

Identification of nighttime urban flood inundation extent using deep learning

Jiaquan Wan^{1, 2, 3}, Xing Wang⁴, Yannian Cheng⁵, Cuiyan Zhang⁴, ~~Yufang Shen^{1, 2, 3}~~, Fengchang Xue⁵, Tao Yang^{1, 2, 3}, [Fei Tong⁶](#) and Quan J. Wang^{6,7}

¹ College of Hydrology and Water Resources, Hohai University, Nanjing, 210098, China

² The National Key Laboratory of Water Disaster Prevention, Hohai University, Nanjing, 210098, China

³ Yangtze Institute for Conservation and Development, Hohai University, Jiangsu, 210098, China

⁴ School of Computer Engineering, ~~NanJing~~[Nanjing](#) Institute of Technology, Nanjing 211167, China

⁵ School of Remote Sensing and Surveying Engineering, Nanjing University of Information Science & Technology, Nanjing 210044, China

⁶ [China Institute of Water Resources and Hydropower, Beijing, 100048, China](#)

^{6,7} Department of Infrastructure Engineering, Faculty of Engineering and Information Technology, The University of Melbourne, Victoria 3010, Australia

Correspondence to: Xing Wang (jwangxing0719@163.com)

Abstract. With the acceleration of urbanization, the disaster of urban flooding has had a serious impact on urban socio-economic activities and has become one of the important factors restricting social development in China. Accurate and timely identification of urban flooding extents is crucial for decision-making. Traditional remote sensing technologies are often limited by environmental factors, making them less suitable for application in complex urban terrains. [With the increase in urbanization and the development of emerging technologies, video imagery has become a significant data source with great potential for urban flood identification.](#) However, existing research has primarily focused on flood extent identification in daytime scenarios, often neglecting the nighttime, a period of high flood occurrence. [In this study, we propose an efficient model \(NWseg\) to identify flood extents in nighttime scenes.](#)~~The development of emerging technologies and the increase in urbanisation have led to a significant increase in the number of surveillance devices within cities, while the development of deep learning techniques has led to their widespread application in various fields. Deep learning methods using video images as a data source provide a new technical methods for intra-urban waterlogging recognition. However, current research mainly focuses on waterlogging recognition in daytime scenes, often ignoring nighttime, a time of high waterlogging incidence. To address these challenges faced by flooding recognition in the nighttime, this study proposes a deep learning model—NWseg—to achieve the recognition of the extent of waterlogging at night.~~ Initially, we constructed [a nighttime flood inundation dataset consisting of 4,000 images.](#)~~a dataset of 4,000 images of nighttime urban flooding.~~ Subsequently, MobilenetV2 and ResNet101 networks were used to replace the DeepLabv3+ backbone network and compared with the NWseg model. Next, the NWseg model [was](#) compared with ResNet50-FCN, LRASPP and U-Net models to evaluate the performance of different models in nighttime urban flooding [extent](#) identification. Finally, [we verified the applicability and performance differences of each model in specific environments. Overall, this study successfully](#)

demonstrates the effectiveness of the NWseg model for nighttime urban flooding extent identification, providing new insights for nighttime flood monitoring in cities. ~~the applicability and performance differences of each model in specific environments were verified. In conclusion, this study successfully demonstrates the effectiveness of the NWseg model for nighttime urban flooding identification and provides new insights for nighttime urban flooding identification.~~

Keywords: Deep learning, Nighttime flooding [extent](#) identification, Urban flooding, NWseg

1 Introduction

In recent years, extreme rainfall events have been occurring frequently in the context of complex climate change ([Burn and Whitfield, 2023](#); [Kim et al., 2024](#)). Concurrently, with the acceleration of urbanization processes, the proportion of impervious surfaces has been continuously expanding, resulting in serious urban flooding issues in many cities worldwide ([Ghosh et al., 2024](#); [Liu et al., 2023](#); [Kundzewicz et al., 2019](#))~~([Xue et al., 2023](#))~~. Urban flooding often coincides with multiple compounded disasters and may even trigger secondary calamities, posing serious threats to the safety of urban residents, the normal operation of city functions, and sustainable development. This exacerbates the vulnerability of urban socio-economic system ([Gu et al., 2025](#); [Visser, 2014](#); [Zheng et al., 2014](#))~~([Luo et al., 2020](#))~~. Therefore, achieving real-time and effective [identification](#) ~~monitoring~~ of urban flooding [extent](#) has become a critical issue that urgently needs to be addressed.

Remote sensing technology has made significant advancements in the field of urban flood [identification](#)~~monitoring~~, providing new perspectives for flood disasters identification through high spatial, temporal and spectral resolution data ([Bofana et al., 2022](#))~~([Hao., 2022](#))~~. However, despite its excellent performance at the macro scale, remote sensing technology has limitations in urban area monitoring. [Due to the limitations in temporal resolution and the impact of cloud cover and atmospheric variations, remote sensing technology struggles to capture the dynamic changes of urban flooding, making real-time identification of rapidly evolving flood events challenging](#) ([Mason et al., 2012](#))~~.Due to insufficient temporal resolution as well as the influence of cloud cover and changing atmospheric conditions, remote sensing techniques have difficulty in capturing subtle topographic changes within cities, and are unable to monitor fast-changing flooding events in real-time~~ ([Gao., 2023](#)). In addition, the complexity of the urban environment, especially the dynamic changes of small-scale water bodies and localized waterlogging, further increases the difficulty of remote sensing technology in urban [flooding](#) ~~flood-extent identification~~~~monitoring~~. Therefore, an intelligent and real-time urban flood monitoring method is urgently needed to achieve more precise flood identification.

With technological advancements, the emerging fields of deep learning and computer vision have matured and engaged in interdisciplinary collaborations, achieving [significant performance](#)~~remarkable results~~ that offer new technical approaches for urban flood identification ([Choi and Yoo, 2023](#)). Particularly in image [segmentation](#)~~recognition~~, deep learning's advantages

in extracting global features and contextual information make it highly promising for inundation detection (Liu et al., 2020)(Liao, 2023). Simultaneously, the increasing level of urbanization has led to the widespread deployment of video surveillance devices across urban areas, particularly in highly urbanized areas (Muhadi et al., 2024; Hao et al., 2022) along city roads, where they are ubiquitous. During rainfall, these cameras can fully record the flooding process, providing real-time reflections of road inundation changes (Wang et al., 2024)(Wang et al., 2024; Yang et al., 2022; Cheng et al., 2018). Therefore, combining deep learning with traffic cameras can effectively achieve the identification of urban flooding extent.

Existing research has demonstrated that deep learning excels in segmenting inundated areas. Sarp et al. (2020) (Bai et al., 2021) utilized the YOLOv2 object detection model to extract water accumulation features from images collected by Xi'an University of Science and Technology, achieving an average recognition accuracy of over 85% through multiple model training sessions, demonstrating the precision of this method for inundation area extraction. (Wang et al., 2021) classified road images into four categories—background, dry surface, inundated area, and slippery surface—and used the Res-UNet+ semantic segmentation network to handle different lighting and scene conditions, achieving an Mean Intersection over Union (MIoU) of 90.07%, outperforming traditional classification methods. (Sarp et al., 2020) applied the Mask R-CNN model to automatically detect and segment floodwaters in urban, suburban, and natural scenes, achieving 99% accuracy in the detection phase and 93% in the segmentation phase. (Sazara et al., 2019) Sazara et al. (2019) used a deep learning approach to detect standing water on urban roads, in which a pre-trained VGG-16 network was used in the classification phase and a full convolutional neural network was used in the segmentation task, and compared it with the traditional classifier and extraction algorithms with manually-designed features, and the results showed that the deep learning approach has a more obvious advantage in both the recognition and segmentation of standing water. Wang et al. (2024) used a deep convolutional neural network (DCNN) for urban flood extent recognition based on video images acquired from surveillance cameras. Zeng et al. (2024) proposed a DeepLabv3+ based flood image recognition method, which effectively improves the model performance through image enhancement and the introduction of the super-resolution generative adversarial network. However, current research focuses on daytime scenes, and the existing datasets lack diversity to cover flooding scenes at night or under complex weather conditions. Meanwhile, some algorithms underperform when processing images in low-light or adverse conditions, making flood extent identification at night or in challenging weather a technical challenge. This limitation underscores the urgent need for accurate nighttime flood extent identification monitoring—and the necessity for algorithm improvements and dataset expansion.

To address the above challenges, this study proposes an efficient method for nighttime urban flood extent identification. First, an urban flood inundation dataset for nighttime scenes is constructed to provide sufficient sample support for model training. Subsequently, a NWseg model for nighttime image segmentation is proposed, which combines a Content-Light Splitter with a Dual-Feature Integrator to enhance the model's performance in identifying flooding extent in low-light environments. Meanwhile, given that the data are mainly sourced from urban road surveillance systems, the method is

particularly suitable for street (Street) and local area (District) scale flood detection. Finally, the robustness and performance advantages of the NWseg model in nighttime urban flood recognition are verified through experimental comparison with mainstream segmentation models. This study not only promotes the development of nighttime urban flood recognition technology but also provides theoretical support and practical experience for future deep learning research in low-light environments using nighttime. ~~For this specific scenario, we propose the NWseg model for waterlogging recognition in nighttime, inspired by the method introduced by (Wei et al., 2023). The problem of insufficient model recognition accuracy in nighttime scenes is effectively solved by two core components, Semantic-Oriented Disentanglement (SOD) and Illumination-Aware Parser (IAParser) (Wei et al., 2023). On this basis, this study constructs an urban flooding dataset for nighttime scenarios, based on which the model is trained to improve the model's ability to recognise the extent of urban flooding in nighttime environments.~~

Totally, the main contributions of this paper are as follows: ~~This study aims to enhance urban flood extent recognition in nighttime scenes by utilizing advanced semantic segmentation techniques and a comprehensive all-nighttime dataset, addressing the current limitations in both datasets and methodologies. More specific, our aims are as follows:~~

(1) Contributed a method for nighttime urban flooding extent identification ~~assessing urban flooded areas~~ based on urban surveillance cameras, aiming at realizing efficient assessment of nighttime urban flooding areas and filling the gaps of research in this field at this stage. ~~in response to common challenges in the field of nighttime urban flooding identification.~~

(2) To support the generalization ability of the model in complex nighttime environments, this study constructs a nighttime flood inundation dataset covering a variety of nighttime scenarios (e.g., different weather, illumination intensity, and urban structure), which provides diverse sample resources required for training and testing. ~~A comprehensive and representative nighttime urban flooding dataset is constructed. It covers a wide range of nighttime scenes, including different weather conditions and city layouts, providing a rich resource for training and testing semantic segmentation models.~~

(3) Replace the original DeepLabv3+ model network backbone with MobilenetV2 and ResNet101 networks and verify the effect of different network backbones on the performance of the Deeplabv3+ model. ~~Replacement of the original DeepLabv3+ model network backbone with MobileNetV2 and Resnet101 networks is used to verify the performance impact of different network backbones on the DeeplabV3+ model through ablation experiments.~~

(4) An urban flood identification ~~A waterlogging recognition~~ model NWseg for nighttime scenarios is proposed ~~contributed~~, and the significant advantages of the model in terms of robustness, effectiveness and practicality are verified by comparing with other existing models, which advances the research and development of nighttime urban flooding extent identification ~~flood recognition~~.

2 Model

2.1 Nighttime Urban Segmentation Model

Flood segmentation faces significant challenges in nighttime scenes. Insufficient illumination and interference from complex artificial light sources such as streetlights and headlights result in blurring of texture, edge, and color information in flooded regions, further exacerbating the difficulty of the segmentation task and severely affecting the robustness and accuracy of the model. To address this challenge, this study proposed a flood extent recognition model specifically designed for low-light nighttime scenes - the NWseg model, which aims to alleviate the impact of low illumination and complex lighting conditions on segmentation performance. As shown in Figure 1, the NWseg model consists of two key modules: Content-Light Splitter (CLS) and Dual-Feature Integrator (DFI).

The design of the CLS module is based on the Retinex theory, which states that an image can be decomposed into a pixel-by-pixel product between a light-independent reflectance component (reflectance) and a light-related illumination component (illumination) (Land, 1977). Based on this principle, the CLS module decomposes the night image into a "reflectance map" and an "illumination map," which represent the inherent semantic information of the flood area and the lighting distribution in the scene, respectively (Wei et al., 2023). Subsequently, a semantic guidance mechanism is introduced to optimize the semantic segmentation loss during training (i.e., the difference between predicted pixel-level class labels and true labels), enabling the reflectance map to learn clearer boundaries and stronger semantic expression, thereby achieving accurate identification of the true contours of the flood areas. In addition, to addressing the interference from artificial light sources (such as car headlights and traffic lights), NWseg further designs the DFI module to enhance segmentation performance by adaptively fusing reflectance and illumination features. The DFI module first encodes the reflectance and illumination features and then constructs an attention mechanism that learns the degree of dependency between each pixel and the two feature types, enabling adaptive feature-weighted fusion at the pixel level (Li et al., 2024). This process adopts a pixel-wise weighting strategy, effectively enhancing the model's ability to recognize light-dominated categories. Finally, the DFI module introduces a dual semantic supervision mechanism: it not only applies semantic segmentation supervision to the fused output but also imposes semantic loss on the illumination channel separately, to enhance its independent discriminative ability and improve the model's overall generalization capability (Wei et al., 2023).

In summary, NWseg, through the collaborative design of the CLS and DFI modules, demonstrates superior semantic understanding and segmentation ability in complex nighttime lighting scenarios. It shows significant robustness and recognition advantages, particularly in high-reflection, low-contrast, and locally overexposed areas.

~~Nighttime scenes are typically characterized by low-temperature illumination and complex artificial light sources, which lead to changes in object appearance due to variations in lighting conditions. This reinforces the entanglement between light-invariant reflectance and light-specific illumination, making it challenging to extract discriminative features for semantic~~

segmentation. Based on this background, proposed a nighttime waterlogging recognition model—NWseg, specifically designed to cope with the problem of degraded segmentation performance due to insufficient illumination and complexity in nighttime scenes (Wei et al., 2023).

The paradigm consists of two core steps: decoupling and parsing. The inference is shown in Figure 1. In the decoupling phase, NWseg decomposes the input image into light-invariant reflectance and illumination-specific components. The designed SOD framework decomposes the image into illumination-independent reflectance components and light-specific components by semantically supervising the training of the de-entanglement module. It utilises Retinex theory to ensure that stable light-invariant reflectance is extracted under complex illumination conditions, which enhances the semantic recognition in the subsequent parsing phase. The parsing phase then extracts illumination features using an Illumination-Aware Parser (IAParser), which quantitatively evaluates the semantic information contained in the illumination by using a pyramid-pooling module and a convolutional layer to construct an attention mask. The final segmentation result is obtained by combining reflectance and illumination features. The model effectively copes with the complex and variable lighting challenges in nighttime scenes through the dual mechanism of decoupling and parsing, and significantly improves the performance of semantic segmentation (Wei et al., 2023).

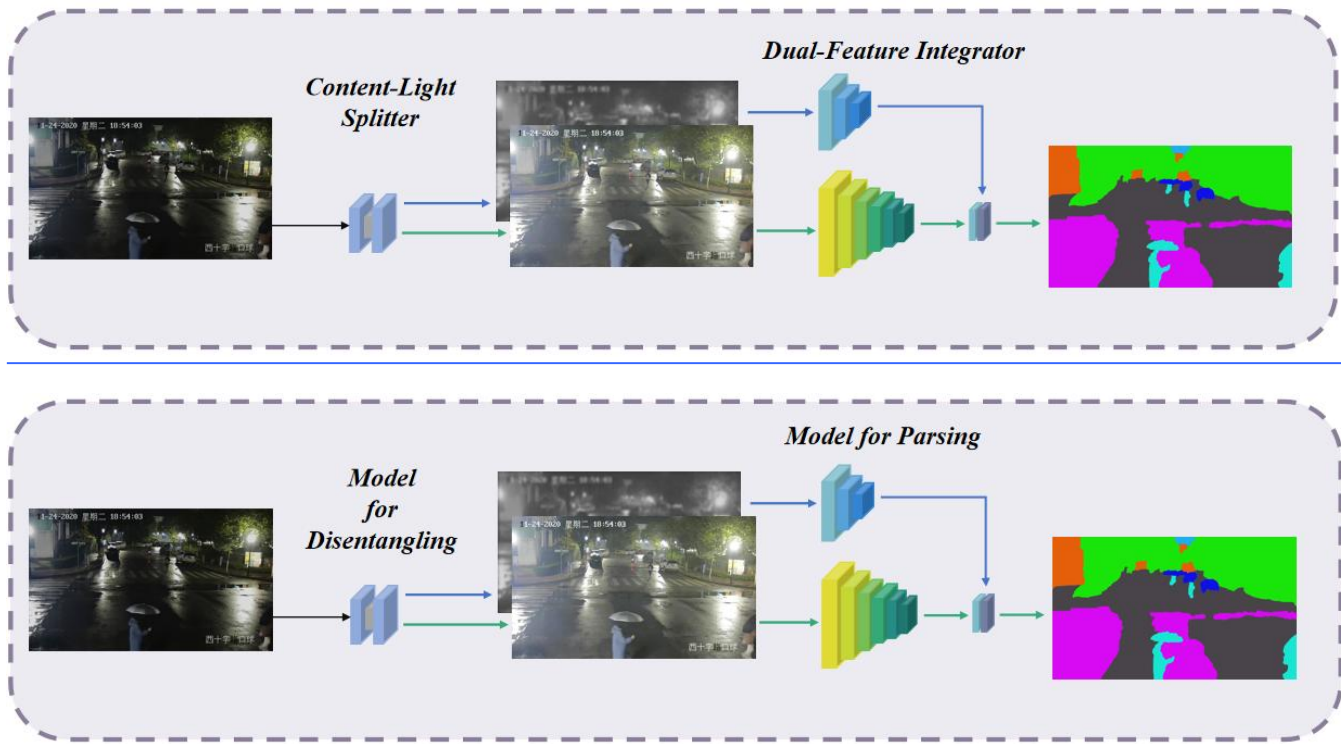


Figure 1: NWseg model inference process

2.2 Typical semantic segmentation model

DeepLabv3+ is an advanced model in the field of image segmentation, which significantly improves the accuracy and detail processing ability of image segmentation by introducing an encoder-decoder structure (Bai et al., 2023; Peng et al., 2023). The encoder part is responsible for extracting the high-level features of the image, while the decoder focuses on recovering

the details of the image, thus realizing a more fine-grained segmentation effect (Fu et al., 2021). The model also employs the techniques of void convolution and Atrous Spatial Pyramid Pooling (ASPP), which can effectively capture the multi-scale information of the image and improve the processing capability of complex scenes and object boundaries (Wang et al., 2024; Peng et al., 2024). Cavity convolution enables the model to capture a larger range of image information without increasing the computational effort by introducing voids in the convolution kernel. It is particularly helpful in capturing the relationships between distant objects in an image (Yu et al., 2017). Atrous Spatial Pyramid Pooling (ASPP), on the other hand, enhances the recognition of objects of different sizes by using different scales of null convolution to extract multiple levels of image features, which helps the model to focus on both detailed and global information (He et al., 2014). In addition, DeepLabv3+ uses Xception as the backbone network, combined with depth-separable convolution to improve computational efficiency. Depth separable convolution divides the traditional convolution operation into two steps: first, each image feature is processed independently, and then the results are combined (Zhang et al., 2023). This approach effectively reduces computation and storage requirements, allowing the model to operate more efficiently while maintaining high accuracy. However, due to its relatively complex network structure, DeepLabv3+ is still slow in the inference stage.

~~The DeepLab network series is an improved set of models based on fully convolutional neural networks (FCNs). These methods effectively enhance the receptive field of convolutional kernels to acquire multi-scale feature information, thereby optimizing the spatial accuracy of segmentation results (Feng et al., 2023; Chen et al., 2024). The network models mainly utilize techniques such as atrous convolution and atrous spatial pyramid pooling (ASPP) to extract multi-scale features and capture contextual information from images. The series includes DeepLabV1, DeepLabV2, DeepLabV3, and DeepLabV3+. DeepLabV3+ is the latest version in the DeepLab series (Li et al., 2024; Peng et al., 2024; Ma et al., 2024; Zhang et al., 2023); it introduces an encoder-decoder structure by adopting DeepLabV3 as the encoder and adding a decoder to form a new model. The Xception model is applied to the segmentation task, extensively using depthwise separable convolutions within the model. However, this network still has limitations in modeling long-range dependencies, insufficient handling of class-imbalanced data, and higher latency for real-time applications. While DeepLabV3+ combines the spatial pyramid pooling module and encoder-decoder structure in deep neural networks to achieve fine segmentation of object boundaries, it remains constrained in modeling long-range dependencies, dealing with class imbalance, and reducing latency for real-time applications (Li et al., 2023; Zhang et al., 2024; Tao et al., 2023).~~

To enhance the segmentation performance of DeepLabv3+ in urban flood scenes, this study designs a series of controlled experiments, systematically modifying or removing network components to verify the effectiveness of different backbone networks (i.e., ablation studies) and compares the results with the NWseg model.~~designed ablation experiments to verify the effectiveness of different backbone networks and compared them with the NWseg model.~~ First, experiments were conducted on the original, unmodified DeepLabv3+ network as a baseline model. Then, we replaced the original DeepLabv3+ backbone network with the lightweight MobilenetV2, constructing an improved model (denoted as

210 MobilenetV2-DeepLabv3+). [MobilenetV2 is structurally optimized to compress the feature information while retaining the](#)
211 [key information as much as possible, thus achieving a lightweight model while maintaining high accuracy \(Jin et al.,](#)
212 [2023\).](#)~~MobilenetV2 optimizes the number of model parameters by introducing a linear bottleneck layer and inverted residual~~
213 ~~structures, ensuring a lightweight model while maintaining high accuracy (Jin et al., 2023).~~ Finally, we replaced the
214 backbone network of DeepLabv3+ with the residual neural network ResNet101 to form another improved model (denoted as
215 ResNet101-DeepLabv3+). [ResNet101 adopts the residual connection mechanism so that part of the feature information can](#)
216 [bypass the intermediate layer and be transmitted directly, which avoids the problems of “learning stagnation” or “training](#)
217 [instability” during the training process](#) ~~ResNet101 leverages a residual learning mechanism, allowing input information to~~
218 ~~bypass certain layers, addressing gradient vanishing and explosion issues during deep network training. This enhances the~~
219 ~~model’s ability to capture spatial depth and details, ultimately improving the accuracy and robustness of flood area~~
220 ~~recognition (Wang et al., 2024)(Yang et al., 2023; Wang et al., 2024).~~

221 [The Fully Convolutional Network \(FCN\) is a deep learning model that divides images into different regions by assigning a](#)
222 [specific label to each pixel. Traditional deep learning models typically provide only an overall classification result for an](#)
223 [entire image. In contrast, FCNs improve upon these models by replacing the fully connected layers with convolutional](#)
224 [operations, enabling the network to handle input images of any size and produce detailed, pixel-level predictions \(Yang et al.,](#)
225 [2017\). FCNs progressively compress the spatial dimensions of the image to extract essential information and then restore the](#)
226 [original size to achieve precise localization of different regions \(Zhao et al., 2018\).](#) Additionally, FCNs combine information
227 [from both shallow and deep layers, further enhancing segmentation accuracy in complex areas, such as flood boundaries. In](#)
228 [this study, ResNet50 was selected as the backbone network for FCN, referred to as ResNet50-FCN. ResNet50 is a deep](#)
229 [neural network that effectively alleviates the gradient vanishing problem during training, improving stability and efficiency.](#)
230 [By combining the depth of ResNet50 with the flexibility of FCN, the proposed model enhances the accurate detection of](#)
231 [inundated areas in complex environments.](#)

232 ~~The Fully Convolutional Network (FCN) is an architecture specifically designed for semantic segmentation by replacing~~
233 ~~the fully connected layers of traditional Convolutional Neural Networks (CNNs) with convolutional layers. This allows~~
234 ~~FCNs to process input images of arbitrary sizes and perform accurate pixel-wise classification. FCNs extract features~~
235 ~~through convolutional layers, reduce feature dimensionality via pooling layers, and restore feature map sizes using~~
236 ~~upsampling layers, achieving precise pixel-level segmentation. Techniques such as bilinear interpolation are employed to~~
237 ~~preserve image details (Zhao et al., 2018).~~ Additionally, skip connections in FCNs effectively fuse shallow and deep feature
238 ~~information, improving segmentation accuracy. In this study, ResNet50 is adopted as the backbone network for FCN,~~
239 ~~denoted as ResNet50-FCN. ResNet50 utilizes a residual learning mechanism to address gradient vanishing issues during~~
240 ~~deep model training, maintaining training stability and efficiency while enabling greater depth. The multiple residual blocks~~
241 ~~in ResNet50 capture rich multi-scale features, adapting to structures from coarse to fine. Its skip connections preserve the~~

detailed information that can be lost during upsampling, ensuring high-precision semantic segmentation. Combining the depth of ResNet50 with the flexibility of FCN, this model enhances the accurate detection of inundated areas in complex environments.

LRASPP (Lightweight Refine Atrous Spatial Pyramid Pooling) is a lightweight model designed for image segmentation tasks. The model adopts MobileNetV3 as the backbone network for extracting the base features of an image and fuses shallow features to enhance the retention of detailed information. To further enhance the inference efficiency, LRASPP reduces the number of convolutional operations in the structural design and streamlines the feature channels to effectively reduce the computational complexity. Ultimately, the model restores the feature map to the same size as the input image through the upsampling operation to achieve accurate prediction of each pixel category. ~~The LRASPP network is a lightweight semantic segmentation model designed for efficient operation on resource-constrained devices such as mobile and embedded systems. It simplifies the classic ASPP (Atrous Spatial Pyramid Pooling) module, retaining its ability to capture multi-scale contextual information while significantly reducing computational complexity and memory usage. By leveraging depthwise separable convolutions to reduce the number of parameters and incorporating detailed information from lower-level features, LRASPP achieves a balance between model efficiency and accuracy. The model employs MobileNetV3 as the lightweight backbone to extract image features and generate multi-scale feature maps. It also simplifies the original ASPP module by capturing multi-scale contextual information through atrous convolutions and fusing low-level detailed features to improve segmentation accuracy. By reducing convolutional layers and the number of channels, the network significantly lowers computational complexity. The final output is upsampled to match the input image size, ensuring both efficiency and accuracy in segmentation tasks (Tang et al., 2024).~~

U-Net is a deep learning model commonly used for image segmentation tasks, and its structure is mainly composed of two parts: encoder and decoder (Siddique et al., 2021). The encoder is responsible for gradually reducing the image size and extracting key features, while the decoder recovers the detailed information by gradually enlarging the feature map, thus realizing the accurate classification of each pixel. In addition, to compensate for the information lost during the process of reducing the image size, U-Net introduces a jump-join mechanism, which passes the features extracted at different stages in the encoder directly to the corresponding decoder stage (Sengupta et al., 2025). This design enables the model to better preserve the detailed features in the image while maintaining overall semantic understanding. ~~U-Net is a classic network architecture for image segmentation, built on fully convolutional networks (FCNs). It utilizes skip connections to directly concatenate features from downsampling and upsampling layers along the channel dimension, effectively integrating information from different layers. U-Net features a symmetrical encoder-decoder structure, with a left downsampling path, a right upsampling path, and intermediate skip connections. The downsampling path resembles traditional CNN architectures, consisting of alternately stacked convolutional and pooling layers, while the upsampling path uses transposed convolution to progressively restore the feature maps to the original image resolution (Zhang et al., 2023). Shallow features primarily~~

~~capture fine-grained information such as flood area edges, texture, and pixel position distribution, while deeper features extract more abstract, coarse-grained semantic information, helping solve the final pixel-level classification problem. U-Net's structural characteristics enable it to effectively handle detailed information in low-light environments, making it particularly suitable for nighttime flood detection and other low-light image segmentation tasks~~ (Yadabendra et al., 2022).

We conducted comparative experiments on the FCN, LRASPP, U-Net, and NWseg models, evaluating their performance using metrics such as Precision, Recall, Mean Intersection over Union (MIoU), and F1 Score. All models were initialized with pretrained weights for their backbone networks and trained on the nighttime urban flooding dataset. The models were then evaluated on the test set, with relevant metrics calculated to determine the most suitable model for nighttime urban flood recognition.

3 Design of experiments

3.1 Construction of dataset

In this study, we employed web crawler technology using Google Chrome to construct a comprehensive nighttime urban waterlogging dataset by searching with the keyword "nighttime urban flooding." This dataset contains 4,000 images that capture a wide range of nighttime waterlogging scenes, varying in extent and shape. To enhance the dataset's robustness and comprehensiveness, we included images of complex scenes, such as strong lighting conditions and splashes caused by vehicles, ensuring its applicability to diverse nighttime flooding situations. During the data selection process, careful attention was given to the representativeness and balance of waterlogged areas across different scales, ranging from localized ponding to large-scale flood events, to ensure broad coverage of possible urban flooding conditions (Du et al., 2025).

In addition, we employed Labelme, an open-source image annotation tool widely used in the field of computer vision, to manually annotate the flooded regions in the images. Through its graphical interface, annotators can polygonally map the inundated areas in an image and assign corresponding category labels to each area, thus generating high-quality semantic segmentation data that can be used for deep learning model training (Zhang et al., 2023). Using this tool, we precisely labeled the inundated areas in a total of 4,000 images. To ensure the accuracy and consistency of the annotations, three graduate students with research backgrounds in urban flooding were recruited to independently perform the annotation work. Specifically, each flood image was annotated separately by all three annotators, followed by a cross-review process to identify potential discrepancies in the flood boundaries. In cases of inconsistency, the annotators engaged in multiple rounds of collaborative discussion and iterative refinement, optimizing the boundaries based on image details. This process ensured the overall quality and reliability of the dataset.~~performed the labeling work on the 4000 images in the dataset using the Labelme tool, which accurately extracted the waterlogged regions in each image. To further improve the accuracy of the annotations, we specifically assigned three graduate students to rigorously review and calibrate the boundary annotations for~~

quality assurance. The annotation results are saved as labeled images. Figure 2 presents a comparison between the original images and the labeled images, where the inundated ~~waterlogged~~ areas are marked in white and the non-inundated ~~waterlogged~~ areas are marked in black.

307

308

309

310

311

312

313

314

315

316

317

318

319

320

321

322

323



Figure 2: Samples and the flood area labels in the dataset, the white marked range is the flood extent.~~Data Samples–~~

3.2 Evaluation metrics

In validation and testing, mean Intersection over Union (MIOU), F1score, precision and recall were used to assess the performance of the semantic segmentation models ([Munawar et al., 2021](#))([Jin et al., 2024](#)).

The MIOU value is defined as the ratio of the intersection area of the predicted bounding box and the real bounding box to the concatenation area, and is calculated by averaging the results for each category. It is used to evaluate the accuracy of the location information of the predicted results of the target detection task. The larger the overlap area between the real and the presumed area of the object, the larger the calculated value of MIOU, and the more accurate the presumed target area. The calculation of the MIOU value follows the following formula:

$$MIOU = \frac{1}{k+1} \sum_{i=0}^k \frac{TP}{TP + FP + FN} \quad (1)$$

Precision, which is the proportion of samples predicted to be positive that are actually positive, is also known as the check rate, and can be expressed by the following formula:

$$Precision = \frac{TP}{TP + FP} \quad (2)$$

Recall, which is the proportion of actual positive samples that are predicted to be positive, is also known as the check all rate, and is given by the following formula:

$$Recall = \frac{TP}{TP + FN}$$

(3)

F1_score is the reconciled mean of precision and recall. The formula for each precision evaluation metric is as follows:

$$F1score = \frac{2 \times Precision \times Recall}{Precision + Recall}$$

(4)

In the above formula, TP is the number of actual situations that are true and predicted to be true; FP is the number of actual situations that are false and predicted to be true; FN is the number of actual situations that are true and predicted to be false; and TN is the number of actual situations that are false and predicted to be false.

3.3 Experimental configuration

All experiments were conducted using an operating system of Windows 10, a CPU model of Intel(R)Core(TM)i712700F@2.10GHz, a GPU model of [NVIDIA GeForce RTX 3080](#), 32GB of operating memory, a programming language of Python 3.13, and a deep learning framework of PyTorch1.13, [GPU acceleration is enabled during model training with CUDA 11.7 and cuDNN 8.4.1 to improve training efficiency.](#) ~~GPU acceleration libraries are CUDA11.7, CUDNN8.4.1. the~~ The input image resolution is 512*512 pixels, the training optimizer type is Adam, the weight decay index is 0.0001, and the initialized learning rate is 0.005. Parameters are shown in Table 1.

Table 1. Configuration table of the experiment

Project	Model
Operating System	Windows 10
Programming Language	Python3.13
GPU	NVIDIA GeForce RTX3080
GPU memory	32GB

4 Result

4.1 Ablation study

Table 2. NWseg and DeepLabv3+ series model training results

Models	P/%	R/%	F1_score/%	MIoU/%	Params/M
Mobilenetv2-DeepLabv3+	67.46	50.64	57.85	46.15	5.81
ResNet101-DeepLabv3+	67.74	57.24	62.05	51.98	59.34
DeepLabv3+	53.34	50.61	51.94	46.07	54.70
NWseg	95.99	94.8	95.39	91.46	122.6

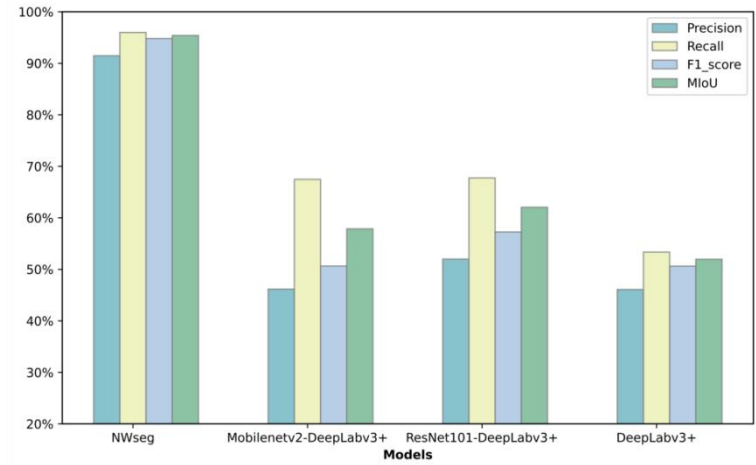
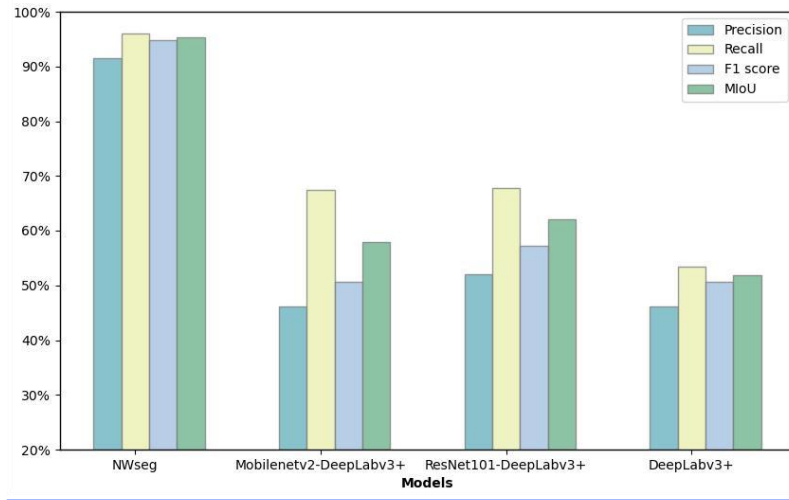


Figure 3. Comparison of experimental results between NWseg and DeepLabv3+ series of models

In this section, we present a comparative analysis of the DeepLabv3+ model with different backbone networks and compare it with the NWseg model. As shown in Table 2 and Figure 3, all evaluation metrics are improved after replacing the original backbone network of DeepLabv3+ with Mobilenetv2 and ResNet101, respectively. Notably, ~~replacing the original backbone of DeepLabv3+ with MobileNetV2 resulted in improvements across all evaluation metrics. Precision and F1 score increased significantly by 14.12% and 5.91%, respectively, while Recall and MIoU saw marginal improvements of 0.03% and 0.08%.~~ When ~~when~~ ResNet101 was used as the backbone, the model's performance improved even more, with Precision, F1 score, Recall, and MIoU increasing by 14.4%, 10.11%, 6.63%, and 5.91%, respectively, compared to the baseline model. However, all DeepLabv3+ variants still exhibited a significant performance gap when compared to NWseg. The NWseg model achieved 95.99% in Precision, 94.80% in Recall, 95.39% in F1 score, and 91.46% in MIoU, demonstrating its superior capability in nighttime urban flood extent recognition. Although NWseg has a relatively large number of parameters, it delivers outstanding accuracy and robustness. ~~Despite these improvements, all three DeepLabV3+ models still exhibited a noticeable performance gap compared to the NWseg model. The NWseg model significantly outperforms the other models by achieving 95.99%, 94.8%, 95.39%, and 91.46% in Precision, Recall, F1 score, and MIoU, respectively.~~

4.2 Model performance experiments

Table 3. NWseg and other segmentation model training results

Models	P/%	R/%	F1_score/%	MIoU/%	Params/M
NWseg	95.99	94.8	95.39	91.46	122.6
ResNet50-FCN	85	77.23	80.93	82.7	35.31
LRASPP	80.17	25.39	38.57	59.21	3.22
U-Net	94.7%	83.57	88.24%	80.5%	43.93

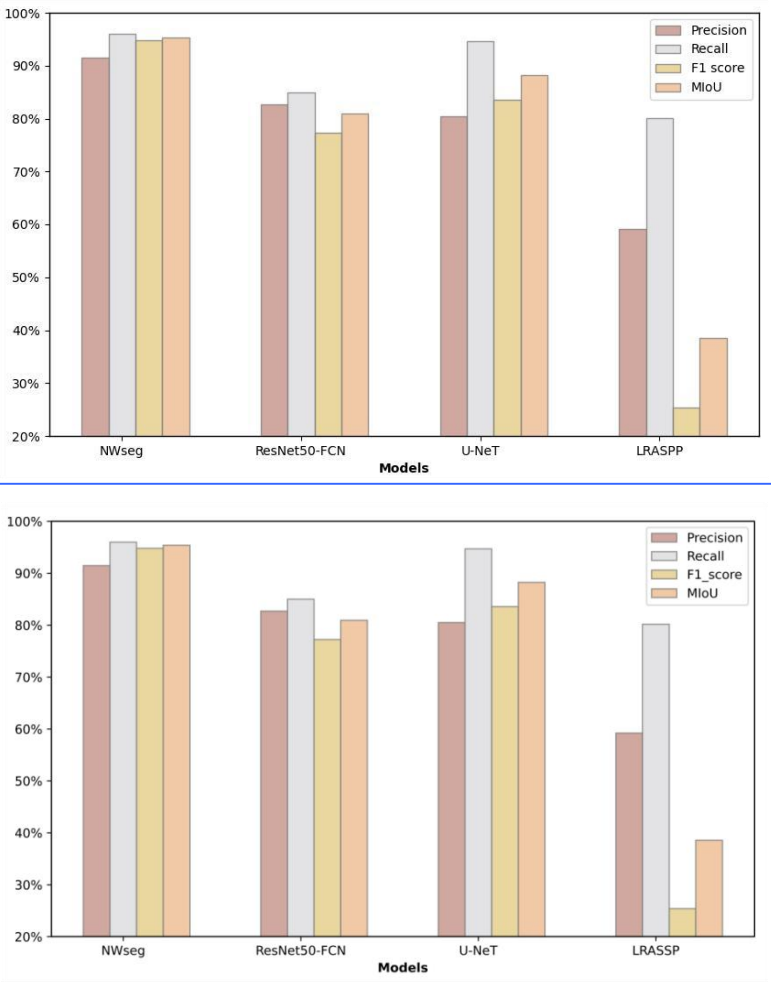


Figure 4. Comparison of experimental results between NWseg and other segmentation models

In this section, we present a comparative analysis of the experimental results of the NWseg model against other segmentation models. As shown in Table 3 and Figure 4, the NWseg model achieved optimal results on the test set of the [nighttime flood inundation](#)~~social inundation~~ dataset, with a Precision of 95.99%, Recall of 94.8%, F1_score of 95.39%, and MIoU of 91.46%. These metrics are significantly higher than those of the other models, demonstrating [superior](#) ~~exceptional~~ accuracy and recall rates. Compared to the ResNet50-FCN model, the NWseg model exhibits superior performance across all indicators, with increases of 10.99% in Precision, 17.57% in Recall, 14.46% in F1_score, and an 8.76% improvement in MIoU. When compared with the U-Net model, while the NWseg's Precision is similar, it outperforms in other metrics, with Recall, F1_score, and MIoU higher by 11.23%, 7.15%, and 10.96% respectively. Additionally, compared to the lightweight LRASPP model, the NWseg model shows more pronounced advantages, with Precision increased by 15.82%, Recall

372 significantly increased by 69.41%, F1_-score improved by 56.82%, and MIoU enhanced by 32.25%. [Although NWseg is](#)
373 [higher in the number of model parameters than the other comparative models, it still demonstrates significant advantages in](#)
374 [several evaluation metrics. Future research will aim to further optimize the structure of the model while maintaining its](#)
375 [performance to achieve a higher degree of lightweighting.](#)~~The lightweight design of LRASPP limits its ability to precisely~~
376 ~~capture details and edges, resulting in lower overall recognition accuracy.~~

377 Overall, the NWseg model demonstrates superior performance across all evaluation metrics and also shows strong
378 performance in real scenario tests. In contrast, although the ResNet50-FCN model performs well in precision and detail
379 processing, it lacks efficacy in handling edge regions, leading to slightly insufficient performance in complex scenes. While
380 LRASPP offers advantages in computational efficiency due to its lightweight design, it has limitations in the precise capture
381 of details and boundaries. The U-Net model is comparable to NWseg in accurately detecting target areas but is somewhat
382 less robust and consistent when processing complex scenes.

383 4.3 Real-world scenes prediction comparison

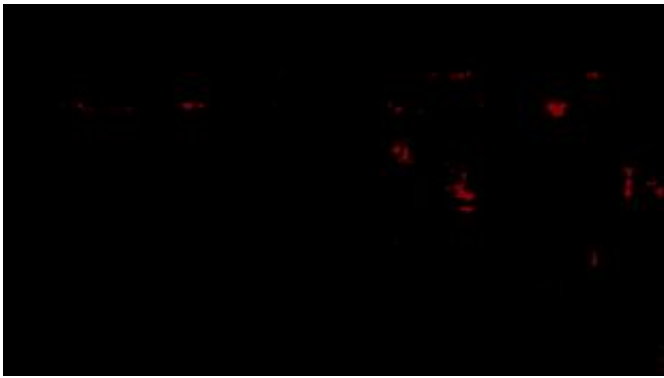
384 To validate the effectiveness and stability of each model under challenging scenes, we conducted tests on seven models
385 using nighttime rainfall scenes and nighttime strong illumination scenes ([Wan et al., 2025](#))([Liang et al., 2023](#)). As shown in
386 Figure 5(a) presents the original scene where streetlights at night generate strong reflections and halos on the water surface.
387 Additionally, the intense lighting affects the detailed features of the ground. By comparing the recognition results of each
388 model, it is evident that the NWseg, ResNet50-FCN, and U-Net models accurately detected the flooding conditions in the
389 scene. Notably, the NWseg model exhibited a more refined recognition ability in identifying water accumulation in road
390 depressions. However, both ResNet50-FCN and U-Net showed certain false detections when recognizing the overall flooded
391 areas. In contrast, the Mobilenetv2-DeepLabv3+, DeepLab, and LRASPP models could only sporadically identify small
392 flooded regions and exhibited varying degrees of false detections. Although the ResNet101-DeepLabv3+ model recognized a
393 larger flooded area, a comparison with the original image reveals a relatively high false detection rate, indicating deviations
394 in prediction accuracy. Overall, the NWseg model outperformed the others in this scene recognition task, demonstrating
395 superior capability in recognizing flooded areas under complex lighting conditions.



397

(A) The Original image

(B) Segmentation of Mobilenetv2-DeepLabv3+

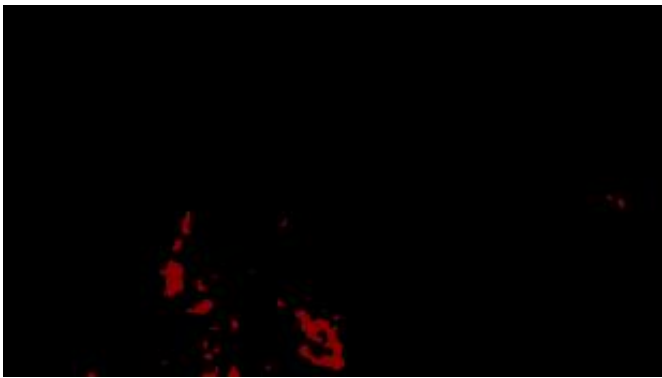


398

399

(C) Segmentation of ResNet101-DeepLabv3+

(D) Segmentation of DeepLabv3+



400

401

(E) Segmentation of LRASPP

(F) Segmentation of U-Net



402

403

(G) Segmentation of ResNet50-FCN

(H) Segmentation of NWseg

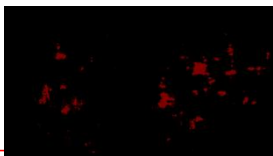
404

405

(a) Original image

(b) Mobilenetv2-DeepLabv3+(c) ResNet101-DeepLabv3+

(d) Deeplab



406

407

(a) Original image

(b) Mobilenetv2-DeepLabv3+(c) ResNet101-DeepLabv3+

(d) Deeplab



408

Figure 5. Scene with nighttime strong illumination: (A) the original scene; (B) the segmentation result of Mobilenetv2-DeepLabv3+;
(C) the segmentation result of ResNet101-DeepLabv3+; (D) the segmentation result of DeepLabv3+; (E) the segmentation result of
LRASPP; (F) the segmentation result of U-Net; (G) the segmentation result of ResNet50-FCN; (H) the segmentation result of
NWseg;**Examination of nighttime strong illumination scenes**

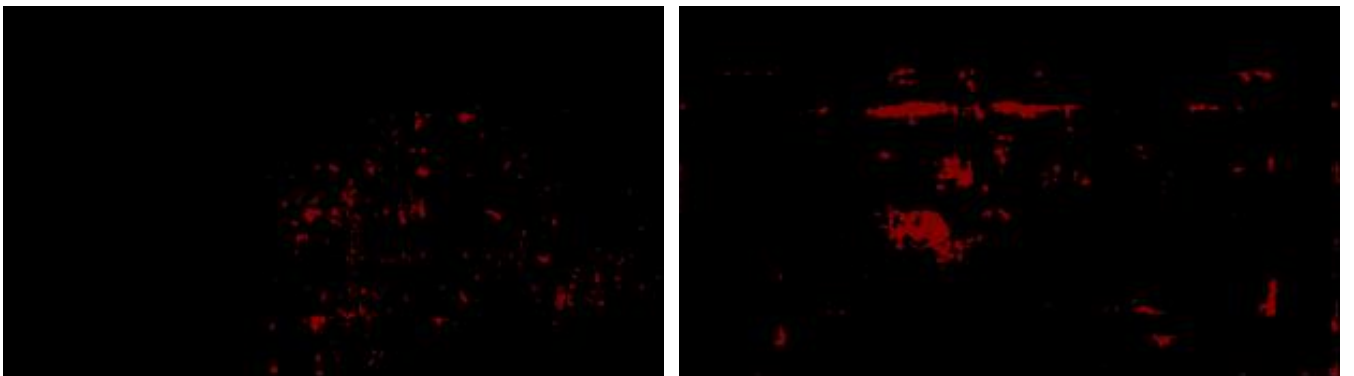
Furthermore, in the nighttime rainfall scene tests, we evaluated each model's performance to simulate urban flood recognition under real-world conditions ~~(Tan et al., 2021)~~. In such scenes, reflections from rainwater, slippery road surfaces, and interference from raindrops on the camera lens can adversely affect image clarity and the models' recognition accuracy ~~(Zhao et al., 2025)~~. As shown clearly in Figure 6, the NWseg, ResNet50-FCN, and U-Net models were able to correctly identify the flooded areas in the images, with the NWseg model providing the most detailed performance by accurately capturing the edges of the flooded regions. While ResNet50-FCN and U-Net also identified the extent of flooding relatively well, they were somewhat insufficient in recognizing the flood boundaries and exhibited some false detections.

In contrast, the other four models performed relatively poorly. Specifically, the LRASPP and Mobilenetv2-DeepLabv3+ models were almost unable to detect the flooding, indicating weaker recognition capabilities in nighttime rainfall scenes. Although ResNet101-DeepLabv3+ and DeepLab could detect some flooded areas, comparison with the original images revealed that the regions identified did not accurately reflect the actual flooding conditions and had high false detection rates. Through comparative analysis, we further confirmed the challenges posed by nighttime rainfall environments for urban flood recognition and demonstrated the superior performance of the NWseg model in handling complex conditions such as nighttime rainfall.



(A) The Original image

(B) Segmentation of Mobilenetv2-DeepLabv3+



431

(C) Segmentation of ResNet101-DeepLabv3+

(D) Segmentation of DeepLabv3+

432

433

(E) Segmentation of LRASPP

(F) Segmentation of U-Net

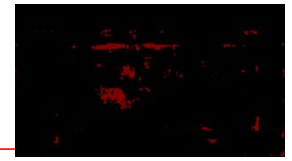
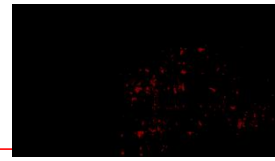
434

435

(G) Segmentation of ResNet50-FCN

(H) Segmentation of NWseg

436



437

(a) Original image (b) MobileNetV2-DeepLabv3+ (c) ResNet101-DeepLabv3+ (d) DeepLabv3+

438



439

440

(e) NWseg (f) ResNet50-FCN (g) LRASPP (h) U-Net

441

442

443

444

Figure 6. Scene with nighttime rainfall: (A) the original scene; (B) the segmentation result of MobileNetV2-DeepLabv3+; (C) the segmentation result of ResNet101-DeepLabv3+; (D) the segmentation result of DeepLabv3+; (E) the segmentation result of LRASPP; (F) the segmentation result of U-Net; (G) the segmentation result of ResNet50-FCN; (H) the segmentation result of NWseg;

445

Figure 6. Examination of nighttime rainfall scenes

446

5 Discuss

447

448

In this study, a state-of-the-art model named NWseg is proposed to address the challenges of nighttime urban flood extent identification. Through a series of experimental validations, the NWseg model demonstrates superior performance with

449 [95.99%, 94.8%, 95.39%, and 91.46% in Precision, Recall, F1 score, and MIoU, respectively. In the prediction comparison of](#)
450 [real scenarios, the model also shows high accuracy and robustness, and effectively recognizes flooded areas in complex](#)
451 [nighttime environments. In addition, NWseg achieves an inference speed of 37.8 FPS \(i.e., approximately 26.5 milliseconds](#)
452 [per image\) under the NVIDIA GeForce RTX 3080 environment, demonstrating its potential for real-time applications in](#)
453 [high-performance computing platforms. This study bridges the current research gap in flood extent recognition in nighttime](#)
454 [scenarios, providing a technical reference for flood monitoring and emergency response.](#)

455 [Nevertheless, this study still has some limitations. First, the overall structure of NWseg is relatively complex, and the](#)
456 [model parameters are large in scale, which limits its deployment capability on resource-constrained edge devices. On the](#)
457 [other hand, in nighttime scenarios with extremely low illumination or even complete power outage \(e.g., the case of city](#)
458 [blackout triggered by heavy rainfall\), the model has difficulty in extracting effective edge and texture information, which](#)
459 [leads to a significant degradation of the recognition performance. In the future, we will further optimize the network](#)
460 [structure to reduce the computational complexity of the model and improve deployment flexibility. In addition, we consider](#)
461 [combining infrared thermal imaging, low-light image enhancement, or multimodal fusion methods to improve the robustness](#)
462 [and generalization ability of the model under extreme low-light conditions.](#)

463 **5.6 Conclusions**

464 This study [successfully verified the excellent performance of the NWseg model in nighttime urban flood monitoring \(Wan et](#)
465 [al., 2024\), which provides a new idea for multi-scene flood extent identification and helps to promote the flood monitoring](#)
466 [system towards all-weather and all-scene intelligent identification.](#)~~addresses the technical challenges of nighttime urban~~
467 ~~flood detection by evaluating the performance of seven different models (Wan et al., 2024).~~ First, we constructed a
468 representative dataset comprising 4,000 images of nighttime urban flooding scenes, covering various nighttime environments
469 and diverse urban backgrounds. Second, a model for nighttime waterlogging recognition, NWseg, is proposed to address the
470 limitations in nighttime waterlogging recognition due to insufficient lighting and complex lighting conditions. Furthermore,
471 we replaced the backbone networks of the DeepLabv3+ model with MobilenetV2 and ResNet101 and conducted ablation
472 experiments to validate the performance of DeepLabv3+ with different backbones in nighttime flood recognition. We also
473 performed a comparative analysis between these DeepLabv3+ models and the NWseg model, as well as systematically
474 analyzed the NWseg, ResNet50-FCN, U-Net, and LRASPP models. Based on this, we reached the following empirical
475 findings:

476 (1) Within the DeepLab series, the DeepLabv3+ model using ResNet101 as the backbone outperformed other variants in
477 capturing water surface edges and shadow details. However, when compared to the NWseg model, there remains a
478 considerable performance gap.

479 (2) The NWseg, U-Net, and ResNet50-FCN models demonstrated excellent performance in recognizing large-scale
480 flooded areas, effectively capturing the overall contours of flood zones and exhibiting strong generalization capabilities.
481 Specifically, NWseg shows higher accuracy and robustness in complex scene tests, while ResNet50-FCN and U-Net have
482 some deficiencies and false detections in detecting edge details. In contrast, the lightweight LRASPP model showed limited
483 ability to recognize flooded areas in nighttime scenes, resulting in relatively poor performance.

484 (3) Through examining each model in complex scenes, we validated the NWseg model's effectiveness and stability in
485 diverse environments and conditions.

486 This study successfully demonstrates the superior performance of the NWseg model in nighttime urban flood
487 detection, [filling the research gap in nighttime flood range identification. Our work not only promotes the development of](#)
488 [the field of nighttime urban flood identification but also provides a reference for future deep learning applications under](#)
489 [extreme lighting conditions](#) (Wan et al., 2024). However, the model's decoupling and parsing process involves complex
490 decomposition of lighting components and adaptive fusion, leading to high computational resource demands, which
491 may impact its practical usability. Future work will focus on reducing the model's parameters and computational costs
492 while maintaining accuracy. Additionally, further optimization of the dataset and model improvements will be pursued
493 to enhance the overall performance of the NWseg model, broadening its potential applications.

494 [Acknowledgments. This research was funded by the National Natural Science Foundation of China \(NSFC\)](#)
495 [\(No.42405140\), the China Postdoctoral Foundation \(No. 2024M761383\), the Open Grants of China Meteorological](#)
496 [Administration Radar Meteorology Key Laboratory \(No. 2024LRM-A02\), and the Talent Startup project of NJIT](#)
497 [\(YKJ.202315\).](#)

498 **Data availability.** Data will be made available on request.

499 **Author contributions.** ~~Xing-Wang, Jiaquan-Wan, Yannian-Cheng, Cuiyan-Zhang~~: Writing – original draft, Validation,
500 Software, Methodology, Investigation. ~~Xing-Wang, Jiaquan-Wan, Yannian-Cheng~~: Writing – review & editing,
501 Validation. ~~Tao-Yang~~: Writing – review & editing, Supervision. ~~Fengchang-Xue~~: Formal analysis, Validation. ~~Yufang-~~
502 ~~Shen~~: Data curation, Validation. ~~FT~~: Data curation, Validation. ~~Quan-J-Wang~~: Data curation, Validation.

503 **Competing interests.** The contact author has declared that none of the authors has any competing interests.

504 References

- 505 [Bai, Y., Li, J., Shi, L., Jiang, Q., Yan, B., and Wang, Z.: DME-DeepLabV3+: a lightweight model for diabetic macular edema extraction](#)
506 [based on DeepLabV3+architecture, Frontiers in Medicine, 10, 11, <http://doi.org/10.3389/fmed.2023.1150295>, 2023.](#)
- 507 [Bofana, J., Zhang, M., Wu, B., Zeng, H., Nabil, M., Zhang, N., Elnashar, A., Tian, F., Da Silva, J.M., Botˆ ao, A., Atumane, A., Mushore,](#)
508 [T.D., and Yan, N.: How long did crops survive from floods caused by Cyclone Idai in Mozambique detected with multi-satellite data,](#)
509 [Remote Sens. Environ., 269, 112808, <https://doi.org/10.1016/j.rse.2021.112808>, 2022.](#)
- 510 [Burn, D. H., Whitfield, P. H.: Climate related changes to flood regimes show an increasing rainfall influence, Journal of Hydrology, 617,](#)
511 [13, <http://doi.org/10.1016/j.jhydrol.2023.129075>, 2023.](#)
- 512 ~~Bai, G., Hou, J., Han, H., Xia, J., Li, B., Zhang, Y., and Wei, Z.: Intelligent monitoring method for road inundation based on deep learning,~~
513 ~~—Water Resources Protection, 37, 75-80, 2021.~~
- 514 [Choi, Y.H., and Yoo, S.J.: Quantized-state-based decentralized neural network control of a class of uncertain interconnected nonlinear](#)
515 [systems with input and interaction time delays, Eng. Appl. Artif. Intel., 125, 106759, <https://doi.org/10.1016/j.engappai.2023.106759>,](#)
516 [2023.](#)
- 517 ~~Cheng, Y., Gu, Q., Wang, Z., and Li, Z.: Wood Defect Image Segmentation Based on Deep Learning, FORESTRY MACHINERY &~~
518 ~~WOODWORKING EQUIPMENT, 46, 33-36, <http://doi.org/10.13594/j.cnki.mejgix.2018.05.004>, 2018.~~
- 519 ~~Chen, C., Hao, X., Long, H., and Sun, X.: Asphalt Road Crack Detection Method Based on Improved DeepLabv3+ Network,~~
520 ~~SEMICONDUCTOROPTOELECTRONICS, 45, 493-500, <http://doi.org/10.16818/j.issn1001-5868.2023103002>, 2024.~~
- 521 [Du, W., Qian, M., He, S., Xu, L., Zhang, X., Huang, M., and Chen, N.: An improved ResNet method for urban flooding water depth](#)
522 [estimation from social media images, Measurement, 242, 12, <http://doi.org/10.1016/j.measurement.2024.116114>, 2025.](#)
- 523 [Fu, H., Meng, D., Li, W., and Wang, Y.: Bridge Crack Semantic Segmentation Based on Improved Deeplabv3+, Journal of Marine](#)
524 [Science and Engineering, 9\(6\), 14, <http://doi.org/10.3390/jmse9060671>, 2021.](#)
- 525 [Ghosh, P., Sudarsan, J. S., and Nithiyantham, S.: Nature-Based Disaster Risk Reduction of Floods in Urban Areas, Water Resources](#)
526 [Management, 38 \(6\), 1847-1866, <http://doi.org/10.1007/s11269-024-03757-4>, 2024.](#)
- 527 [Gu, Q., Chai, F., Zang, W., Zhang, H., Hao, X., and Xu, H.: A Two-Level Early Warning System on Urban Floods Caused by Rainstorm,](#)
528 [Sustainability, 17\(5\), 16, <http://doi.org/10.3390/su17052147>, 2025.](#)
- 529 [Hao, X., Lyu, H., Wang, Z., Fu, S., and Zhang, C.: Estimating the spatial-temporal distribution of urban street ponding levels from](#)
530 [surveillance videos based on computer vision, Water Resources Management, 36\(6\), 1799-1812,](#)
531 [https://doi.org/10.1007/s11269-022-03107-2, 2022.](#)
- 532 [He, K., Zhang, X., Ren, S. and Sun, J.: Spatial pyramid pooling in deep convolutional networks for visual recognition. In Proceedings of](#)
533 [the 13th European Conference on Computer Vision \(ECCV\), 8691, 346-361, \[https://doi.org/10.1007/978-3-319-10578-9_23\]\(https://doi.org/10.1007/978-3-319-10578-9_23\), 2014.](#)
- 534 ~~Feng, T., Guo, Y., Huang, X., and Qiao, Y.: Cattle Target Segmentation Method in Multi-Scenes Using Improved DeepLabV3+ Method,~~
535 ~~Animals, 13, 2521, <http://doi.org/10.3390/ani13152521>, 2023.~~

536 Gao, F.: Urban Flood Disaster Risk Factor Identification Based on Deep Learning MBA thesis, Changzhou University, Changzhou,
537 ~~—<http://doi.org/10.27739/d.cnki.gjsgy.2023.000080>, 2023.~~

538 Hao, Y.: Urban Flood Disaster Risk Warning Based on Deep Learning MBA thesis, Changzhou University, Changzhou,
539 ~~<http://doi.org/10.27739/d.cnki.gjsgy.2022.000104>, 2022.~~

540 Jin, K., Zhang, J., Wang, Z., Zhang, J., Liu, N., Li, M., and Ma, Z.: Application of deep learning based on thermal images to identify the
541 water stress in cotton under film-mulched drip irrigation, Agricultural Water Management, 299, 108901,
542 <http://doi.org/10.1016/j.agwat.2024.108901>, 2024.

543 Jin, N.: Research on Passible Area Detection Based on Deep Learning for All Types of Roads in both Fog and Sunny Conditions MEng-
544 thesis Sichuan University, Chengdu, ~~<http://doi.org/10.27342/d.cnki.gsedu.2023.001503>, 2023.~~

545 Luo, H.: A Study on Risk Assessment Method of Urban Flood Disaster and Its Applications, MSe thesis, South China University of
546 Technology, Guangzhou, ~~<http://doi.org/10.27151/d.cnki.ghnlu.2020.002269>, 2020.~~

547 Liao, Y.: Research on recognition method of urban road waterlogging information based on deep learning MEng thesis, South China
548 University of Technology, Guangzhou, ~~<http://doi.org/10.27151/d.cnki.ghnlu.2023.004803>, 2023.~~

549 Kim, H., Villarini, G., Wasko, C., Trambly, Y.: Changes in the Climate System Dominate Inter-Annual Variability in Flooding Across
550 the Globe, Geophysical Research Letters, 51(6), <http://doi.org/10.1029/2023gl107480>, 2024.

551 Kundzewicz, Z. W., Su, B., Wang, Y., Xia, J., Huang, J., and Jiang, T.: Flood risk and its reduction in China. Adv. Water Resour., 130, 37
552 – 45, <https://doi.org/10.1016/j.advwatres.2019.05.020>, 2019.

553 Land, E. H.: The retinex theory of color vision, Scientific American, 237(6), 108-28, <https://doi.org/10.1038/scientificamerican1277-108>,
554 1977.

555 Li, X., Wang, W., Feng, X., and Li, M.: Deep parametric Retinex decomposition model for low-light image enhancement, Computer
556 Vision and Image Understanding, 241, 14, <https://doi.org/10.1016/j.cviu.2024.103948>, 2024

557 Liu, W., Feng, Q., Engel, B. A., Yu, T., Zhang, X., and Qian, Y.: A probabilistic assessment of urban flood risk and impacts of future
558 climate change, Journal of Hydrology, 618, 11, <http://doi.org/10.1016/j.jhydrol.2023.129267>, 2023.

559 Liu, X., Song, L., Liu, S., and Zhang, Y.: A review of deep-learning-based medical image segmentation methods, Sustainability, 13(3),
560 1224. <https://doi.org/10.3390/su13031224>, 2020.

561 Mason, D. C., Davenport, I. J., Neal, J. C., Schumann, G. J. P., and Bates, P. D.: Near Real-Time Flood Detection in Urban and Rural
562 Areas Using High-Resolution Synthetic Aperture Radar Images, Ieee Transactions on Geoscience and Remote Sensing, 50(8),
563 3041-3052, <http://doi.org/10.1109/tgrs.2011.2178030>, 2012.

564 Muhadi, N. A., Abdullah, A. F., Bejo, S. K., Mahadi, M. R., Mijic, A., and Vojinovic, Z.: Deep learning and LiDAR integration for
565 surveillance camera-based river water level monitoring in flood applications, Natural Hazards, 120(9), 8367-8390,
566 <http://doi.org/10.1007/s11069-024-06503-6>, 2024.

567 [Munawar, H.S., Ullah, F., Qayyum, S., and Heravi, A.: Application of deep learning on UAV-based aerial images for flood detection,](#)
568 [Smart Cities, 4 \(3\), 1220 – 1242. <https://doi.org/10.3390/smartcities4030065>, 2021.](#)

569 ~~Li, F., Shi, J., Liang, X., Li, Y., Liu, P., and Chen, X.: Research on Hilly Field Road Image Segmentation Method Based on Improved~~
570 ~~DeepLabV3+, Journal of Southwest University (Natural Science Edition), 46, 172-183,~~
571 ~~<http://doi.org/10.13718/j.cnki.xdzk.2024.08.016>, 2024.~~

572 ~~Li, B.: Research on forestland classification based on improved DeepLabV3+, MSc thesis, Central South University of Forestry and~~
573 ~~Technology, Changsha, <http://doi.org/10.27662/d.cnki.gznle.2024.000886>, 2024.~~

574 ~~Liang, Y., Li, X., Tsai, B., Chen, Q., and Jafari, N.: V-FloodNet: A video segmentation system for urban flood detection and quantification,~~
575 ~~Environmental Modelling & Software, 160, 105586, <http://doi.org/10.1016/j.envsoft.2022.105586>, 2023.~~

576 ~~Ma, J., Guo, Z., Ma, Z., Ma, X., and Li, J.: Remote sensing image land feature segmentation method based on lightweight DeepLabV3+,~~
577 ~~Chinese Journal of Liquid Crystals and Displays, 39, 1001-1013, 2024.~~

578 Peng, H., Xiang, S., Chen, M., Li, H., and Su, Q.: DCN-Deeplabv3+: A Novel Road Segmentation Algorithm Based on Improved
579 Deeplabv3+, Ieee Access, 12, 87397-87406, <http://doi.org/10.1109/access.2024.3416468>, 2024.

580 [Peng, H., Zhong, J., Liu, H., Li, J., Yao, M., and Zhang, X.: ResDense-focal-DeepLabV3+enabled litchi branch semantic segmentation for](#)
581 [robotic harvesting, Computers and Electronics in Agriculture, 206, 12, <http://doi.org/10.1016/j.compag.2023.107691>, 2023.](#)

582 Sarp, S., Kuzlu, M., Cetin, M., Sazara, C., and Guler, O.: Detecting floodwater on roadways from image data using Mask-R-CNN, 2020
583 International Conference on INnovations in Intelligent SysTems and Applications, Novi Sad, Serbia, 24-26,
584 <http://doi.org/10.1109/INISTA49547.2020.9194655>, 2020.

585 Sazara, C., Cetin, K., and Iftekharuddin, K.: Detecting floodwater on roadways from image data with handcrafted features and deep
586 transfer learning, 2019 IEEE Intelligent Transportation Systems Conference, Auckland, New Zealand, 27-30,
587 <http://doi.org/10.1109/ITSC.2019.8917368>, 2019.

588 [Sengupta, S., Chyrmang, G., Bora, K., Das, H. S., Li, A., Lemos, B., and Mallik, S.: Assessment of different U-Net backbones in](#)
589 [segmenting colorectal adenocarcinoma from H&E histopathology, Pathology Research and Practice, 266, 11,](#)
590 <https://doi.org/10.1016/j.prp.2025.155820>, 2025.

591 [Siddique, N., Paheding, S., Elkin, C. P., and Devabhaktuni, V.: U-Net and Its Variants for Medical Image Segmentation: A Review of](#)
592 [Theory and Applications, Ieee Access, 9, 82031-82057, <https://doi.org/10.1109/access.2021.3086020>, 2021.](#)

593 ~~Tao, L.: Research on Road Scene Semantic Segmentation Method Based on DeepLabV3+ Model, MEng thesis, Dalian Jiaotong University,~~
594 ~~Dalian, <http://doi.org/10.26990/d.cnki.gslte.2023.000111>, 2023.~~

595 ~~Tan, X., Xu, K., Cao, Y., Zhang, Y., Ma, L., and Lau, RWH.: Night Time Scene Parsing With a Large Real Dataset Ieee Transactions on~~
596 ~~Image Processing, 30, 9085-9098, <http://doi.org/10.1109/tip.2021.3122004>, 2021.~~

597 Tang, Y., Tan, D, Li, H., Zhu, M., Li, X., Wang, X., Wang, J., Wang, Z., Gao, C., Wang, J., and Han, A.: RTC_TongueNet: An improved
598 tongue image segmentation model based on DeepLabV3 DIGITAL HEALTH, 10, 20552076241242773,
599 <http://doi.org/10.1177/20552076241242773>, 2024.

600 [Visser, F.: Rapid mapping of urban development from historic Ordnance Survey maps: an application for pluvial flood risk in Worcester. J.
601 Maps., 10 \(2\), 276 – 288. <https://doi.org/10.1080/17445647.2014.893847>, 2014.](#)

602 Wan, J., Qin, Y., Shen, Y., Yang, T., Yan, X., Zhang, S., Yang, G., Xue, F., and Wang, Q.: Automatic detection of urban flood level with
603 YOLOv8 using flooded vehicle dataset, Journal of Hydrology, 639, 131625, <http://doi.org/10.1016/j.jhydrol.2024.131625>, 2024.

604 [Wan, J., Xue, F., Shen, Y., Song, H., Shi, P., Qin, Y., Yang, T., and Wang, Q. J.: Automatic segmentation of urban flood extent in video
605 image with DSS-YOLOv8n, Journal of Hydrology, 655, 12, <https://doi.org/10.1016/j.jhydrol.2025.132974>, 2025.](#)

606 Wang, Y., Shen, Y., Salahshour, B., Cetin, M., Iftikharuddin, K., Tahvildari, N., Huang, G., Harris, D., Ampofo, K., and Goodall, J.:
607 Urban flood extent segmentation and evaluation from real-world surveillance camera images using deep convolutional neural network,
608 Environmental Modelling & Software, 173, 105939, <http://doi.org/10.1016/j.envsoft.2023.105939>, 2024.

609 [Wang, Y., Yang, L., Liu, X., and Yan, P.: An improved semantic segmentation algorithm for high-resolution remote sensing images based
610 on DeepLabv3+, Scientific Reports, 14, 15, <http://doi.org/10.1038/s41598-024-60375-1>, 2024.](#)

611 Wei, Z., Chen, L., Tu, T., Ling, P., Chen, H., and Jin, Y.: Disentangle then Parse: Night-time Semantic Segmentation with Illumination
612 Disentanglement, Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV), 21593-21603,
613 <https://doi.org/10.48550/arXiv.2307.09362>, 2023.

614 ~~Wang, Z.: Semantic segmentation-based waterlogged region recognition model and application research, MEng thesis, Institute of Disaster
615 Prevention, Sanhe, <https://doi.org/10.27899/d.cnki.gfzkj.2024.000003>, 2024.~~

616 ~~Xue, F., Lv, X., and Chen, X.: Urban flood monitoring method based on improved DeeplabV3+, Science of Surveying and Mapping, 48,
617 216-224, <http://doi.org/10.16251/j.cnki.1009-2307.2023.10.022>, 2023.~~

618 ~~Yang, Y.: Research on Multispectral Retinal Image Segmentation Based on Convolutional Neural Network, MEng thesis, Hebei
619 University, Baoding, <http://doi.org/10.1016/10.27103/d.cnki.ghebu.2023.001067>, 2023.~~

620 Yadavendra, and Chand, S.: Semantic segmentation of human cell nucleus using deep U-Net and other versions of U-Net models,
621 Network-Computation in Neural Systems, 33, 167-186, <http://doi.org/10.1080/0954898x.2022.2096938>, 2022.

622 ~~Yang, K., Zhang, S., Yang, X., and Wu, N.: Flood Detection Based on Unmanned Aerial Vehicle System and Deep Learning, Complexity,
623 2022, 6155300, <http://doi.org/10.1155/2022/6155300>, 2022.~~

624 [Yang, Y., Zhuang, Y., Bi, F., Shi, H., and Xie, Y.: M-FCN: Effective Fully Convolutional Network-Based Airplane Detection Framework,
625 Ieee Geoscience and Remote Sensing Letters, 14, 1293-1297, <https://doi.org/10.1109/lgrs.2017.2708722>, 2017.](#)

626 [Yu, F., Koltun, V. and Funkhouser, T.: Dilated residual networks. In Proceedings of the 30th IEEE/CVF Conference on Computer Vision
627 and Pattern Recognition \(CVPR\), 636 – 644. <https://doi.org/10.1109/cvpr.2017.75>, 2017.](#)

628 [Zhang, L., Han, G., Qiao, Y., Xu, L., Chen, L., and Tang, J.: Interactive Dairy Goat Image Segmentation for Precision Livestock Farming,](#)
629 [Animals, 13\(20\), 17, <http://doi.org/10.3390/ani13203250>, 2023.](#)

630 [Zhao, J., Wang, X., Zhang, C., Hu, J., Wan, J., Cheng, L., Shi, S., and Zhu, X.: Urban Waterlogging Monitoring and Recognition in](#)
631 [Low-Light Scenarios Using Surveillance Videos and Deep Learning, Water, 17\(5\), 19, <http://doi.org/10.3390/w17050707>, 2025.](#)

632 Zhao, W., Zhang, H., Yan, Y., Fu, Y., and Wang, H.: A Semantic Segmentation Algorithm Using FCN with Combination of BSLIC,
633 Applied Sciences-Basel, 8, 500, <http://doi.org/10.3390/app8040500>, 2018.

634 [Zheng, F., Westra, S., Leonard, M., Sisson, S.A.: Modeling dependence between extreme rainfall and storm surge to estimate coastal](#)
635 [flooding risk, Water Resour. Res., 50 \(3\), 2050 – 2071, <https://doi.org/10.1002/2013WR014616>, 2014.](#)

636 [Zeng, Y., Chang, M., and Lin, G.: A novel AI-based model for real-time flooding image recognition using super-resolution generative](#)
637 [adversarial network, Journal of Hydrology, 638, 12, <https://doi.org/10.1016/j.jhydrol.2024.131475>, 2024.](#)