

Concept-based explanations of Segmentation and Detection models in Natural Disaster Management

Samar Heydari¹, Jawher Said¹, Galip Ümit Yolcu¹, Evgenii Kortukov¹, Elena Golimblevskaia¹, Evgenios Vlachos², Vasileios Mygdalis², Ioannis Pitas², Sebastian Lapuschkin^{1,*} and Leila Arras^{1,*}

¹Department of Artificial Intelligence, Fraunhofer Heinrich Hertz Institute, Berlin, Germany

²Department of Informatics, Aristotle University of Thessaloniki, Thessaloniki, Greece

Abstract

Deep learning models for flood and wildfire segmentation and object detection enable precise, real-time disaster localization when deployed on embedded drone platforms. However, in natural disaster management, the lack of transparency in their decision-making process hinders human trust required for emergency response. To address this, we present an explainability framework for understanding flood segmentation and car detection predictions on the widely used PIDNet and YOLO architectures. More specifically, we introduce a novel redistribution strategy that extends Layer-wise Relevance Propagation (LRP) explanations for sigmoid-gated element-wise fusion layers. This extension allows LRP relevances to flow through the fusion modules of PIDNet, covering the entire computation graph back to the input image. Furthermore, we apply Prototypical Concept-based Explanations (PCX) to provide both local and global explanations at the concept level, revealing which learned features drive the segmentation and detection of specific disaster semantic classes. Experiments on a publicly available flood dataset show that our framework provides reliable and interpretable explanations while maintaining near real-time inference capabilities, rendering it suitable for deployment on resource-constrained platforms, such as Unmanned Aerial Vehicles (UAVs).

Keywords

Concept-based Explanations, Semantic Segmentation, Object Detection, Natural Disaster Management.

1. Introduction

Natural disasters such as flood and wildfire pose a serious threat to public safety and the global economy [1], requiring advanced monitoring systems. Deep neural networks (DNNs) for the semantic segmentation and the detection of objects from drone and satellite images can provide crucial information for situational awareness. However, their “black-box” nature limits the trust and confidence required for emergency responses. We address this challenge by introducing an end-to-end explainable framework for DNN-based segmentation and object detection in natural disaster management (NDM) building upon Layer-wise Relevance Propagation (LRP) [2], an explainable AI attribution technique that provides feature-level importance scores (aka relevances) to explain individual predictions of neural networks. Together with its concept-based extensions, Concept Relevance Propagation (CRP) [3, 4] and Prototypical Concept-based Explanations (PCX) [5], LRP provides an understanding of DNN model behavior both globally (i.e., dataset-wide) and in terms of human-comprehensible concepts.

LRP and its concept-based extensions already demonstrated their usefulness in computer vision, but they have not yet been applied to a PIDNet architecture nor to the NDM domain. Additionally, Prototypical Concept-based Explanations (PCX) [5] were so far confined to image *classification* models: we present its first extension to semantic *segmentation* and object *detection* models.

Contributions. In summary, our novel contributions are as follows:

Late-breaking work, Demos and Doctoral Consortium, colocated with the 4th World Conference on eXplainable Artificial Intelligence: July 01–03, 2026, Fortaleza, Brazil

*Corresponding author.

✉ sebastian.lapuschkin@hhi.fraunhofer.de (S. Lapuschkin); leila.arras@hhi.fraunhofer.de (L. Arras)

🆔 0009-0002-2247-2124 (E. Vlachos); 0000-0001-5473-5262 (V. Mygdalis); 0009-0006-7555-8641 (I. Pitas); 0000-0002-0762-7258 (S. Lapuschkin)



© 2026 Copyright for this paper by its authors. Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).

- We apply LRP [2] and concept-based explanations [3, 4] for the first time in NDM.
- We extend LRP [2] to a PIDNet architecture.
- We apply PCX [5] for the first time on segmentation and detection models.

2. Related Work

2.1. Deep learning for Natural Disaster Management

Deep learning methods are increasingly common in Natural Disaster Management (NDM) [6, 7] for real-time monitoring and response planning. In the context of wildfire, deep learning has been employed for satellite-based [8] and aerial-based image analysis [9]. Recent solutions leverage high-performance architectures like PIDNet [10], with specialized extensions for fusing infrared and RGB spectra such as RoboFireFuseNet [11], and Neural Architecture Search (NAS) methods [12] for optimizing the number of trainable parameters. Regarding flood segmentation, new learning processes utilize Self-Knowledge Distillation (Self-KD) to encode foreground information while suppressing background noise [13]. Core enabling resources include publicly available datasets for wildfire/burned area segmentation such as Blaze [14], datasets and benchmarks for flood segmentation [15], as well as synthetic data [16, 17] and weakly supervised learning frameworks [18, 19] to address data scarcity and improve model generalization. Recent efforts regarding object detection in NDM employed transfer learning on YOLO models [20] and visibility-enhanced models, such as VE-DINO [21], to handle occlusion-prone environments. Other research focused on combining deep learning vision models with social media analysis [22] and sentiment analysis [23] to give first responders a holistic understanding of the situation, which can help them allocate resources and better predict the progression of the disaster.

2.2. Explainable AI

Explainable AI (XAI) methods aim to achieve fidelity (explanations that align with the actual logic of the model) and comprehensibility for human operators. One popular class of XAI methods includes local feature attributions that explain individual predictions by assigning importance scores to the input and latent features. These can be roughly divided into three groups: 1) perturbation/surrogate-based (e.g., SHAP[24]), 2) gradient-based (e.g., Gradient and Grad-CAM[25, 26]), and 3) decomposition-based methods (e.g., LRP [2]). Group 1 and 2 present the advantage of being model-agnostic. While group 1 are computationally expensive, group 2 are cheap but typically noisy and prone to gradient shattering. Group 3 methods can be made as efficient as a gradient backward pass [27], however they require a careful design for new types of neural network layers. Evaluation in controlled environments w.r.t. ground truths has demonstrated the superior quality of the latter group of methods [28, 29] in computer vision networks.

2.3. Concept-based explanations

Concept Relevance Propagation (CRP) [3, 4] breaks down decisions into human-interpretable concepts by considering feature maps inside a convolutional neural network (CNN) as semantic concept detectors and conditioning the LRP backward pass on these concepts, making it possible to visualize concept-conditioned heatmaps in the input space, as well as retrieve samples that maximize the relevance of a concept. Prototypical Concept-based Explanations (PCX) [5] generalize this to global XAI by summarizing the model’s prediction behavior dataset-wide into prototypes. Another concept-based explanation is Testing with Concept Activation Vectors (TCAV) [30] which measures the gradient-based sensitivity of the model’s latent features w.r.t. pre-defined concept directions that are obtained by training a linear classifier to separate latent features of samples with and without that concept. Such concept-based explanations have been shown to foster human trust in automated medical decisions [31].

3. Methods

3.1. Layer-wise Relevance Propagation (LRP)

Layer-wise Relevance Propagation (LRP) [2] is a post-hoc, model-specific explanation technique that performs a conservative backward decomposition of the model’s prediction, for instance, an object detection score or a semantic segmentation logit, backward through all the layers of the network until the input. The distribution of the relevance is controlled by local propagation rules that preserve the total relevance at each layer, i.e., $\sum_i R_i^{(l)} = \sum_j R_j^{(l+1)}$, where $R_i^{(l)}$ is the relevance of neuron i in layer l . The basic LRP ϵ -rule for a linear layer with neurons i and j in consecutive layers is given by:

$$R_i^{(l)} = \sum_j \frac{z_{ij}}{\sum_k z_{kj} + \epsilon \cdot \text{sign}(\sum_k z_{kj})} R_j^{(l+1)} \quad (1)$$

In this equation, z_{ij} is the contribution of neuron i to neuron j in the forward pass (typically equal to the neuron’s activation x_i multiplied by the connection weight w_{ij} , i.e., $z_{ij} = x_i \cdot w_{ij}$), and ϵ is a small numerical stabilizer. This conservative redistribution guarantees that the total evidence for a detection or segmentation is taken into account throughout the whole computation graph. Other LRP rules for CNNs include the z^+ -rule and γ -rule. For an overview of LRP rules we refer to Montavon et al. [32].

3.2. Extending LRP to PIDNet

The PIDNet architecture introduces special layers which require propagation rules consistent with the LRP conservative backward decomposition. In particular, residual summations and bilinear interpolations are treated as linear layers and explained with the ϵ -rule from Eq. (1).

Additionally, for element-wise multiplications of branches, as they occur in the PIDNet Pixel-Attention-Guided (Pag) and Boundary-Attention-Guided (Bag) fusion layers, of the form $y = x \odot \sigma(g)$, where σ is the sigmoid activation function, we propose to follow the signal-take-all redistribution strategy introduced for gated interactions in LSTMs [33], i.e., the relevance is assigned entirely to the signal branch while the gating branch does not receive any relevance:

$$R_x = R_y, \quad R_g = 0, \quad (2)$$

reflecting the interpretation that the gated input $\sigma(g)$ acts only as a modulator of the signal x in the forward pass; its effect is therefore already reflected in the relevance R_y .

3.3. Concept-based explanations for Segmentation and Detection

To leverage LRP latent feature attributions into concept-based explanations for segmentation and detection, building upon Concept Relevance Propagation (CRP) [3, 4] and Prototypical Concept-based Explanations (PCX) [5], we proceed in the following way. For each prediction of a segmentation mask or of an object bounding box, we start by generating latent LRP relevances of feature maps inside convolutional layers. These relevances are summed up across spatial dimensions to obtain *concept relevance vectors* with one value per feature map. Then, in a second step, we cluster these vectors for all training samples using Gaussian Mixture Model (GMM) clustering. Each resulting cluster then represents a model prediction *strategy*. Then, at inference time, when a new test prediction is made, we compare it to the nearest cluster centroid, also called a *prototype*, in terms of concept usage. This enables us to quantify how similar or dissimilar a new test prediction is w.r.t. prototypical decisions, and whether the prediction shall be labeled as ordinary or as an outlier by PCX.

Besides, we visualize the semantic of concepts through the retrieval of concept maximizing reference samples over the training data, together with generating concept-conditioned heatmaps of the prediction. Example PCX prototypes, concepts and heatmaps in NDM will be provided in Section 4.3 and 4.4.

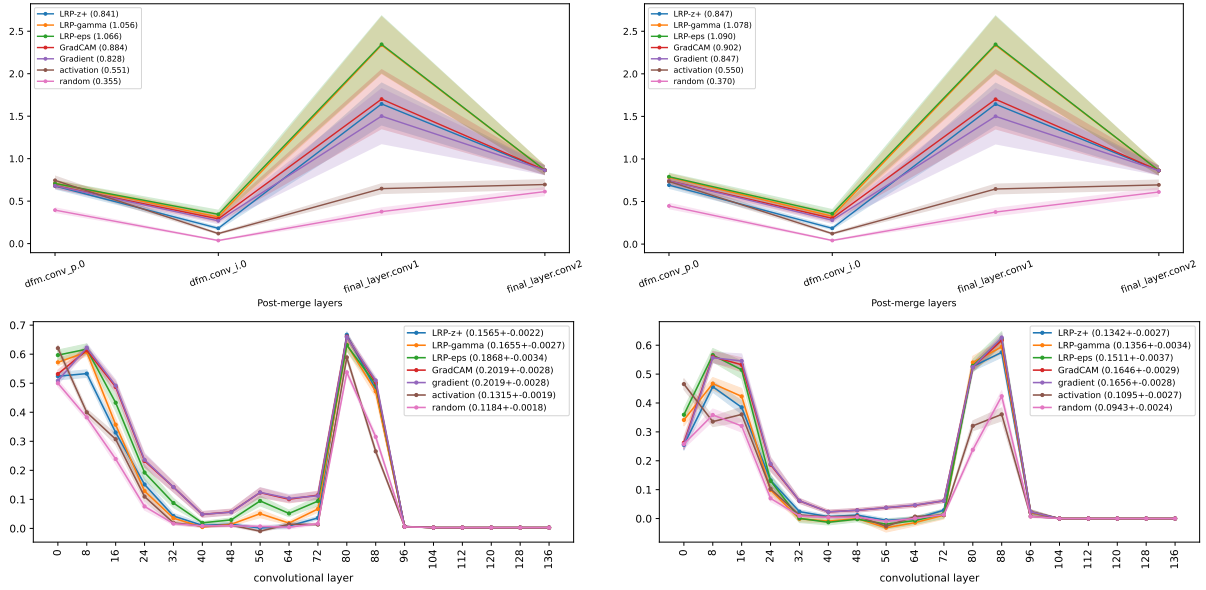


Figure 1: Perturbation-based evaluation of concept-based explanations. *Top:* PIDNet flood segmentation, *Bottom:* YOLOv6s6 car detection. *Left:* AOC for concept deletion, *Right:* AUC for concept insertion. The higher the AOC/AUC, the better. Mean across samples are shown as bold lines, with the Standard Error of the Mean in semi-transparent shading. Average over all layers are given in parenthesis.

4. Experimental results

4.1. Setting: data and models

We test our XAI framework on two state-of-the-art DNNs trained for flood segmentation and car/person detection in a flood scenario using the TEMA AIIA Ahrtal flood dataset¹ captured from UAVs. The code to reproduce our experiments is made publicly available².

PIDNet-small [10]: This segmentation DNN is made of a three-branch architecture inspired by control theory whose Pixel-Attention-Guided (Pag) and Boundary-Attention-Guided (Bag) fusion modules aggregate information from the detail (P-branch), context (I-branch), and boundary (D-branch) branches. It consists of 27 layers and 8.06M parameters. The prediction performance over 244 validation images is: 0.833 mIoU, 0.91 Pixel Accuracy. The number of training images is 1321.

YOLOv6s6 [34]: This object detection DNN employs an EfficientRep backbone for feature extraction and a Rep-PAN neck for multi-scale feature fusion, enabling the detection of objects across various scales in a single stage. It consists of 136 layers and 41.32M parameters. The prediction performance over 179 validation images is: 0.428 F1@0.5, 0.255 mAP@0.5. Since the performance on the car class alone is higher: 0.571 F1@0.5, resp. 0.425 mAP@0.5, we focus on car detection in our experiments. The number of training images is 435 (pre-training was performed on the VisDrone Dataset).

4.2. Evaluation of concept-based explanations

In order to quantitatively assess concept-based explanations, we measure the Area Over the Curve (AOC), resp. Area Under the Curve (AUC), in the prediction logit *change* when deleting, resp. inserting, feature maps (deletion means setting their activation to 0) according to their relevance over 100 validation samples (the feature map relevance is obtained by summing up LRP latent relevances across spatial dimensions). This is analog to pixel-flipping for input relevance evaluation [35]. For PIDNet we perform this perturbation-based evaluation only on the last 4 convolutional layers (i.e., after the 3 branches of PIDnet are merged, since it is unclear how to evaluate concepts in parallel layers), while for YOLOv6s6

¹Dataset available at: <https://doi.org/10.5281/zenodo.18377521>

²Code available at: <https://github.com/samarheydari/Segmentation-and-Object-detection-PCX>

we consider all 136 convolutional layers of the model (in steps of 8). Results are provided in Fig. 1. We find that explanations based on LRP- ϵ [2], Gradient [25] and Grad-CAM [26] are superior, while explanations using standard Activation are only slightly better than Random, which is consistent with previous work [4]. Among the LRP rules, we note that although LRP- ϵ performs better in the present evaluation, other evaluations taking into account localization or ground truth masks have shown that other rules such as the z^+ -rule should generally be preferred in computer vision CNNs [28, 36].

4.3. Visualizing prototypes and concepts

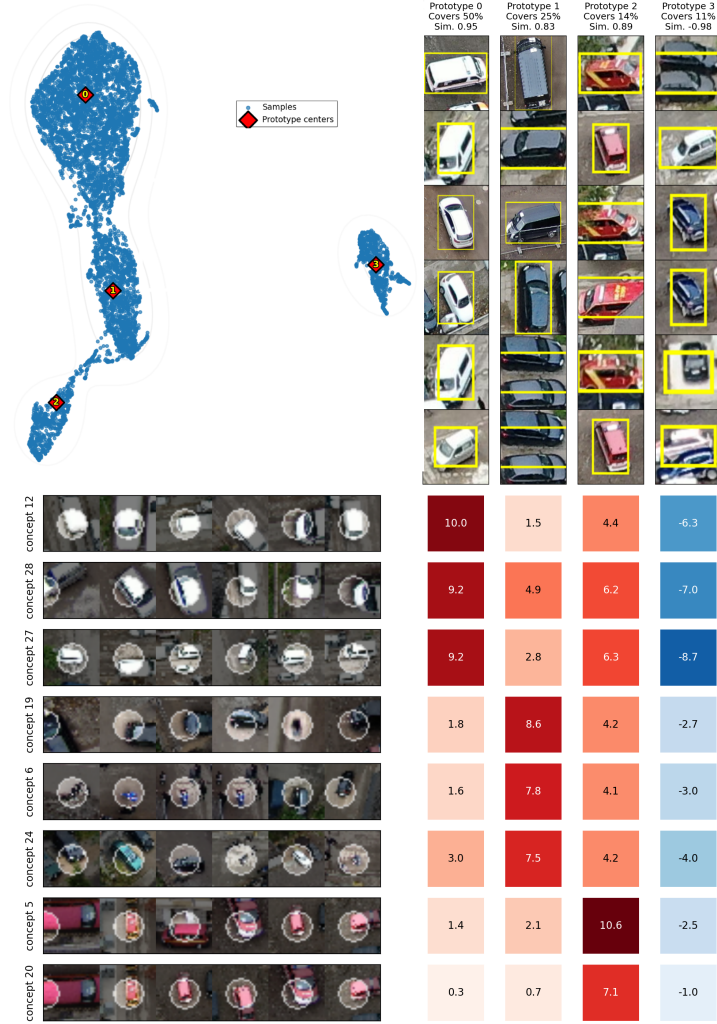


Figure 2: PCX prototypes, and their concept contributions, for car detection with YOLOv6s6.

We visualize PCX prototypes, and most relevant concepts, for car detection on the YOLOv6s6 training data in Fig. 2 (using layer backbone .stem.rbr_dense.conv and 4 clusters for GMM clustering). We observe that the model has developed 4 different strategies to detect cars in a flood scenario, although it was trained to recognize all types of cars altogether. Prototypes 0 and 1 correspond to white and dark passenger cars resp., prototype 2 consists of red ambulances, while the remaining prototype 3 groups various blurry cars. Accordingly, predictions of each prototype rely on concepts of vehicle parts of the same color, except for prototype 3 where common concepts are under-used. In Fig. 2 we further indicate the coverage of each prototype in the data, and the cosine similarity of the *concept relevance vector* of each cluster w.r.t. the full data. Prototype 3 covers 11% of the training samples, and has a negative similarity to the car’s mean of -0.98, reinforcing a data quality issue for this subset of samples. These PCX insights can serve to warn the end-user when a new test prediction is made which gets assigned to

prototype 3, indicating that the prediction is less reliable and that the car type is atypical for the model.



Figure 3: PCX prototypes for flood segmentation with PIDNet.

In Fig. 3 we visualize the prototypes for flood segmentation on the PIDNet training data (using layer `final_layer.conv1` and 10 clusters). We observe that the model has learned to distinguish different flood patterns. For example, prototypes 0, 3 and 9 correspond to linear flood structures, while prototype 8 represents wide-area plain inundation and prototype 6 captures small-scale, street-level flooding. Prototypes 1 and 4 with lower similarity to the flood’s mean likely indicate outlier clusters with atypical flood patterns. Concepts in flood segmentation mainly correspond to water colors and floods occurring near vegetation, roads, or habitation (we refrain from retrieving such concepts due to space constraints).

4.4. Understanding an individual prediction

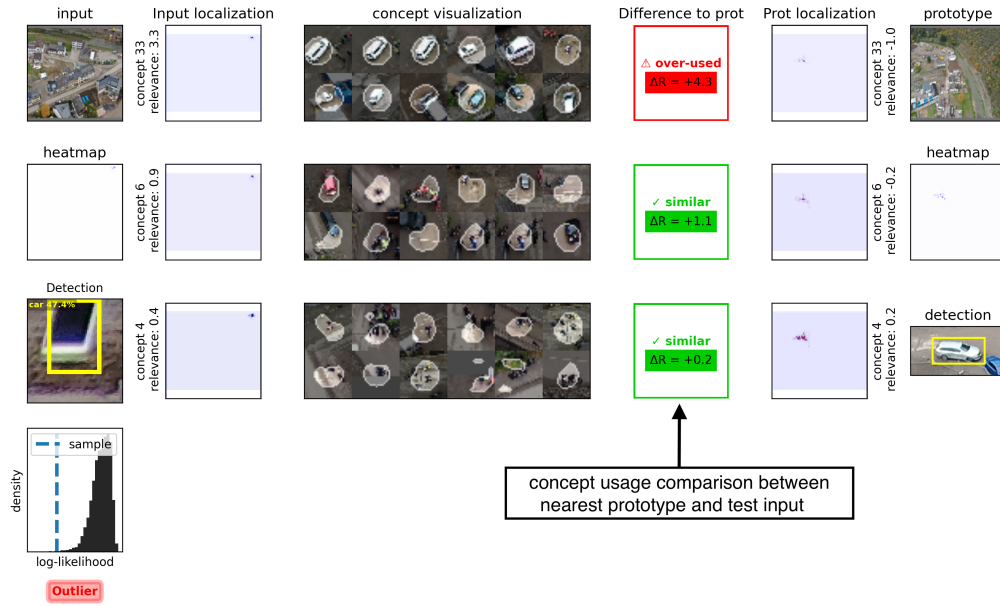


Figure 4: PCX explanation of an outlier prediction for car detection with YOLOv6s6.

In Fig. 4 we showcase an example prediction from the validation data explained with PCX. The top left and top right images are the test input and the nearest prototype from the training data. The 3 middle rows correspond to the most relevant concepts used for the test prediction, which are visualized by retrieving reference images maximizing each concept (further the full heatmap and concept-conditioned heatmaps are provided for each concept). The most important result can be found in the "difference to prototype" column, where concept usage between the test input and the prototype are compared. In particular the top concept, which corresponds to "white vehicle side windows" was over-used in the test input. Hence PCX labeled the prediction as an outlier. Indeed the model mis-detected a window on a roof as a white car, and PCX was able to identify this mistake through the unusual concept profile.

5. Conclusion

In this pilot study we demonstrated, both quantitatively and qualitatively, that concept-based explanations can be extended to the NDM domain on two DNN models and tasks in a flood scenario: flood segmentation with PIDNet and car detection with YOLOv6s6. In particular we highlighted that PCX's prototypes can help to identify and inspect the model's prediction strategies, and that it can successfully detect an outlier prediction. We believe this showcases the usefulness of concept-based XAI for supporting the transparency of DNN decisions in real-world scenarios such as in natural disasters.

Acknowledgments

We thank Maximilian Dreyer for helpful discussions. This work was supported by the European Union's Horizon Europe research and innovation programme's TEMA project, grant number 101093003.

Declaration on Generative AI

The authors have not employed any Generative AI tools.

References

- [1] H. Ritchie, P. Rosado, M. Roser, Natural disasters, Our World in Data (2022). URL: <https://ourworldindata.org/natural-disasters>.
- [2] S. Bach, A. Binder, G. Montavon, F. Klauschen, K.-R. Müller, W. Samek, On pixel-wise explanations for non-linear classifier decisions by layer-wise relevance propagation, *PLOS ONE* 10 (2015) 1–46.
- [3] R. Achibat, M. Dreyer, I. Eisenbraun, S. Bosse, T. Wiegand, W. Samek, S. Lapuschkin, From attribution maps to human-understandable explanations through Concept Relevance Propagation, *Nature Machine Intelligence* 5 (2023) 1006–1019.
- [4] M. Dreyer, R. Achibat, T. Wiegand, W. Samek, S. Lapuschkin, Revealing Hidden Context Bias in Segmentation and Object Detection through Concept-specific Explanations, in: *Conference on Computer Vision and Pattern Recognition Workshops*, 2023, pp. 3828–3838.
- [5] M. Dreyer, R. Achibat, W. Samek, S. Lapuschkin, Understanding the (Extra-)Ordinary: Validating Deep Model Decisions with Prototypical Concept-based Explanations, in: *Conference on Computer Vision and Pattern Recognition Workshops*, 2024, pp. 3491–3501.
- [6] A. Akhyar, M. Asyraf Zulkifley, J. Lee, T. Song, J. Han, C. Cho, S. Hyun, Y. Son, B.-W. Hong, Deep artificial intelligence applications for natural disaster management systems: A methodological review, *Ecological Indicators* 163 (2024) 112067.
- [7] V. Linardos, M. Drakaki, P. Tzionas, Y. L. Karnavas, Machine learning in disaster management: Recent developments in methods and applications, *Mach. L. and Knowl. Extract.* 4 (2022) 446–473.
- [8] A. W. Ali, S. Kurnaz, Optimizing deep learning models for fire detection, classification, and segmentation using satellite images, *Fire* 8 (2025).
- [9] M. Li, Y. Zhang, L. Mu, J. Xin, Z. Yu, S. Jiao, H. Liu, G. Xie, Y. Yingmin, A real-time fire segmentation method based on a deep learning approach, *IFAC-PapersOnLine* 55 (2022) 145–150.
- [10] J. Xu, Z. Xiong, S. P. Bhattacharyya, PIDNet: A Real-time Semantic Segmentation Network Inspired by PID Controllers, in: *CVPR*, 2023, pp. 19529–19539.
- [11] D. Fotiou, V. Mygdalis, I. Pitas, RoboFireFuseNet: Robust Fusion of Visible and Infrared Wildfire Imaging for Real-Time Flame and Smoke Segmentation, *TechRxiv* (2025).
- [12] E. Vlachos, C. Papaioannidis, I. Pitas, Neural architecture search and knowledge distillation for semantic image segmentation on big wildfire datasets, in: *EUSIPCO*, 2025, pp. 750–754.
- [13] P. Montesidis, V. Mygdalis, I. Pitas, Improve real-time flood segmentation by encoding and distilling foreground information, in: *Int. Conference on Image Processing*, 2025, pp. 1840–1845.

- [14] M. Siavrakas, C. Papaioannidis, I. Pitas, Blaze: A Dataset For Wildfire And Burnt Area UAV Image Classification And Segmentation, in: Int. Conference on Image Processing, 2025, pp. 1960–1965.
- [15] A. Gerontopoulos, D. Papaioannou, C. Papaioannidis, I. Pitas, Real-time flood water segmentation with deep neural networks, in: IEEE Int. S. Cluster, Cloud Int. Comp. Workshops, 2025, pp. 85–91.
- [16] M. Arlovic, F. Hrzic, M. Patel, T. Bednarz, J. Balen, Evaluation of synthetic data impact on fire segmentation models performance, Scientific Reports 15 (2025).
- [17] E. Spatharis, C. Papaioannidis, V. Mygdalis, I. Pitas, Unrealfire: A Synthetic Dataset Creation Pipeline for Annotated Fire Imagery in Unreal Engine, in: Int. Conference on Image Processing Workshops, 2025, pp. 610–615.
- [18] M. Tzimas, V. Mygdalis, C. Papaioannidis, I. Pitas, Extreme weakly supervised binary semantic image segmentation via one-pixel supervision, Pattern Recognition (2026) 113048.
- [19] A. Apostolidis, V. Mygdalis, M. Tzimas, I. Pitas, MEWS: Semantic image segmentation with multiclass extreme weak supervision, Neurocomputing (2026) 133290.
- [20] Y. Pi, N. D. Nath, A. H. Behzadan, Convolutional neural networks for object detection in aerial imagery for disaster response and recovery, Advanced Engineering Informatics 43 (2020) 101009.
- [21] Z.-A. Zhao, S. Wang, M.-X. Chen, Y.-J. Mao, A. C.-H. Chan, D. K.-H. Lai, D. W.-C. Wong, J. C.-W. Cheung, Enhancing Human Detection in Occlusion-Heavy Disaster Scenarios: A Visibility-Enhanced DINO (VE-DINO) Model with Reassembled Occlusion Dataset, Smart Cities 8 (2025).
- [22] M. Wieland, S. Schmidt, B. Resch, A. Abecker, S. Martinis, Fusion of geospatial information from remote sensing and social media to prioritise rapid response actions in case of floods, Natural Hazards 121 (2025) 8061–8088.
- [23] S. Alqithami, Integrating sentiment analysis and reinforcement learning for equitable disaster response: A novel approach, Sustainability 17 (2025).
- [24] S. M. Lundberg, S.-I. Lee, A unified approach to interpreting model predictions, in: Advances in Neural Information Processing Systems, 2017, pp. 4765–4774.
- [25] K. Simonyan, A. Vedaldi, A. Zisserman, Deep inside convolutional networks: Visualising image classification models and saliency maps, in: ICML, 2014.
- [26] R. R. Selvaraju, M. Cogswell, A. Das, R. Vedantam, D. Parikh, D. Batra, Grad-CAM: Visual explanations from deep networks via gradient-based localization, in: ICCV, 2017.
- [27] L. Arras, B. Puri, P. Kahardipraja, S. Lapuschkin, W. Samek, A close look at decomposition-based XAI-methods for Transformer language models, arXiv:2502.15886 (2025).
- [28] L. Arras, A. Osman, W. Samek, CLEVR-XAI: A benchmark dataset for the ground truth evaluation of neural network explanations, Information Fusion 81 (2022) 14–40.
- [29] A. Mamalakis, E. A. Barnes, I. Ebert-Uphoff, Investigating the fidelity of explainable artificial intelligence methods for applications of convolutional neural networks in geoscience, Artificial Intelligence for the Earth Systems 1 (2022).
- [30] B. Kim, M. Wattenberg, J. Gilmer, C. J. Cai, J. Wexler, F. B. Viégas, R. Sayres, Interpretability beyond feature attribution: Quantitative testing with concept activation vectors (TCAV), in: ICML, 2018.
- [31] E. Poeta, G. Ciravegna, E. Pastor, T. Cerquitelli, E. Baralis, Concept-based Explainable Artificial Intelligence: A Survey, ACM Computing Surveys (2025).
- [32] G. Montavon, A. Binder, S. Lapuschkin, W. Samek, K.-R. Müller, Layer-wise relevance propagation: An overview, in: Explainable AI: Interpreting, Explaining and Visualizing Deep Learning, volume 11700 of LNCS, 2019, pp. 193–209.
- [33] L. Arras, G. Montavon, F. Klauschen, K.-R. Müller, W. Samek, Explaining recurrent neural network predictions in sentiment analysis, EMNLP Workshop on Computational Approaches to Subjectivity, Sentiment and Social Media Analysis (2017).
- [34] C. Li, L. Li, H. Jiang, K. Weng, Y. Geng, et al., YOLOv6: A Single-Stage Object Detection Framework for Industrial Applications, arXiv.2209.02976 (2022).
- [35] W. Samek, A. Binder, G. Montavon, S. Lapuschkin, K.-R. Müller, Evaluating the visualization of what a deep neural network has learned, IEEE TNNLS 28 (2017) 2660–2673.
- [36] M. Kohlbrenner, A. Bauer, S. Nakajima, A. Binder, W. Samek, S. Lapuschkin, Towards best practice in explaining neural network decisions with LRP, in: Int. Joint Conf. on Neur. Netw., 2020, pp. 1–7.