

# Application of Improved Light-weight Feedback Attention Networks for Multi-Class Segmentation in Urban Scenes

Jiayi Yang  
*ECE-Department (MEng)*  
*Concordia University*  
Montreal, Canada  
jiayi.yang@mail.concordia.ca

Yuhang Chen  
*ECE-Department (MEng)*  
*Concordia University*  
Montreal, Canada  
yuhang.chen@mail.concordia.ca

Qian Zhang  
*ECE-Department (MEng)*  
*Concordia University*  
Montreal, Canada  
qian.zhang@mail.concordia.ca

## Abstract

Image segmentation is one of the most important techniques for computer vision. It is not only the foundation for biomedical imaging, but also essential in urban planning and autonomous driving. In this project, we extend the Feedback Attention Network (FANet), originally designed for precise biomedical segmentation, to address the complex demands of urban scenes through the Cityscapes dataset. The original FANet was designed for binary segmentation and applying it to multi-class segmentation is a non-trivial task. To address this challenge, we propose a new competitive layer for inference. This enhancement, coupled with model light-weighting and the incorporation of spatial and channel squeeze and excitation (scSE) modules, not only reduces the computational complexity but also maintains or even surpasses the original model performance. Through rigorous testing and evaluation, the modified FANet demonstrates robust adaptability and strong performance across diverse urban environments. The results highlight the model's versatility, confirming its extended applicability beyond its original medical domain and its potential as a powerful tool for a broad range of image segmentation tasks.

## Keywords

Feedback Attention Network, FANet, multi-class segmentation, urban scenes, model light-weighting

## I. INTRODUCTION

In the rapidly evolving field of computer vision, image segmentation is essential for interpreting and processing visual data across various applications. From detailed medical diagnostics to complex urban scene analysis, the ability to precisely delineate objects against diverse backgrounds is critical. Among the array of technologies developed for image segmentation, deep learning frameworks, particularly the Feedback Attention Network (FANet), have achieved notable advancements. Initially designed for binary segmentation of medical images, FANet effectively utilized feedback mechanisms across training epochs to set new benchmarks for segmentation precision.

However, FANet's potential is not confined to medical imaging alone. As image capture methods become increasingly complex, leading to an explosion of data complexity, there is a pressing need to adapt and optimize such deep learning frameworks for new domains. This paper discusses the expansion of FANet's capabilities from its traditional role in medical image segmentation to address multi-class segmentation challenges in urban environments. To facilitate this transition, we have introduced several methodological enhancements aimed at decoding the intricate structure of urban scenes, thereby bridging the gap between the precise needs of biomedical imaging and the varied components of natural scenes such as vehicles, pedestrians, and roadways.

The adaptation includes integrating a competitive layer during the prediction phase to effectively distinguish between multiple categories such as vehicles, pedestrians, and roadways, enhancing the network's ability to manage complex urban scenes. Additionally, to address the challenges of model light-weighting which is crucial for processing efficiency. Meanwhile, we replaced the original SE (Squeeze-and-Excitation) modules with scSE (concurrent spatial and channel squeeze and excitation) modules. This enhancement enhances the accuracy of the model amidst reductions in its complexity.

These improvements ensure that FANet's effectiveness extends beyond the confines of medical imaging to meet the multi-faceted challenges presented by our selected urban dataset. Through rigorous experimental simulation and analysis, we validate the effectiveness of our improved FANet approach and its implications for urban scene applications. This exploration underscores FANet's adaptability and marks a significant step towards its evolution into a versatile framework capable of tackling the next generation of image segmentation challenges.

## II. PROBLEM STATEMENT

The pursuit of advanced image segmentation technologies has predominantly focused on enhancing precision and accuracy within well-defined, controlled imaging environments. In the field of biomedical imaging, success largely hinges on the stability and predictability of medical images. However, moving beyond the structured realm of biomedical imaging to the more complex and unpredictable urban scene imaging introduces a series of unique and uncharted challenges.

In the processing of urban scenes, variations in environmental conditions such as lighting and weather, along with the presence of a higher density of overlapping objects—such as vehicles, pedestrians, and diverse types of roadways—introduce complexities not typically encountered in medical image segmentation. Consequently, the problem at hand is twofold. First, there is a need for a segmentation model that is not only sensitive to the subtleties of urban imagery, particularly to overlapping areas, but also robust against its inherent variabilities. Second, and more critically, this model must adaptively refine its learning process, utilizing iterative feedback to enhance its precision with each new data exposure.

To address these challenges, we have expanded FANet’s capabilities by integrating a competitive layer during the prediction phase. This addition facilitates multi-class segmentation by enabling the network to effectively distinguish between categories such as vehicles, pedestrians, and roadways, thus enhancing the model’s ability to manage the complexities of urban scenes. Moreover, to maintain high performance amidst the increased demand on computational resources, the network has undergone light-weighting modifications. We have strategically reduced the complexity of certain network modules without compromising the overall segmentation quality.

Additionally, to address the potential loss of accuracy due to the model’s light-weighting, the original SE (Squeeze-and-Excitation) modules were replaced with more advanced scSE (spatial and channel squeeze and excitation) modules. These modules are specifically designed to enhance feature recalibration capabilities both spatially and across channels, ensuring that segmentation accuracy is maintained or even improved despite the model’s reduced scale.

This necessitates a reevaluation of the existing feedback mechanisms and attention modules, originally designed for the homogeneity of medical datasets, repurposing them to accommodate the heterogeneity of urban scenes. As FANet transitions to this new application, assessing its performance and iteratively refining its architecture to ensure robust segmentation across multiple object classes in complex urban environments is crucial.

This paper will elaborate on the enhancements required for FANet to tackle these emerging challenges. Subsequent sections will detail the specific issues within the selected dataset, describe the adaptations to the network, and present an analysis of the outcomes to demonstrate how the revised FANet architecture rises to the complex task of multi-domain image segmentation.

## III. LITERATURE REVIEW

The field of image segmentation has witnessed significant advancements through the integration of deep learning techniques, specifically in adapting to complex imaging environments beyond traditional medical settings. Here, we highlight image segmentation for urban scenes, along with key developments and methodological innovations that have influenced current practice.

### A. Image Segmentation Techniques for Urban Environments

The adaptation of segmentation models to urban environments has necessitated handling a higher degree of variability and unpredictability in scene composition. One seminal work in this area is by Cordts et al. (2016)[1], who demonstrated the use of Convolutional Neural Networks (CNNs) for semantic understanding of urban scenes. Their research highlighted the ability of CNNs to accurately recognize and categorize critical urban elements like roads, vehicles, and pedestrians, even in densely populated settings.

### B. Advancements in Deep Learning for Urban Scene Analysis

Deep learning frameworks have continuously evolved to address the dynamic aspects of urban environments. Neuhold et al. (2017)[2] extended the capabilities of deep learning in urban scene segmentation by introducing a method that leverages large-scale video data to improve the temporal consistency and accuracy of segmentation models. This approach is particularly beneficial for applications such as autonomous driving, where understanding the temporal dynamics of urban scenes is crucial.

### C. Utilization of Attention Mechanisms in Urban Scene Segmentation

The application of attention mechanisms in image segmentation has provided significant improvements in the discriminative capabilities of models focused on urban scenes. A noteworthy contribution was made by Zhao et al. (2017)[3], who incorporated pyramid scene parsing networks that apply attention mechanisms to aggregate context from different scales, thus enhancing the segmentation of complex urban landscapes. This methodology has improved the accuracy and efficiency of segmenting detailed urban components.

#### *D. Context-Aware Segmentation Approaches*

Recent advancements have also focused on context-aware segmentation techniques that adapt to the specific challenges presented by urban environments. For instance, Zhang et al. (2018)[4] developed a context encoding module that effectively captures contextual relationships between different urban elements, significantly boosting the performance of segmentation tasks in heterogeneous urban settings.

### IV. PROBLEM FORMULATION

Modify the original FANet approach to suit the complexities of the new datasets, including changes in input data characteristics and expected segmentation outcomes. This section involves multiple components and will outline the theoretical modifications proposed for adapting FANet to handle diverse data types effectively:

#### *A. Dataset Replacement*

To adapt the FANet model, originally utilized in the domain of medical imaging for binary classification tasks, to the new and complex urban scenarios presented by the Cityscapes dataset, it was essential to modify the network structure. This adaptation involved introducing a competitive layer during the prediction phase of the model, transitioning from binary to multi-class segmentation. This adjustment enables effective categorization of various urban elements such as vehicles, pedestrians, and roadways. Such an adaptation necessitates a comprehensive reevaluation of the network architecture to cope with the increasing complexity and diversity of urban scenes. This step ensures that the model can efficiently manage and accurately process the details of city landscapes.

#### *B. Introduction of a Competitive Layer*

To facilitate the multi-class segmentation, a competitive layer has been integrated during the prediction phase. This layer functions to discern and classify multiple overlapping urban elements, enhancing the network's ability to differentiate between closely situated and often overlapping objects in urban environments.

#### *C. Model Light-weighting*

Given the computational intensity associated with processing high-resolution urban imagery, reducing the model's size and complexity was imperative. This light-weighting was achieved by scaling down certain network modules, thus decreasing the computational demand and improving the model's responsiveness and usability in resource-constrained environments.

#### *D. Accuracy Preservation Amidst Model Reduction*

A critical aspect of the model adaptation was to ensure that the reduction in complexity did not compromise the segmentation accuracy. To address the potential loss of accuracy due to model light-weighting, the original SE (Squeeze-and-Excitation) modules were replaced with scSE (concurrent spatial and channel squeeze and excitation) modules. These modules are designed to enhance feature recalibration capabilities both spatially and per channel, thereby maintaining, and sometimes even improving, segmentation accuracy despite the reduced model size.

#### *E. Optimization Goal*

The overarching goal is to develop an optimized version of FANet that not only addresses the increased complexity and class variability found in the Cityscapes dataset but also operates efficiently within the constraints of reduced computational resources. The objective is to achieve this without sacrificing the accuracy required for reliable urban scene analysis.

### V. ANALYSIS

#### *A. Original FANet Model: Principles and Application*

The original Feedback Attention Network (FANet) is designed for high-precision biomedical image segmentation, leveraging the principles of deep learning to enhance segmentation accuracy. The core mechanism of FANet centers on feedback loops that utilize errors from previous predictions to refine future outputs. This innovative approach integrates attention mechanisms to focus computational resources on critical areas of the image that are prone to segmentation errors. The specific network architecture is as follows:

•**SE-Residual Block:** The SE-Residual Block(Fig. 1), central to FANet's architecture, consists of two 3x3 convolutions, each followed by batch normalization (BN) and ReLU activation. An identity mapping links the input and output of these convolutions, helping mitigate the effects of vanishing or exploding gradients common in deep networks.

Adopting techniques from Hu et al.[4], the SE layer within the block dynamically recalibrates channel-wise feature responses using a global average pooling that compresses feature maps into a channel descriptor. This descriptor is then processed through a neural network that scales feature channels, enhancing relevant features and suppressing less significant ones.

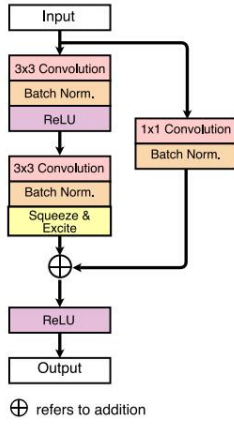


Fig. 1: SE-Residual Block

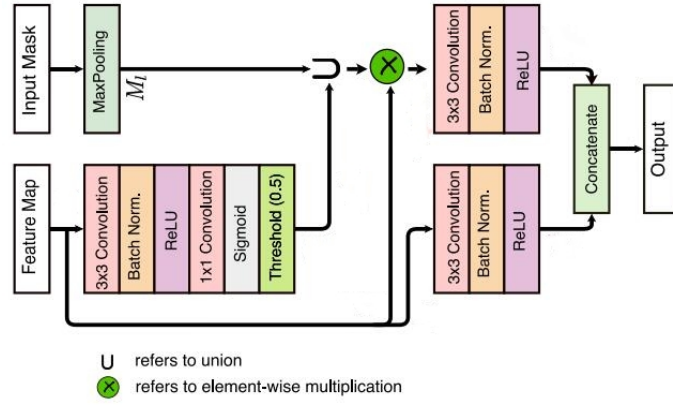


Fig. 2: MixPool Block

•**MixPool Block:** The MixPool block(Fig. 2) in FANet introduces hard attention via binary masks that selectively process features, focusing computational efforts on crucial attributes while ignoring the irrelevant. This selective enhancement and suppression optimize computational efficiency and enhance the model’s interpretability.

•**FANet Encoder-Decoder Structure:** FANet employs an encoder-decoder framework(Fig. 3), augmented with SE-residual and MixPool blocks, to capture deep semantic information and reconstruct accurate segmentation maps. The architecture features recurrent learning mechanisms that allow continuous refinement of predictions, leveraging feedback from prior outputs to improve future segmentation tasks. This streamlined architecture highlights FANet’s ability to dynamically learn and adapt to complex segmentation challenges, offering a marked improvement over traditional methods by incorporating iterative refinement and targeted attention mechanisms.

FANet employs a convolutional neural network (CNN) architecture enhanced with attention modules, which help to emphasize salient features while suppressing less relevant information. The feedback mechanism is a distinctive feature that differentiates FANet from other segmentation models, allowing for continuous improvement of the segmentation results as the network learns from its previous predictions.

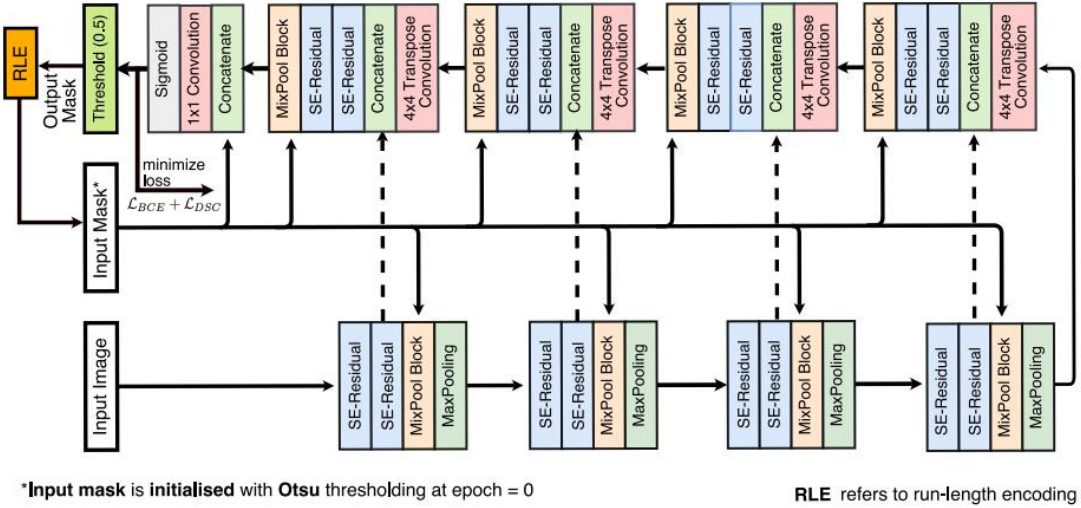


Fig. 3: FANet Framework

## B. Improved FANet Model with New Dataset "Cityscapes"

1) *Dataset Description:* In this study, we use a publicly available dataset named Cityscapes applied in the field of autonomous driving. Cityscapes is a semantic understanding image dataset about urban street scenes. It mainly contains street scenes from 50 different cities and has 3475 high-quality pixel-level annotated images of driving scenes in urban environments. The resolution size of the dataset images is 2048x1024. We have divided the dataset into 3127 and 348 images. Among them, 3127 images



are used as the training set, and the rest are used as the verification set and the test set. The samples of the Cityscapes Dataset are shown in Figure 4 and 5. In order to adapt the dataset to the network structure of FANet, we have segmented the dataset with a total of four categories labeled as vehicles, pedestrians, roads, and other backgrounds, which are shown in Figure 6.



Fig. 4: The samples of Cityscapes Dataset

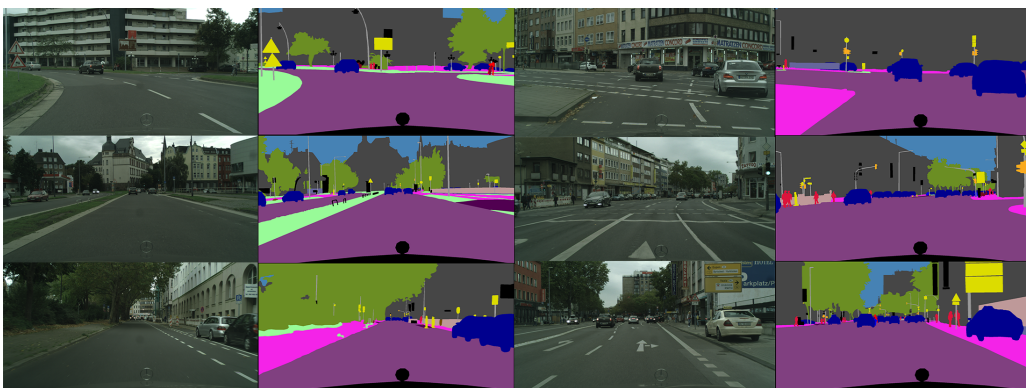


Fig. 5: The samples of labels

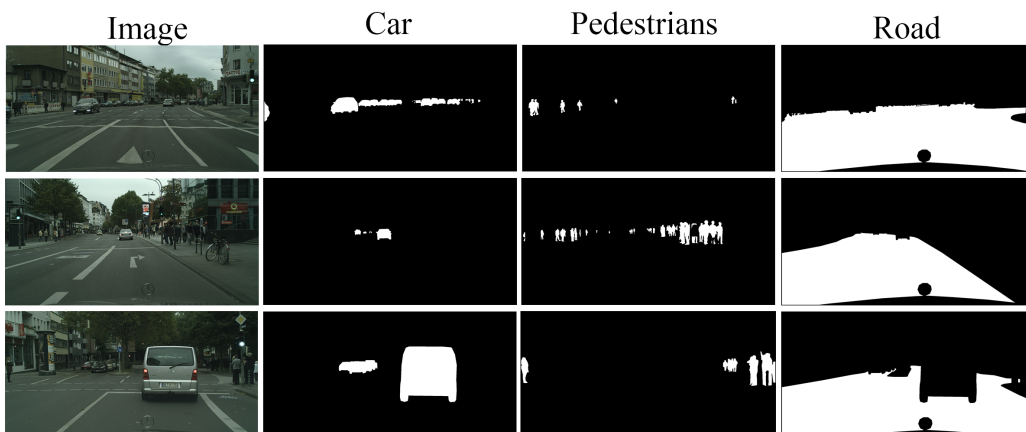


Fig. 6: The segment of labels

2) *Adaptation to Cityscapes Dataset by Introducing a Competitive Layer*: The transition to the Cityscapes dataset, which captures complex urban scenes, required significant modifications to the original FANet model. The primary challenge was the adaptation from a binary classification system, which segments images into foreground and background, to a multi-class system capable of identifying multiple distinct categories such as vehicles, pedestrians, and various types of road infrastructure.

To address this challenge, a competitive layer was introduced in the prediction phase of the network. This layer operates by evaluating the activations corresponding to different class predictions and enforcing competition among them. This mechanism ensures that the most relevant features for each class are enhanced while others are suppressed, making it particularly effective for scenes with overlapping objects where clear delineation is crucial. In the experiments of this paper, we use different labels to train the network separately, which can create different weights. Then during testing, we use the different models generated during training to extract semantic information from a single input image. In this step, we get several different feature maps. Finally, we add this data to the competitive layer, which finds the most probable category in each pixel and finally outputs a multi-categorized semantic segmented image

3) *Model Light-weighting*: Adapting FANet for the complex Cityscapes dataset required not only an expansion of capabilities to handle multi-class segmentation but also necessitated significant modifications to reduce the network's size and complexity. This process, known as model light-weighting, is crucial for enhancing the network's operational efficiency, particularly in resource-constrained environments.

The lightweight of the FANet architecture involved the selective reduction of network modules and layer dimensions, which effectively reduced the computational load while maintaining essential performance capabilities. This approach ensures quicker response times and a more user-friendly experience, as the streamlined model demands fewer computational resources. Here we reduce the dimensions of the two corresponding groups of network modules and layers, as shown in Fig 9

4) *Optimization with Attention Mechanism Modules*: Inspired by the Spatial and Channel Squeeze & Excitation Attention Module[6], we propose a novel residual block named scSE-Residual. After we lightweight the network structure, the result of the accuracy of image segmentation will be decreased to some extent. Therefore we used Spatial and Channel Squeeze & Excitation Attention Block (scSE) to combine with the residual block of the network. This attention block not only possesses less number of parameters, but also enhances the feature extraction ability of the network. The structure of the scSE attention module is shown in Fig.7.

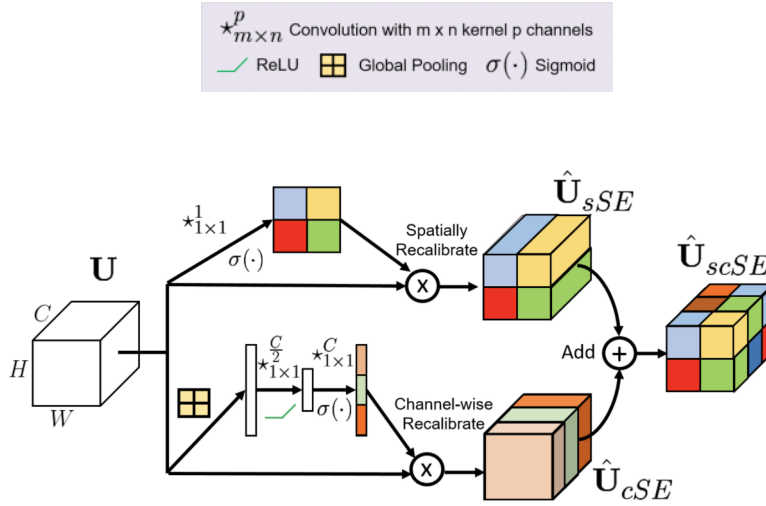


Fig. 7: Spatial and Channel Squeeze Excitation(scSE) Block

As can be seen from the figure, scSE attention is a combination of two SE attention mechanisms. The upper branch of the module is a Channel Squeeze and Spatial Excitation Block (sSE). The lower branch is a Spatial Squeeze and Channel Excitation Block (cSE). In the sSE module, the incoming feature layer  $U$  ( $H \times W \times C$ ) will be input, where  $H$  represents the height of the feature map,  $W$  represents the width of the feature map, and  $C$  represents the number of channels in the feature map. The input feature map is immediately processed using a  $1 \times 1 \times 1$  convolutional layer to output a  $H \times W \times 1$  spatial attention map. Finally, activated by a sigmoid function, the weights of each spatial location are generated and the feature layer  $\hat{U}_{sSE}$  is output. This processing is shown in expression 1, where  $f$  represents the convolution operation,  $U$  represents the input feature map, and  $kernel$  represents the size of the convolution kernel. The sSE attention module is a spatial attention mechanism that works by focusing on spatial locations in the feature map. By emphasizing important local regions in the image, the spatial

attention mechanism enables the model to better capture critical spatial information, such as features of key parts or specific shape features.

$$\hat{U}_{sSE} = \sigma(f(U, kernel)) \quad (1)$$

The bottom branch of the scSE attention module is a cSE attention module. cSE is a channel attention module whose main role is to re-assign weights to the information content of each channel of the input feature map. By learning the importance of different channels, the model can select the more important feature channels while suppressing those that are not important. This helps the model to focus on those features that are more critical to the task at hand. The feature map undergoes two processes through the cSE attention module, the Squeeze operation and the Excitation operation. First the feature map  $U$  goes through the Squeeze operation first, which is a global average pooling of the feature map  $U$  to obtain a  $1 \times 1 \times C$  vector, where the average of each channel represents the global response of that channel. The Squeeze operation is shown by expression 2, where  $x_{i,j,c}$  are the values of the  $c^{th}$  channel of the feature map at position  $i,j$ . Next, the Excitation operation is performed on  $Z_c$  obtained after the Squeeze operation to learn the weights of each channel from the global information of that channel using two fully connected layers. In the first fully connected layer, the number of channels is dimensionalized down to  $\frac{1}{2}$  of the original number of channels, and then restored to the original number of channels  $C$  in the second fully connected layer. This “bottleneck” structure allows the network to extract the most useful information from the features obtained by global average pooling, and then enhance the importance of this critical information by dimensionalizing back to a higher dimensional space. Eventually, the weights of each channel are fixed to 0-1 by a sigmoid activation function, and the Excitation operation is shown by expression 3, where  $W1$  represents the weights of the first full connection and  $W2$  represents the weights of the second full connection.

$$Z_c = \frac{1}{H \times W} \sum_H \sum_W x_{i,j,c} \quad (2)$$

$$\hat{U}_{cSE} = \sigma(W1, W2) \quad (3)$$

Followed by, the combination of the sSE module with the cSE module through element-wise addition forms the scSE attention mechanism module. scSE attention module not only possesses the advantage of the channel attention mechanism (cSE) for the selection of the important features of each channel, but also possesses the ability of the spatial attention mechanism (sSE) to capture the key spatial information. sSE with the cSE fusion method is shown by expression 4.

$$\hat{U}_{scSE} = \hat{U}_{sSE} + \hat{U}_{cSE} \quad (4)$$

Eventually, the scSE attention mechanism module is combined with the residual block to form our proposed scSE-Residual Block. The scSE-Residual Block’s structure is shown in Fig. 8. The network structure diagram of scSE-Residual Block introduced after lightweight improvement of FANet is shown in Fig. 9.

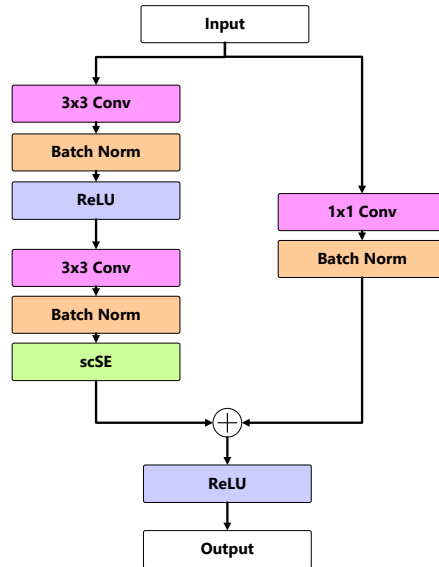


Fig. 8: Structure of scSE-Residual

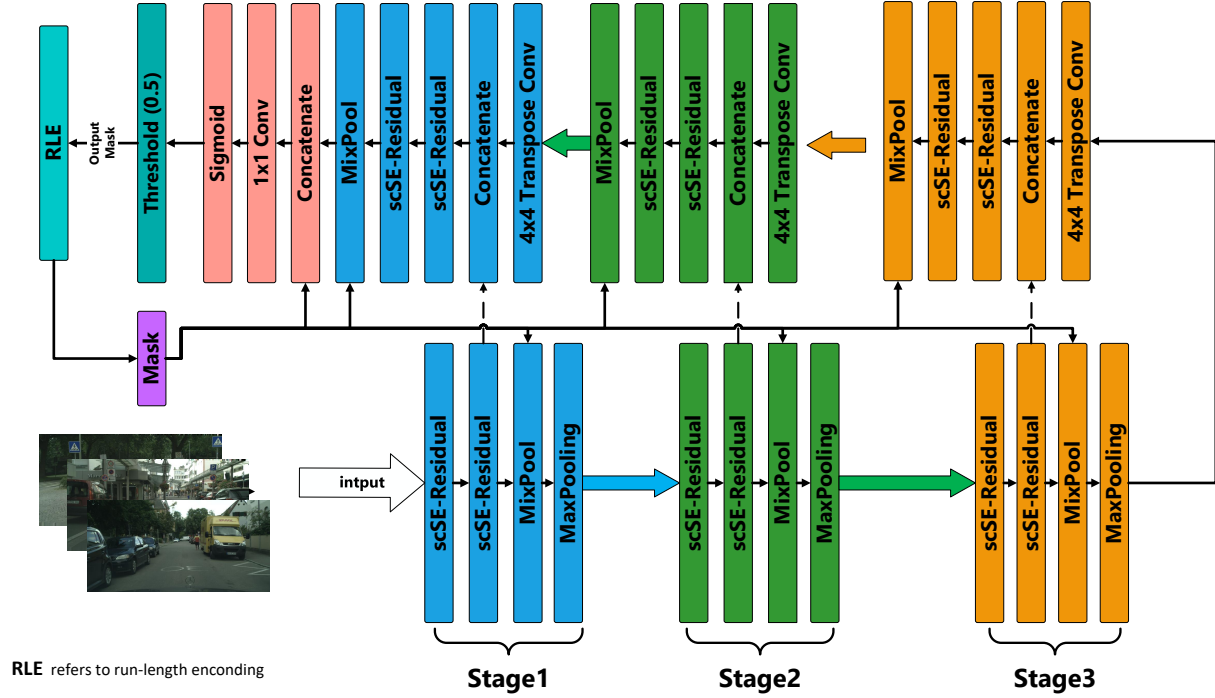


Fig. 9: Structure of Lightweight FANet With Attention Mechanism

As shown in the Fig. 9, we replaced the SE-Residual Block with scSE-Residual Block, and we reduced upsampling and downsampling block to three block respectively.

## VI. EXPERIMENTAL RESULTS

### A. Implementation Details

All the training is performed on a T4 GPU using the PyTorch 2.2.1 framework. For test inference, we have used an NVIDIA RTX 4060 GPU for our method. Our model is trained for 50 epochs (empirically set) using an Adam optimizer with a learning rate of  $1 * \exp^{-3}$  for all the experiments, batch size was set as 12.

### B. Results and Analysis

As shown in Figs. 10, the three categories of vehicles, people, and roads were trained using 50 epochs to achieve the fitted state. As shown in Table 1, the model size of our proposed method is reduced to 25% of the size of the original method.

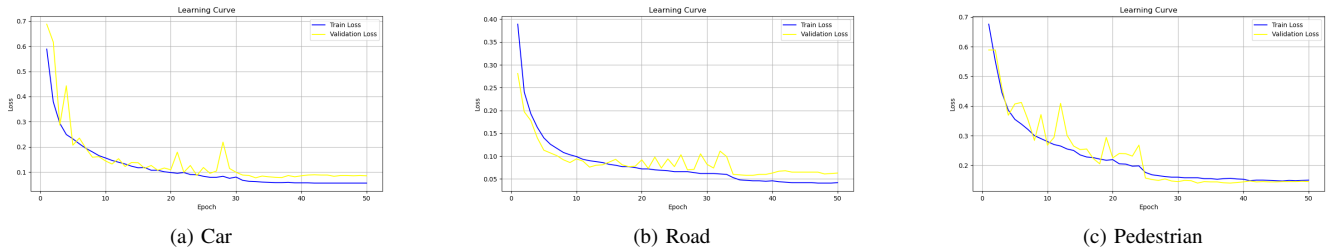


Fig. 10: Learning curves for vehicles (a), roads (b) and people (c)

Although we have lightened the network, which can lead to a decrease in accuracy, we have added a residual module with a scSE attention mechanism to compensate for this significant decrease in accuracy. The accuracy results of our approach in

semantic segmentation of vehicles and roads are almost close. The results in pedestrian category segmentation are better than the original method, Jaccard value rises from 0.4528 to 0.5413, F1 value rises from 0.5533 to 0.6495, Recall rises from 0.572 to 0.6709, Precision value rises by 11.32% to 77.6%. f2 value rises from 0.549 to 0.645. mIoU value went up by 0.0452 to 0.7292.

| Methods              | Category | Jaccard       | F1            | Recall        | Specificity   | Accuracy        | F2              | mIoU          | Size(MB) |
|----------------------|----------|---------------|---------------|---------------|---------------|-----------------|-----------------|---------------|----------|
| Original             | car      | 0.0263        | 0.0322        | <b>0.1682</b> | 0.9046        | 0.915596        | 0.037733        | <b>0.4614</b> | 30       |
| Light-weighting      | car      | 0.0246        | 0.029         | 0.1641        | 0.905         | 0.905954        | 0.034053        | 0.4606        | 7.6      |
| Light-weighting+scSE | car      | <b>0.0267</b> | <b>0.0329</b> | 0.1669        | <b>0.9047</b> | <b>0.915748</b> | <b>0.038513</b> | 0.4612        | 7.8      |
| Original             | people   | 0.4528        | 0.5533        | 0.572         | 0.9176        | 0.995089        | 0.549075        | 0.684         | 30       |
| Light-weighting      | people   | 0.4939        | 0.6054        | <b>0.6823</b> | 0.9169        | 0.995032        | 0.62651         | 0.7048        | 7.6      |
| Light-weighting+scSE | people   | <b>0.5413</b> | <b>0.6495</b> | 0.6709        | <b>0.9184</b> | <b>0.996241</b> | <b>0.645008</b> | <b>0.7292</b> | 7.8      |
| Original             | path     | <b>0.0006</b> | <b>0.0012</b> | <b>0.1094</b> | 0.6625        | 0.655846        | <b>0.002224</b> | 0.3286        | 30       |
| Light-weighting      | path     | 0.0005        | 0.001         | 0.1082        | <b>0.6719</b> | 0.644969        | 0.001973        | <b>0.3326</b> | 7.6      |
| Light-weighting+scSE | path     | 0.0004        | 0.0008        | 0.1088        | 0.6638        | <b>0.656963</b> | 0.001557        | 0.3282        | 7.8      |

TABLE I: Result

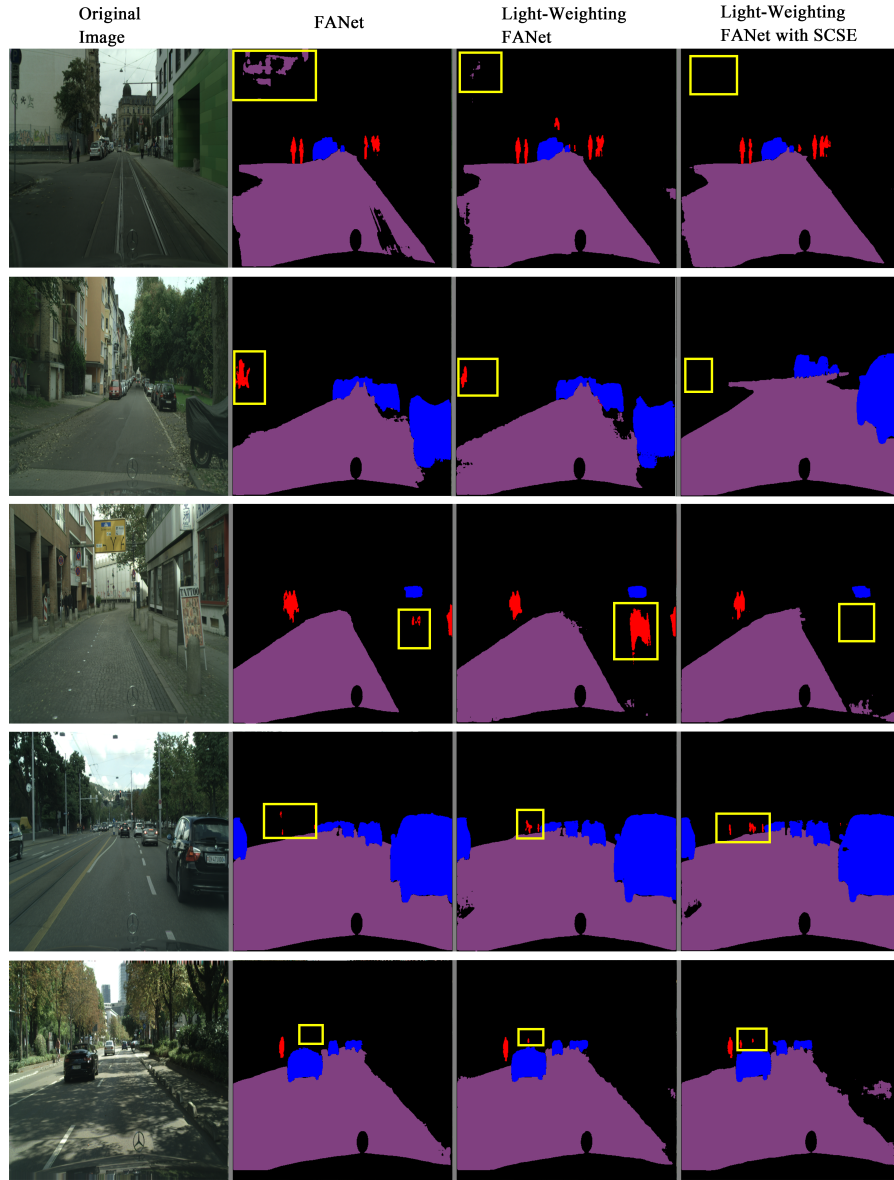


Fig. 11: Visulization Results

It can also be seen from the visualization results in Fig.11 that although our method still has some misclassifications and missed detections. In comparison to the original method, which in the first set of results recognizes buildings as roads as well, the method we have adopted avoids misclassification to a certain extent. In the second set of results, the original method recognizes walls as pedestrians, and our method effectively avoids this misclassification. And in this set of results, the segmentation accuracy of the original method for roads is not as high as that of our proposed method. In the fourth set of experimental results, it is shown that our method can effectively mitigate the high rate of missed detection for pedestrians. In the original method, only one pedestrian can be segmented and the segmentation accuracy is not particularly high, while our method can segment four pedestrians and the accuracy is significantly higher than the original method.

## VII. DISCUSSION

These results highlight the potential of applying advanced segmentation techniques, originally developed for medical imaging, to broader applications such as urban scene analysis. The successful adaptation of FANet to urban landscapes underscores the flexibility and scalability of the original architecture.

The integration of scSE modules has proven crucial not only in maintaining segmentation accuracy but often in enhancing it, despite the model's reduced complexity after light-weighting. However, challenges remain in handling edge cases and extremely crowded urban scenarios. Future work will focus on further enhancing the model's adaptability and exploring additional techniques for feature recalibration.

## VIII. CONCLUSION

The adaptation of the Feedback Attention Network (FANet) from its original application in biomedical imaging to the complex dynamics of urban scene segmentation with the Cityscapes dataset marks a substantial advancement in the field of image segmentation. This adaptation not only extended FANet's capabilities to tackle multi-class segmentation in densely populated urban settings but also demonstrated the model's scalability and flexibility. Strategic enhancements, including the integration of a competitive layer and scSE modules, have not only upheld but in many instances improved the segmentation accuracy. Additionally, efforts to lighten the network have significantly increased its suitability for deployment in resource-constrained environments, enhancing its applicability for real-time critical applications such as urban planning and autonomous driving.

While the adaptation of FANet to urban scene segmentation has shown promising results, there are several avenues for further development to enhance its effectiveness and applicability in real-world scenarios. Moving forward, we aim to refine the accuracy and robustness of our segmentation algorithm by integrating more advanced deep learning models, which will better handle the intricate and variable aspects of urban environments. To address the fluctuating conditions encountered in autonomous driving, such as diverse weather and lighting scenarios, we plan to expand our dataset to include these complexities. This expansion will allow for more comprehensive training and testing, ensuring the model's robust performance under varied conditions.

Moreover, we recognize the critical need for real-time processing capabilities in autonomous driving systems. Our ongoing efforts will focus on optimizing the algorithm's speed and efficiency to deliver quick responses essential for the safety and reliability of self-driving vehicles. Ultimately, we will conduct extensive on-road testing to validate the performance of our updated models, ensuring they meet the high safety standards required for deployment in autonomous driving applications.



## REFERENCES

- [1] M. Cordts et al., “The Cityscapes Dataset for Semantic Urban Scene Understanding,” in 2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR), Las Vegas, NV, USA: IEEE, Jun. 2016, pp. 3213–3223. doi: 10.1109/CVPR.2016.350.
- [2] G. Neuhold, T. Ollmann, S. R. Bulò, and P. Kotschieder, “The Mapillary Vistas Dataset for Semantic Understanding of Street Scenes,” in 2017 IEEE International Conference on Computer Vision (ICCV), Venice: IEEE, Oct. 2017, pp. 5000–5009. doi: 10.1109/ICCV.2017.534.
- [3] H. Zhao, J. Shi, X. Qi, X. Wang, and J. Jia, “Pyramid Scene Parsing Network,” in 2017 IEEE Conference on Computer Vision and Pattern Recognition (CVPR), Honolulu, HI: IEEE, Jul. 2017, pp. 6230–6239. doi: 10.1109/CVPR.2017.660.
- [4] H. Zhang et al., “Context Encoding for Semantic Segmentation,” in 2018 IEEE/CVF Conference on Computer Vision and Pattern Recognition, Salt Lake City, UT, USA: IEEE, Jun. 2018, pp. 7151–7160. doi: 10.1109/CVPR.2018.00747.
- [5] J. Hu, L. Shen, S. Albanie, G. Sun, and E. Wu, “Squeeze-and-Excitation Networks,” IEEE Trans. Pattern Anal. Mach. Intell., vol. 42, no. 8, pp. 2011–2023, Aug. 2020, doi: 10.1109/TPAMI.2019.2913372.
- [6] A. G. Roy, N. Navab, and C. Wachinger, “Recalibrating Fully Convolutional Networks With Spatial and Channel ‘Squeeze and Excitation’ Blocks,” IEEE Trans. Med. Imaging, vol. 38, no. 2, pp. 540–549, Feb. 2019, doi: 10.1109/TMI.2018.2867261.
- [7] N. K. Tomar et al., “FANet: A Feedback Attention Network for Improved Biomedical Image Segmentation,” IEEE Trans. Neural Netw. Learning Syst., vol. 34, no. 11, pp. 9375–9388, Nov. 2023, doi: 10.1109/TNNLS.2022.3159394.
- [8] J. Long, E. Shelhamer, and T. Darrell, “Fully convolutional networks for semantic segmentation,” in 2015 IEEE Conference on Computer Vision and Pattern Recognition (CVPR), Boston, MA, USA: IEEE, Jun. 2015, pp. 3431–3440. doi: 10.1109/CVPR.2015.7298965.
- [9] G. Han et al., “Improved U-Net based insulator image segmentation method based on attention mechanism,” Energy Reports, vol. 7, pp. 210–217, Nov. 2021, doi: 10.1016/j.egy.2021.10.037.