection using deep neural networks. In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR). 2014. p. 2147-2154.
9. Redmon J, Farhadi A. YOLO9000: Better, faster, stronger. In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR). 2017. p. 6517-6525.
10. Shih YF, et al. Deep co-occurrence feature learning for visual object recognition. In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR). 2017.
11. He K, Zhang X, Ren S, Sun J. Deep residual learning for image recognition. In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR). 2016. p. 770-778.
12. Huang J, Rathod V, Sun C, Zhu M, Korattikara A, Fathi A, et al. Speed/accuracy trade-offs for modern convolutional object detectors. In: Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition (CVPR). 2017. p. 7310-7311.
13. Krizhevsky A, Sutskever I, Hinton GE. ImageNet classification with deep convolutional neural networks. In: Advances in Neural Information Processing Systems (NeurIPS). 2012. p. 1097-1105.
14. Pathak AR, Pandey M, Rautaray S. Application of deep learning for object detection. Procedia Comput Sci. 2018;132:1706-1717.
15. Zhang Q, Yang LT, Chen Z, Li P. A survey on deep learning for big data. Inf Fusion. 2018;42:146-157.
16. Abadi M, Barham P, Chen J, Chen Z, Davis A, Dean J, et al. TensorFlow: A system for large-scale machine learning. In: Proceedings of the 12th USENIX Symposium on Operating Systems Design and Implementation (OSDI). 2016. p. 265-283.
17. Shevlin H, Vold K, Crosby M, Halina M. The limits of machine intelligence: Despite progress in machine intelligence, artificial general intelligence is still a major challenge. EMBO Rep. 2019;20:e49177.
18. Shorten C, Khoshgoftaar TM. A survey on image data augmentation for deep learning. J Big Data. 2019;6(1). doi:10.1186/s40537-019-0197-0.

19. Brazil G, Yin X, Liu X. Illuminating pedestrians via simultaneous detection and segmentation. arXiv. 2017;arXiv:1706.08564.
20. Radford A, Metz L, Chintala S. Unsupervised representation learning with deep convolutional generative adversarial networks. arXiv. 2015;arXiv:1511.06434.
21. Lin Y, Han S, Mao H, Wang Y, Dally WJ. Deep gradient compression: Reducing the communication bandwidth for distributed training. arXiv. 2017;arXiv:1712.01887.
22. Wang S, Cheng J, Liu H, Tang M. PCN: Part and context information for pedestrian detection with CNNs. arXiv. 2018;arXiv:1804.04483.
23. Sultania F, Sufian A, Dutta P. A review of object detection models based on convolutional neural networks. arXiv. 2019;arXiv:1905.01614.
24. Shrivastava A, Sukthankar R, Malik J, Gupta A. Beyond skip connections: Top-down modulation for object detection. arXiv. 2016;arXiv:1612.06851.
25. Lee K, Choi J, Jeong J, Kwak N. Residual features and unified prediction network for single-stage detection. arXiv. 2017;arXiv:1707.05031.
26. Shen Z, Shi H, Feris R, Cao L, Yan S, Liu D, et al. Learning object detectors from scratch with gated recurrent feature pyramids. arXiv. 2017;arXiv:1712.00886.
27. Wong A, Shafiee MJ, Li F, Chwyl B. Tiny SSD: A tiny single-shot detection deep convolutional neural network for real-time embedded object detection. arXiv. 2018;arXiv:1802.06488.
28. Sermanet P, Eigen D, Zhang X, Mathieu M, Fergus R, LeCun Y. OverFeat: Integrated recognition, localization and detection using convolutional networks. arXiv. 2013;arXiv:1312.6229.
29. Sevo I, Avramovic A. Convolutional neural network-based automatic object detection on aerial images. IEEE Geosci Remote Sens Lett. 2016;13(5):740-744.
30. Zhu X, Wang Y, Dai J, Yuan L, Wei Y. Flow-guided feature aggregation for video object detection. In: Proceedings of the IEEE International Conference on Computer Vision (ICCV). 2017. p. 408-417.

Creative Commons License

This article is an open-access article distributed under the terms and conditions of the Creative Commons Attribution–Non-commercial–No Derivatives 4.0 International (CC BY-NC-ND 4.0) License. This license permits users to copy and redistribute the material in any medium or format for non-commercial purposes only, provided that appropriate credit is given to the original author(s) and the source. No modifications, adaptations, or derivative works are permitted.

Disclaimer: The views, opinions, statements, and conclusions expressed in the papers, abstracts, presentations, and other scholarly contributions included in this conference are solely those of the respective authors. The organisers and publisher shall not be held responsible for any loss, harm, damage, or consequences — direct or indirect — arising from the use, application, or interpretation of any information, data, or findings published or presented in this conference. All responsibility for the originality, authenticity, ethical compliance, and correctness of the content lies entirely with the respective authors.