



UNIVERSITY OF LEEDS

This is a repository copy of *PIP-Net: Pedestrian Intention Prediction in the Wild*.

White Rose Research Online URL for this paper:

<https://eprints.whiterose.ac.uk/228796/>

Version: Accepted Version

Article:

Azarmi, M. orcid.org/0000-0003-0737-9204, Rezaei, M. orcid.org/0000-0003-3892-421X and Wang, H. (2025) PIP-Net: Pedestrian Intention Prediction in the Wild. IEEE Transactions on Intelligent Transportation Systems, 26 (7). pp. 9824-9837. ISSN 1524-9050

<https://doi.org/10.1109/tits.2025.3570794>

This is an author produced version of an article published in IEEE Transactions on Intelligent Transportation Systems, made available under the terms of the Creative Commons Attribution License (CC-BY), which permits unrestricted use, distribution and reproduction in any medium, provided the original work is properly cited.

Reuse

This article is distributed under the terms of the Creative Commons Attribution (CC BY) licence. This licence allows you to distribute, remix, tweak, and build upon the work, even commercially, as long as you credit the authors for the original work. More information and the full terms of the licence here:

<https://creativecommons.org/licenses/>

Takedown

If you consider content in White Rose Research Online to be in breach of UK law, please notify us by emailing eprints@whiterose.ac.uk including the URL of the record and the reason for the withdrawal request.



eprints@whiterose.ac.uk
<https://eprints.whiterose.ac.uk/>

PIP-Net: Pedestrian Intention Prediction in the Wild

Mohsen Azarmi , Mahdi Rezaei , He Wang 

Abstract—Accurate pedestrian intention prediction (PIP) by Autonomous Vehicles (AVs) is one of the current research challenges in this field. In this article, we introduce PIP-Net, a novel framework designed to predict pedestrian crossing intentions by AVs in real-world urban scenarios. We offer two variants of PIP-Net designed for different camera mounts and setups. Leveraging both kinematic data and spatial features from the driving scene, the proposed model employs a recurrent and temporal attention-based solution, outperforming state-of-the-art performance. To enhance the visual representation of road users and their proximity to the ego vehicle, we introduce a categorical depth feature map, combined with a local motion flow feature, providing rich insights into the scene dynamics. Additionally, we explore the impact of expanding the camera’s field of view, from one to three cameras surrounding the ego vehicle, leading to an enhancement in the model’s contextual perception. Depending on the traffic scenario and road environment, the model excels in predicting pedestrian crossing intentions up to 4 seconds in advance, which is a breakthrough in current research studies in pedestrian intention prediction. Finally, for the first time, we present the Urban-PIP dataset, a customised pedestrian intention prediction dataset, with multi-camera annotations in real-world automated driving scenarios.

Index Terms—Autonomous vehicles, pedestrian behaviour, pedestrian crossing prediction, computer vision, deep neural networks.

I. INTRODUCTION

PEDESTRIANS are the most vulnerable road users and face a high risk of fatal accidents [1]. Ensuring pedestrian safety in automated driving, particularly in mixed AV-pedestrian traffic scenarios, heavily relies on the AV’s capability in “pedestrian intention prediction (PIP)” [2]. A PIP system determines if a pedestrian is likely to cross the road shortly (within the next few seconds). This study aims to investigate the critical visual clues that pedestrians exhibit when they intend to cross the road, and then provide a model which predicts crossing behaviour, a few seconds in advance.

Anticipating pedestrian crossing behaviour is a difficult task due to various environmental factors that affect human intention [3]–[5]. Even in the simulated scenarios in which the majority of parameters are under control, crossing prediction is a challenging endeavour [6]. Factors like interactions with other pedestrians, traffic signs, road congestion, and vehicle



Fig. 1. Pedestrians’ crossing intention prediction in complex urban scenarios via contextual data analysis and a multi-camera perception setup.

speed can influence pedestrians’ tendency to cross the road in front of AVs [7].

Computer vision plays a crucial role in enabling AVs to perceive their surrounding environment by analysing the visual data captured via multiple sensors, such as cameras, LiDAR, Radar, etc. Learning-based models, in particular deep neural networks (DNNs), have shown remarkable success in various computer vision tasks, including scene understanding, semantic segmentation [8], road users classification, localisation [9], and motion prediction [10]. Figure 1 illustrates some of the perceivable factors such as depth, pedestrian pose, and surrounding objects, that an autonomous vehicle should consider to interpret the scene and estimate the pedestrians’ intention. DNNs are particularly effective at learning complex patterns and features from visual data, making them a natural fit for tasks that involve analysing images or videos to comprehend pedestrian behaviour [11]–[13]. They also offer significant capabilities in multi-modal integration by providing a neural-based mechanism to process and fuse all the perceived information from diverse sensors. This integration may enhance the overall understanding of the environment and help to make more accurate and safer decisions [14].

Several datasets, such as JAAD [15], PIE [16], and STIP [17], use onboard camera recordings and their data are publicly released for the study of pedestrians’ behaviour before and during road crossing. However, most of the current research works suffer from supplying a multi-camera setup to leverage the benefits of sensor fusion and multi-modal perception. In addition to the above-mentioned datasets, some baseline approaches [18] have also been established for analysing the visual cues and signals that pedestrians emit through their body language and positioning. The approaches highlight the benefits of combining these features with contextual information. Contextual information may include factors such as the road’s location, the time of day, weather conditions, the presence of traffic signals or crosswalks, the type of road (urban, suburban, rural), and the position and behaviour of other vehicles near

Received 19 February 2024; revised 18 March 2025; accepted 5 May 2025. This work was supported by European Union’s Horizon 2020 Research and Innovation Programme under Grant 101006664. The associate Editor for this article was V. Chamola. (Corresponding author: Mohsen Azarmi)

Mohsen Azarmi and Mahdi Rezaei are with the Institute for Transport Studies, Computer Vision and Machine Learning Group, University of Leeds, LS2 9JT Leeds, U.K. (e-mail: tsmaz@leeds.ac.uk and m.rezaei@leeds.ac.uk)

He Wang is with UCL Centre for Artificial Intelligence, Department of Computer Science, University College London, WC1E 6BT London, U.K. (e-mail: he_wang@ucl.ac.uk)

Digital Object Identifier 1-1109/TITS.2025.3570794

the scene [19]. To the best of our knowledge, no extensive research has been conducted to understand and interpret such contextual details and their effects on pedestrian’s decision-making.

In this study, we propose a customised DNN-based framework, called “PIP-Net” that takes various features of pedestrians, the environment, and the ego-vehicle state into account, to learn the context of a crossing scenario and consequently predict the crossing intention of pedestrians in real-world AV urban driving scenarios. The main contributions of this research are highlighted as follows:

- A novel feature fusion model is presented to integrate AV’s surrounding cameras and combine visual and non-visual modalities, as well as a hybrid feature map that incorporates depth and instance semantic information of each road user to comprehend the latent dynamics in the scene.
- Introducing the multi-camera “Urban-PIP dataset”, which includes various real-world scenarios of pedestrian crossing for autonomous driving in urban areas.
- We examine the effectiveness of the various input features, temporal prediction expansion, and the worthiness of expanding the vehicle field of view from one camera to three cameras based on the latest Waymo car camera setup [20] to ensure the developed model will be in line with the current technology developments in the AV industry.
- Finally, we evaluate the effectiveness of the proposed model on the widely utilised PIE dataset and the introduced Urban-PIP dataset, outperforming the state-of-the-art (SOTA) for crossing prediction.

II. RELATED WORKS

Recently, pedestrian crossing intention prediction research has surged and gained significant attention within the autonomous driving research community [11]–[13]. Most current methods mainly address the problem by taking two aspects into account: Discovering influential factors and features for interpreting road users’ interactions [4], [6], [7], and designing the analytic model to predict the pedestrians’ crossing intention [15]–[19]. Both research directions mainly utilise advanced learning-based techniques. Deep learning methods have been fostered on multiple features of pedestrians and the environment, whether derived from annotations, visual information from videos or their combinations [2], [21]. The following two subsections are dedicated to introducing approaches that utilise DNN-based architectures for spatio-temporal analysis and feature selection/fusion.

A. Spatio-temporal Analysis

Recently, there has been a shift from still image analysis to the incorporation of temporal information into the prediction models. Rather than relying on individual images, most contemporary methods utilise sequences of input images for decision-making by their prediction models. This adaptation recognises the significance of temporal data in enhancing the

prediction task, resulting in what is known as spatio-temporal modelling.

Spatio-temporal modelling can be achieved through a two-step process. Initially, visual (spatial) features per frame can be extracted using a visual encoder such as 2D convolutional neural networks (CNNs) [22], vision transformers (ViTs) [23], or graph convolution networks (GCNs) [17]. Subsequently, these extracted features are then fed into a temporal encoder like recurrent neural networks (RNNs) [24], using long short-term memory (LSTM) [25]–[27] or the gated recurrent unit (GRU) modules [28]–[30]. For instance, in [24], [26], [31], 2D convolutions are employed to extract visual features from image sequences, while RNNs encode the temporal relationships among these features. These sequentially encoded visual features are then inputted into a fully-connected layer to generate the ultimate crossing prediction.

An alternative approach is extracting both spatial and temporal features involves the utilisation of 3D CNNs (Conv3D) [32] or Transformer architecture [33]. 3D CNNs can directly capture spatio-temporal features by substituting the 2D kernels within the convolution and pooling layers of a 2D CNN with their 3D equivalents. For instance, in works such as [34]–[37], a framework based on a 3D CNN is employed to directly extract sequential visual features from sequences of pedestrian images. The ultimate prediction is then made using a fully-connected layer. Transformer architecture uses self-attention mechanisms [38] to capture long-range relationships, both within a single frame (spatial) and between different frames (temporal). This helps in thoroughly analysing pedestrian dynamics [39]–[42]. Furthermore, hybrid models combining CNNs, RNNs, and attention mechanisms have been explored [43]–[45] to leverage the strengths of these approaches.

B. Feature Selection and Integration

Reliance on a single pedestrian feature for crossing prediction, such as using only the pedestrian’s pose kinematics [46] or bounding box [41], has resulted in context-free crossing predictions that neglect other scene-specific modalities and miss traffic and situational awareness. Instead, it is possible to treat various types of information, such as the pedestrian’s image, body pose keypoints, bounding box, vehicle dynamics, and the broader contextual backdrop as distinct input channels for the prediction model [13].

Studies such as [47]–[49] have incorporated human poses or skeletons into pedestrian crossing prediction tasks alongside various features such as pedestrian’s local image, bounding box, and vehicle speed, to construct the intention classifier. This approach has shown improved prediction accuracy but often neglects other important features or lacks proper attention to feature integration.

The investigation into the types of features, such as pedestrians and environmental context, is still ongoing. For instance, pedestrian-to-vehicle distance is considered one of the most influential factors in pedestrians’ decisions to cross [6]. This feature is typically estimated as a single measure between the target pedestrian and the ego vehicle [4], [50], while relative velocity and closing speed, which refers to how quickly the

gap between the ego vehicle and the pedestrian is decreasing, are also investigated [51]. Alternatively, a depth map of the scene (see Figure 2e) can be used to assess the distance from other road users and possibly reveal the underlying dynamics [52]. However, the depth map is susceptible to noise due to rough estimation, which can lead to inaccuracies in scenarios involving multiple pedestrians crossing [53].

On the other hand, some studies specifically concentrated on feature integration. For instance, vision and non-vision branches fusion [43] suggest how to efficiently combine diverse data modalities at different stages of a DNN model to surge the intention prediction accuracy. Another study [11] is conducted to merge two visual and three non-visual elements of the pedestrian, scene, and subject vehicle in a multi-stream network. From a different perspective, in studies such as [12], [45], local and global contextual information has been weighted by an attention mechanism and fused together to apply a prediction on *Joint Attention in Autonomous Driving* (JAAD) [15] and *Pedestrian Intention Estimation* (PIE) [16] datasets.

C. Research Gaps

Despite advancements in pedestrian crossing intention prediction, several key gaps persist, as outlined in recent surveys [2], [5], [21]. In this study, we focus on addressing four research gaps:

1) Pedestrian crossing intention highly relies on the distance of the AV to the pedestrian and the relative distance of the pedestrian to other road users, which may fall into various categories of instance segmentation (e.g., cars, other pedestrians, etc.). None of the reviewed research has considered the simultaneous impact of both features on pedestrian intention.

2) To the best of our knowledge, no prior study has considered instance segmentation to smooth and normalise the distance measurement of road user instances. In this article, we propose a new concept of *Categorical Depth* which integrates the classic noisy depth measurement with instance segmentation to gain more accurate depth information and provide context-aware spatial data.

3) The reviewed models, often have limited generalisability and are incapable of performing in the wild and real-world automated driving scenarios, as they have normally been tested with a human-driven vehicle [54]. Our study focuses on the real-world Waymo dataset, which is collected from an AV's field of view. Our added pedestrian intention annotations to this dataset help address the issue of inadequate annotated ground-truth data availability in this field of research.

4) There is a shortage of dedicated neural network architectures capable of effectively accommodating and extracting maximal multi-camera information from around the AV for context recognition, hence an accurate model for predicting pedestrian crossing intentions. Camera integration is proposed in this study to cope with the limited field of view.

III. METHODOLOGY

We propose the PIP-Net model, which is based on deep neural networks for predicting pedestrian crossing intention.

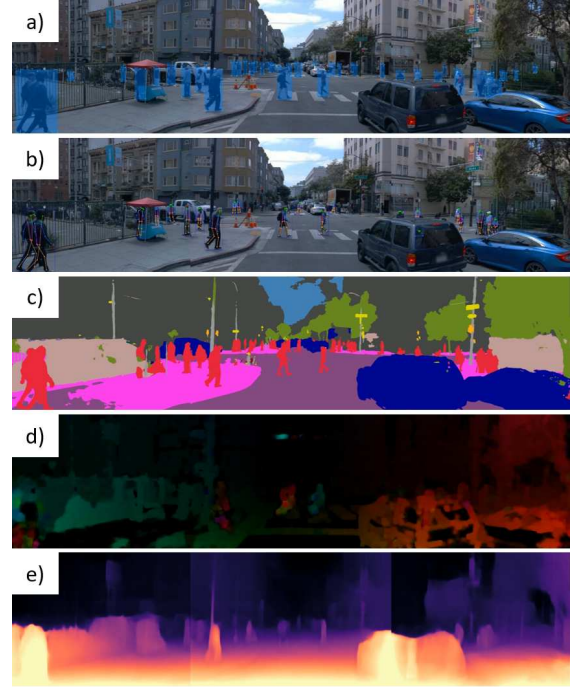


Fig. 2. Analysing pedestrians' features and traffic scene dynamics through: a) Pedestrian localisation, highlighted with bounding boxes, b) Gesture understanding with pose estimation, c) Object categorisation through segmentation, d) Global motion patterns using optical flow, and e) Estimating distance via a global depth heatmap.

The model incorporates spatial-temporal features such as road users' positioning, pose, and dynamic movements, along with a hybrid feature map, the categorised semantic and depth information as input to the network. A multi-camera stitching and integration model is developed to facilitate panoramic viewing, enabling a synchronised pedestrian ID assignment and tracking across the entire multi-view scene, thus enhancing the PIP-Net model's understanding of spatial characteristics and contextual information.

An overview of the proposed architecture is illustrated in Figure 3. The input features are categorised into spatial kinematic data and contextual data, and they are passed to the model through distinct pipelines based on their data type. Finally, the recurrent module and attention module are utilised to improve temporal data processing.

A. Spatial Kinematics

Kinematic input data includes the positioning of the pedestrian in the scene with reference to the detected pedestrian bounding box P_{bb} , pedestrian body pose keypoints P_{bp} , and the ego-vehicle speed V_s .

The data is arranged in a gated recurrent unit (GRU) layer [28], beginning with the *Bounding Box* feature P_{bb} . It indicates the location of the pedestrians which is detected through the customised You-Only-Look-Once algorithm for road user detection [55]. This feature is defined as:

$$P_{bb} = \{b_i^{t-m}, b_i^{t-m+1}, \dots, b_i^t\}, \quad (1)$$

where $b_i = [x_1, y_1, x_2, y_2] \in \mathbb{R}^4$ represents the coordinates of a pedestrian bounding box. It consists of the top-left ($[x_1, y_1]$) and the bottom-right coordinates ($[x_2, y_2]$). The dimension of

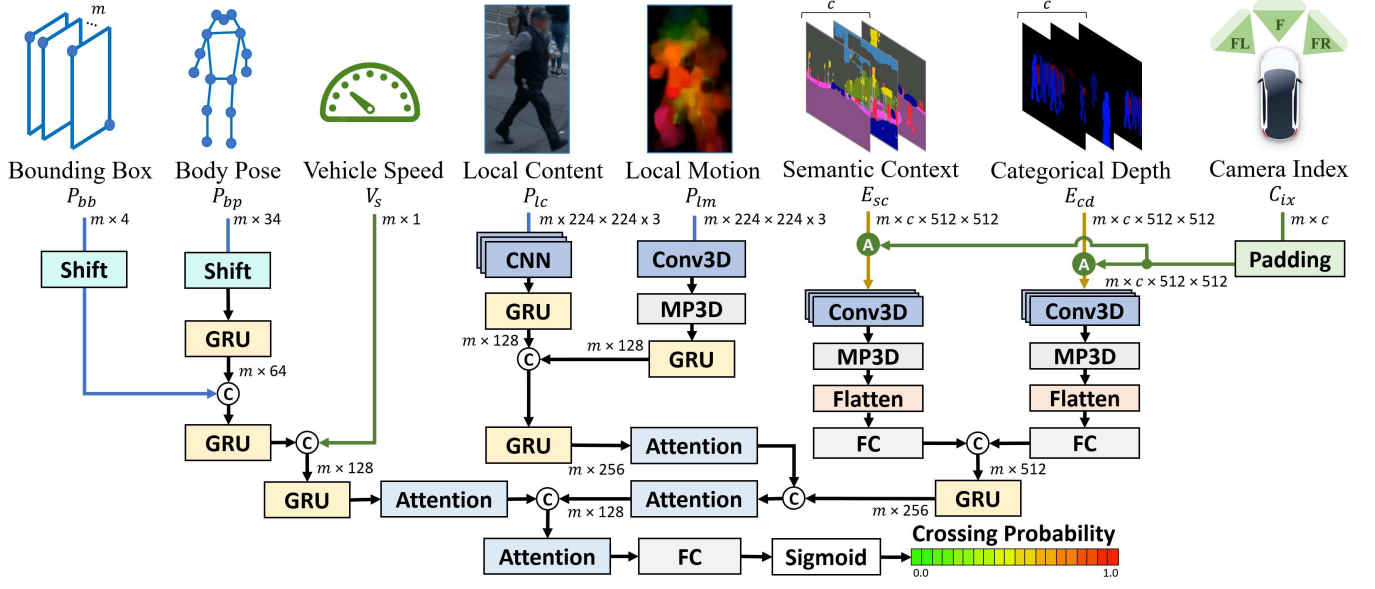


Fig. 3. The overview of the proposed DNN-based framework. PIP-Net receives *Spatial Kinematic* and *Spatial Context* data, and applies multi-modal feature fusion and multi-camera integration with temporal and attention-based analysis to predict the pedestrians' crossing intentions. Shift and Padding units adjust the location of the target pedestrian with respect to the corresponding camera. Aggregation module (A) fuses the cameras' features as shown in Figure 5.

the bounding box matrix P_{bb} is determined as $m \times 4$, where m is the observation sequence length which indicates the number of frames that are observed to predict the pedestrian crossing. We define t as the *decisive moment*, 0.5 ~ 4 seconds before the crossing event.

The *Body Pose* feature P_{bp} is defined as:

$$P_{bp} = \{p_i^{t-m}, p_i^{t-m+1}, \dots, p_i^t\}, \quad (2)$$

where the pose keypoints are obtained using YOLO-Pose [56], which estimates the pose of a person by detecting 17 keypoints joints, including the shoulders, elbows, wrists, hips, knees, ankles, eyes, ears, and nose. The keypoints are represented by a 34-dimensional vector, p_i , which contains the 2D coordinates of each joint for the i -th pedestrian at time t .

The *Vehicle Speed* is also defined as:

$$V_s = \{s_i^{t-m}, s_i^{t-m+1}, \dots, s_i^t\}, \quad (3)$$

where s_i refers to the exact speed of the ego-vehicle in *km*.

B. Spatial Context

Contextual input data includes pedestrian features, such as a pedestrian-bounded image (*Local Content*, P_{lc}) and corresponding motion flow analysis of the pedestrian (*Local Motion*, P_{lm}), the environment features like semantic segmentation of the scene (*Semantic Context*, E_{sc}), as well as our proposed hybrid feature map (*Categorical Depth*, E_{cd}), which refines depth information for specific pedestrians and vehicles in the scene.

The *Local Content* feature P_{lc} is defined as:

$$P_{lc} = \{lc_i^{t-m}, lc_i^{t-m+1}, \dots, lc_i^t\}, \quad (4)$$

where lc_i denotes the feature vector that is output by applying the CNN backbone to an RGB image. An ImageNet pre-trained VGG19 network is used as the CNN backbone, with

a maximum pooling layer as suggested in [43]. Subsequently, a GRU is applied recursively to process the temporal dimension of the feature. The RGB image contains an individual pedestrian, cropped based on the bounding box location and subsequently warped to dimensions of $224 \times 224 \times 3$ pixels, which is the optimum spatial size in the network [23].

The *Local Motion* feature P_{lm} is derived from the dense optical flow analysis within the pedestrian-bounded image. This analysis is more consistent than examining the entire scene, which can be affected by ego-vehicle motions. We opt for a more advanced optical flow approach using FlowNet2 [57]. This deep learning-based method offers improved accuracy and faster run-time performance. The *Local Motion* is defined as:

$$P_{lm} = \{lm_i^{t-m}, lm_i^{t-m+1}, \dots, lm_i^t\}, \quad (5)$$

where lm_i is considered as the localised i -th pedestrian motion descriptor. A Conv3D layer is used to extract a feature vector of size $(m, 512)$, where m represents the observation sequence length. The feature vector is then passed through a 3D max-pooling layer (MP3D) with a kernel size of 4×4 and a GRU module. This process yields a $(m, 128)$ vector, this vector is then inputted into a GRU layer, resulting in an $(m, 128)$ vector, suitable for concatenation with the *Local Content* feature vector.

The *Semantic Context* feature E_{sc} is defined as:

$$E_{sc} = \{sc^{t-m}, sc^{t-m+1}, \dots, sc^t\}, \quad (6)$$

where sc refers to the semantic segmentation of objects within the entire scene encompassing road structure and users. This feature ensures that the model considers the spatial distribution of classes for both moving and static objects within the scene. The semantic information is extracted by Slot-VPS [58] model, which is a panoptic video segmentation algorithm that not only offers semantic segments but also a unique ID for each instance of the objects in the scene. The segmented

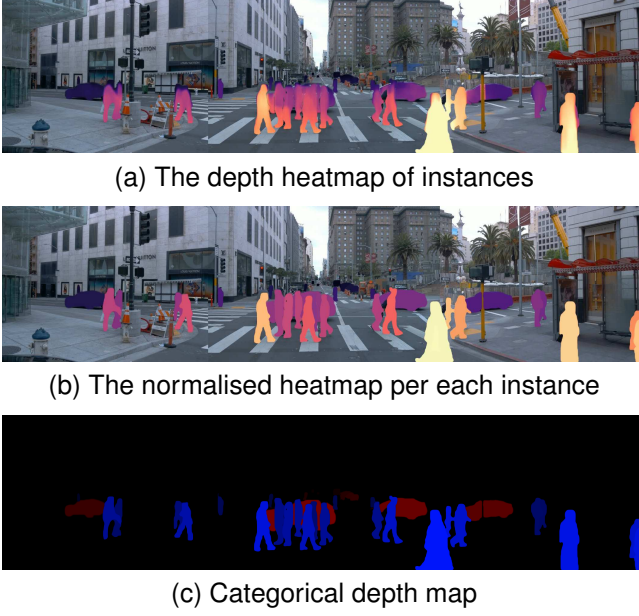


Fig. 4. **Colour-coded visualisation of road users' distance to the ego-vehicle after the tripe camera stitching.** 4a is the proportion of the global depth heatmap which pedestrians and vehicles occupy; 4b is the normalised heatmap values per each instance; 4c is a visualised E_{cd} feature in an RGB format in which pedestrians and vehicles integrated into the blue and red channel, respectively, and high-intensity values indicate closer users. A categorical depth map has a positive effect on removing the camera stitching effect and increasing the clarity and saliency of pedestrians and vehicles in the scene.

classes include 8 dynamic classes (person, rider, car, truck, bus, train, motorcycle, and bicycle) and 11 static classes (traffic light, fire hydrant, stop sign, parking meter, bench, handbag, road, sidewalk, sky, building, and vegetation). The instance segments corresponding to pedestrians and vehicles are integrated into the *Categorical Depth* feature, denoted as E_{cd} , using the following formulation:

$$E_{cd} = \{cd^{t-m}, cd^{t-m+1}, \dots, cd^t\}, \quad (7)$$

where cd represents the hybrid feature map containing depth information for pedestrians and vehicles within the scene. The depth data are initially estimated using the ManyDepth model [59] and encoded in a heatmap representation, resulting in a global depth heatmap. As illustrated in Figure 2e, high-intensity spots (white and oranges) indicate proximity to the ego vehicle, while low-intensity spots (navy blue and black) represent greater distances. However, our experiments revealed that the global depth heatmap is unreliable due to inconsistencies in providing clear object boundaries. To address this, the pedestrian and vehicle instances are cropped using instance masks obtained from the Slot-VPS, as shown in Figure 4a. Subsequently, the intensity of pixels within each instance is normalised by averaging, yielding a normalised heatmap as seen in Figure 4b. This process ensures a clear and consistent depth estimation for each class, as depicted in Figure 4c.

Both inputs, E_{sc} and E_{cd} , undergo extraction via three Conv3D layers to be assessed for the spatio-temporal analysis. The feature dimensions are gradually reduced ($512 \rightarrow 256 \rightarrow 128 \rightarrow 64$) by repeatedly subjecting them to max-pooling layers. This process not only selects the most important information from the local neighbourhood of each pooling window

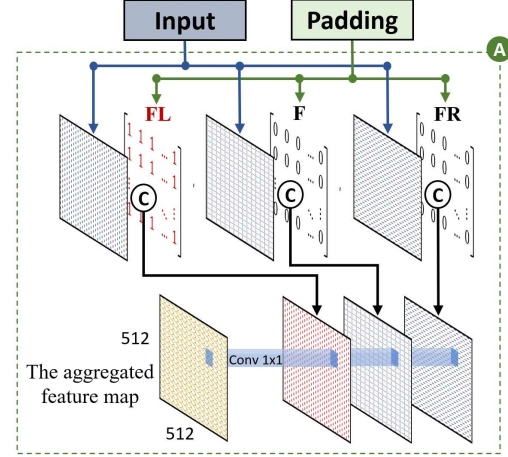


Fig. 5. **The multi-camera feature aggregation module.** In this example, the front-left camera (FL) is the sentinel camera, where the target pedestrian is observed. The padding module expands the *Camera Index* value and provides a camera indicator mask for combining with *Sense Context* and *Categorical Depth* features.

but also reduces the spatial dimensions (width and height) of the feature maps and the computational complexity of the network. Then the features are organised using a flattened layer, resulting in a one-dimensional array that is suitable for concatenation and can be fed into a fully-connected layer (FC). Finally, the data is passed through a GRU module. The outputs of the three GRUs are combined and concatenated into a single output, which is then passed through an attention mechanism.

C. Cameras Features Integration

The incorporation of multiple cameras might be beneficial for capturing complex traffic scenarios, such as intersections, thanks to providing a surrounding field of view. In these scenarios, pedestrians may approach the road from the sides rather than directly in front of the vehicle. They may also choose to cross the road while a vehicle is changing lanes or making a turn. By incorporating left and right-side cameras, we can gather critical information about pedestrians in adjacent lanes or at the side of the vehicle. The Waymo dataset is one of the best options with three cameras ($c = 3$) that also offer a diversity of real-world pedestrian crossing scenarios. The cameras are named front-left (FL), front (F), and front-right (FR) positioned from the AV's left to right angles (as shown in Figure 2). The synchronised videos provided have an approximately 11% overlap along the edges. These overlapping areas can introduce redundancy in data and make it challenging to precisely determine the pedestrian's position, movement, and intention when it moves from one camera scene to another. Therefore, we merge the cameras using the Panoptic stitching over time approach [20], excluding the overlapping regions and giving higher priority to the front-view camera. Figure 2 illustrates an example of different types of features that have been extracted from three cameras, and then stitched together to constitute a single wide image.

We define *Sentinel Camera* as a variable that indicates the index of the camera (C_{ix}) on which the target pedestrian has been observed. Using the *Camera Index*, we can adjust the pose and bounding box coordinates with respect to the sentinel camera. This task will be accomplished by the Shift

unit, as shown in Figure 3 in cyan, which extends the global coordinates from the leftmost camera to the rightmost one, and applies these adjustments to the inputs. In this context, the Padding unit is responsible for generating a zero binary mask of size $c \times 512 \times 512$, where the c -th dimension corresponds to the sentinel camera being set to one. Here, c represents the number of cameras. As the aggregation module (A) shown in Figure 5, the binary masks are combined with the input vectors i.e., E_{sc} and E_{cd} . Subsequently, a pointwise convolution (Conv 1×1) operator aggregates all the features across cameras. This process combines features from different channels (cameras) at each spatial location, allowing a weighted combination of input features.

D. Temporal and Attention Module

To account for the temporal context of input features, GRU is employed. When describing the recursion for the GRU equation, the variables at the j -th level of the stack can be outlined as follows:

$$z_j^t = \sigma(x_j^t W_j^{xz} + h_j^t W_j^{hz} + b_j^z), \quad (8)$$

$$r_j^t = \sigma(x_j^t W_j^{xr} + h_j^t W_j^{hr} + b_j^r), \quad (9)$$

$$\tilde{h}_j^t = \tanh(x_j^t W_j^x + (r_j^t \odot r_j^{t-1}) W_j^h + b), \quad (10)$$

$$h_j^t = (1 - z_j^t) \odot h_j^{t-1} + z_j^t \odot \tilde{h}_j^t, \quad (11)$$

where $\sigma(\cdot)$ denotes the logistic sigmoid function x_j^t is the input feature at time step t . The reset and update gates at time step t are denoted as r_j^t and z_j^t , respectively, and the weights between the two units are represented by W . The hidden state at the previous time step and the current time step are represented by h_j^{t-1} and h_j^t , respectively.

To assess the significance of the processed features during network training, the attention mechanism [38]) is utilised to focus on specific segments of features, thereby enhancing the effectiveness of feature analysis. The resulting vector from the attention module is defined as follows:

$$A = \tanh(W_c[h_c : h_m]), \quad (12)$$

$$h_c = \sum_{s^t} \alpha^t h_{s^t}, \quad (13)$$

where W_c represents a weight matrix, h_c denotes the cumulative sum of all attention-weighted hidden states, h_m signifies the final hidden state of the encoder, h_{s^t} corresponds to the preceding hidden state of the encoder, and α^t denotes the vector of attention weights, which is defined as follows:

$$\alpha_t = \frac{\exp(\text{score}(h_m, \tilde{h}_s))}{\sum_{s^t=1}^{t_m} \exp(\text{score}(h_m, \tilde{h}_{s^t}))}, \quad (14)$$

$$\text{score}(h_m, h_{s^t}) = h_m^T W_p h_{s^t} \quad (15)$$

where t_m is the input sequence length at time t . h_m^T represents the transpose of the h_m vector, and W_p is a weight matrix that can be estimated through the training phase of the network.

In the tile of the network, the outputs of the attention modules are concatenated and then forwarded through the last attention module and an FC layer. The ultimate output, normalised to a range between zero and one using the Sigmoid function, represents the predicted probability of a pedestrian crossing.

IV. EXPERIMENTS

In this section, we conduct four distinct experiments to thoroughly assess the robustness of the proposed framework. Each experiment is designed to provide unique insights into various aspects of the model's performance. We compare our model against the PIE dataset in four different intervals ranging from 1 to 4 seconds, allowing us to scrutinise the model's predictive capabilities for an in-time response in different driving scenarios. Additionally, we evaluate the model's performance under varying temporal resolution and observation time. The second experiment examines the impact of the introduced Categorical Depth and Local Motion features in enhancing the framework's prediction accuracy. The third experiment evaluates the model generalisability and reliability on a diverse dataset from Waymo's self-driving vehicles. This investigation ensures that the framework performs effectively across different real-world scenarios. Lastly, we investigate the scalability of the model by assessing the framework's ability to handle one to three cameras simultaneously, expanding its view angles. This provides insights into the model's efficiency when processing information from multiple cameras. This multi-faceted evaluation allows us to gain a detailed understanding of the model's effectiveness in comparison to a singular experiment.

A. Datasets

The JAAD [15] and STIP [17] datasets suffer from no annotation in terms of ego-vehicle speed values. Furthermore, there is a slight bias in these datasets, as the majority of annotations indicate the cases of crossing which may not result in effective training of deep models. Therefore, the evaluations were conducted on the *Pedestrian Intention Estimation* (PIE) [16] dataset, which is extensively employed in the majority of prior studies. Additionally, we utilised our custom dataset named Urban-PIP, specifically annotated for pedestrian crossing behaviour, built upon the Waymo [20] dataset. Waymo is a widely used dataset for traffic perception by AVs thanks to the diversity of the video data from urban and rural environments under various driving conditions and situations. The specifications of datasets are briefly mentioned in Table I.

The evaluation metrics are based on [18], including *accuracy* (Acc), *precision*, and *recall rate*, which quantify the model's ability to accurately predict the binary classification task. Additionally, the *area under the ROC curve* (AUC), indicates the model's proficiency in distinguishing between different classes, and *F1 score*, represents the harmonic mean of precision and recall rate.

1) PIE dataset: The dataset is recorded on a sunny clear day for 6 hours in HD format (1920×1080). Each video segment lasts approximately 10 minutes, resulting in a total of 6 sets. We utilised approximately 50% (880 samples) of the dataset for training, 40% (719) for testing, and 10% (243) for validation as per the same split proportion as [12]. Regarding occlusion levels, partial occlusion is defined when an object is obstructed between 25% and 75%, while full occlusion occurs when the object is obstructed by 75% or more. The

TABLE I
DATASET SPECIFICATIONS

Specification	PIE	Urban-PIP
Autonomous Driving	No	Yes
Number of Cameras	1	3
Auxiliary Sensors	OBD	LiDAR, Radar, IMU
Lighting Conditions	Daylight	Daylight, Nighttime
Demographic Information	Age, Gender	None
Video clip lengths	10 min	16 sec
Total Number of Frames	909,000	32,790
Total Number of Annotated Frames	293,000	32,790
Total Number of Pedestrians	1,842	1,481
Crossed Pedestrians	519	409
Not Crossed Pedestrians	1,323	1,072

dataset includes the vehicle speed, heading direction, and GPS coordinates.

2) **Urban-PIP dataset:** The dataset captures diverse traffic conditions, ranging from busy intersections to quieter residential streets, with dynamic environmental changes. Pedestrian behaviour may be influenced by the actions of nearby vehicles or other pedestrians, making intention prediction more challenging. It is recorded under various weather conditions, including sunny, cloudy, rainy, and foggy, in three geographical locations using a multi-sensor setup. This multi-modal dataset is collected via a combination of LiDAR, cameras, radar, and IMU sensors mounted on the ego-vehicle. LiDAR provides 360° field of view with approximately a 300-meter range by beaming out millions of laser pulses per second and measuring the time of laser beam flight from the sensor to the surface of an object, then reflecting from the object to the sensor on the ego vehicle. The radar system has a continuous 360° view to track the presence and speed of road users in front, behind and sides of the vehicle. The multi-camera perspective (left, front, and right) offers a richer contextual understanding of pedestrian interactions compared to single-camera datasets like PIE and JAAD. The cameras simultaneously capture the traffic scene videos in HD format (1920 × 1080), The IMU module uses accelerometers and gyroscopes with input from GPS, maps, wheel speeds, as well as laser and radar measurements to provide position, velocity, and heading information to the vehicle.

The Urban-PIP dataset captures both daylight and nighttime scenarios, making it more diverse and challenging than datasets such as PIE, JAAD, and STIP, which are limited to daylight conditions. The inclusion of low-light environments introduces additional challenges, where the visibility of pedestrians is reduced, and sensor data may become noisy, leading to challenges in detecting and tracking pedestrians’ behaviour accurately.

In this study, the experiments are conducted using camera sensors as they provide rich visual information, including detailed information about pedestrian behaviour, their body language, and contextual information that can be crucial for predicting crossing intentions. Also, the affordability of camera sensors has made them a practical choice for current research on crossing prediction. We annotated 1,481 pedestrian crossing intentions from the front cameras including 448 in the front-left camera, 541 from the front camera, and 492 from the front-right camera.

3) **Frontal Urban-PIP dataset:** To assess models limited to a single camera, we introduce a subset of the Urban-PIP dataset, named *Frontal Urban-PIP* (FU-PIP), focusing on pedestrians observed only by the front camera. This subset, featuring 55 pedestrians with crossing intentions and 129 without, ensures a fair comparison will be conducted with similar methods which are limited to a single camera only.

B. Hardware and Implementation Settings

All proposed models were executed on a CUDA parallel computing platform with an Nvidia Quadro RTX A5000 GPU, 64GB of RAM and an Intel Core i9 13900K 24-core processor and the Torch environment.

PIP-Net was trained using the RMSProp optimiser. 256 hidden units were used for the GRUs, and the sigmoid (σ) activation function was applied to the GRUs for handling spatial kinematic data. To mitigate overfitting, a dropout rate of 0.5 was introduced after the attention block. Additionally, an L2 regularisation term of 0.0001 was incorporated into the last fully connected layer.

PIP-Net- α : This model is designed for a single camera setup. The model is trained on the PIE dataset and evaluated against the PIE test set and the FU-PIP datasets. The model doesn’t have the Camera Index pipeline within its architecture. Thereby, the outputs of the Max-pooling 3D (MP3D in Figure 3) are directly forwarded to the flattened block, and the Shift blocks are deactivated. A learning rate of 5×10^{-5} was used for 300 epochs with a batch size 10. The model was tested on the FU-PIP dataset to evaluate its generalisability.

PIP-Net- β : This model is designed for multi-camera setups to be evaluated against Urban-PIP with three cameras. The training of this model is performed using the Urban-PIP dataset with various observation times, a learning rate of 4×10^{-5} across 400 epochs, and a batch size of 6. The split ratio for training and testing samples is 80% (1,181 samples) and 20% (296 samples) of the dataset, respectively.

C. Temporal Resolution and Observation Time

Temporal resolution plays a crucial role in pedestrian crossing prediction, as it determines the frequency at which visual observations are sampled over time. To analyse its impact, we conducted experiments by varying the stride length while keeping the observation sequence length fixed at $m = 15$ (processing 15 frames), as followed the benchmark [18]. This approach allows us to expand the temporal window without increasing the model’s parameters, enabling a trade-off between capturing long-term dependencies and maintaining computational efficiency. We experimented with five different temporal resolutions over PIE videos by varying the stride from 1 (processing every frame) to 5 (processing every fifth frame), leading to observation times ranging from 0.5 seconds to 2.5 seconds. Figure 6 presents the performance of our proposed model, PIP-Net- α , across different temporal resolutions. The results indicate that increasing the stride initially improves performance, as observed with stride 1, which achieves the best results (Acc: 0.915, AUC: 0.897, F1: 0.846). However, beyond stride 2, the performance degrades,



Fig. 6. Performance of PIP-Net- α across different temporal resolutions: As the stride increases, the sampling frequency decreases, expanding the observation time from 0.5 to 2.5 seconds, influencing the model’s ability to predict pedestrian crossings.

suggesting that excessive temporal gaps between frames may lead to loss of fine-grained motion cues critical for precise crossing prediction. This aligns with other studies that suggest the intuition that the closer to crossing events, the more visual information clues will be revealed [29], [42], [44].

TABLE II
PERFORMANCE RESULTS OF PIP-Net- α , MEMORY USAGE, INFERENCE TIME, AND OBSERVATION TIME FOR DIFFERENT TEMPORAL STRIDES AND SEQUENCE LENGTHS ON PIE DATASET.

s	m	OT	Params	M	IT	Acc	AUC	F1
1	10	0.3s	5.5M	1.4GB	55ms	0.913	0.907	0.835
2	10	0.5s				0.921	0.935	0.882
1	15	0.5s	5.9M	1.8GB	58ms	0.915	0.897	0.846
2	15	1.0s				0.907	0.891	0.836
1	20	0.6s	6.6M	2.5GB	62ms	0.908	0.911	0.853
2	20	1.3s				0.896	0.874	0.810

s: Temporal Stride; m: Sequence Length (frame); OT: Observation Time (second); Params: Number of model parameters (million); M: Memory usage (gigabyte); IT: Inference Time (millisecond).

To further investigate the impact of temporal dependencies, we conducted experiments by varying the observation sequence length m while keeping the temporal stride fixed at 1 and 2. The objective was to determine whether increasing the number of frames to process improves prediction performance or introduces redundancy or irrelevant information. Table II presents the results, including the number of model parameters, memory usage, and inference time for different configurations. Our findings indicate that a shorter sequence length ($m = 10$) with a temporal stride of 2 ($s = 2$) achieves the highest prediction performance (AUC = 0.938, F1 = 0.882) while maintaining computational efficiency. This configuration ($s = 2, m = 10$) keeps the observation time to 0.5 seconds like ($s = 1, m = 15$) but with fewer frames processing. Furthermore, this setting requires fewer learnable parameters (5.5M) and lower memory usage (1.4GB), leading to a more efficient inference time of 55ms. These results suggest that increasing sequence length beyond a certain threshold may introduce redundant or irrelevant information about pedestrians’ crossing intentions.

D. Comparative Results

Table III highlights the performance of our method on the PIE dataset. The observation time (m) has been set to 15 frames, the same as the previous methods, to ensure a fair comparison with other methods. PIP-Net- α , achieves

TABLE III
PERFORMANCE COMPARISON ON THE PIE DATASET.

Model	Acc	AUC	F1	Precision	Recall
ATGC [15]	0.59	0.55	0.39	0.33	0.47
Multi-RNN [30]	0.83	0.80	0.71	0.69	0.73
SingleRNN [31]	0.81	0.75	0.64	0.67	0.61
SFRNN [29]	0.84	0.82	0.72	0.75	0.80
PCPA [18]	0.87	0.86	0.77	0.75	0.79
CAPformer [23]	0.88	0.80	0.71	0.69	0.74
PPCI [43]	0.89	0.86	0.80	0.79	0.81
GraphPlus [49]	0.89	0.90	0.81	0.83	0.79
MCIP* [11]	0.89	0.87	0.81	0.81	0.81
CIPF* [12]	0.91	0.89	0.84	0.85	0.83
PIT* [42]	0.91	0.92	0.82	0.84	0.81
VMIGI [44]	0.92	0.91	0.87	0.86	0.88
TrEP [39]	0.92	0.93	0.86	0.87	0.84
PIP-Net-α (Ours)	0.92	0.94	0.88	0.89	0.88

All models are evaluated at ETC = 0.5s. Results marked with * are reported from the original article due to unavailable source code.

the highest values in all metrics, showcasing its prowess in capturing the crossing intention classification.

In comparison to MultiRNN, SingleRNN, and SFRNN, which rely on CNN encoders for visual features and RNN-based encoder-decoder structures, PIP-Net- α achieves significant improvements by integrating a more diverse set of input features, capturing both pedestrian kinematics and environmental context more effectively. Unlike transformer-based (CAPformer) and graph-based (GraphPlus) architectures, which struggle with motion generalisation and real-time inference, PIP-Net- α leverages an RNN-based architecture to maintain temporal continuity and robust sequential reasoning, resulting in higher predictive performance. Furthermore, the proposed feature fusion approach, combining three kinematic and four contextual features, enables a more comprehensive scene understanding, improving AUC by +4% compared to CIPF, which integrates eight pedestrian-vehicle interaction features.

While recent models such as PIT, VMIGI, and TrEP have introduced advanced techniques for long-range dependencies, relational learning, and uncertainty estimation, their effectiveness remains constrained by the inherent limitations of their respective architectures. PIT, incorporating self-attention mechanisms and Temporal Fusion Blocks, enhances sequence modelling but requires heavy computational resources, limiting real-time feasibility. VMIGI, which applies a multimodal fusion strategy with GCN-based relational learning, strengthens interaction reasoning but struggles with fast motion scenarios due to graph structure limitations. TrEP, adopting a transformer-based evidential prediction framework, explicitly models uncertainty by correlating it with human annotator disagreement, yet its reliance on pure transformer encoding leads to weaker sequential motion modelling. In contrast, PIP-Net- α effectively balances multimodal integration, temporal consistency, and computational efficiency, making it the most well-rounded solution for pedestrian crossing prediction in dynamic urban environments, modelling uncertainty and correlating it with human annotator disagreement.

1) **Crossing Time Prediction:** Depending on the traffic scenario, the model’s prediction performance may vary. The model can predict the pedestrians’ estimated time to cross (ETC), 1 to 4 seconds in advance. For example, an ETC=2

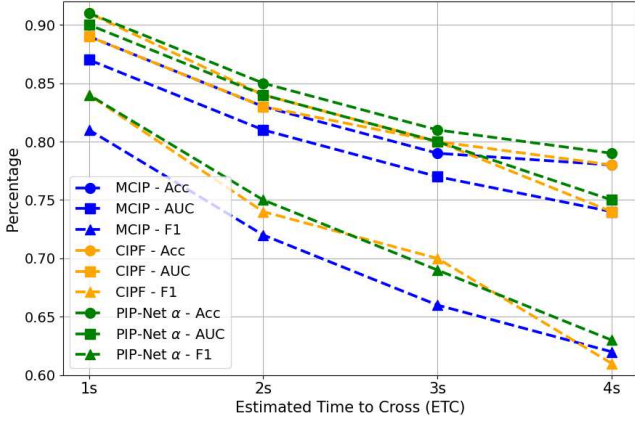


Fig. 7. Comparative analysis of prediction performance is conducted for future time points in the study. Models are evaluated in terms of their predictive accuracy from 1 second to 4 seconds ahead.

means the model expects or predicts the target pedestrian crosses in 2 seconds.

We evaluated the performance of the proposed PIP-Net- α model across various ETCs from 1 to 4 seconds. Figure 7 presents a comparison of the Acc, AUC, and F1 performance of PIP-Net- α with two recent prediction models, CIPF and MCIP, using the PIE dataset. All three models exhibited a decline in performance across all metrics as the Estimated Time to Cross (ETC) increased, i.e. when the models tried to have a longer-term prediction. Notably, the most significant drop in AUC occurred between the ETC = 1-second and ETC = 2, with MCIP and CIPF models decreasing by 6.8% and 6.7%, respectively. This is while our model shows a 6.6% decrease in the AUC from ETC = 1 to 2. The Accuracy also shows gradual decreases from the 3-second to the 4-second interval, with all models experiencing a decrease of approximately 1% (MCIP) and 2% (CIPF), respectively. As can be seen, the proposed model with green dashed lines consistently outperformed the other models for all ETCs.

2) **Ablation Study:** To investigate the inner workings of the proposed models, incorporating interpretability methods such as Grad-CAM [60] can enhance understanding by visualising which parts of the image influence predictions. However, in complex models that process heterogeneous input modalities and sequential data, these techniques may not always provide clear or actionable insights. Although VIMGI [44] successfully visualised the visual encoder, such methods often focus on individual components, missing feature interactions and kinematic data that affect predictions. Our approach, similar to studies such as CIPF [12], MCIP [11], TrEP [39] and CAPformer [23], emphasises performance-driven modelling, prioritising accuracy and efficiency. We have conducted comprehensive ablation studies to address interpretability concerns, providing valuable insights into our model’s decision-making processes.

The baseline model, α^0 , comprises the primary input features recognised in the task of pedestrian intention prediction, including *Bounding Box*, *Body Pose*, *Local Content*, *Semantic Context*, and *Vehicle Speed* [11]–[13], [23], [51]. Ablation experiments were conducted on the best-performing configuration, as presented in Table II, with $s = 2$ and $m = 10$. The results demonstrate that excluding the *Bounding Box* input leads

TABLE IV
PIP-Net- α VARIANTS USING DIFFERENT INPUT FEATURES.

Variant	GM	LM	MD	CD	Acc	AUC	F1	IT
α^0	-	-	-	-	0.893	0.907	0.830	25
α^1	✓	-	-	-	0.904↑	0.913↑	0.827↓	+10
α^2	-	✓	-	-	0.898↑	0.919↑	0.839↑	+3
α^3	-	-	✓	-	0.884↓	0.902↓	0.836↑	+12
α^4	-	-	-	✓	0.914↑	0.924↑	0.867↑	+27
α^5	✓	-	✓	-	0.883↓	0.903↓	0.827↓	+22
α	-	✓	-	✓	0.921↑	0.938↑	0.882↑	+30

Features: **GM** is Global Motion, **LM** is Local Motion, **MD** is ManyDepth, and **CD** is Categorical Depth, **IT**: Inference Time (milliseconds). The up arrow indicates an improvement, and the down arrow indicates a decrement in comparison with the baseline variant.

to an 8.6% decrease in accuracy compared with the baseline, whereas omitting the *Body Pose* parameter results in a lower accuracy reduction of 3.5%. This represents less importance of body pose compared to the bounding box which may seem counter-intuitive. However, our further investigations confirm that the bounding box data is notably more important and useful as it includes the pedestrian’s moving trajectory and tracking history over time, providing valuable spatio-temporal information. While body poses spatio-temporal information is not very important. The body pose only seems to matter when the pedestrian is about to cross the road in the last few frames before crossing. Otherwise, the body pose in previous moments, such as when the pedestrian is on the sidewalk, is redundant.

The impact of removing *Local Content* causes a decrease in accuracy by 1.7%. It appears to lack comprehensive features on pedestrians’ intentions, given the wide variety in appearance and accessories they may have. Excluding *Semantic Context* leads to a 3.4% accuracy decrease as it includes details about road layout such as sidewalk positioning and drivable zones.

Interestingly, removing the *Vehicle Speed* feature results in a 3.8% drop in accuracy, making it the second most important input. This aligns with the findings of the study by [23], which states that a model trained with the ego-vehicle speed tends to focus on the ego-vehicle speed adjustment (e.g. deceleration) to learn the pedestrian intention, rather than learning to predict the intention from the pedestrian behaviour.

Table IV shows a comparison of the two input features used in the proposed model, including scene motion and depth information, along with their impact on inference time. The α^2 variant demonstrates that the *local motion* feature can exhibit superior performance when compared to the global motion feature (2d), which includes optical flow analysis of the entire scene. While optical flow is typically sensitive to any movement between consecutive frames, it appears that local motion can provide a coarse-grained feature and treat more concisely to account for pedestrians’ velocity and direction of movement, regardless of other irrelevant objects in the surroundings. Additionally, local motion introduces only a modest increase in inference time (+3ms), making it an efficient feature choice.

Regarding depth information features, the categorical depth feature proposed stands out as the most effective standalone feature, as evidenced by the results of variant α^4 in Table IV,

highlighting the significance of pedestrians’ distance to the ego-vehicle and their interactions with other road users in the traffic scene. However, this improvement comes at the cost of a notable increase in inference time (+27ms). Conversely, α^3 performs the weakest among the variants when compared with the baseline, which utilises the global depth heatmap of the entire scene (2e). This underperformance may be attributed to the unstable estimation of depth for irrelevant surrounding objects. This is addressed in the categorical depth by focusing only on pedestrians and vehicles, and then applying per-instance normalisation (as shown in Figure 4).

Finally, the optimal outcome is attained by taking into account both the local motion and categorical depth map in α , leading to enhancements in the baseline regarding Acc, AUC, and F1 score by 3.14%, 3.1%, and 6.2%, respectively. This improvement, however, comes with the highest inference time (+30ms), emphasising the trade-off between model complexity and real-time applicability. The full model (α) achieves 55 milliseconds (≈ 18 FPS) inference time while maintaining high prediction performance, ensuring its suitability for real-time autonomous systems.

3) **Generalisation:** The PIP-Net- α evaluation is presented against the SOTA methodologies on the FU-PIP dataset in Table V. Notably, the models have never seen the scenarios in their training phase. The outcomes of other methods were generated using the pre-trained weights they provided. Overall, models exhibit limited generalisation capability on Urban-PIP, suggesting that our dataset is a more challenging benchmark compared to PIE. Given its environmental variety and inclusion of nighttime scenes, we believe our findings are largely transferable to PIE and JAAD, which mainly comprise daylight conditions and may not fully capture real-world complexities.

Most models show improvements over ATGC across various metrics, with each demonstrating unique strengths. However, we witnessed a decrement in terms of performance for PCPA, PPCI, CAPformer, and GraphPlus. As far as the research curiosity demand, we investigated the architecture of these models, and it turned out they are suffering from low-quality global context (see Figure 8) and body pose features. This is caused by the feature extractor algorithms, i.e., the semantic segmentation and pose estimator they have used, which hinder the classifier from judging based on precise features.

In our model, we enhanced semantic context representation by employing panoptic segmentation instead of conventional semantic segmentation. Unlike semantic segmentation, which classifies each pixel into a category for a single image, panoptic segmentation processes video sequences to generate temporally consistent object segments, distinguishing between object instances and semantic categories simultaneously. However, failure cases arise in challenging scenarios, particularly in low-light conditions or occluded scenes. Figure 8 illustrates qualitative comparisons between Slot-VPS (used in our model) and DeepLabV3 [61] (used in previous works). During daytime, Slot-VPS [58] provides sharper instance boundaries and more stable object tracking. However, at night, we observe increased misclassification and temporal inconsistencies, where pedestrians may be incorrectly segmented as part of the background or falsely merged with nearby objects. These

TABLE V
PERFORMANCE COMPARISON ON THE FRONTAL URBAN-PIP DATASET.

Model	Acc	AUC	F1	Precision	Recall
ATGC	0.52	0.51	0.35	0.32	0.44
Multi-RNN	0.64	0.63	0.49	0.51	0.48
SingleRNN	0.65	0.64	0.54	0.57	0.53
SFRNN	0.65	0.65	0.55	0.58	0.53
PCPA	0.62	0.60	0.50	0.52	0.47
PPCI	0.63	0.61	0.51	0.55	0.47
CAPformer	0.64	0.60	0.55	0.58	0.54
GraphPlus	0.64	0.61	0.57	0.59	0.56
VMIGI	0.65	0.61	0.57	0.58	0.55
TrEP	0.64	0.62	0.58	0.59	0.57
PIP-Net-α (Ours)	0.73	0.71	0.69	0.70	0.68

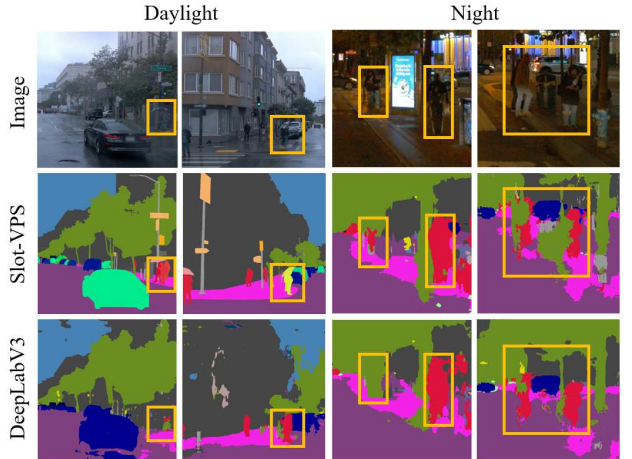


Fig. 8. Comparison of semantic and panoptic segmentation, where panoptic segmentation offers more precise scenes in daylight by distinguishing object instances and semantic categories, leading to improved global context representation.

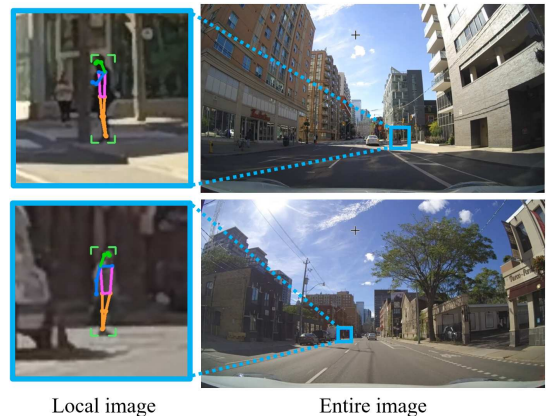


Fig. 9. Body pose estimation on local image rather than on the entire image enables more precise and reliable pose extraction.

segmentation errors propagate to downstream predictions, potentially leading to incorrect pedestrian intention estimates, especially when multiple pedestrians are present.

We observed more precise body pose information by applying pose estimation on the local image region containing the pedestrian rather than the entire scene image. This method ensures a more focused and accurate pose representation, reducing noise and ambiguity caused by irrelevant background elements. Figure 9 illustrates how our approach successfully retrieves clear and precise pose data in scenarios where applying pose estimation on the entire scene fails to obtain reliable pose information. Applying the pose estimation algorithm to

the local image region results in more robust and consistent pose features.

While these enhancements contribute to better generalisation, performance variations in challenging conditions emphasise the need for further research into feature extraction under adverse lighting and occlusions to ensure model reliability across all operational scenarios.

4) **View Angle Expansion:** We explored the enhancement of the field of view using three cameras to enable the autonomous vehicle to perceive a larger portion of its surroundings. For this purpose, we train PIP-Net- β with three different observation times (m) of 10, 20, and 30 frames. Subsequently, we examined how the prediction performance evolves as expand the ETC prediction horizon from 1 to 4 seconds. As depicted in Figure 10, the accuracy of crossing intention prediction decreases as the ETC prediction expands. However, the accuracy often increases with the enlargement of observation time. Intriguingly, when ETC = 4 and $m = 30$ the accuracy was lower compared to $m = 20$. This discrepancy arises from the model predicting pedestrians to be crossing based on long-term observations when, in reality, they did not cross. This highlights the observation that a pedestrian’s previous actions do not always accurately indicate their future intentions, as they can change their mind and act in an instant [3].

5) **Computational Cost of Multi-Camera Integration:** To assess the computational cost of integrating three cameras, we evaluated the inference time, memory usage, and parameter count of PIP-Net- β with stitched multi-camera inputs. Our results indicate that while computational demands increase, the model remains suitable for real-time execution. Specifically, the three-camera version required 9.9M parameters, 2.9GB of memory, and an inference time of 127ms (≈ 7.8 FPS). This represents an approximate 1.8 $\times$ increase in model size, 2.1 $\times$ in memory usage, and 2.3 $\times$ in inference time compared to the single-camera setup (PIP-Net- α). The feature integration module efficiently mitigates redundancy, ensuring that multi-camera integration remains feasible for real-world deployment.

E. Observational Results

We present the qualitative results of the PIP-Net- α in Figure 11 for FU-PIP datasets. The crossing probability is represented by a probability bar, where higher values (reddish colour) indicate a high probability of the pedestrian crossing. Pedestrians with no intention to cross are depicted with greenish bounding boxes and lower values on the probability bar.

Figure 12 presents the qualitative results of PIP-Net- α on 12 sample scenarios in the PIE and FU-PIP dataset. The model successfully predicted cases in rows (a) to (c), demonstrating robustness across varying daylight ambient lighting conditions. In row (c) of FU-PIP, the model accurately inferred the crossing intention of a pedestrian despite partial occlusion. However, in row (d) of PIE, the model struggled to make a correct prediction, when the pedestrian was nearly fully occluded.

In row (d) of FU-PIP, it seemed the model correctly predicted the intention of the pedestrian who stepped onto

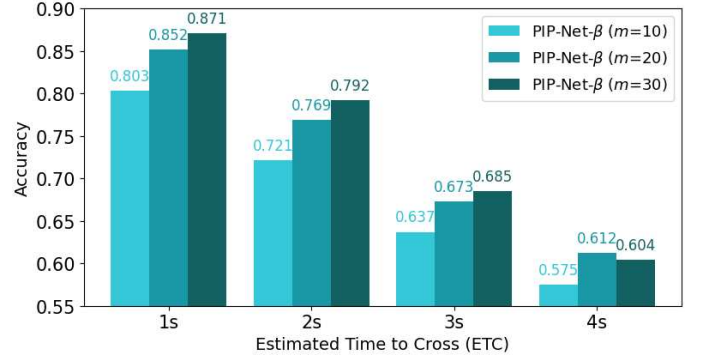
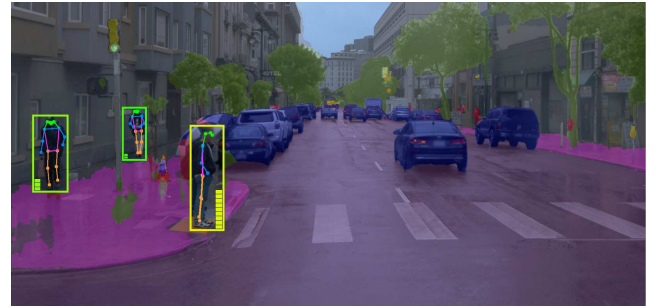


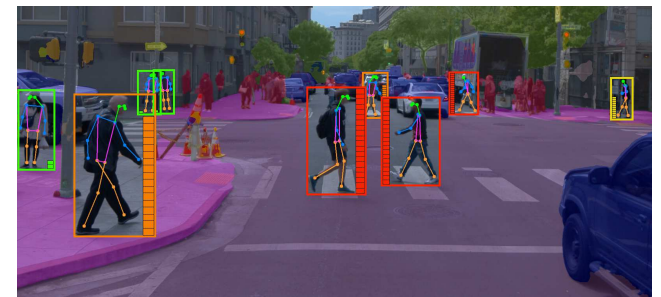
Fig. 10. Comparison of the crossing intention accuracy (Acc) for PIP-Net- β models (in three cameras mode) across various observation sequence lengths (m). The results are based on the Urban-PIP dataset.



(a) A non-signalised junction



(b) A signalised junction with a blue light



(c) A signalised junction with a red light

Fig. 11. A representation of the proposed model’s analysis on traffic video scenes of FU-PIP dataset. Three distinct images showcase intersections and traffic light conditions with probability bars positioned in front of pedestrians, aligning with their body pose direction. These bars visually represent the model’s predicted probability of each detected pedestrian crossing within the next 2 to 4 seconds into the future.

the road for crossing. However, the pedestrian subsequently abandoned the crossing and stepped back. In such rare cases, where the pedestrian ultimately does not cross, the ground truth is labelled as “not crossing”, consistent with the annotation criteria of the PIE and JAAD datasets.

In row (e) of PIE, the ego vehicle was accelerating, and

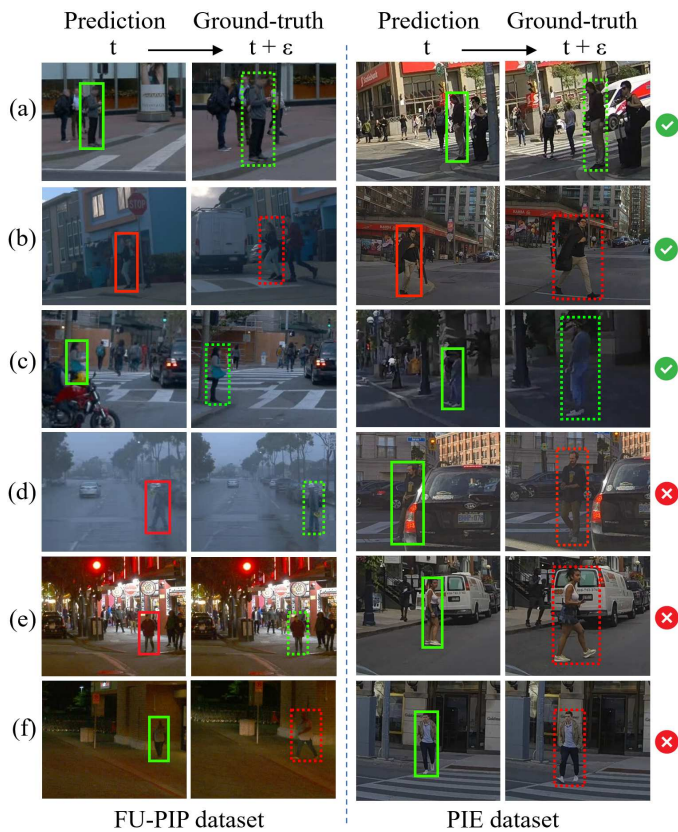


Fig. 12. Qualitative results on 12 sample scenarios. Rows with blue ticks represent successful predictions by the model and red crosses indicate wrong predictions. t demonstrates the prediction at the decisive moment and $(t + \epsilon)$ shows the ground truth (dashed bounding boxes). Red bounding boxes denote crossing, while blue denotes not crossing.

a hesitation in crossing was observed in the pedestrian’s behaviour. Interestingly, the model predicted a 43% probability of crossing, but this is below the 50% probability threshold to be accepted as a crossing event. However, we hypothesised that the acceleration of the ego vehicle might have a dominant effect over visual observations. To investigate this, we applied the permutation feature importance approach, as suggested in [51], and replaced the speed value with a scenario in which the ego vehicle was decelerating. Notably, the model then predicted an 88% probability of crossing intention for this case. This finding aligns with previous studies suggesting that, in certain scenarios, the model prediction heavily relies on ego-vehicle speed.

In row (f) of PIE, the ego vehicle was preparing to turn left while the pedestrian was inclined to cross the intersecting road. This suggests that incorporating additional motion-related information about the ego vehicle, such as the heading angle/steering angle, could help to address such cases.

We also evaluated the model’s performance in nighttime scenarios, as shown in rows (e) and (f) of FU-PIP. In these cases, the model failed to classify pedestrian crossing intentions accurately. This poor performance can be attributed to two primary factors: (1) the reduced accuracy of input features extracted from panoptic segmentation and depth estimation under low-light conditions, and (2) the limited number of nighttime samples in the training dataset, which likely led to model underfitting in such scenarios.

V. CONCLUSION AND FUTURE WORK

This paper presented a framework called PIP-Net for predicting pedestrian crossing intentions in real-world urban self-driving situations. Two variants, α and β , were introduced to support different camera setups. By utilising both kinematic data and spatial features of the driving scene, the proposed model employed a recurrent and temporal attention-based methodology to predict pedestrians’ future crossing intentions accurately. Through quantitative and qualitative experiments on the PIE dataset, the proposed model achieved state-of-the-art performance with a 94% AUC and an 88% F1 score.

The Urban-PIP was introduced as a new dataset for the pedestrian crossing prediction task, including various AV driving scenarios and comprehensive annotations on a multi-sensory setup, thereby enabling a better future investigation of crossing behaviour studies. Our model demonstrated a generalisation capability when applied to the Urban-PIP dataset by +9%, +10%, and 12% improvement compared with other models in terms of accuracy, AUC, and F1 score, respectively. This was underlined by the scene feature extractors employed in training our model.

To enhance the visual encoding of road users and their relative distances to the ego vehicle, we introduced a categorical depth feature map. This, combined with the sliced motion flow feature, provided salient information about the dynamics of the scene. Our results reveal that they cumulatively enhanced the accuracy and F1 score of the baseline model by +3.1% and +6.2%, respectively. Additionally, we investigated the impact of expanding the view angle using three cameras and enlarging prior observation frames.

Our algorithm achieved 85.4% accuracy in predicting pedestrian crossing intentions 2 seconds in advance and 79.3% accuracy for predictions between 2 and 4 seconds in advance. However, the algorithm is sensitive to the quality and precision of the input features, specifically, scene context and body pose information.

Overall, this study paved the way for developing multi-modal solutions for pedestrian crossing prediction and provided insights to effectively protect vulnerable road users by foreseeing the crossing behaviour of nearby pedestrians. Following the approach taken by seminal works in the field this study also focused on RGB cameras; however, we extended the use of a single camera to multiple cameras (left, front, and right) to evaluate the effectiveness of contextual information expansion. As one of the possible future work, we recognise the potential benefits of integrating data from other sensors, such as LiDAR and radar, to further enhance the model robustness in adverse weather or challenging lighting conditions. Specifically, LiDAR can offer more precise and reliable depth estimation, enhancing the Categorical Depth representation beyond what can be achieved with monocular-based depth estimation algorithms.

DECLARATION OF COMPETING INTEREST

The authors declare that they have no known competing financial interests or personal relationships that could have appeared to influence the work reported in this paper.

ACKNOWLEDGEMENTS

The authors would like to thank all partners within the Hi-Drive project for their cooperation and valuable contribution. This research has received funding from the European Union's Horizon 2020 research and innovation programme, under grant Agreement No 101006664. The article reflects only the author's view based on limited datasets and experiments, and neither the European Commission nor CINEA is responsible for any real-world use that may be made of the information this document contains.

REFERENCES

- [1] G. Yannis, D. Nikolaou, A. Laiou, Y. A. Stürmer, I. Buttler, and D. Jankowska-Karpa, "Vulnerable road users: Cross-cultural perspectives on performance and attitudes," *IATSS research*, vol. 44, no. 3, pp. 220–229, 2020.
- [2] N. Sharma, C. Dhiman, and S. Indu, "Pedestrian intention prediction for autonomous vehicles: A comprehensive survey," *Neurocomputing*, 2022.
- [3] A. Najmi, T. Waller, M. Memarpour, D. Nair, and T. H. Rashidi, "A human behaviour model and its implications in the transport context," *Transportation research interdisciplinary perspectives*, vol. 18, p. 100800, 2023.
- [4] Z. Zhou, Y. Liu, B. Liu, M. Ouyang, and R. Tang, "Pedestrian crossing intention prediction model considering social interaction between multi-pedestrians and multi-vehicles," *Transportation Research Record*, p. 03611981231187643, 2023.
- [5] J. Fang, F. Wang, J. Xue, and T.-S. Chua, "Behavioral intention prediction in driving scenes: A survey," *IEEE Transactions on Intelligent Transportation Systems*, 2024.
- [6] A. H. Kalantari, Y. Yang, J. G. de Pedro, Y. M. Lee, A. Horrobin, A. Solernou, C. Holmes, N. Merat, and G. Markkula, "Who goes first? a distributed simulator study of vehicle–pedestrian interaction," *Accident Analysis & Prevention*, vol. 186, p. 107050, 2023.
- [7] B. Yang, W. Zhan, P. Wang, C. Chan, Y. Cai, and N. Wang, "Crossing or not? context-based recognition of pedestrian crossing intention in the urban environment," *IEEE Transactions on Intelligent Transportation Systems*, vol. 23, no. 6, pp. 5338–5349, 2021.
- [8] K. Muhammad, T. Hussain, H. Ullah, J. Del Ser, M. Rezaei, N. Kumar, M. Hijji, P. Bellavista, and V. H. C. de Albuquerque, "Vision-based semantic segmentation in scene understanding for autonomous driving: Recent achievements, challenges, and outlooks," *IEEE Transactions on Intelligent Transportation Systems*, 2022.
- [9] L. Chen, S. Lin, X. Lu, D. Cao, H. Wu, C. Guo, C. Liu, and F.-Y. Wang, "Deep neural network based vehicle and pedestrian detection for autonomous driving: A survey," *IEEE Transactions on Intelligent Transportation Systems*, vol. 22, no. 6, pp. 3234–3246, 2021.
- [10] M. Gulzar, Y. Muhammad, and N. Muhammad, "A survey on motion prediction of pedestrians and vehicles for autonomous driving," *IEEE Access*, vol. 9, pp. 137 957–137 969, 2021.
- [11] J.-S. Ham, K. Bae, and J. Moon, "MCIP: Multi-stream network for pedestrian crossing intention prediction," in *Computer Vision–ECCV 2022 Workshops: Tel Aviv, Israel, October 23–27, 2022, Proceedings, Part I*. Springer, 2022, pp. 663–679.
- [12] J.-S. Ham, D. H. Kim, N. Jung, and J. Moon, "CIPF: Crossing intention prediction network based on feature fusion modules for improving pedestrian safety," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2023, pp. 3665–3674.
- [13] R. Ni, B. Yang, Z. Wei, H. Hu, and C. Yang, "Pedestrians crossing intention anticipation based on dual-channel action recognition and hierarchical environmental context," *IET Intelligent Transport Systems*, vol. 17, no. 2, pp. 255–269, 2023.
- [14] T. Zhang, X. Chen, Y. Wang, Y. Wang, and H. Zhao, "Mutr3d: A multi-camera tracking framework via 3d-to-2d queries," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2022, pp. 4537–4546.
- [15] A. Rasouli, I. Kotseruba, and J. K. Tsotsos, "Are they going to cross? a benchmark dataset and baseline for pedestrian crosswalk behavior," in *Proceedings of the IEEE International Conference on Computer Vision Workshops*, 2017, pp. 206–213.
- [16] A. Rasouli, I. Kotseruba, T. Kunic, and J. K. Tsotsos, "PIE: A large-scale dataset and models for pedestrian intention estimation and trajectory prediction," in *Proceedings of the IEEE/CVF International Conference on Computer Vision*, 2019, pp. 6262–6271.
- [17] B. Liu, E. Adeli, Z. Cao, K.-H. Lee, A. Shenoi, A. Gaidon, and J. C. Nibbles, "Spatiotemporal relationship reasoning for pedestrian intent prediction," *IEEE Robotics and Automation Letters*, vol. 5, no. 2, pp. 3485–3492, 2020.
- [18] I. Kotseruba, A. Rasouli, and J. K. Tsotsos, "Benchmark for evaluating pedestrian action prediction," in *Proceedings of the IEEE/CVF Winter Conference on Applications of Computer Vision*, 2021, pp. 1258–1268.
- [19] F. Schneemann and P. Heinemann, "Context-based detection of pedestrian crossing intention for autonomous driving in urban environments," in *2016 IEEE/RSJ International Conference on Intelligent Robots and Systems (IROS)*. IEEE, 2016, pp. 2243–2248.
- [20] J. Mei, A. Z. Zhu, X. Yan, H. Yan, S. Qiao, L.-C. Chen, and H. Kretschmar, "Waymo open dataset: Panoramic video panoptic segmentation," in *European Conference on Computer Vision*. Springer, 2022, pp. 53–72.
- [21] C. Zhang and C. Berger, "Pedestrian behavior prediction using deep learning methods for urban scenarios: A review," *IEEE Transactions on Intelligent Transportation Systems*, vol. 24, no. 10, pp. 10 279–10 301, 2023.
- [22] H. Razali, T. Mordan, and A. Alahi, "Pedestrian intention prediction: A convolutional bottom-up multi-task approach," *Transportation research part C: emerging technologies*, vol. 130, p. 103259, 2021.
- [23] J. Lorenzo, I. Parra, R. Izquierdo, A. L. Ballardini, Á. Hernández-Saz, D. F. Llorca, and M. Á. Sotelo, "CAPformer: Pedestrian crossing action prediction using transformer," *Sensors (Basel, Switzerland)*, vol. 21, 2021.
- [24] J. Lorenzo, I. Parra, F. Wirth, C. Stiller, D. F. Llorca, and M. A. Sotelo, "RNN-based pedestrian crossing prediction using activity and pose-related features," in *2020 IEEE Intelligent Vehicles Symposium (IV)*. IEEE, 2020, pp. 1801–1806.
- [25] N. Sharma, C. Dhiman, and S. Indu, "Intelligent pedestrian intention prediction framework," in *2022 IEEE International Conference on Service Operations and Logistics, and Informatics (SOLI)*. IEEE, 2022, pp. 1–5.
- [26] K. M. Abughalieh and S. G. Alawneh, "Predicting pedestrian intention to cross the road," *IEEE Access*, vol. 8, pp. 72 558–72 569, 2020.
- [27] A. Alofi, R. Greer, A. Gopalkrishnan, and M. Trivedi, "Pedestrian safety by intent prediction: A lightweight lstm-attention architecture and experimental evaluations with real-world datasets," in *2024 IEEE Intelligent Vehicles Symposium (IV)*. IEEE, 2024, pp. 77–84.
- [28] K. Cho, B. Van Merriënboer, C. Gulcehre, D. Bahdanau, F. Bougares, H. Schwenk, and Y. Bengio, "Learning phrase representations using rnn encoder-decoder for statistical machine translation," in *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, 2014, pp. 1724–1734.
- [29] A. Rasouli, I. Kotseruba, and J. K. Tsotsos, "Pedestrian action anticipation using contextual feature fusion in stacked rnns," *ArXiv*, vol. abs/2005.06582, 2020.
- [30] A. Bhattacharyya, M. Fritz, and B. Schiele, "Long-term on-board prediction of people in traffic scenes under uncertainty," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2018, pp. 4194–4202.
- [31] I. Kotseruba, A. Rasouli, and J. K. Tsotsos, "Do they want to cross? understanding pedestrian intention for behavior prediction," in *2020 IEEE Intelligent Vehicles Symposium (IV)*. IEEE, 2020, pp. 1688–1693.
- [32] D. Tran, L. Bourdev, R. Fergus, L. Torresani, and M. Paluri, "Learning spatiotemporal features with 3d convolutional networks," in *Proceedings of the IEEE international conference on computer vision*, 2015, pp. 4489–4497.
- [33] A. Vaswani, "Attention is all you need," *Advances in Neural Information Processing Systems*, 2017.
- [34] K. Saleh, M. Hossny, and S. Nahavandi, "Real-time intent prediction of pedestrians for autonomous ground vehicles via spatio-temporal densenet," in *2019 International Conference on Robotics and Automation (ICRA)*, 2019, pp. 9704–9710.
- [35] —, "Spatio-temporal densenet for real-time intent prediction of pedestrians in urban traffic environments," *Neurocomputing*, vol. 386, pp. 317–324, 2020.
- [36] A. Singh and U. Suddamalla, "Multi-input fusion for practical pedestrian intention prediction," in *Proceedings of the IEEE/CVF International Conference on Computer Vision*, 2021, pp. 2304–2311.
- [37] Y. Yan, M. Zhou, C.-c. Feng, L. Lv, and H. Ding, "Three-dimensional cnn-based model for fine-grained pedestrian crossing behavior recognition in automated vehicles," *Journal of Transportation Engineering, Part A: Systems*, vol. 151, no. 2, p. 04024116, 2025.

- [38] H. Zhao, J. Jia, and V. Koltun, "Exploring self-attention for image recognition," in *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, 2020, pp. 10 076–10 085.
- [39] Z. Zhang, R. Tian, and Z. Ding, "Trep: Transformer-based evidential prediction for pedestrian intention with uncertainty," in *Proceedings of the AAAI Conference on Artificial Intelligence*, vol. 37, no. 3, 2023, pp. 3534–3542.
- [40] N. Osman, G. Camporese, and L. Ballan, "Tamformer: Multi-modal transformer with learned attention mask for early intent prediction," in *ICASSP 2023-2023 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP)*. IEEE, 2023, pp. 1–5.
- [41] L. Achaji, J. Moreau, T. Fouqueray, F. Aioun, and F. Charpillet, "Is attention to bounding boxes all you need for pedestrian action prediction?" in *2022 IEEE Intelligent Vehicles Symposium (IV)*. IEEE, 2022, pp. 895–902.
- [42] Y. Zhou, G. Tan, R. Zhong, Y. Li, and C. Gou, "Pit: Progressive interaction transformer for pedestrian crossing intention prediction," *IEEE Transactions on Intelligent Transportation Systems*, 2023.
- [43] D. Yang, H. Zhang, E. Yurtsever, K. A. Redmill, and Ü. Özgüner, "Predicting pedestrian crossing intention with feature fusion and spatio-temporal attention," *IEEE Transactions on Intelligent Vehicles*, vol. 7, no. 2, pp. 221–230, 2022.
- [44] N. Sharma, C. Dhiman, and S. Indu, "Visual-motion-interaction guided pedestrian intention prediction framework," *IEEE Sensors Journal*, 2023.
- [45] M. Azarmi, M. Rezaei, T. Hussain, and C. Qian, "Local and global contextual features fusion for pedestrian intention prediction," in *Artificial Intelligence and Smart Vehicles*, 2023, pp. 1–13.
- [46] J. Gesnouin, S. Pechberti, B. Stanciulescu, and F. Moutarde, "Trous-pinet: Spatio-temporal attention on parallel atrous convolutions and u-grus for skeletal pedestrian crossing prediction," in *2021 16th IEEE International Conference on Automatic Face and Gesture Recognition (FG 2021)*. IEEE, 2021, pp. 01–07.
- [47] Z. Fang and A. M. López, "Is the pedestrian going to cross? answering by 2d pose estimation," in *2018 IEEE intelligent vehicles symposium (IV)*. IEEE, 2018, pp. 1271–1276.
- [48] Z. Fang and A. M. López, "Intention recognition of pedestrians and cyclists by 2d pose estimation," *IEEE Transactions on Intelligent Transportation Systems*, vol. 21, no. 11, pp. 4773–4783, 2020.
- [49] P. R. G. Cadena, Y. Qian, C. Wang, and M. Yang, "Pedestrian graph+: A fast pedestrian crossing prediction model based on graph convolutional networks," *IEEE Transactions on Intelligent Transportation Systems*, vol. 23, no. 11, pp. 21 050–21 061, 2022.
- [50] L. Xu, S. You, G. He, and Y. Li, "Pedestrian-vehicle information modulation for pedestrian crossing intention prediction," *IEEE Transactions on Intelligent Vehicles*, 2024.
- [51] M. Azarmi, M. Rezaei, H. Wang, and A. Arabian, "Feature importance in pedestrian intention prediction: A context-aware review," *arXiv preprint arXiv:2409.07645*, 2024.
- [52] S. Neogi, M. Hoy, K. Dang, H. Yu, and J. Dauwels, "Context model for pedestrian intention prediction using factored latent-dynamic conditional random fields," *IEEE transactions on intelligent transportation systems*, vol. 22, no. 11, pp. 6821–6832, 2020.
- [53] D. Zhang, F. Shi, Y. Meng, Y. Xu, X. Xiao, and W. Li, "Pedestrian intention prediction via depth augmented scene restoration," in *2021 5th CAA International Conference on Vehicular Control and Intelligence (CVCI)*. IEEE, 2021, pp. 1–6.
- [54] J. Gesnouin, S. Pechberti, B. Stanciulescu, and F. Moutarde, "Assessing cross-dataset generalization of pedestrian crossing predictors," in *2022 IEEE Intelligent Vehicles Symposium (IV)*. IEEE, 2022, pp. 419–426.
- [55] M. Rezaei, M. Azarmi, and F. M. P. Mir, "3D-Net: Monocular 3D object recognition for traffic monitoring," *Expert Systems with Applications*, vol. 227, p. 120253, 2023.
- [56] D. Maji, S. Nagori, M. Mathew, and D. Poddar, "YOLO-Pose: Enhancing YOLO for multi person pose estimation using object keypoint similarity loss," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2022, pp. 2637–2646.
- [57] E. Ilg, N. Mayer, T. Saikia, M. Keuper, A. Dosovitskiy, and T. Brox, "FlowNet 2.0: Evolution of optical flow estimation with deep networks," in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2017, pp. 2462–2470.
- [58] Y. Zhou, H. Zhang, H. Lee, S. Sun, P. Li, Y. Zhu, B. Yoo, X. Qi, and J.-J. Han, "Slot-vps: Object-centric representation learning for video panoptic segmentation," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2022, pp. 3093–3103.
- [59] J. Watson, O. Mac Aodha, V. Prisacariu, G. Brostow, and M. Firman, "The temporal opportunist: Self-supervised multi-frame monocular depth," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2021, pp. 1164–1174.
- [60] R. R. Selvaraju, M. Cogswell, A. Das, R. Vedantam, D. Parikh, and D. Batra, "Grad-cam: Visual explanations from deep networks via gradient-based localization," in *Proceedings of the IEEE international conference on computer vision*, 2017, pp. 618–626.
- [61] L.-C. Chen, Y. Zhu, G. Papandreou, F. Schroff, and H. Adam, "Encoder-decoder with atrous separable convolution for semantic image segmentation," in *Proceedings of the European conference on computer vision (ECCV)*, 2018, pp. 801–818.

VI. BIOGRAPHY

Mohsen Azarmi is a Ph.D. student at the Institute for Transport Studies, University of Leeds, UK. He holds a master's degree in Artificial Intelligence and Robotics. His research focuses on the intersection of computer vision, deep neural networks, and multi-sensor data fusion, with particular interests in pedestrian behaviour understanding, intelligent transportation systems, traffic safety, 3D scene modelling, and multimodal large language models.



Mahdi Rezaei is an Associate Professor in Computer Vision and Machine Learning and the team leader of the Computer Vision Research Group at the University of Leeds, Institute for Transport Studies. With a background in Artificial Intelligence and 18 years of research experience in academia and industry, he is a PI and lead CO-I in various EU and UK flagship projects, such as EPSRC DTP, Hi-Drive, MAVIS, and L3Pilot, with a focus on Driver Monitoring Systems, AV-Pedestrian interactions, XAI, and Intelligent Transportation Systems.



He Wang is an Associate Professor at the Department of Computer Science, University College London (UCL) and a Visiting Professor at the University of Leeds. He serves as an Associate Editor of Computer Graphics Forum, an Academic Advisor at the Commonwealth Scholarship Council, and a former Fellow of the Alan Turing Institute. His current research interest is mainly in computer graphics, computer vision and machine learning.

