

Robust Multi-Modality Multi-Object Tracking

Wenwei Zhang¹, Hui Zhou², Shuyang Sun³, Zhe Wang², Jianping Shi², Chen Change Loy¹

¹Nanyang Technological University, ²SenseTime Research, ³University of Oxford

{wenwei001, ccloy}@ntu.edu.sg, {zhouhui, wangzhe, shijianping}@sensetime.com,
 shuyang.sun@eng.ox.ac.uk

Abstract

Multi-sensor perception is crucial to ensure the reliability and accuracy in autonomous driving system, while multi-object tracking (MOT) improves that by tracing sequential movement of dynamic objects. Most current approaches for multi-sensor multi-object tracking are either lack of reliability by tightly relying on a single input source (e.g., center camera), or not accurate enough by fusing the results from multiple sensors in post processing without fully exploiting the inherent information. In this study, we design a generic sensor-agnostic multi-modality MOT framework (mmMOT), where each modality (i.e., sensors) is capable of performing its role independently to preserve reliability, and could further improving its accuracy through a novel multi-modality fusion module. Our mmMOT can be trained in an end-to-end manner, enables joint optimization for the base feature extractor of each modality and an adjacency estimator for cross modality. Our mmMOT also makes the first attempt to encode deep representation of point cloud in data association process in MOT. We conduct extensive experiments to evaluate the effectiveness of the proposed framework on the challenging KITTI benchmark and report state-of-the-art performance. Code and models are available at <https://github.com/ZwwWayne/mmMOT>.

1. Introduction

Reliability and accuracy are the two fundamental requirements for autonomous driving system. Dynamic object perception is vital for autonomous driving. To improve its reliability, multi-modality sensors can be employed to provide loosely coupled independent clues to prevent failure showed in Figure 1 (a). To improve accuracy, sequential information from multiple object tracking can be incorporated, and better multi-sensor information can reinforce the final score as in Figure 1 (b). In this paper, we propose the *multi-modality Multi-Object Tracking* (mmMOT) framework, which preserves reliability by a novel fusion

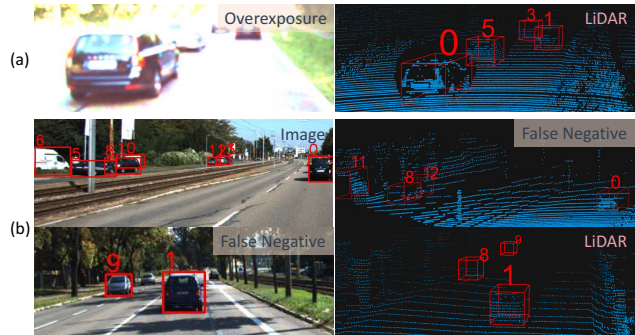


Figure 1. Figure (a) For reliability: camera is disabled when overexposure or crashed in transmission. Figure (b) For accuracy: multi-sensor information could reinforce the perception ability. The image is cropped and best viewed in color and zoomed in.

module for multiple sensors and improves accuracy with attention guided multi-modality fusion mechanism.

It is non-trivial for traditional methods to design a multi-modality (i.e., multi-sensor) MOT framework and preserve both reliability and accuracy. A majority of traditional methods [1, 9, 12, 25] use camera, LiDAR or radar with hand-crafted features fused by Kalman Filter or Bayesian framework. Their accuracy is bounded by the expression ability of hand-crafted features. Another stream of methods uses deep feature extractors [11], which significantly improve the accuracy. Nevertheless, they focus on image level deep representation to associate object trajectories and use LiDAR only in detection stage. Such a binding method cannot preserve reliability if the camera is down.

In this work, we design a multi-modality MOT (mmMOT) framework that is extendable to camera, LiDAR and radar. Firstly, it obeys a loose coupling scheme to allow high reliability during the extraction and fusion of multi-sensor information. Specifically, multi-modality features are extracted from each sensor independently, then a fusion module is applied to fuse these features, and pass them to an adjacency estimator, which is capable of performing inference based on each modality. Second, to enable the network to learn to infer from different modalities simultaneously, our mmMOT is trained in an end-to-end manner, so that

the multi-modality feature extractor and cross-modality adjacency estimator are jointly optimized. Last but not least, we make the first attempt of using deep representation of point cloud in the data association process for MOT and achieve competitive results.

We conduct extensive experiments on the fusion module and evaluate our framework on the KITTI tracking dataset [13]. Without bells and whistles, we achieve state-of-the-art results on KITTI tracking benchmark [13] under the online setting, purely relying on image and point cloud, and our results with single modality (under sensor failure condition) by the same model are also competitive (only 0.28% worse).

To summarize, our contributions are as follows:

1. We propose a multi-modality MOT framework with a robust fusion module that exploits multi-modality information to improve both reliability and accuracy.
2. We propose a novel end-to-end training method that enables joint optimization of cross-modality inference.
3. We make the first attempt to apply deep features of point cloud for tracking and obtain competitive results.

2. Related Work

Multi-Object Tracking Framework. Recent research of MOT primarily follows the tracking-by-detection paradigm [6, 11, 38, 50], where object of interests is first obtained by an object detector and then linked into trajectories via data association. The data association problem could be tackled from various perspectives, e.g., min-cost flow [11, 20, 37], Markov decision processes (MDP) [48], partial filtering [6], Hungarian assignment [38] and graph cut [44, 49]. However, most of these methods are not trained in an end-to-end manner thus many parameters are heuristic (e.g., weights of costs) and susceptible to local optima.

To achieve end-to-end learning within the min-cost flow framework, Schuster et al. [37] applies bi-level optimization by smoothing the linear programming and Deep Structured Model (DSM) [11] exploits the hinge loss. Their frameworks, however, are not designed for cross-modality. We solve this problem by adjacency matrix learning.

Apart from different data association paradigms, correlation features have also been explored widely to determine the relation of detections. Current image-centric methods [11, 35, 38, 50] mainly use deep features of image patches. Hand-crafted features are occasionally used as auxiliary inputs, including but not limited to bounding box [15], geometric information [27], shape information [38] and temporal information [45]. 3D information is also beneficial and thus exploited by prediction from 3D detection [11] or estimation from RGB image with either neural networks [36] or geometric prior [38]. Osep et al. [25] fuses the information from RGB images, stereo, visual odometry, and optionally scene flow, but it cannot be trained in an end-to-end

manner. All the aforementioned methods must work with camera thus lack of reliability. By contrast, our mmMOT extracts feature from each sensor independently (both deep image features and deep representation of point cloud), and each sensor plays an equally important role and they can be decoupled. The proposed attention guided fusion mechanism further improves accuracy.

Deep Representation of Point Cloud. A traditional usage of point cloud for tracking is to measure distances [31], provide 2.5D grid representation [2, 10] or to derive some hand-crafted features [42]. None of them fully exploit the inherent information of the point cloud for the data association problem. Recent studies [3, 7, 24] have demonstrated the value of using 3D point cloud as perception features in autonomous driving. To learn a good deep representation for point cloud, PointNet [29] and PointNet++ [30] process raw unstructured point clouds using symmetric functions. We adopt this effective method in our framework. Other studies such as PointSIFT [17] proposes an orientation-encoding unit to learn SIFT-like features of point cloud, and 3DSmoothNet [14] learns a voxelized smoothed density value representation. There are also methods [46, 47] which project the point cloud to a sphere thus 2D CNN can be applied for the segmentation task.

Object Detection. An object detector is also a vital component in the tracking by detection paradigm. Deep learning approaches for 2D object detection have improved drastically [23, 32, 43] since Faster R-CNN [33]. 3D object detection receives increasing attention recently. To exploit both image and point cloud, some methods [8, 18] aggregate point cloud and image features from different views, while F-PointNet [28] obtains frustum proposal from an image, and then applies PointNet [29] for 3D object localization with the point cloud. There exist state-of-the-art methods [19, 39, 51] that use point cloud only. One-stage detectors [19, 51] usually apply CNN on the voxelized representation, and two-stage detectors such as Point RCNN [39] generates proposals first by segmentation, which are refined in the second stage. Our mmMOT is readily adaptable for both 2D and 3D object detectors.

3. Multi-Modality Multi-Object Tracking

We propose a multi-modality MOT (mmMOT) framework, which preserves reliability via independent multi-sensor feature extraction and improves accuracy via modality fusion. It generally follows the widely adopted tracking-by-detection paradigm from the min-cost flow perspective. Specifically, our framework contains four modules including an object detector, feature extractor, adjacency estimator and min-cost flow optimizer, as shown in Fig. 2 (a), (b), (c), (d), respectively. First, an arbitrary object detector is used to localize objects of interests. We use PointPillar [19] for convenience. Second, the feature extractor extracts

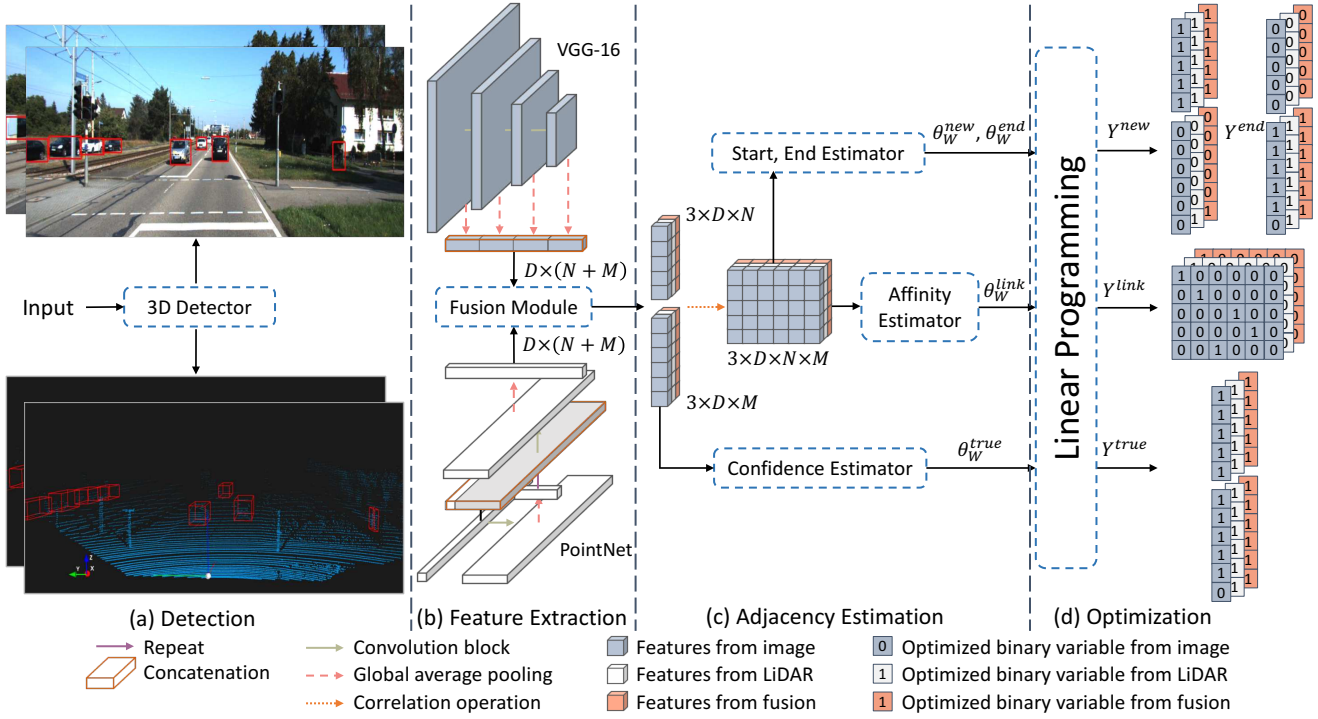


Figure 2. The pipeline of mmMOT. The feature extractors first extract features from image and LiDAR, and the robust fusion module fuses the multi-sensor features. Next, the correlation operator produces the correlation features for each detection pair, by which the adjacency estimator predicts the adjacency matrix. All the predicted scores are optimized to predict the binary variable Y .

features from each sensor independently for each detection (Section 3.2), after which a fusion module is applied to fuse and pass the single modality feature to the adjacency estimator (Section 3.3). The adjacency estimator is modality agnostic. It infers the scores necessary for the min-cost flow graph computation. The structure of the adjacency estimator and the associated end-to-end learning method will be demonstrated in Section 3.4. The min-cost flow optimizer is a linear programming solver that finds the optimal solution based on the predicted scores (Section 3.5).

3.1. Problem Formulation

Our mmMOT follows the tracking-by-detection paradigm to define the data association cost, which is solved as a min-cost flow problem [11, 20, 37]. Take the online MOT setting for example, assume there are N and M detections in two consecutive frames i and $i + 1$, denoted by $X^i = \{x_j^i | j = 1, \dots, N\}$ and $X^{i+1} = \{x_k^{i+1} | k = 1, \dots, M\}$, respectively. Each detection is associated to four types of binary variables in this paradigm. We introduce them following the notation of Deep Structured Model (DSM) [11]. First, for any x_j , a binary variable y_j^{true} indicates whether the detection is a true positive. Second, binary variables y_{jk}^{link} indicates if the j -th detection in the first frame and k -th detection in the second frame belong to the same trajectory, and

all these y_{jk}^{link} form an adjacency matrix $A^i \in R^{N \times M}$, where $A_{jk}^i = y_{jk}^{link}$. The other two variables y_j^{new} , y_j^{end} represents whether the detection is the start or the end of a trajectory, respectively. For convenience, we flatten the adjacency matrix into a vector Y^{link} , and gather all the binary variables having the same type as Y^{true} , Y^{new} , Y^{end} , then all these variables are collapsed into a vector $\mathbf{Y} = [Y^{true}, Y^{link}, Y^{new}, Y^{end}]$, which comprises all states of edges in the network flow. For each binary variable in Y^{true} , Y^{link} , Y^{new} , Y^{end} , the corresponding scores are predicted by the confidence estimator, affinity estimator, start and end estimator, respectively. These estimators form the adjacency estimator, and we solve them in a multi-task learning network as shown in Figure 2.

3.2. Single Modality Feature Extractor

In an online setting, only detections in two consecutive frames are involved. To estimate their adjacency, their deep representations are first extracted from the respective image or point cloud. The features of each single modality form a tensor with a size of $1 \times D \times (N + M)$, where $D = 512$ is the vector length, and $N + M$ is the total number of detections in the two frames.

Image Feature Extractor. Upon obtaining 2D bounding boxes from either a 2D or 3D detector, the image patches associated to each detection are cropped and resized to a square with a side length of 224 pixels to form a batch. All

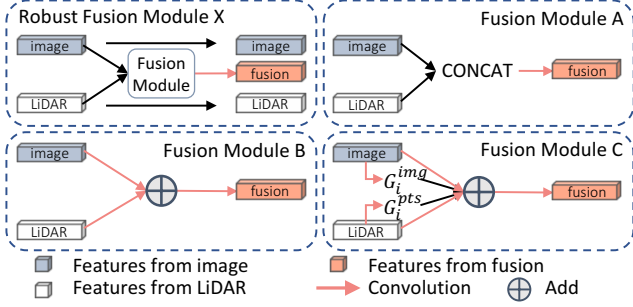


Figure 3. The robust fusion module and three multi-modality fusion modules. The robust fusion module can apply any one of the fusion modules A, B and C to produce the fused modality. Unlike the conventional fusion modules, the robust fusion module produces both the single modalities and the fused modality as an output. Fusion module A concatenates the multi-modality features, module B fuses them with a linear combination, module C introduces attention mechanism to weights the importance of sensor’s feature adaptively.

these patches form a 4D tensor with a size of $(N + M) \times 3 \times 224 \times 224$. We use VGG-Net [40] as the image feature extractor’s backbone. To exploit features at different level, we modify the skip-pooling [4] so as to pass different level’s feature to the top, as shown in the VGG-Net depicted in Figure 2. The details of skip-pooling are provided in the supplementary material.

Point Cloud Feature Extractor. One of our contributions is to apply deep representation of point cloud in data association process for MOT. While the LiDAR point cloud associated to a single detection could be easily obtained by a 3D bounding box, it remains a challenge if only a 2D bounding box is provided. It is possible to obtain a 3D bounding box using F-PointNet [28], or directly estimated the 3D bounding box with other geometric information and priori [36, 38]. In this study, we choose not to locate the detection in 3D space because we observed more errors. Rather, inspired by F-PointNet [28], we exploit all the point clouds in the frustum projected by the 2D bounding box. This leads to high flexibility and reliability, and save computation from obtaining 3D bounding box.

The point cloud forms a tensor with a size $1 \times D \times L$, where L is the total number of all the points in all bounding boxes, and $D = 3$ is the dimension of the point cloud information. We empirically found the reflectivity of point cloud provides only with marginal improvement, thus we only used the coordinates in 3D space. We modify the vanilla PointNet [30] to extract features from point cloud for each detection as shown in the PointNet depicted in Figure 2. To enhance the global information of points in each bounding box, we employ the global feature branch originally designed for the segmentation task in PointNet [30], and we found that average pooling works better than max pooling in PointNet for tracking. During pooling, only the

feature of points belonging to the same detection are pooled together. The feature vector of point cloud has a length of 512 for each detection.

3.3. Robust Multi-Modality Fusion Module

In order to better exploit multi-sensor features while maintaining the ability to track with each single sensor, our robust fusion module is designed to have the capability of fusing features of multiple modalities as well as handing original features from just a single modality.

Robust Fusion Module. The operations in the adjacency estimator is batch-agnostic, thus we concatenate single modalities and the fused modality in the batch dimension to ensure that the adjacency estimator could still work as long as there is an input modality. This design enables the proposed robust fusion module to skip the fusion process or fuse the remaining modalities (if there are still multiple sensors) during sensor malfunctioning, and pass them to the adjacency estimator thus the whole system could work with any sensor combination. Formally, we denote the feature vectors of different modalities as $\{F_i^s\}_{s=0}^S$, where the number of sensors is S , and the fused feature is denoted as F_i^{fuse} . In our formulation, the features of fused modality has the same size as a single modality. The robust fusion module concatenates $\{F_i^s\}_{s=0}^S$ and F_i^{fuse} along the batch dimension and feeds them to the adjacency estimator. They form a tensor of size $(S + 1) \times D \times (N + M)$.

The robust fusion module could employ arbitrary fusion module, and we investigate three fusion modules as shown in Figure 3. Take two sensors’ setting as an example, the fusion module A naively concatenates features of multiple modalities; the module B add these features together; the module C introduces attention mechanism.

Fusion Module A. A common approach is to concatenate these features, and use point-wise convolution with weight W to adapt the length of the output vector to be the same as a single sensor’s feature as follows:

$$F_i^{fuse} = W \otimes \text{CONCAT}(F_i^0, \dots, F_i^S), \quad (1)$$

where $\otimes$ denotes a convolution operation, and $\text{CONCAT}(\cdot)$ denotes a concatenation operation.

Fusion Module B. Another intuitive approach is to fuse these two features with addition, we reproject the features of each modality and add them together as follows:

$$F_i^{fuse} = \left(\sum_{s=0}^S W^s \otimes F_i^s \right), \quad (2)$$

where W^s denotes the corresponding convolution kernels to the s -th sensor’s feature. By addition the module gathers information from each sensor, and correlation feature of fused modality is also more like that of single sensor. It is favorable for the adjacency estimator to handle different

modality since the correlation operation is multiplication or subtraction.

Fusion Module C. The module C introduces an attention mechanism for guiding the information fusion from different sensors, since the significance of a sensor's information might vary in different situations, e.g., the point cloud feature might be more important when the illumination condition is bad, and the image feature might be more important when the point cloud is affected in rainy days. The attention map G_i^s for each sensor is first calculated as follows:

$$G_i^s = \sigma(W_{att}^s \otimes F_i^s), \quad (3)$$

where W_{att}^s is the convolution parameter and σ is a sigmoid function. We expect the W_{att}^s to learn predict the importance conditioned on the feature itself, and the sigmoid function normalizes the attention map to be in the range from 0 to 1. Then the information is fused as follows:

$$F_i^{fuse} = \frac{1}{\sum_{s=0}^S G_i^s} \sum_{s=0}^S G_i^s \odot (W^s \otimes F_i^s), \quad (4)$$

where $\odot$ denotes element-wise multiplication, and the summation of G_i^s is taken as a denominator for normalization.

3.4. Deep Adjacency Matrix Learning

Given the extracted multi-modality features, the adjacency estimator infers the confidence, affinity, start and end scores in the min-cost flow graph [11, 37] based on each modality. These features are shared for each branch of the adjacency estimator, namely the confidence estimator, affinity estimator, start and end estimator. It is straightforward to learn a model for confidence estimator by taking it as a binary classification task. We focus on the design of the two other branches.

Correlation Operation. To infer the adjacency, the correlation of each detection pair is needed. The correlation operation is batch-agnostic thus it can handle cross-modality, and the operation applied channel by channel to take advantage of the neural network. The commutative property is theoretically favorable for learning paired data, since it is agnostic of the order of F_j^i and F_k^{i+1} . In this work, we compare three simple yet effective operators as follows:

- Element-wise multiplication,: $F_{jk} = F_j^i \odot F_k^{i+1}$,
- Subtraction: $F_{jk} = F_j^i - F_k^{i+1}$,
- Absolute subtraction: $F_{jk} = |F_j^i - F_k^{i+1}|$.

The element-wise multiplication is equivalent to a depth-wise correlation filter [21], where the filter size is 1×1 . The subtraction measures the distance of two vectors. By taking the absolute value of subtraction, the operation becomes commutative and agnostic to the chronology of detection, which makes the network more robust.

Affinity Estimator. The obtained F_{jk} is then used by the affinity estimator to predict the adjacency matrix A^i .

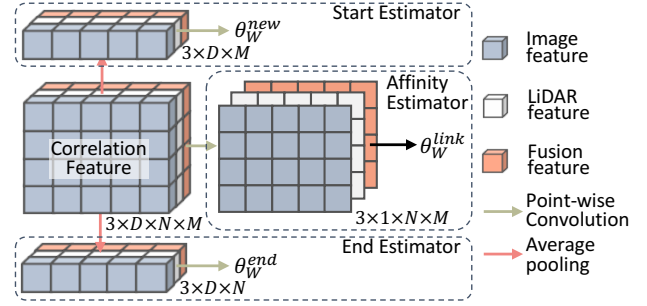


Figure 4. The structure of the affinity estimator and start and end estimator. The affinity estimator estimates the adjacency using point-wise convolution. The start and end estimator gathers the correlation feature of each detection to check whether a detection is linked to make prediction more robust.

Since the correlation operation handles multi-modality in the batch dimension and is performed on each detection pair between two frames, the correlation feature map has a size of $3 \times D \times N \times M$. We use 2D point-wise convolution as shown in Figure 4. This makes the network handle each correlation feature separately since it only needs to determine whether F_{jk} indicates a link. Since the convolution is batch-agnostic, it could work on any combination of modality, and the output adjacency matrix will have a size of $3 \times 1 \times N \times M$. Because these three predictions have the same target, we apply supervision signal to each of them, which enables joint optimization of feature extractor for each modality and affinity estimator for cross modality. During inference, the affinity estimator needs no modification if the sensor combination is changed, which allows both flexibility and reliability.

Start and End Estimator. The start and end estimator estimate whether a detection is linked, thus their parameters are shared for efficiency. Given the correlation feature F_{jk} , after gathering all the correlation information of one detection in each row or column by average pooling, the estimator also uses point-wise convolution to infer whether one detection is linked as shown in Figure 4. Since the pooling layer is batch-agnostic, the start and end estimator is also flexible for different sensor settings. During inference, we simply pad zeros for new score of detection in the first frame and end score of detection in the last frame, since they cannot be estimated from the correlation feature map.

Ranking Mechanism. We denote the raw output of the neural network's last layer as o_{jk}^i , and we found that a_{jk} should also be the greatest value among a_{js} , $s = 1, \dots, M$ and a_{tk} , $t = 1, \dots, N$, but directly take $A_{jk}^i = o_{jk}^i$ does not exploit this global information, thus we design a ranking mechanism to handle this problem. Specifically, we apply a softmax function for each row and each column in the output matrix, and gather these two matrices to get the final adjacency matrix. In this work, we investigate three operations to combine the two softmax feature maps: max,

multiplication and average. Taking the multiplication for example, the ranking mechanism is introduced as follows:

$$a_{jk}^i = \frac{e^{o_{jk}^i}}{\sum_{s=0}^N e^{o_{js}^i}} \times \frac{e^{o_{jk}^i}}{\sum_{t=0}^M e^{o_{tk}^i}}. \quad (5)$$

Loss Function. The whole framework can be learnt in an end-to-end manner in a multi-task learning framework. We adopt the cross entropy loss for the classification branch and the L2 loss for the other two, thus the overall loss function can be written as follows:

$$L = L_{link} + \alpha L_{start} + \gamma L_{end} + \beta L_{true}, \quad (6)$$

where α , γ and β indicates the weight of loss for each task. We empirically set $\alpha = \gamma = 0.4$ and $\beta = 1.5$ in all the experiments in this paper.

3.5. Linear Programming

After obtaining the prediction score from the neural networks, the framework needs to find an optimal solution from the min-cost flow graph. There are several facts that could be exploited as linear constraints among these binary variables in $\mathbf{Y}$. Firstly, if a detection is a true positive, it has to be either linked to another detection in the previous frame or the start of a new trajectory. Therefore, for one detection in current frame and all detections in its previous frame, a linear constraint can be defined in the form as follows:

$$\forall k, y_k^{true} = \sum_{j=0}^N y_{jk}^{link} + y_k^{start}. \quad (7)$$

Symmetrically, for one detection in previous frame and all detections in current frame, a linear constraint can be defined as follows:

$$\forall j, y_j^{true} = \sum_{k=0}^M y_{jk}^{link} + y_j^{end}. \quad (8)$$

These two constraints can be collapsed in a matrix form to yield $\mathbf{CY} = 0$, which has already encodes all valid trajectories. Then the data association problem is formulated as an integer linear program as follows:

$$\begin{aligned} \arg \max_y &= \Theta(\mathbf{X})^\top \mathbf{Y} \\ \text{s.t. } &\mathbf{CY} = 0, \mathbf{Y} \in \{0, 1\}^{|\mathbf{Y}|}, \end{aligned} \quad (9)$$

where $\Theta(\mathbf{X})$ is a flattened vector comprising all the predicted scores by the adjacency estimator.

4. Experiments

Dataset. Our method is evaluated on the challenging KITTI Tracking Benchmark [13]. This dataset contains 21 training sequences and 29 test sequences. We select 10 sequences

from the training partition as the training set and the remaining 11 sequences as the validation set. The train/validation set split is entirely based on frame number of these sequences to make the total frame number of training set (3975) close to that of validation set (3945). We submit our test-set result with the model trained only on split training set for fair comparison [36].

Each vehicle in the dataset is annotated with 3D and 2D bounding boxes with a unique ID across different frames, and this allows us to obtain the ground truth adjacency matrix for each detection predicted by the detector. We calculate the Intersection over Union (IoU) between each detection and ground truth (GT) bounding boxes, and assign the ID of one GT box to a detection if one has an IoU greater than 0.5 and has the greatest IoU among other detections. This setting is consistent with the test setting of KITTI Benchmark. The KITTI Benchmark [13] assesses the performance of tracking algorithms relying on standard MOT metrics, CLEAR MOT [5] and MT/PT/ML [22]. This set of metrics measures recall and precision of detection, and counts the number of identity switches and fragmentation of trajectories. It also counts the mostly tracked or mostly lost objects, and provides an overall tracking accuracy (MOTA).

Implementation Details. We first produce detections using the official code of PointPillar¹ [19]. The whole tracking framework is implemented with PyTorch [26]. The image appearance model's backbone is VGG-16 [40] with Batch Normalization [16] pretrained on ImageNet-1k [34]. For linear programming, we use the mixed integer programming (MIP) solver provided by Google OR-Tools². We train the model for 40 epochs using ADAM optimizer with a learning rate of $6e^{-4}$ and the super convergence strategy [41]. We manually set the score to be -1 if the confidence score falls below 0.2, this forces any detection having low confidence to be ignored during linear programming.

4.1. Ablation Study

To evaluate the proposed approach and demonstrate the effectiveness of the key components, we conduct an ablation study on the KITTI benchmark [13] under the online setting, with the state-of-the-art detector PointPillar [19]. We found that PointPillar detector produces large amount of false positive detections with low prediction score, so we discard detections with a score below 0.3. This does not hurt the mAP of detection, but saves a lot of memory in training. **Competency of Point Cloud for Tracking.** We set a 2D tracker as our baseline, which only employs 2D image patches as cues and use multiplication as correlation operator during data association, without the ranking mechanism. We first compare the effectiveness of image and Li-

¹<https://github.com/nutonomy/second.pytorch>

²<https://developers.google.com/optimization>

Table 1. Comparison of different modalities. ‘Frustum’ indicates using point cloud in the frustum. Robust Modules X indicates using fusion module X in the robust fusion module.

Method	Modality	MOTA↑	ID-s↓	FP↓	FN↓
Baseline	Image	74.88	454	951	1387
	Frustum	75.50	387	918	1418
	Point Cloud	75.70	362	946	1393
	Ensemble	77.54	158	949	1388
Robust Module A	Image	75.40	396	951	1387
	Point Cloud	76.13	317	948	1392
	Fusion	77.57	177	910	1406
Robust Module B	Image	75.17	421	951	1387
	Point Cloud	74.55	490	951	1387
	Fusion	77.62	193	850	1444
Robust Module C	Image	74.86	456	951	1387
	Point Cloud	74.94	452	946	1398
	Fusion	78.18	129	895	1401
Module A	Fusion	77.31	176	934	1412
Module B	Fusion	77.31	212	913	1396
Module C	Fusion	77.62	142	945	1400

DAR point cloud, and evaluate two approaches to employ the point cloud: using point cloud in the frustum or in the bounding box. From the row of baseline in Table 1, it is observed that using the point cloud in the frustum yields competitive results as using that in the bounding box. The results suggest the applicability of point cloud even with 2D detections (as discussed in Section 3.2), thus the proposed framework is adaptable for 2D or 3D detector with arbitrary modality. More surprisingly, all point cloud methods perform better than the image baseline, which suggests the efficacy of point cloud’s deep representation, and indicates that the system could still work when camera is failed.

Robust Multi-Modality Fusion Module. We compare the effectiveness of the robust fusion modules A, B, and C. Baselines comprise of trackers using a single sensor, i.e, camera or LiDAR; we train and evaluate each modality separately. To form a stronger baseline, we ensemble the image model (MOTA 74.88) and point cloud model in bounding box (MOTA 75.70), and yields much better result (MOTA 77.54). As shown in Table 1, only robust fusion module C with attention mechanism surpasses the ensemble result remarkably, although all fusion methods surpass single-sensor baselines. The results suggest the non-triviality of finding a robust fusion module for multi-sensor inputs.

Since each methods with robust fusion module also provides prediction of single sensor, we compare the single sensor results of each robust fusion module in Table 1. As can be observed, while the proposed Robust Module is capable of fusing multi-modality effectively, it can maintain competitive performance on the single modality in comparison to baselines (wherein dedicated training on single modality is conducted). Such kind of reliability in fusion is new in the literature.

Table 2. Comparison of 2D trackers with further modification.

Modification	MOTA↑	ID-s↓	FP↓	FN↓
Multiplication	74.88	454	951	1387
Subtraction	75.27	410	951	1387
Absolute subtraction	77.76	143	941	1387
Softmax w mul	75.08	431	951	1387
Softmax w max	76.24	313	940	1387
Softmax w add	77.40	234	891	1387

Table 3. Further improvement on fusion results. ‘Correlation’ indicates using absolute subtraction as correlation operation, ‘Ranking’ indicates using softmax with addition in ranking mechanism.

Correlation	Ranking	MOTA↑	ID-s↓	FP↓	FN↓
		78.18	129	895	1401
✓		79.18	23	873	1418
✓	✓	80.08	13	790	1411

Fusion Module. We further compare the results of normal fusion modules, which only outputs fused modality to the adjacency estimator, thus the tracker cannot perform tracking with single modality under multi-modality setting. The results in the last row of Table 1 shows that the proposed robust module outperforms the baseline modules A, B, and C consistently, with the additional capability of handling single modality. The results suggests that by preserving reliability, mmMOT gets more supervision signal which is favorable, and thus further improves the accuracy.

4.2. Further Analysis

Correlation Operator. We further conduct experiments on correlation function discussed in Section 3.4, and compare the effectiveness of three different correlation functions on the 2D baseline. As shown in Table 2, the subtraction variant always performs better than the multiplication variant, and with commutative property the absolute subtraction performs the best.

Ranking Mechanism. We also examine the effectiveness of the ranking mechanism, and investigate three different variants: the Softmax w mul, Softmax w max, Softmax w add, which indicate combining the softmax output by multiplication, argmax, addition, respectively. From Table 2, we can see that the ranking mechanism could improve MOTA by 0.2 at least, and adding the softmax output could yield improvement of about 2.5 in MOTA.

Best Results with 3D Detection. We further improve the results of fusion model. Following the conclusion in Table 2, we use the absolute subtraction for correlation operation, and softmax activation by addition for ranking mechanism. We compare the efficacy of each modification in Table 3. The absolute subtraction correlation improves the fusion model’s MOTA by 1, and the softmax activation with addition further improves 1 in MOTA and decreases the count of ID switches to 13, which is a remarkable improvement.

Table 4. Comparison on the testing set of KITTI tracking benchmark. Only published online methods are reported.

Method	MOTA↑	MOTP↑	Prec.↑	Recall↑	FP↓	FN↓	ID-s↓	Frag↓	MT↑	ML↓
DSM [11]	76.15	83.42	98.09	80.23	578	7328	296	868	60.00	8.31
extraCK [15]	79.99	82.46	98.04	84.51	642	5896	343	938	62.15	5.54
PMBM [36]	80.39	81.26	96.93	85.01	1007	5616	121	613	62.77	6.15
JCSTD [45]	80.57	81.81	98.72	83.37	405	6217	61	643	56.77	7.38
IMMDP [48]	83.04	82.74	98.82	86.11	391	5269	172	365	60.62	11.38
MOTBeyondPixels [38]	84.24	85.73	97.95	88.80	705	4247	468	944	73.23	2.77
mmMOT-normal	84.77	85.21	97.93	88.81	711	4243	284	753	73.23	2.77
mmMOT-lose image	84.53	85.21	97.93	88.81	711	4243	368	832	73.23	2.77
mmMOT-lose point cloud	84.59	85.21	97.93	88.81	711	4243	347	809	73.23	2.77

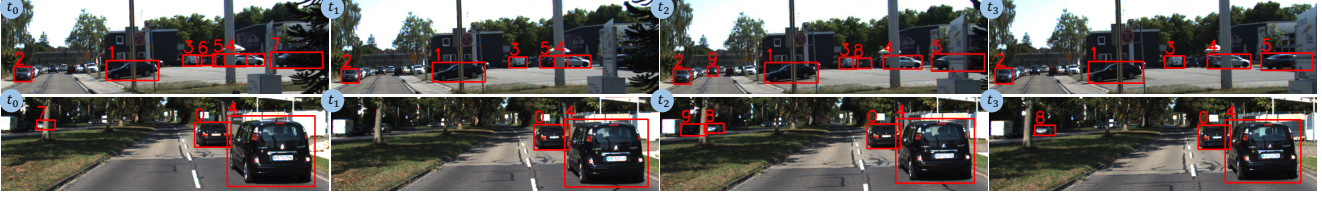


Figure 5. Failure case analysis.

4.3. KITTI Results

We achieve state-of-the-art and competitive results using 2D detection from RRC-Net [32] provided by MOT-BeyondPixels [38]. We use PointNet [30] to process point clouds in frustum, and VGG-16 [40] for image patches. More details are provided in the supplementary material. Table 4 compares our method with other published state-of-the-art online methods. We first test mmMOT using all the modalities, namely mmMOT-normal. Then we simulate the sensor failure case by only passing single modality to the same model, named mmMOT-lose image/point cloud. Under both conditions our mmMOT surpass all the other published state-of-the-art online methods on MOTA.

The proposed method by modality fusion surpasses the previous best method MOTBeyondPixels [38] by far fewer ID switches (184 fewer) with the same detection method. It is noteworthy that our single modality results still perform better, and we did not use bounding box and shape information of detections while MOTBeyondPixels does. PMBM [27], JCSTD [45], and IMMDP [48] exhibit fewer ID switches but miss approximately one to two thousand detections. Those missed detections are hard examples not only for detection but also for tracking, so it is likely that they would exhibit higher number of ID switches than our method if they use the same detections. Our method with each of the modalities surpasses the DSM [11] and extraCK [15] with fewer False Negatives and ID switches, i.e., our method makes fewer mistakes even when more hard examples are given.

4.4. Failure Case Analysis

We observe several conditions that could cause failure in our mmMOT. The statistical results are provided in the supplementary material, and the examples are shown in Figure

5, where each row includes four consecutive frames in a video. First, for objects far away, early errors caused by 2D detector will lead to false negative detection as shown by the car with ID 9 in the first row. The error could also cause ID switches if the car is missed but recovered, as shown by the car with ID 6 in the first row and the car with ID 7 in the second row. Second, the illumination also affects the performance, e.g., the black car in the shade with ID 9 in the second row. Third, the occlusion also causes difficulties, e.g., the detector missed the car with ID 7 in the first row. And partial observation makes the cars hard to be distinguished, e.g., the cars with ID 5 and 7 in the first row both only have black rears observed, thus are inferred to be the same. To further address the challenge caused by the occlusion, illumination and long distance, one may further exploit multi-modality in detection to prevent early errors, or exploit more information (e.g., temporal information) in data association to reinforce the prediction.

5. Conclusion

We have presented the **mmMOT**: a **multi-modality Multi-object Tracking** framework. We make the first attempt to avoid single-sensor instability while keeping multi-modality effective via a deep end-to-end network. Such a function is crucial for safe autonomous driving and has been overlooked by the community. Our framework is learned in an end-to-end manner with adjacency matrix learning, thus could learn to infer from arbitrary modality well in the same time. In addition, this framework is the first to introduce deep representation of LiDAR point cloud into data association problem, and enhances the multi stream framework's robustness against sensor malfunctioning.

Acknowledgment This work is supported by SenseTime Group Limited, Singapore MOE AcRF Tier 1 (M4012082.020), NTU SUG, and NTU NAP.

References

- [1] Alireza Asvadi, Pedro Girão, Paulo Peixoto, and Urbano Nunes. 3D object tracking using RGB and LIDAR data. In *ITSC*, 2016. 1
- [2] Alireza Asvadi, Paulo Peixoto, and Urbano Nunes. Detection and tracking of moving objects using 2.5D motion grids. In *ITSC*, 2015. 2
- [3] Min Bai, Gellért Mátyus, Namdar Homayounfar, Shenlong Wang, Shrinidhi Kowshika Lakshmikanth, and Raquel Urtasun. Deep multi-sensor lane detection. In *IROS*, 2018. 2
- [4] Sean Bell, C. Lawrence Zitnick, Kavita Bala, and Ross B. Girshick. Inside-outside net: Detecting objects in context with skip pooling and recurrent neural networks. In *CVPR*, 2016. 4
- [5] Keni Bernardin and Rainer Stiefelhausen. Evaluating multiple object tracking performance: The CLEAR MOT metrics. *EURASIP J. Image and Video Processing*, 2008. 6
- [6] Michael D. Breitenstein, Fabian Reichlin, Bastian Leibe, Esther Koller-Meier, and Luc J. Van Gool. Online multi-person tracking-by-detection from a single, uncalibrated camera. *IEEE TPAMI*, 2011. 2
- [7] Sergio Casas, Wenjie Luo, and Raquel Urtasun. IntentNet: Learning to predict intention from raw sensor data. In *CoRL*, 2018. 2
- [8] Xiaozhi Chen, Huimin Ma, Ji Wan, Bo Li, and Tian Xia. Multi-view 3D object detection network for autonomous driving. In *CVPR*, 2017. 2
- [9] Hyunggi Cho, Young-Woo Seo, B. V. K. Vijaya Kumar, and Ragunathan Rajkumar. A multi-sensor fusion system for moving object detection and tracking in urban driving environments. In *ICRA*, 2014. 1
- [10] Jaebum Choi, Simon Ulbrich, Bernd Lichte, and Markus Maurer. Multi-target tracking using a 3D-lidar sensor for autonomous vehicles. In *ITSC*, 2013. 2
- [11] Davi Frossard and Raquel Urtasun. End-to-end learning of multi-sensor 3D tracking by detection. In *ICRA*, 2018. 1, 2, 3, 5, 8
- [12] Ricardo Omar Chávez García and Olivier Aycard. Multiple sensor fusion and classification for moving object detection and tracking. *IEEE TITS*, 2016. 1
- [13] Andreas Geiger, Philip Lenz, and Raquel Urtasun. Are we ready for autonomous driving? the kitti vision benchmark suite. In *CVPR*, 2012. 2, 6
- [14] Zan Gojcic, Caifa Zhou, Jan D. Wegner, and Andreas Wieser. The perfect match: 3D point cloud matching with smoothed densities. *CoRR*, abs/1811.06879, 2018. 2
- [15] Gultekin Gunduz and Tankut Acarman. A lightweight online multiple object vehicle tracking method. In *IV*, 2018. 2, 8
- [16] Sergey Ioffe and Christian Szegedy. Batch normalization: Accelerating deep network training by reducing internal covariate shift. *CoRR*, abs/1502.03167, 2015. 6
- [17] Mingyang Jiang, Yiran Wu, and Cewu Lu. PointSIFT: A SIFT-like network module for 3D point cloud semantic segmentation. *CoRR*, abs/1807.00652, 2018. 2
- [18] Jason Ku, Melissa Mozifian, Jungwook Lee, Ali Harakeh, and Steven L. Waslander. Joint 3D proposal generation and object detection from view aggregation. In *IROS*, 2018. 2
- [19] Alex H. Lang, Sourabh Vora, Holger Caesar, Lubing Zhou, Jiong Yang, and Oscar Beijbom. PointPillars: Fast encoders for object detection from point clouds. *CoRR*, abs/1812.05784, 2018. 2, 6
- [20] Philip Lenz, Andreas Geiger, and Raquel Urtasun. Followme: Efficient online min-cost flow tracking with bounded memory and computation. *ICCV*, 2015. 2, 3
- [21] Bo Li, Wei Wu, Qiang Wang, Fangyi Zhang, Junliang Xing, and Junjie Yan. SiamRPN++: Evolution of siamese visual tracking with very deep networks, 2018. 5
- [22] Yuan Li, Chang Huang, and Ram Nevatia. Learning to associate: Hybridboosted multi-target tracker for crowded scene. In *CVPR*, 2009. 6
- [23] Tsung-Yi Lin, Piotr Dollár, Ross B. Girshick, Kaiming He, Bharath Hariharan, and Serge J. Belongie. Feature pyramid networks for object detection. In *CVPR*, 2017. 2
- [24] Wenjie Luo, Bin Yang, and Raquel Urtasun. Fast and furious: Real time end-to-end 3D detection, tracking and motion forecasting with a single convolutional net. In *CVPR*, 2018. 2
- [25] Aljoša Ošep, Wolfgang Mehner, Markus Mathias, and Bastian Leibe. Combined image- and world-space tracking in traffic scenes. In *ICRA*, 2017. 1, 2
- [26] Adam Paszke, Sam Gross, Soumith Chintala, Gregory Chanan, Edward Yang, Zachary DeVito, Zeming Lin, Alban Desmaison, Luca Antiga, and Adam Lerer. Automatic differentiation in pytorch. In *NeurIPS-W*, 2017. 6
- [27] Horst Possegger, Thomas Mauthner, Peter M. Roth, and Horst Bischof. Occlusion geodesics for online multi-object tracking. In *CVPR*, 2014. 2, 8
- [28] Charles Ruizhongtai Qi, Wei Liu, Chenxia Wu, Hao Su, and Leonidas J. Guibas. Frustum PointNets for 3D object detection from RGB-D data. *CoRR*, abs/1711.08488, 2017. 2, 4
- [29] Charles Ruizhongtai Qi, Hao Su, Kaichun Mo, and Leonidas J. Guibas. PointNet: Deep learning on point sets for 3D classification and segmentation. In *CVPR*, 2017. 2
- [30] Charles Ruizhongtai Qi, Li Yi, Hao Su, and Leonidas J. Guibas. PointNet++: Deep hierarchical feature learning on point sets in a metric space. In *NeurIPS*, 2017. 2, 4, 8
- [31] Akshay Ranges and Mohan M. Trivedi. No blind spots: Full-surround multi-object tracking for autonomous vehicles using cameras & lidars. *CoRR*, abs/1802.08755, 2018. 2
- [32] Jimmy Ren, Xiaohao Chen, Jianbo Liu, Wenxiu Sun, Jiahao Pang, Qiong Yan, Yu-Wing Tai, and Li Xu. Accurate single stage detector using recurrent rolling convolution. *CVPR*, 2017. 2, 8
- [33] Shaoqing Ren, Kaiming He, Ross Girshick, and Jian Sun. Faster R-CNN: Towards real-time object detection with region proposal networks. *IEEE TPAMI*, 2017. 2
- [34] Olga Russakovsky, Jia Deng, Hao Su, Jonathan Krause, Sanjeev Satheesh, Sean Ma, Zhiheng Huang, Andrej Karpathy, Aditya Khosla, Michael Bernstein, Alexander C. Berg, and Li Fei-Fei. ImageNet Large Scale Visual Recognition Challenge. *IJCV*, 2015. 6
- [35] Amir Sadeghian, Alexandre Alahi, and Silvio Savarese. Tracking the untrackable: Learning to track multiple cues with long-term dependencies. *ICCV*, 2017. 2

- [36] Samuel Scheidegger, Joachim Benjaminsson, Emil Rosenberg, Amrit Krishnan, and Karl Granström. Mono-camera 3D multi-object tracking using deep learning detections and PMBM filtering. In *IV*, 2018. 2, 4, 6, 8
- [37] Samuel Schuster, Paul Vernaza, Wongun Choi, and Manmohan Chandraker. Deep network flow for multi-object tracking. 2017. 2, 3, 5
- [38] Sarthak Sharma, Junaid Ahmed Ansari, J. Krishna Murthy, and K. Madhava Krishna. Beyond pixels: Leveraging geometry and shape cues for online multi-object tracking. In *ICRA*, 2018. 2, 4, 8
- [39] Shaoshuai Shi, Xiaogang Wang, and Hongsheng Li. PointRCNN: 3D object proposal generation and detection from point cloud. *CoRR*, abs/1812.04244, 2018. 2
- [40] Karen Simonyan and Andrew Zisserman. Very deep convolutional networks for large-scale image recognition. *CoRR*, abs/1409.1556, 2014. 4, 6, 8
- [41] Leslie N. Smith and Nicholay Topin. Super-convergence: Very fast training of neural networks using large learning rates, 2017. 6
- [42] Shiyang Song, Zhiyu Xiang, and Jilin Liu. Object tracking with 3D lidar via multi-task sparse learning. In *ICMA*, 2015. 2
- [43] Shuyang Sun, Jiangmiao Pang, Jianping Shi, Shuai Yi, and Wanli Ouyang. FishNet: A versatile backbone for image, region, and pixel level prediction. In *NeurIPS*, 2018. 2
- [44] Siyu Tang, Bjoern Andres, Mykhaylo Andriluka, and Bernt Schiele. Subgraph decomposition for multi-target tracking. In *CVPR*, 2015. 2
- [45] Wei Tian, Martin Lauer, and Long Chen. Online multi-object tracking using joint domain information in traffic scenarios. *IEEE TITS*, 2019. 2, 8
- [46] Bichen Wu, Alvin Wan, Xiangyu Yue, and Kurt Keutzer. SqueezeSeg: Convolutional neural nets with recurrent CRF for real-time road-object segmentation from 3D lidar point cloud. In *ICRA*, 2018. 2
- [47] Bichen Wu, Xuanyu Zhou, Sicheng Zhao, Xiangyu Yue, and Kurt Keutzer. SqueezeSegV2: Improved model structure and unsupervised domain adaptation for road-object segmentation from a lidar point cloud. *CoRR*, abs/1809.08495, 2018. 2
- [48] Yu Xiang, Alexandre Alahi, and Silvio Savarese. Learning to track: Online multi-object tracking by decision making. In *ICCV*, 2015. 2, 8
- [49] Amir Roshan Zamir, Afshin Dehghan, and Mubarak Shah. GMCP-Tracker: Global multi-object tracking using generalized minimum clique graphs. In *ECCV*, 2012. 2
- [50] Hui Zhou, Wanli Ouyang, Jian Cheng, Xiaogang Wang, and Hongsheng Li. Deep continuous conditional random fields with asymmetric inter-object constraints for online multi-object tracking. *IEEE TCSVT*. 2
- [51] Yin Zhou and Oncel Tuzel. VoxelNet: End-to-end learning for point cloud based 3D object detection. In *CVPR*, 2018. 2