

DAMAGED ROAD EXTRACTION BASED ON SIMULATED POST-DISASTER REMOTE SENSING IMAGES

Yansong Huang, Haocai Wei, Junli Yang*, Ming Wu

Beijing University of Posts and Telecommunications
yangjunli@bupt.edu.cn

ABSTRACT

Damaged road extraction is a challenging task in the field of remote sensing. Some existing methods include the step to extract road from pre- and post-disaster remote sensing images of the same area. In practice, it often occurs that one of these two images is missing. To solve this problem, we use CoCosNet, the model for exemplar-based image translation, to translate pre-disaster images to simulated post-disaster ones. Then we use D-LinkNet, the state-of-the-art method in road extraction, to extract road from the pre- and post-disaster images of the same area. We extract damaged road area by comparing pre-disaster road masks with post-disaster ones and output the damage level by calculating the proportion of the damaged road area. Finally, we evaluate the damaged road extraction accuracy. Experimental results on simulated post-disaster images prove the effectiveness of the simulation method and the framework for damaged road extraction and damage level evaluation.

Index Terms— damaged road extraction, exemplar-based image translation, remote sensing image, post-disaster

1. INTRODUCTION

Road extraction is a basis task in the field of remote sensing. Damaged road extraction is a specific application of road extraction. It is especially important after disasters like earthquake and flood because disaster assessment and management highly depend on road damage information. Some of the existing methods for damaged road extraction contain a similar step, that is extracting road from both pre- and post-disaster remote sensing images of the same area[1]. By comparing these two extracted road masks, information of damaged road is obtained, including the area and length of damaged road. However, in practice, this condition is hard to satisfy. It often occurs that the pre-disaster images are missing when the disaster has taken place, or the pre-disaster images exists but disaster has never taken place in this area.

To solve this problem, we resort to exemplar-based image translation technique to simulate high quality post-disaster

remote sensing images automatically. Zhang et al. proposed CoCosNet[2] in 2020, which is a state-of-the-art method for image translation. Given the exemplar image, CoCosNet translates it to target image based on the target semantic segmentation masks. CoCosNet establishes the reliable dense cross-domain correspondence[2] between the exemplar image and target semantic segmentation masks. This cross-domain correspondence enables CoCosNet to outperform other state-of-the-art methods in terms of image quality by a large margin. In this task, for we already have the pre-disaster images as exemplar images, we erase part of road masks in the original segmentation masks associated with pre-disaster images to generate target semantic segmentation masks. In this way, CoCosNet can synthesize simulated post-disaster images, namely target images, from pre-disaster images and altered segmentation masks.

After generating the post-disaster images, we apply D-LinkNet, the state-of-the-art method in road extraction, to pre- and post-disaster images to extract road respectively. By comparing these two road masks, we obtain damaged road masks.

Finally, we assess the result of damaged road extraction by comparing the extracted damaged road masks with the erased road masks, namely ground truth of damaged road. The evaluation in IoU demonstrates that damaged road is extracted precisely, which verifies the effectiveness of the simulation method and the framework for damaged road extraction and damage level evaluation.

2. APPROACH

2.1. Framework Description

The framework will be described in this section. The framework diagram is shown in Fig.1. The details of the steps are listed as follows:

(1) Preprocess: In this step, we edit the semantic segmentation mask of the pre-disaster images. The detail will be given in the next subsection. This is the pre-process for generating post-disaster images.

(2) Generate simulated post-disaster remote sensing images: In this step, we use CoCosNet to synthesize post-disaster images from pre-disaster images and altered seman-

*Corresponding author

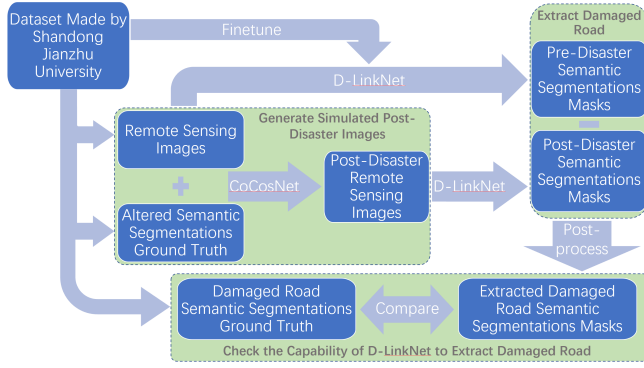


Fig. 1. Framework Diagram

tic segmentation masks. The details of CoCosNet will be given in the later subsection. The simulated post-disaster images will have similar style with corresponding pre-disaster images while part of road is removed, indicating the effect of disaster.

(3) Finetune D-LinkNet: To improve the performance of D-LinkNet, we first pretrain it on the DeepGlobe Road Extraction dataset and then finetune it on the SJD Roads Dataset. The details will be given in the next section.

(4) Extract road: In this step, we use the finetuned D-LinkNet model to extract road from pre- and post-disaster images respectively.

(5) Extract damaged road: In this step, we compare pre- and post-disaster road network pixel-by-pixel to obtain the damaged road. If a pixel is labelled as road in pre-disaster images but non-road in post-disaster images, it will be labeled as damaged road.

(6) Postprocess: The extracted masks of the damaged road are subjected to denoising using a simulated annealing algorithm[3].

(7) Assess the performance of damaged road extraction: In this step, we compare the extracted damaged road to the ground truth and analyze the result.

Following subsections are descriptions of main processes in detail.

2.2. Preprocess

To simulate the phenomenon of damaged road, we turn part of pixels whose label is road to pixels whose label is non-road. Firstly, we select pixels in a square area, which has fixed length but random location. The number of pixels whose label is road in this square area is summed up. If it is larger than a default value, all the pixels in this area will be set as non-road. Otherwise, the location of the square area will be randomly changed. The steps above will be repeated until the label of pixels is altered. The diagram of this process is shown in Fig.2.

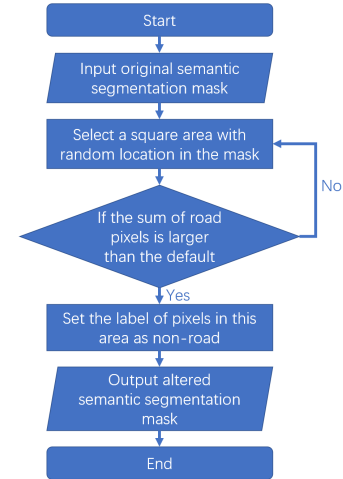


Fig. 2. Flow diagram of editing semantic segmentation masks.

2.3. CoCosNet

In this step, we use CoCosNet to generate simulated post-disaster remote sensing images.

CoCosNet is proposed by Zhang et al. It was presented as an oral presentation at Computer Vision and Pattern Recognition, 2020. CoCosNet is an exemplar-based image translation technique. Combining the exemplar image and the target semantic segmentation mask, it can generate the image that has the style (e.g., color, texture) in consistency with the semantically corresponding objects in the exemplar. If the target semantic segmentation mask is similar to that of the exemplar image, the image quality of the generated image will be much better.

2.4. D-LinkNet

D-LinkNet is an effective model in the field of road extraction from remote sensing images. It was the winner of CVPR DeepGlobe 2018 Road Extraction Challenge[4]. The model itself is a binary semantic segmentation model for extracting roads from high resolution satellite images. D-LinkNet model has the advantages of skip connection, residual block and encoder-decoder architecture. And it has additional Dilated convolution in the central part, which can increase the receptive field of feature points without decreasing the resolution of the feature map.

2.5. Postprocess

Comparing image pixel-by-pixel to extract damaged road will cause randomly distributed noisy pixel points on the binary image. We know that the density of noise pixel is small and there should be a strong correlation between neighboring pixels in our binary road images. Based on these properties, we

use the simulated annealing algorithm combined with gradient descent in an iterative process to find a global optimal solution. We implemented this intelligent algorithm in python and successfully removed most of the peripheral noises.

3. EXPERIMENTS

3.1. Datasets

ADE20K[5]: consists of about 20k training images, each image associated with a 150-class segmentation mask.

DeepGlobe Road Extraction dataset[4]: consists of 6226 training images. The image size is 1024×1024 .

SJU Road Extraction dataset: made by Shandong Jianzhu University, which consists of 493 training images and 132 test images. The image size is 1280×1280 . All pixels in the images of the two datasets are labeled as road or non-road.

3.2. Implementation details

CoCosNet and D-LinkNet are both based on PyTorch as the deep learning framework. They are trained on 2 32G NVIDIA V100.

In the step of generating simulated post-disaster images, we used CoCosNet pretrained on ADE20k.

SJU Road Extraction dataset contains pre-disaster remote sensing images and associated segmentation masks. We first erased part of road masks in the original segmentation masks in SJU Road Extraction dataset. CoCosNet took pre-disaster images and altered segmentation masks as input, then output simulated post-disaster images. 132 pre-disaster images were translated. The image size of pre-disaster images and altered segmentation masks is 1280×1280 pixels. The image size of simulated post-disaster images is 256×256 pixels.

In the step of damaged road extraction, we use D-LinkNet pretrained on DeepGlobe Road Extraction dataset.

We finetuned our model in the training image of Shandong Jianzhu University dataset with the same DA (Data augmentation), TTA (Test time augmentation) method and loss function[6] as the original D-LinkNet model did. To correspond to the size of the simulated post-disaster image, we crop the size of the dataset from 1280×1280 to 256×256 for training. It took about half an hour for our finetuned model to converge.

For noise reduction of images, we use the simulated annealing algorithm combined with gradient descent. By continuously adjusting the algorithm parameters, the noise reduction effect reached a relatively good point.

3.3. Results and analysis

We first generated post-disaster images with CoCosNet, part of results is shown in Fig.3.

Then we evaluated the IoU of masks predicted by the models. The capability of pretrained and finetuned models

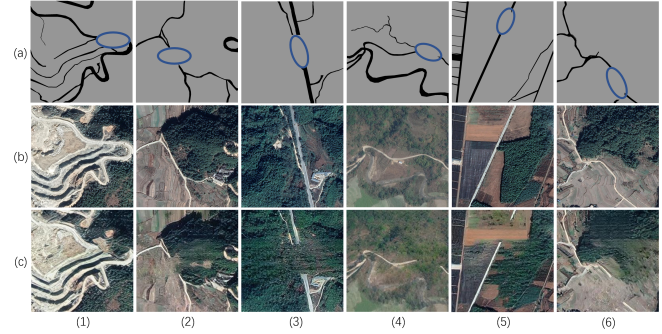


Fig. 3. Simulated post-disaster images. (a) Altered segmentation masks. Black pixels stand for road and gray pixels stand for background. The road masks in blue ellipses are erased. (b) Pre-disaster remote sensing images. (c) Simulated post-disaster images.

were evaluated separately. The performances of the two models are shown in Tab.1.

Table 1. Results on the test images of Shandong Jianzhu University dataset.

Model	IoU on test set
D-LinkNet34(pretrained)	0.1756
D-LinkNet34(finertuned)	0.4132

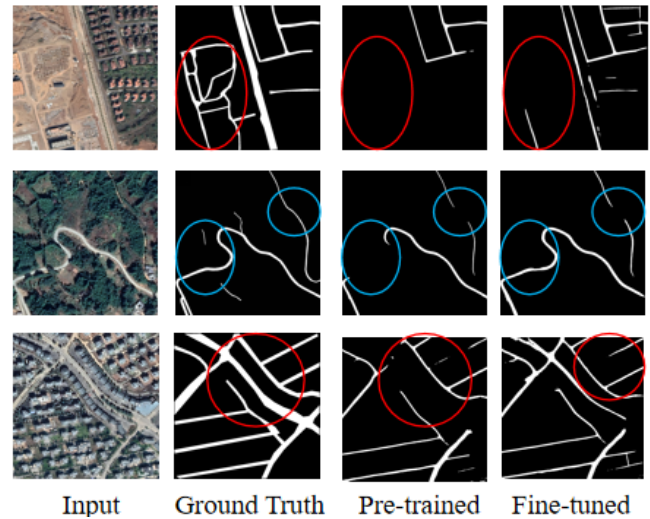


Fig. 4. Example results of pretrained and fine-tuned D-LinkNet34 models.

We found that the finetuned model obtains an improvement of 23.76% in IoU. The images in Fig.4 shows that the model has a good extraction ability for urban roads and a weak extraction ability for dirt roads as shown in the red circled region, and the connectivity at the end of the road is unsatisfied, as shown in the blue circled region.

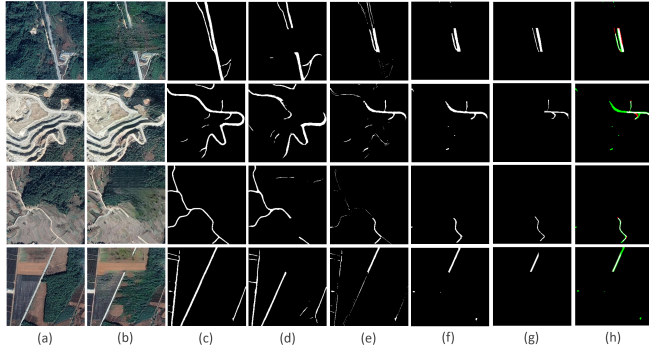


Fig. 5. Example results of our framework for damaged road extraction. (a) Pre-disaster image (b) Simulated post-disaster image (c) Extracted mask of pre-disaster image (d) Extracted mask of post-disaster image (e) Initial extracted damaged road (f) Denoised of images in (e) (g) Ground truth of damaged road (h) Assessment of results of damaged road extraction, white pixels represent correctly predicted damaged road areas, green pixels represent areas incorrectly predicted as damaged road, namely false positive, red pixels represent areas that are not predicted to be damaged road, namely false negative.

Then, we tested the ability of our model to extract damaged road regions. Compared the final damage road extraction result with the ground truth. The IoU of predicted damage area is 32.5%. Some examples of the results are shown in the Fig.5.

The visualization of assessment, namely (h) in Fig. 5 has following characteristics:

(1) Proportion of red pixels is small: Red pixels represent areas that are not predicted to be damaged road. Small proportion of red pixels means most of damaged road area is extracted.

(2) Green pixels beside white pixels: This phenomenon is shown in the first row of images in Fig. 5. Green pixels represent areas incorrectly predicted as damaged road and white pixels represent correctly predicted damaged road areas. The reason for these green pixels is that CoCosNet may displace the pixels representing the road when generating the simulated post-disaster remote sensing images, thus causing some displacement of the road extracted by D-LinkNet on the remote sensing images.

(3) Blocks of green pixels: This phenomenon is shown in the second and fourth row of images in Fig. 5. This is because of deficiencies in road extraction by D-LinkNet for simulated post-disaster remote sensing images.

Since damaged road extraction is performed on the pre-disaster images and simulated post-disaster images, the assessment of damaged road extraction proves the effectiveness of the simulation method and the framework for damaged road extraction and damage level evaluation.

4. CONCLUSION

In this paper, we propose a simulation method to synthesize simulated post-disaster images. Then we perform the task of damaged road extraction based on the pre- and post-disaster images. The IoU of predicted damage area achieves 32.5%. Assessment of the result proves the effectiveness of the simulation method and the framework for damaged road extraction and damage level evaluation.

Acknowledgement

This work is funded by Research Innovation Fund for College Students (No.202011065) and Teaching Reform Project No.2019JYA05) of Beijing University of Posts and Telecommunications.

5. REFERENCES

- [1] Moslem Ouled Sghaier and Richard Lepage, "Road damage detection from vhr remote sensing images based on multiscale texture analysis and dempster shafer theory," in *2015 IEEE international geoscience and remote sensing symposium (IGARSS)*. IEEE, 2015, pp. 4224–4227.
- [2] Pan Zhang, Bo Zhang, Dong Chen, Lu Yuan, and Fang Wen, "Cross-domain correspondence learning for exemplar-based image translation," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2020, pp. 5143–5153.
- [3] GY Chen, Tien D Bui, and A Krzyżak, "Image denoising with neighbour dependency and customized wavelet and threshold," *Pattern recognition*, vol. 38, no. 1, pp. 115–124, 2005.
- [4] Ilke Demir, Krzysztof Koperski, David Lindenbaum, Guan Pang, Jing Huang, Saikat Basu, Forest Hughes, Devis Tuia, and Ramesh Raskar, "Deepglobe 2018: A challenge to parse the earth through satellite images," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition Workshops*, 2018, pp. 172–181.
- [5] Bolei Zhou, Hang Zhao, Xavier Puig, Sanja Fidler, Adela Barriuso, and Antonio Torralba, "Scene parsing through ade20k dataset," in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2017, pp. 633–641.
- [6] Lichen Zhou, Chuang Zhang, and Ming Wu, "D-linknet: Linknet with pretrained encoder and dilated convolution for high resolution satellite imagery road extraction," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition Workshops*, 2018, pp. 182–186.