

Learned Two-Plane Perspective Prior based Image Resampling for Efficient Object Detection

Anurag Ghosh N. Dinesh Reddy Christoph Mertz Srinivasa G. Narasimhan
 Carnegie Mellon University
 {anuraggh, dnapur, cmertz, srinivas}@cs.cmu.edu

Abstract

Real-time efficient perception is critical for autonomous navigation and city scale sensing. Orthogonal to architectural improvements, streaming perception approaches have exploited adaptive sampling improving real-time detection performance. In this work, we propose a learnable geometry-guided prior that incorporates rough geometry of the 3D scene (a ground plane and a plane above) to re-sample images for efficient object detection. This significantly improves small and far-away object detection performance while also being more efficient both in terms of latency and memory. For autonomous navigation, using the same detector and scale, our approach improves detection rate by $+4.1 AP_S$ or $+39\%$ and in real-time performance by $+5.3 sAP_S$ or $+63\%$ for small objects over state-of-the-art (SOTA). For fixed traffic cameras, our approach detects small objects at image scales other methods cannot. At the same scale, our approach improves detection of small objects by 195% ($+12.5 AP_S$) over naive-downsampling and 63% ($+4.2 AP_S$) over SOTA.

1. Introduction

Visual perception is important for autonomous driving and decision-making for smarter and sustainable cities. Real-time efficient perception is critical to accelerate these advances. For instance, a single traffic camera captures half a million frames every day or a commuter bus acting as a city sensor captures one million frames every day to monitor road conditions [10] or to inform public services [22]. There are thousands of traffic cameras [24] and nearly a million commuter buses [53] in the United States. It is infeasible to transmit and process visual data on the cloud, leading to the rise of edge architectures [43]. However, edge devices are severely resource constrained and real-time inference requires down-sampling images to fit both latency and memory constraints severely impacting accuracy.

On the other hand, humans take visual shortcuts [19] to

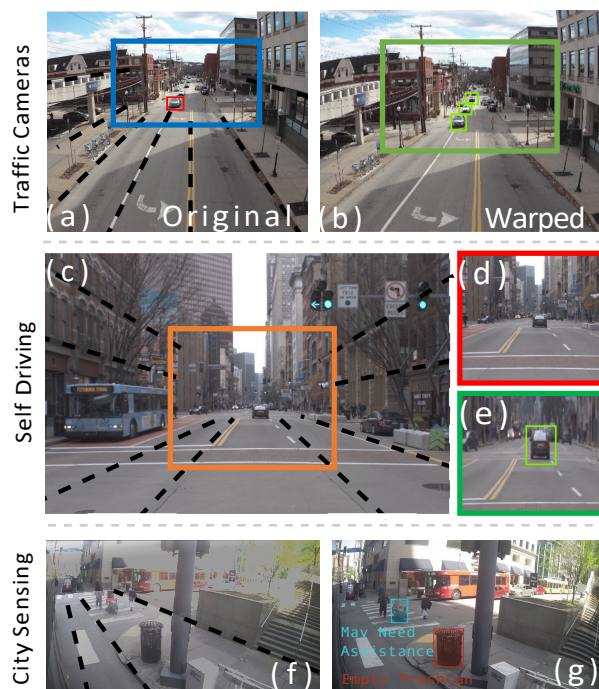


Figure 1. Geometric cues (black dashed lines) are implicitly present in scenes. Our Perspective based prior exploits this geometry. Our method (a) takes an image and (b) warps them, and performs detection on warped images. Small objects which are (d) not detected when naively downsampled but (e) are detected when enlarged with our geometric prior. Our method (f) uses a geometric model to construct a saliency prior to focus on relevant areas and (g) enables sensing on resource-constrained edge devices.

recognize objects efficiently and employ high-level semantics [19, 52] rooted in scene geometry to focus on relevant parts. Consider the scene in Figure 1 (c), humans can recognize the distant car despite its small appearance (Figure 1 (d)). We are able to contextualize the car in the 3D scene, namely (1) it's on the road and (2) is of the right size we'd expect at that distance. Inspired by these observations, can we incorporate semantic priors about scene geometry in our neural networks to improve detection?

In this work, we develop an approach that enables object detectors to “zoom” into relevant image regions (Figure 1 (d) and (e)) guided by the geometry of the scene. Our approach considers that most objects of interests are present within two planar regions, either on the ground plane or within another plane above the ground, and their size in the image follow a geometric relationship. Instead of uniformly downsampling, we sample the image to enlarge far away regions more and detect those smaller objects.

While methods like quantization [17], pruning [18], distillation [6] and runtime-optimization [16] improve model efficiency (and are complementary), approaches exploiting spatial and temporal sampling are key for enabling efficient real-time perception [21, 27]. Neural warping mechanisms [23, 36] have been employed for image classification and regression, and recently, detection for self-driving [50]. Prior work [50] observes that end-to-end trained saliency networks fail for object detection. They instead turn to heuristics such as dataset-wide priors and object locations from previous frames, which are suboptimal. We show that formulation of learnable geometric priors is critical for learning end-to-end trained saliency networks for detection. We validate our approach in a variety of scenarios to showcase the generalizability of geometric priors for detection in self-driving on Argoverse-HD [27] and BDD100K [57] datasets, and for traffic-cameras on WALT [37] dataset.

- On Argoverse-HD, our learned geometric prior improves performance over naive downsampling by **+6.6 AP** and **+2.7 AP** over SOTA using the same detection architecture. Gains from our approach are achieved by detecting small far-away objects, improving by **9.6 AP_S** (or 195%) over naive down-sampling and **4.2 AP_S** (or 63%) over SOTA.
- On WALT, our method detects small objects at image scales where other methods perform poorly. Further, it significantly improves detection rates by **10.7 AP_S** over naive down-sampling and **3 AP_S** over SOTA.
- Our approach improves object tracking (**+4.8% MOTA**) compared to baseline. It also improves tracking quality, showing increase of **+7.6% MT%** and reduction of **-6.7% ML%**.
- Our approach can be deployed in resource constrained edge devices like Jetson AGX to detect 42% more rare instances while being **2.2X** faster to enable real-time sensing from buses.

2. Related Work

We contextualize our work with respect to prior works modelling geometry and also among works that aim to make object detection more accurate and efficient.

Vision Meets Geometry: Geometry has played a crucial role in multiple vision tasks like detection [8, 20, 48, 55],

segmentation [28, 46], recognition [14, 47] and reconstruction [25, 35, 44]. Perspective Geometric constraints have been used to remove distortion [58], improve depth prediction and semantic segmentation [26] and feature matching [51]. However, in most previous works [25, 44, 55] exploiting these geometric constraints have mainly been concentrated around improving 3D understanding. This can be attributed to a direct correlation between the constraints and the accuracy of reconstruction. Another advantage is the availability of large RGB-D and 3D datasets [5, 15, 39] to learn and exploit 3D constraints. Such constraints have been under-explored for learning based vision tasks like detection and segmentation. A new line of work interpreting classical geometric constraints and algorithms as neural layers [7, 42] have shown considerable promise in merging geometry with deep learning.

Learning Based Detection: Object detection has mostly been addressed as an learning problem. Even classical-vision based approaches [12, 54] extract image features and learn to classify them into detection scores. With deep learning, learnable architectures have been proposed following this paradigm [4, 34, 38, 40], occasionally incorporating classical-vision ideas such as feature pyramids for improving scale invariance [29]. While learning has shown large improvements in accuracy over the years they still perform poorly while detecting small objects due to lack of geometric scene understanding. To alleviate this problem, we guide the input image with geometry constraints, and our approach complements these architectural improvements.

Efficient Detection with Priors: Employing priors with learning paradigms achieves improvements with little additional human labelling effort. Object detection has traditionally been tackled as a learning problem and geometric constraints were sparsely used for such tasks, constraints like ground plane [20, 48] were used.

Temporality [13, 50, 56] has been exploited for improving detection efficiently. Some of these methods [13, 50] deform the input image using approach that exploit temporality to obtain saliency. This approach handles cases where object size decreases with time (object moving away from the camera in scene), but cannot handle new incoming objects. None of these methods explicitly utilize geometry to guide detection, which handles both these cases. Our two-plane prior deforms the image while taking perspective into account without biasing towards previous detections.

Another complementary line of works automatically learn metaparameters (like image scale) [9, 16, 45] from image features. However, as they do not employ adaptive sampling accounting for image-specific considerations, performance improvements are limited. Methods not optimized for on-line perception like AdaScale [9] for video object detection do not perform well in real-time situations.

3. Approach

We describe how a geometric model rooted in the interpretation of a 3D scene can be derived from the image. We then describe how to employ this rough 3D model to construct saliency for warping images and improving detection.

3.1. Overview

Object sizes in the image are determined by the 3D geometry of the world. Let us devise a geometric inductive prior considering a camera mounted on a vehicle. Without loss of generality, assume the vehicle is moving in direction of the dominant vanishing point.

We are interested in objects that are present in a planar region (See Figure 2) of width P_1P_2 corresponding to the camera view, of length P_1P_3 defined in the direction of the vanishing point. This is the planar region on the ground on which most of the objects of interest are placed (vehicles, pedestrians, etc) and another planar region $Q_1...Q_4$ parallel to this ground plane above horizon line, such that all the objects are within this region (e.g., traffic lights).

From this simple geometry model, we shall incorporate relationships derived from perspective geometry about objects, i.e., the scale of objects on ground plane is inversely proportional to their depth w.r.t camera [20].

3.2. 3D Plane parameterization from 2D images

We parameterize the planes of our inductive geometric prior. We represent 2D pixel projections $u_1...u_4$ of 3D points $P_1...P_4$. Assume that the dominant vanishing point in the image is $v = (v_x, v_y)$ and let the image size be (w, h) . Consider u_1 (Figure 2 (b)). We can define a point on the edge of the image plane,

$$u_L = (0, v_y + v_x \tan \theta_1) \quad (1)$$

u_1 can be expressed as a linear combination of v and u_L ,

$$u_1 = \alpha_1 u_L + (1 - \alpha_1)v \quad (2)$$

Similarly, for u_2 , we can define u_R in terms of v and θ_2 and α_2 while u_3 and u_4 are defined like Equation 1 to represent any arbitrary plane in this viewing direction. However, for simplicity, for ground plane we fix them as $(0, h)$ and (w, h) respectively. Consider the planar region $Q_1...Q_4$ at height H above the horizon line. We can similarly define θ_3 and θ_4 to represent the angles from the horizon in the opposite direction and define q_1 and q_2 . Again, we set q_3 as $(0, 0)$ and q_4 as $(w, 0)$. We now have 4 points to calculate homographies H_{plane} for both planes.

For now, assume v is known. However, we still do not know the values for θ 's and α , and we shall learn these parameters end-to-end from task loss. These parameters are learned and fixed for a given scenario in our learning paradigm. Our re-parameterization aims to ease learning of these parameters

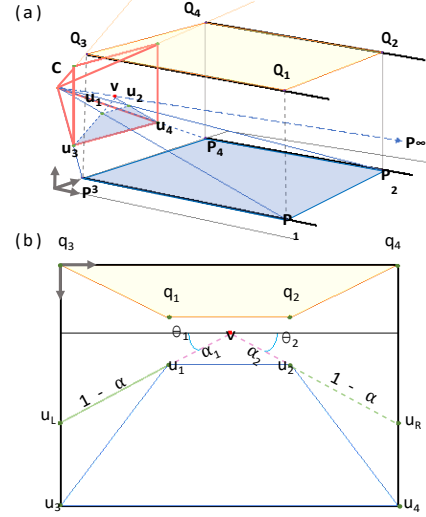


Figure 2. Geometry Of The Two Plane Perspective Prior: (a) describes the single view geometry of the proposed two plane prior. Region on the ground plane defined by $P_1, \dots, P_4$, and rays emanating from camera C to P_i intersect at $u_1...u_4$ on the image plane. The vanishing point v maps to P_∞ . This planar region accounts for small objects on the ground plane. To account for objects that are tall or do not lie on the ground plane, we consider another plane $Q_1...Q_4$ above the horizon line. These two planes encapsulate all the relevant objects in the scene. (b) depicts the re-parameterization of the two planes in the 2D image. Instead of representing the planar points $u_1...u_4$ as pixel coordinates, we instead parameterize them in terms of the vanishing point v , θ 's and α to ease learning.

as we clamp the values of α 's to $[0, 1]$ and θ 's to $[-\frac{\pi}{2}, \frac{\pi}{2}]$. It should be noted that all the operations are differentiable.

3.3. From Planes to Saliency

We leverage geometry to focus on relevant image regions through saliency guided warping [23, 36], and create a saliency map from a parameterized homography using $u_1...u_4$ defined earlier. Looking at ground plane from two viewpoints (Figure 3 (a) and (b)), object size decreases by their distance from the camera [20]. We shall establish a relationship to counter this effect and “sample” far-away objects on the ground plane more than nearby objects. The saliency guided warping proposed by [36] operates using an inverse transformation T_S^{-1} parameterized by a saliency map S as follows,

$$I'(x, y) = W_T(I) = I(T_S^{-1}(x, y)) \quad (3)$$

where the warp W_T implies iterating over output pixel coordinates, using T_S^{-1} to find corresponding input coordinates (non-integral), and bilinearly interpolating output color from neighbouring input pixel grid points. For each

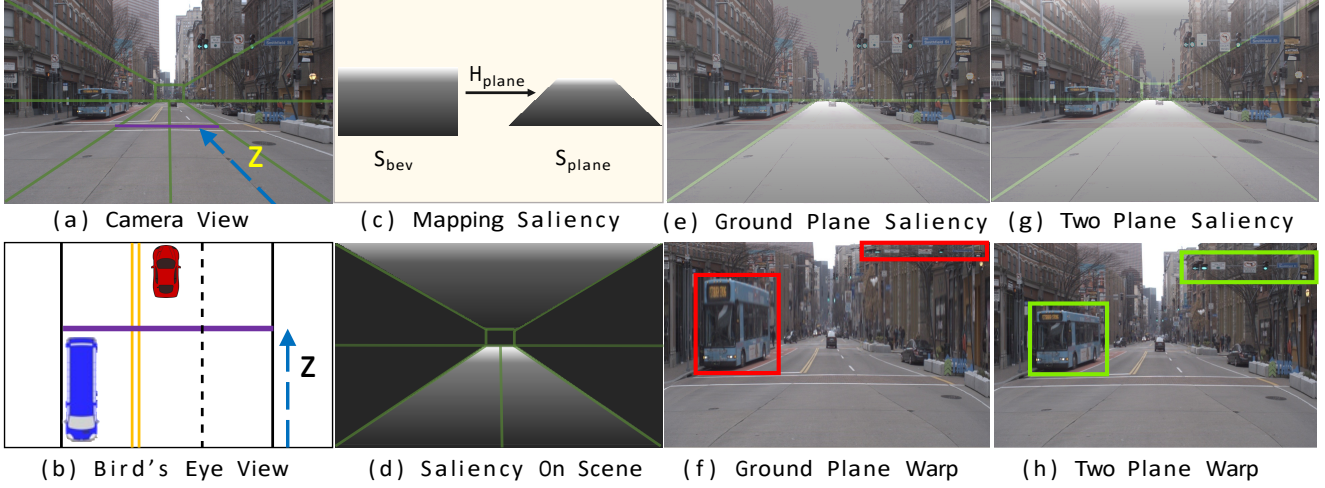


Figure 3. Two-Plane Perspective Prior based Image Resampling: Consider the scene of car, bus and traffic light from (a) camera view and (b) (simplified) bird's eye view. (c) Saliency function that captures the inverse relationship between object size (in camera view) and depth (bird's eye view is looking at X Z plane from above) can be transferred to the camera view (d), by mapping row z using H (marked by blue arrows). (e) and (f) shows that ground plane severely distorts nearby tall objects while squishing traffic light. (g) and (h) shows that additional plane reduces distortion for both tall objects and objects not on ground plane.

input pixel (x, y) , pixel coordinates with higher $S(x, y)$ values (i.e. salient regions) would be sampled more.

We construct a saliency S respecting the geometric properties that we desire. Let H_{plane} be the homography between the camera view (using coordinates $u_1 \dots u_4$) and a bird's eye view of the ground plane assuming plane size to be the original image size (w, h) . In bird's eye view, we propose saliency function for a row of pixels z (assuming bottom-left of this rectangle as $(0, 0)$) as,

$$S_{\text{bev}}(z) = e^{\nu(\frac{z}{h}-1)} \quad (4)$$

with a learnable parameter ν (> 1). ν defines the extent of sampling with respect to depth.

To map this saliency to camera view, we warp S_{bev} via perspective transform W_p and H_{plane} (Figure 3 (c)),

$$S_{\text{plane}} = W_p(H_{\text{plane}}^{-1}, S_{\text{bev}}) \quad (5)$$

We have defined saliency S_{plane} given H_{plane} in a differentiable manner. Our saliency ensures that objects on the ground plane separated by depth Z are sampled by the factor $e^{\nu \frac{Z}{h}}$ in the image.

3.4. Two-Plane Perspective Prior

Ground Plane saliency focuses on objects that are geometrically constrained to be on this plane and reasonably models objects far away on the plane. However, nearby and tall objects, and small objects far above the ground plane are not modelled well. In Fig 3 (f), nearby objects above ground plane (traffic lights), they are highly distorted. Critically, these same objects when further away are rendered

small in size and appear close to ground (and thus modelled well). Objects we should focus more on are thus the former compared to the latter. Thus, another plane is needed, and direction of the saliency function is reversed to $S_{\text{bev}}(z) = e^{\hat{\nu}(((h-z)/h)-1)}$ to account for these objects that would otherwise be severely distorted.

To represent the Two-Plane Prior, we represented the planar regions as saliencies. The overall saliency is,

$$S = S_{\text{ground_plane}} + \lambda S_{\text{top_plane}} \quad (6)$$

where λ is a learned parameter.

3.5. Additional Considerations

Warping via a piecewise saliency function imposes additional considerations. The choice of deformation method is critical, saliency sampler [36] implicitly avoids drastic transformations common in other approaches. For e.g., Thin-plate spline performs worse [36], produces extreme transformations and requires regularization [13].

Fovea [50] observes that restricting the space of allowable warps such that axis alignment is preserved improves accuracy, we adopt the separable formulation of T^{-1} ,

$$T_x^{-1}(x) = \frac{\int_{x'} S_x(x') k(x', x) x'}{\int_{x'} S_x(x') k(x, x')} \quad (7)$$

$$T_y^{-1}(y) = \frac{\int_{y'} S_y(y') k(y', y) y'}{\int_{y'} S_y(y') k(y, y')} \quad (8)$$

where k is a Gaussian kernel. To convert a saliency map S to S_x and S_y we marginalize it along the two axes. Thus entire rows or columns are “stretched” or “compressed”.

Two-plane prior is learnt end to end as a learnable image warp. For object detection, labels need to be warped too, and [36]’s warp is invertible. Like [50], We employ the loss $L(T^{-1}(f_{\phi}(W_T(I))), L)$ where (I, L) is the image-label pair and omit the use of delta encoding for training RPN [40] (which requires the existence of a closed form T), instead adopting GIoU loss [41]. This ensures W_T is learnable, as T^{-1} is differentiable.

We did not assume that the vanishing point is within the field of view of our image, and our approach places no restrictions on the vanishing point. Thus far, we explained our formulation while considering a single dominant vanishing point, however, multiple vanishing points can be also considered. Please see supplementary for more details.

3.6. Obtaining the Vanishing Point

We now describe how we obtain the vanishing point. Many methods exist with trade-offs in accuracy, latency and memory which which inform our design to perform warping efficiently with minimal overheads.

Fixed Cameras: In settings like traffic cameras, the camera is fixed. Thus, the vanishing point is fixed, and we can cache the corresponding saliency S , as all the parameters, once learnt, are fixed. We can define the vanishing point for a camera manually by annotating two parallel lines or any accurate automated approach. Saliency caching renders our approach extremely efficient.

Autonomous Navigation: Multiple assumptions simplify the problem. We assume that there is one dominant vanishing point and a navigating car is often moving in the viewing direction. Thus, we assume that this vanishing point lies inside the image, and directly regress v from image features using a modified coordinate regression module akin to YOLO [33, 38]. This approach appears to be memory and latency efficient. Other approaches, say, using parallel lane lines [1] or inertial measurements [3] might also be very efficient. An even simpler assumption is to employ the average vanishing point, as vanishing points are highly local, we observe this is a good approximation.

Temporal Redundancies: In videos, we exploit temporal redundancies, the vanishing point is computed every n_v frames and saliency is cached to amortize latency cost.

General Case: This is the most difficult case, and many approaches have been explored in literature to find all vanishing points. Classical approaches [49] while fast are not robust, while deep-learned approaches [31, 32, 59] are accurate yet expensive (either in latency or memory).

3.7. Learning Geometric Prior from Pseudo-Labels

Prior work [50] have shown improvements in performance on pre-trained models via heuristics, which didn’t require any training. However, their method still employs domain-specific labels (say, from Argoverse-HD) to generate the

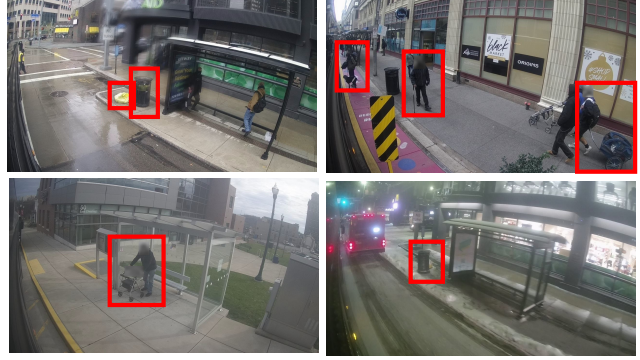


Figure 4. Commuter Bus Dataset: The captured data has a unique viewpoint, the average size of objects is small and captured under a wide variety of lighting conditions. Top-left and bottom-right images depict the same bus-stop and trashcan at different time of day and season.

prior. Our method can’t be used directly with pre-trained models as it learns geometrically inspired parameters end-to-end. However, domain-specific images without labels can be exploited to learn the parameters.

We propose a simple alternative to learn the proposed prior using pre-trained model without requiring additional domain-specific labels. We generate pseudo-labels from our pre-trained model inferred at 1x scale. We fine-tune and learn our warp function (to 0.5x scale) end-to-end using these “free” labels. Our geometric prior shows improvements without access to ground truth labels.

4. Dataset And Evaluation Details

4.1. Datasets

Argoverse-HD [27]: We employ Argoverse-HD dataset for evaluation in the autonomous navigation scenario. This dataset consists of 30fps video sequences from a car collected at 1920×1200 resolution, and dense box annotations are provided for common objects of interest such as vehicles, pedestrians, signboards and traffic lights.

WALT [37]: We employ images from 8 4K cameras that overlook public urban settings to analyze the flow of traffic vehicles. The data is captured for 3-second bursts every few minutes and only images with notable changes are stored. We annotated a set of 4738 images with vehicles collected over the time period of a year covering a variety of day/night/dawn settings, seasonal changes and camera viewpoints. We show our results on two splits (approximately 80% training and 20% testing). The first, All-Viewpoints, images from all the cameras are equally represented in the train and test sets. Alternatively, we split by camera, Split-by-Camera, images from 6 cameras are part of the training set and 2 cameras are held out for testing.

Commuter Bus Dataset: We curated this dataset from a

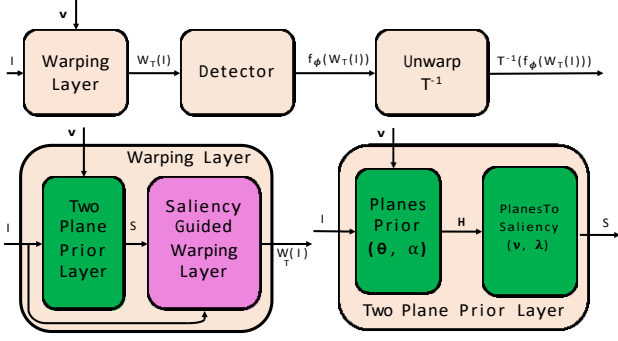


Figure 5. Two-Plane Prior as a Neural Layer: We implemented our approach as a global prior that is learned end-to-end from labelled data. Our prior is dependent on a vanishing point estimate to specify the viewing direction of the camera.

commuter bus running on an urban route. The 720p camera of interest has a unique viewpoint and scene geometry (See Figure 4). Annotated categories are trashcans and garbage bags (to help inform public services) and people with special needs (using wheelchairs or strollers), which are a rare occurrence. The dataset size is small with only 750 annotated images (split into 70% training and 30% testing). This is an extremely challenging dataset due to its unique viewpoint, small object size, rare categories, along with variations in lighting and seasons.

4.2. Evaluation Details

We perform apples-to-apples comparisons on the same detector trained using the datasets with identical training schedule and hardware.

Data: We compare to methods that were trained on fixed training data. In contrast, sAP leaderboards [27] don't restrict data and evaluate on different hardware. We compare with [16, 27] from leaderboard, which follow the same protocols. Other methods on the leaderboard use additional training data to train off-the-self detectors. Our detectors would see similar improvements with additional data.

Real-Time Evaluation: We evaluate using Streaming AP (sAP) metric proposed by [27], which integrates latency and accuracy into a single metric. Instead of considering models via accuracy-latency tradeoffs [21], real-time performance can be evaluated by applying real-time constraints on the predictions [27]. Frames of the video are observed every 33 milliseconds (30 fps) and predictions for every frame must be emitted before the frame is observed (forecasting is necessary). For a fair comparison, sAP requires evaluation on the same hardware. Streaming results are not directly comparable with other work [27, 50, 56] as they use other hardware (say, V100 or 2080Ti), thus we run the evaluation on our hardware (Titan X).

Additional details are in the Supplementary.

Method	Scale	AP	AP _S	AP _M	AP _L	Latency (ms)
Faster R-CNN	0.5x	24.2	4.9	29.0	50.9	78.4 ± 1.8
Fovea (S _D) [50]	0.5x	26.7	8.2	29.7	54.1	83 ± 2.5
Fovea (S _I) [50]	0.5x	28.0	10.4	31.0	54.5	85 ± 2.7
Fovea (L:S _I) [50]	0.5x	28.1	10.3	30.9	54.1	85.4 ± 2.7
Two-Plane Pr. (Pseudo.)	0.5x	27.1	9.8	28.9	50.2	104.5 ± 8.5
Two-Plane Prior	0.5x	30.8	14.5	31.6	52.9	105 ± 8.5
Baseline at higher scales						
Faster R-CNN	0.75x	29.2	11.6	32.1	53.3	142 ± 2.5
Faster R-CNN	1.0x	33.3	16.8	34.8	53.6	220 ± 1.7

Table 1. Evaluation on Argoverse-HD: Two-Plane Prior outperforms both SOTA's dataset-wide and temporal priors in overall accuracy. Our method improves small object detection by +4.1 AP_S or 39% over SOTA.

5. Results and Discussions

Accuracy-Latency Comparisons: On Argoverse-HD, we compare with Faster R-CNN with naive downsampling (Baseline) and Faster R-CNN paired with adaptive sampling from Fovea [50] which proposed two priors, a dataset-wide prior (S_D) and frame-specific temporal priors (S_I & L : S_I; from previous frame detections).

Two-Plane Prior improves (Table 1) upon baseline at the same scale by +6.6 AP and over SOTA by +2.7 AP. For small objects, the improvements are even more dramatic, our method improves accuracy by +9.6 AP_S or 195% over baseline and +4.2 AP_S or 45% over SOTA respectively. Surprisingly, our method at 0.5x scale improves upon Faster R-CNN at 0.75x scale by +1.6 AP having latency improvement of 35%. Our Two-Plane prior trained via pseudo-labels comes very close to SOTA which employ ground truth labels, the gap is only -1 AP and improves upon Faster R-CNN model trained on ground truth by +2.9 AP inferred at the same scale.

WALT dataset [37] comprises images (only images with notable changes are saved) and not videos, we compare with Fovea [50] paired with the dataset-wide prior. We observe similar trends on both splits (Table 3 and Fig 6) and note large improvements over baseline and a consistent improvement over Fovea [50], specially for small objects.

Real-time/Streaming Comparisons: We use Argoverse-HD dataset, and compare using the sAP metric (Described in Section 4.2). Algorithms may choose any scale and frames as long as real-time latency constraint is satisfied.

All compared methods use Faster R-CNN, and we adopt their reported scales (and other parameters). Streamer [27] converts any single frame detector for streaming by scheduling which frames to process and interpolating predictions between processed frames. AdaScale [9] regresses optimal scale from image features to minimize single-frame latency while Adaptive Streamer [16] learns scale choice in

ID	Method	Scale(s)	sAP	sAP _S	sAP _M	sAP _L
1	Streamer [27]	0.5x	18.3	3.1	15.0	40.9
2	1 + Adascale [9]	0.2x-0.6x	13.4	0.2	8.6	37.4
3	Adaptive Streamer [16]	0.2x-0.6x	21.3	4.2	18.8	47.0
4	1 + FOVEA (S _I) [50]	0.5x	24.1	8.4	24.7	48.7
5	1 + Ours (Avg VP)	0.5x	29.9	13.7	31.3	52.2
6	1 + Ours	0.5x	30.0	13.7	31.5	52.2
7	1 + Ours (VP Oracle)	0.5x	30.7	14.5	31.6	52.9

Table 2. Streaming Evaluation on Argoverse-HD: Ours denotes Two-Plane Prior. Every frame’s prediction (streamed at 30FPS) must be emitted before frame is observed [27] (via forecasting). All methods evaluated on Titan X GPU. Underlying detector (Faster R-CNN) is constant across approaches, improvements are solely from spatial sampling mechanisms. Notice improved detection of small objects by $+5.3 \text{ sAP}_S$ or 63% over SOTA.

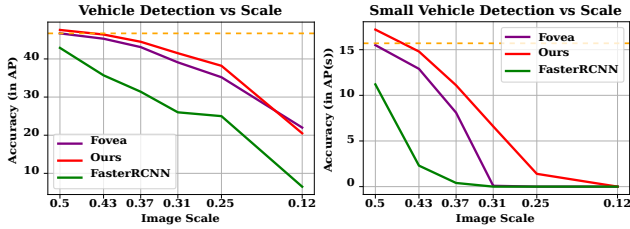


Figure 6. WALT All-Viewpoints Split: Our approach (Two-Plane Prior) shows improved overall performance over naive downsampling and state-of-the-art adaptive sampling technique, specially for small objects at all scales (starting from 0.5x). Horizontal line (orange) indicates performance at maximum possible scale (0.6x) the base detector was trained at (memory constraints).

the streaming setting. Both these methods employ naive-downsampling. State-of-the-art, Fovea [50] employs the temporal prior (S_I). From Table 2, Two-Plane prior outperforms the above approaches by $+16.5 \text{ sAP}$, $+8.6 \text{ sAP}$ and $+5.9 \text{ sAP}$ respectively. Comparison with [9, 16, 27] shows the limitations of naive downsampling, even when “optimal” scale is chosen. Our geometric prior greatly improves small object detection performance by 63% or $+5.3 \text{ sAP}_S$ over SOTA. To consider dependence on accurate Vanishing Point detection (and its overheads), we use NeurVPS [59] as oracle (we simulate accurate prediction with zero delay) to obtain an upper bound, we observe even average vanishing point location’s performance is within 0.8 sAP.

Accuracy-Scale Tradeoffs: We experiment with WALT All-Viewpoints split to observe accuracy-scale trade-offs. The native resolution (4K) of the dataset is extremely large and the gradients don’t fit within 12 GB memory of our Titan X GPU, thus we cropped the skies and other static regions to reduce input scale (1x) to 1500×2000 . Still, the highest scale we were able to train our baseline Faster R-CNN model is 0.6x. So, we use aggressive downsam-

Method	Scale	AP	AP _S	AP _M	AP _L
Faster R-CNN	0.25x	16.7	0	7.2	42.3
FOVEA (S _D) [50]	0.25x	23.2	0	13.9	54.1
Two-Plane Prior	0.25x	25.2	1.0	18.2	54.5
Faster R-CNN	0.5x	29.2	4.9	24.7	55.5
FOVEA (S _D) [50]	0.5x	34.4	8.7	30.5	59.4
Two-Plane Prior	0.5x	36.4	11.6	32.3	59.0
Baseline at higher scales					
Faster R-CNN	0.6x	33.2	9.3	28.6	56.7
Faster R-CNN*	1.0x	34.3	12.6	30.7	54.3

Table 3. WALT Camera-Split: The viewpoints on the test set were not seen, and Two-Plane Prior shows better performance over both naive downsampling and state-of-the-art adaptive sampling as it generalizes better to unseen scenes and viewpoints. *Not trained at that scale due to memory constraints on Titan X.

pling scales {0.5, 0.4375, 0.375, 0.3125, 0.25, 0.125}. The results are presented in Figure 6. We observe a large and consistent improvement over baseline and Fovea [50], specially for small objects. For instance, considering performance at 0.375x scale, our approach is better than baseline by $+13.1 \text{ AP}$ and Fovea by $+1.4 \text{ AP}$ for all objects.

For small objects, we observe dramatic improvement, at scales smaller than 0.375x, other approaches are unable to detect any small objects while our approach does so until 0.125x scale, showing that our approach degrades more gracefully. At 0.375x scale, our approach improves upon Faster R-CNN by $+10.7 \text{ AP}_S$ and Fovea by $+3.0 \text{ AP}_S$.

Generalization to new viewpoints: We use WALT Camera-Split, the test scenes and viewpoint are unseen in training. E.g., the vanishing point of one of the held-out cameras is beyond the image’s field of view. We operate on the same scale factors in the earlier experiments, and results are presented in Table 3. We note lower overall performance levels due to scene/viewpoint novelty in the test sets. Our approach generalizes better due to the explicit modelling of the viewpoint via the vanishing point (See Section 3.2). We note trends similar to previous experiment, we demonstrate improvements of $+8.5 \text{ AP}$ over naive-downsampling and $+2.0 \text{ AP}$ over Fovea [50] at 0.25x scale.

Tracking Improvements: We follow tracking-by-detection and pair IOUTracker [2] with detectors on Argoverse-HD dataset. MOTA and MOTP evaluate overall tracking performance. From Table 4, Our method improves over baseline by $+4.8\%$ and $+0.7\%$. We also focus on tracking quality metrics, Mostly Tracked % (MT%) evaluates the percentage of objects tracked for atleast 80% of their lifespan while Mostly Lost % (ML%) evaluates percentage of objects for less than 20% of their lifespan. In both these cases, our approach improves upon the baseline by $+7.6\%$ and -6.7% respectively. To autonomous navigation, we define two relevant metrics, namely, average lifespan extension (ALE)

Method	MOTA $\uparrow$	MOTP $\uparrow$	MT% $\uparrow$	ML% $\downarrow$	MOS $\downarrow$	ALE% $\uparrow$
Faster RCNN	39.8	82.3	30.7	35.6	37.1	59.3 %
Fovea (S_D) [50]	43.9	81.9	34.1	31.9	34.8	+5.4 %
Fovea (S_I) [50]	44.3	81.8	36.7	28.4	33.8	+8.4 %
Two-Plane Prior	44.6	83.0	38.3	28.9	31.6	+9.9 %

Table 4. Tracking Improvements: We setup a tracking by detection pipeline and replace the underlying detection method and observe improvements if any. All the detectors employ the Faster R-CNN architecture and are executed at 0.5x scale. We observe improvements in tracking metrics due to Two-Plane Prior.

and minimum object size tracked (MOS), whose motivation and definitions can be viewed in supplementary. We observe that our improvements are better than both Fovea (S_D) and (S_I). As we observe, our method improves tracking lifespan and also helps track smaller objects.

Efficient City-Scale Sensing: We detect objects on Commuter Bus equipped with a Jetson AGX. Identifying (Recall) relevant frames is key on the edge. Recall of Faster R-CNN with at 1x scale is 43.3AR (16.9AR_S) (Latency: 350ms; infeasible for real-time execution) but drops to 31.7AR (0.5AR_S) at 0.5x scale when naively down-sampled (Latency: 154ms). Whereas our approach at 0.5x scale improves recall by 42% over full resolution execution to **61.7 AR** (**16.4 AR_S**) with latency of 158 ms (+4 ms).

In supplementary, we perform additional comparisons, provide details on handling multiple vanishing points, more tracking results and edge sensing results. We also present some qualitative results and comparisons.

5.1. Ablation Studies

We discuss some of the considerations of our approach through experiments on the Argoverse-HD dataset.

Ground Plane vs Two-Plane Prior: We discussed the rationale of employing multiple planes in Fig 3, and our results are consistent. From Table 5, Two-Plane Prior outperforms Ground Plane prior considerably (**+1.8 AP**). Ground Plane Prior outperforms Two-Plane Prior on small objects by **+1 AP_S** but is heavily penalized on medium (**-4.1 AP_M**) and large objects (**-4.6 AP_L**). This is attributed to heavy distortion of tall and nearby objects, and objects that are not on the plane (Figure 3). Lastly, this prior was difficult to learn, the parameter space severely distorted the images (we tuned initialization and learning rate). Thus we did not consider this prior further. The second plane acts as a counter-balance and that warping space is learnable.

Vanishing Point Estimate Dependence: From Table 5, dominant vanishing point in autonomous navigation is highly local in nature, and estimating VP improves the result by **+1.2 AP**. Estimating the vanishing point is a design choice, it's important for safety critical applications like autonomous navigation (performance while navigating turns)

Method	Scale	AP	AP _S	AP _M	AP _L
Ground Plane Prior	0.5x	29.1	15.5	27.5	48.3
Two-Plane Prior (Pseudo.)	0.5x	27.1	9.8	28.9	50.2
Two-Plane Prior (Avg VP)	0.5x	29.6	12.7	30.7	52.7
Two-Plane Prior	0.5x	30.8	14.5	31.6	52.9

Table 5. Ablation Study on Argoverse-HD to justify our design choice of using two planes, dependence on accurate vanishing point detection and choice of pseudo labels vs ground truth.

however might be omitted for sensing applications.

Using Pseudo Labels vs Ground Truth: Table 5 shows there is still considerable gap (**-3.7 AP**) between the Two-Plane Prior trained from pseudo labels and ground truth. We observe that the model under-performs on stopsign, bike and truck classes, which are under-represented in the COCO dataset [30] compared to person and car classes. Performance of the pre-trained model on these classes is low even at 1x scale. Hence, we believe that the performance difference is an artifact of this domain gap.

6. Conclusions

In this work, we proposed a learned two-plane perspective prior which incorporates rough geometric constraints from 3D scene interpretations of 2D images to improve object detection. We demonstrated that (a) Geometrically defined spatial sampling prior significantly improves detection performance over multiple axes (accuracy, latency and memory) in terms of both single-frame accuracy and accuracy with real-time constraints over other methods. (b) Not only is our approach is more accurate when adaptively down-sampling at all scales, it degrades much more gracefully for small objects, resulting in latency and memory savings. (c) As our prior is learned end-to-end, we can improve a detector's performance at lower scales for "free". (d) Our approach generalizes better to new camera viewpoints and enables efficient city-scale sensing applications. Vanishing point estimation is the bottleneck of our approach [11, 31, 32, 59] for general scenes, and increasing efficiency of its computation we will see substantial improvements. Investigating geometric constraints to improve other aspects of real-time perception systems as future work, like object tracking and trajectory understanding and forecasting, is promising.

Societal Impact Our approach has strong implications for autonomous-driving and city-scale sensing for smart city applications, wherein efficient data processing would lead to more data-driven decision-making and public policies. However, privacy is a concern, and we shall release the datasets after anonymizing people and license plates.

Acknowledgements: This work was supported in part by an NSF CPS Grant CNS-2038612, a DOT RITA Mobility-21 Grant 69A3551747111 and by General Motors Israel.

References

- [1] Hala Abualsaud, Sean Liu, David B Lu, Kenny Situ, Akshay Rangesh, and Mohan M Trivedi. Laneaf: Robust multi-lane detection with affinity fields. *Robotics and Automation Letters*, 2021. 5
- [2] Erik Bochinski, Volker Eiselein, and Thomas Sikora. High-speed tracking-by-detection without using image information. In *AVSS*, 2017. 7
- [3] Federico Composeco and Marc Pollefeys. Using vanishing points to improve visual-inertial odometry. In *ICRA*, 2015. 5
- [4] Nicolas Carion, Francisco Massa, Gabriel Synnaeve, Nicolas Usunier, Alexander Kirillov, and Sergey Zagoruyko. End-to-end object detection with transformers. In *ECCV*, 2020. 2
- [5] Angel X Chang, Thomas Funkhouser, Leonidas Guibas, Pat Hanrahan, Qixing Huang, Zimo Li, Silvio Savarese, Manolis Savva, Shuran Song, Hao Su, et al. Shapenet: An information-rich 3d model repository. *arXiv preprint arXiv:1512.03012*, 2015. 2
- [6] Guobin Chen, Wongun Choi, Xiang Yu, Tony Han, and Manmohan Chandraker. Learning efficient object detection models with knowledge distillation. *NeurIPS*, 2017. 2
- [7] Hansheng Chen, Pichao Wang, Fan Wang, Wei Tian, Lu Xiong, and Hao Li. Epro-pnp: Generalized end-to-end probabilistic perspective-n-points for monocular object pose estimation. In *CVPR*, 2022. 2
- [8] Xiaozhi Chen, Kaustav Kundu, Yukun Zhu, Andrew G Berneshawi, Huimin Ma, Sanja Fidler, and Raquel Urtasun. 3d object proposals for accurate object class detection. *NeurIPS*, 2015. 2
- [9] Ting-Wu Chin, Ruizhou Ding, and Diana Marculescu. Adascale: Towards real-time video object detection using adaptive scaling. *Machine Learning and Systems*, 1:431–441, 2019. 2, 6, 7
- [10] Kevin Christensen, Christoph Mertz, Padmanabhan Pillai, Martial Hebert, and Mahadev Satyanarayanan. Towards a distraction-free waze. In *HotMobile*, 2019. 1
- [11] James Coughlan and Alan L Yuille. The manhattan world assumption: Regularities in scene statistics which enable bayesian inference. *NeurIPS*, 2000. 8
- [12] Navneet Dalal and Bill Triggs. Histograms of oriented gradients for human detection. In *CVPR*, 2005. 2
- [13] Babak Ehteshami Bejnordi, Amirhossein Habibian, Fatih Porikli, and Amir Ghodrati. Salisa: Saliency-based input sampling for efficient video object detection. In *ECCV*, 2022. 2, 4
- [14] Andreas Geiger, Martin Lauer, Christian Wojek, Christoph Stiller, and Raquel Urtasun. 3d traffic scene understanding from movable platforms. *TPAMI*, 2014. 2
- [15] Andreas Geiger, Philip Lenz, and Raquel Urtasun. Are we ready for autonomous driving? the kitti vision benchmark suite. In *CVPR*, 2012. 2
- [16] Anurag Ghosh, Akshay Nambi, Aditya Singh, Harish YVS, and Tanuja Ganu. Adaptive streaming perception using deep reinforcement learning. *arXiv preprint arXiv:2106.05665*, 2021. 2, 6, 7
- [17] Suyog Gupta, Ankur Agrawal, Kailash Gopalakrishnan, and Prithvi Narayanan. Deep learning with limited numerical precision. In *ICML*, 2015. 2
- [18] Song Han, Jeff Pool, John Tran, and William Dally. Learning both weights and connections for efficient neural network. *NeurIPS*, 2015. 2
- [19] Taylor R Hayes and John M Henderson. Scene semantics involuntarily guide attention during visual search. *Psychonomic Bulletin & Review*, 2019. 1
- [20] Derek Hoiem, Alexei A Efros, and Martial Hebert. Putting objects in perspective. *IJCV*, 2008. 2, 3
- [21] Jonathan Huang, Vivek Rathod, Chen Sun, Menglong Zhu, Anoop Korattikara, Alireza Fathi, Ian Fischer, Zbigniew Wojna, Yang Song, Sergio Guadarrama, et al. Speed/accuracy trade-offs for modern convolutional object detectors. In *CVPR*, 2017. 2, 6
- [22] Bret Hull, Vladimir Bychkovsky, Yang Zhang, Kevin Chen, Michel Goraczko, Allen Miu, Eugene Shih, Hari Balakrishnan, and Samuel Madden. Cartel: a distributed mobile sensor computing system. In *SenSys*, pages 125–138, 2006. 1
- [23] Max Jaderberg, Karen Simonyan, Andrew Zisserman, et al. Spatial transformer networks. *NeurIPS*, 2015. 2, 3
- [24] Sarah Kliff and Soo Oh. America’s 4,150 traffic cameras, in one map. <https://www.vox.com/a/red-light-speed-cameras>, 2015. 1
- [25] Abhijit Kundu, Yin Li, and James M. Rehg. 3d-rcnn: Instance-level 3d object reconstruction via render-and-compare. In *CVPR*, 2018. 2
- [26] Lubor Ladicky, Jianbo Shi, and Marc Pollefeys. Pulling things out of perspective. In *CVPR*, 2014. 2
- [27] Mengtian Li, Yu-Xiong Wang, and Deva Ramanan. Towards streaming perception. In *ECCV*, 2020. 2, 5, 6, 7
- [28] Xin Li, Zequn Jie, Wei Wang, Changsong Liu, Jimei Yang, Xiaohui Shen, Zhe Lin, Qiang Chen, Shuicheng Yan, and Jiashi Feng. Foveanet: Perspective-aware urban scene parsing. In *ICCV*, 2017. 2
- [29] Tsung-Yi Lin, Piotr Dollár, Ross Girshick, Kaiming He, Bharath Hariharan, and Serge Belongie. Feature pyramid networks for object detection. In *CVPR*, 2017. 2
- [30] Tsung-Yi Lin, Michael Maire, Serge Belongie, James Hays, Pietro Perona, Deva Ramanan, Piotr Dollár, and C Lawrence Zitnick. Microsoft coco: Common objects in context. In *ECCV*, 2014. 8
- [31] Yancong Lin, Ruben Wiersma, Silvia L Pintea, Klaus Hildebrandt, Elmar Eisemann, and Jan C van Gemert. Deep vanishing point detection: Geometric priors make dataset variations vanish. In *CVPR*, 2022. 5, 8
- [32] Shichen Liu, Yichao Zhou, and Yajie Zhao. Vapid: A rapid vanishing point detector via learned optimizers. In *ICCV*, 2021. 5, 8
- [33] Yin-Bo Liu, Ming Zeng, and Qing-Hao Meng. D-vpnet: A network for real-time dominant vanishing point detection in natural scenes. *Neurocomputing*, 2020. 5
- [34] Ze Liu, Yutong Lin, Yue Cao, Han Hu, Yixuan Wei, Zheng Zhang, Stephen Lin, and Baining Guo. Swin transformer: Hierarchical vision transformer using shifted windows. In *ICCV*, 2021. 2

- [35] Minh Vo N Dinesh Reddy and Srinivasa G. Narasimhan. Car-fusion: Combining point tracking and part detection for dynamic 3d reconstruction of vehicle. In CVPR, 2018. 2
- [36] Adria Recasens, Petr Kellnhofer, Simon Stent, Wojciech Matusik, and Antonio Torralba. Learning to zoom: a saliency-based sampling layer for neural networks. In ECCV, 2018. 2, 3, 4, 5
- [37] N Dinesh Reddy, Robert Tamburo, and Srinivasa G Narasimhan. Walt: Watch and learn 2d amodal representation from time-lapse imagery. In CVPR, 2022. 2, 5, 6
- [38] Joseph Redmon, Santosh Divvala, Ross Girshick, and Ali Farhadi. You only look once: Unified, real-time object detection. In CVPR, 2016. 2, 5
- [39] Jeremy Reizenstein, Roman Shapovalov, Philipp Henzler, Luca Sbordone, Patrick Labatut, and David Novotny. Common objects in 3d: Large-scale learning and evaluation of real-life 3d category reconstruction. In ICCV, 2021. 2
- [40] Shaoqing Ren, Kaiming He, Ross Girshick, and Jian Sun. Faster r-cnn: Towards real-time object detection with region proposal networks. NeurIPS, 2015. 2, 5
- [41] Hamid Rezatofighi, Nathan Tsoi, JunYoung Gwak, Amir Sadeghian, Ian Reid, and Silvio Savarese. Generalized intersection over union: A metric and a loss for bounding box regression. In CVPR, 2019. 5
- [42] Chris Rockwell, Justin Johnson, and David F Fouhey. The 8-point algorithm as an inductive bias for relative pose prediction by vits. In 3DV, 2022. 2
- [43] Mahadev Satyanarayanan. The emergence of edge computing. IEEE Computer, 2017. 1
- [44] Johannes Lutz Schönberger and Jan-Michael Frahm. Structure-from-motion revisited. In CVPR, 2016. 2
- [45] Gur-Eyal Sela, Ionel Gog, Justin Wong, Kumar Krishna Agrawal, Xiangxi Mo, Sukrit Kalra, Peter Schafhalter, Eric Leong, Xin Wang, Bharathan Balaji, et al. Context-aware streaming perception in dynamic environments. In ECCV, 2022. 2
- [46] Paul Sturgess, Karteek Alahari, Lubor Ladicky, and Philip HS Torr. Combining appearance and structure from motion features for road scene understanding. In BMVC, 2009. 2
- [47] Hang Su, Subhransu Maji, Evangelos Kalogerakis, and Erik Learned-Miller. Multi-view convolutional neural networks for 3d shape recognition. In ICCV, 2015. 2
- [48] Patrick Sudowe and Bastian Leibe. Efficient use of geometric constraints for sliding-window object detection in video. In ICVS, 2011. 2
- [49] Jean-Philippe Tardif. Non-iterative approach for fast and accurate vanishing point detection. In ICCV, 2009. 5
- [50] Chittesh Thavamani, Mengtian Li, Nicolas Cebon, and Deva Ramanan. Fovea: Foveated image magnification for autonomous navigation. In ICCV, 2021. 2, 4, 5, 6, 7, 8
- [51] Carl Toft, Daniyar Turmukhambetov, Torsten Sattler, Fredrik Kahl, and Gabriel J Brostow. Single-image depth prediction makes feature matching easier. In ECCV, 2020. 2
- [52] Antonio Torralba, Aude Oliva, Monica S Castelhana, and John M Henderson. Contextual guidance of eye movements and attention in real-world scenes: the role of global features in object search. Psychological review, 2006. 1
- [53] U.S. DEPARTMENT OF TRANSPORTATION. Bus profile, bureau of transportation statistics. <https://www.bts.gov/content/bus-profile>, 2022. 1
- [54] Paul Viola and Michael Jones. Rapid object detection using a boosted cascade of simple features. In CVPR, 2001. 2
- [55] Yan Wang, Wei-Lun Chao, Divyansh Garg, Bharath Hariharan, Mark Campbell, and Kilian Q Weinberger. Pseudo-lidar from visual depth estimation: Bridging the gap in 3d object detection for autonomous driving. In CVPR, 2019. 2
- [56] Jinrong Yang, Songtao Liu, Zeming Li, Xiaoping Li, and Jian Sun. Real-time object detection for streaming perception. In CVPR, 2022. 2, 6
- [57] Fisher Yu, Haofeng Chen, Xin Wang, Wenqi Xian, Yingying Chen, Fangchen Liu, Vashisht Madhavan, and Trevor Darrell. Bdd100k: A diverse driving dataset for heterogeneous multitask learning. In CVPR, 2020. 2
- [58] Yajie Zhao, Zeng Huang, Tianye Li, Weikai Chen, Chloe LeGendre, Xinglei Ren, Ari Shapiro, and Hao Li. Learning perspective undistortion of portraits. In ICCV, 2019. 2
- [59] Yichao Zhou, Haozhi Qi, Jingwei Huang, and Yi Ma. Neurvps: Neural vanishing point scanning via conic convolution. NeurIPS, 2019. 5, 7, 8