

Article

Automated Visual Recognizability Evaluation of Traffic Sign Based on 3D LiDAR Point Clouds

Shanxin Zhang ^{1,2}, Cheng Wang ^{1,*} , Lili Lin ¹, Chenglu Wen ¹, Chenhui Yang ¹, Zhemin Zhang ¹ and Jonathan Li ³

¹ Fujian Key Laboratory of Sensing and Computing for Smart Cities, School of Information Science and Engineering, Xiamen University, Xiamen, 361005, China; 23020140154352@stu.xmu.edu.cn (S.Z.); 23020181154277@stu.xmu.edu.cn (L.L.); clwen@xmu.edu.cn (C.W.); chyang@xmu.edu.cn (C.Y.); zhangzhemin@xmu.edu.cn (Z.Z.)

² Xizang Key Laboratory of Optical Information Processing and Visualization Technology, Information Engineering College, Xizang Minzu University, Xianyang, 712000, China

³ Department of Geography and Environmental Management University of Waterloo, Waterloo, ON N2L 3G1, Canada; junli@xmu.edu.cn

* Correspondence: cwang@xmu.edu.cn; Tel.: +86-592-2580003

Received: 21 May 2019 ; Accepted: 16 June 2019 ; Published: 19 June 2019



Abstract: Maintaining the high visual recognizability of traffic signs for traffic safety is a key matter for road network management. Mobile Laser Scanning (MLS) systems provide efficient way of 3D measurement over large-scale traffic environment. This paper presents a quantitative visual recognizability evaluation method for traffic signs in large-scale traffic environment based on traffic recognition theory and MLS 3D point clouds. We first propose the Visibility Evaluation Model (VEM) to quantitatively describe the visibility of traffic sign from any given viewpoint, then we proposed the concept of visual recognizability field and Traffic Sign Visual Recognizability Evaluation Model (TSVREM) to measure the visual recognizability of a traffic sign. Finally, we present an automatic TSVREM calculation algorithm for MLS 3D point clouds. Experimental results on real MLS 3D point clouds show that the proposed method is feasible and efficient.

Keywords: traffic sign; visibility; recognizability; mobile laser scanner; point clouds

1. Introduction

Traffic signs are an important kind of transportation facility that present traffic information, such as speed and driver behavior restrictions, road changes ahead, and so forth. A driver's timely visual recognition of traffic signs is critical to ensure safe driving and to avoid accidents [1–3]. However, traffic signs are often partially contaminated due to damage, occlusion, or installed in an unreasonable position, thereby decreasing visual recognition. Figure 1 illustrates some examples of traffic signs with low visibility caused by object/plants occlusion. Maintenance and optimization of the existing infrastructure is a major work for road network management. The most difficult aspect for an engineer is knowing which traffic signs should be repaired efficiently and when. How to accurately and efficiently evaluate the visibility and recognizability of traffic signs in a large-scale traffic environment is a challenging problem.

Visual recognizability of traffic signs is affected by: (1) traffic sign geometric factors, such as placement of the sign, mounting height, tilt, aiming direction, depression angle, shape damage, occlusion, road curvature, fluctuating road surfaces, and so forth; (2) vehicle movement factors, such as vehicle speed, Geometric Field Of View (GFOV) [4], line of sight, and so forth; and (3) other factors, such as weather conditions [5], lighting [6], the age of the drivers and their cognitive burden of

traffic density [7], and so forth. This research focuses on geometric and vehicle movement factors of traffic signs.

Existing research on the evaluation of traffic sign visibility and recognizability is based mainly on simulator and image methods and naturalistic driving experimentation. Simulator based methods [8–10] cannot evaluate the visibility and recognizability of real roads. Image based methods [11–13] are limited by fixed viewpoints and cannot evaluate the visibility and recognizability over the whole road surface. Naturalistic driving experimentation based methods [14,15] cannot obtain recognizability at a given position on a road surface. So far, there is no solution available to evaluate the recognizability distribution of traffic signs in a real traffic environment.



Figure 1. Examples of traffic signs with low visual recognizability.

Mobile Laser Scanning (MLS) systems scan large-scale road environments at normal driving speeds and collect highly accurate 3D point clouds over the area of driving. All the traffic sign geometric factors can be measured in the 3D point cloud and the vehicle movement factors can be calculated in the same coordination. So, the MLS 3D point clouds are an ideal source of data for evaluating the visibility and recognizability of traffic signs.

The recognition of traffic signs depends on the visual judgment of humans. The visibility and recognizability results of simulator based methods are obtained by the recruited volunteers. The visibility and recognizability results of image based methods are obtained by camera view. The underlying rule is that the scene people see is consistent with what the camera sees (in fact, they may not be the same). Like simulator based methods and image based methods, our model is proposed also based on traffic human visual recognition theory. We seek to solve the problem of evaluating recognizability from the perspective of retinal imaging area, imaging location and occlusion.

We present a quantitative Traffic Sign Visual Recognizability Evaluation Model (TSVREM), and propose an automatic TSVREM calculation approach based on MLS 3D point clouds. To the best of our knowledge, this is the first solution for accurate traffic sign visual recognizability evaluation over a large-scale traffic environment.

This paper is organized as follows: Section 2 reviews the previous work; Section 3 defines the TSVREM model; Section 4 gives the automatic calculation of TSVREM on MLS point clouds; Section 5 describes the experiments; and Section 6 concludes the paper.

2. Related Work

2.1. Visibility and Recognizability Evaluation

Simulation-based methods. Sun et al. [16] recruited volunteers for visual cognition time in a driving simulator. The UTDive platform used in Reference [17] investigates driver behavior associated with visual and cognitive distractions. Lyu et al. [8] evaluated traffic safety by analyzing the driving behavior and performance of recruited drivers under a cognitive workload. Motamedi et al. [9] used BIM enabled Virtual Reality (VR) environments to analyze traffic sign visibility. Some researchers [10,18] used eye tracker equipment to determine the visual cognition of traffic signs under simulated driving conditions.

Simulation-based methods are used to gather statistics regarding the visual or cognitive information obtained by the recruited volunteers using a simulator platform. The methods pay attention to only the time required for the recognition of traffic signs and whether or not they can be recognized. The methods do not focus on the quantitative value of recognizability. Most important is that these methods cannot obtain visual or cognitive information on real roads. Even if the methods include wearable devices that can be used on a real road, because recognizability must be measured once at each location, it is very difficult to measure the recognizability of every location on a road surface. Therefore, these methods are not suitable for use in large-scale traffic environments.

Image-based methods. As part of nuisance-free driving safety support systems, Doman et al. [19] proposed a visibility estimation method for traffic signs by using image information. To compute visibility of a traffic sign, they used different contrast ratios and different numbers of pixels in the occluded area of an image. They improved their method by considering temporal environmental changes [20] and integrated both local and global features in a driving environment [12]. Belaroussi et al. [5] investigated the effects of reduced visibility due to fog in traffic sign detection.

Image-based methods, limited by viewpoint position, cannot continuously evaluate visibility or recognizability over an entire road surface. Also, image-based methods are limited by lighting conditions and do not consider the impact of traffic sign placement, road curvature, shape damage and so forth.

Naturalistic driving experimentation-based methods. Naturalistic driving experimentation allows us to recognize driving modes by observing a driver's behavior behind the wheel in natural conditions during long periods of observation. For a deeper understanding of driving behavior, José et al. [14,15] proposed methods for mapping naturalistic driving data with Geographic Information Systems (GIS). Because human cognition takes time, in naturalistic driving experimentation, a driver must stop to obtain recognizability and occlusion results from a given viewpoint. This shortcoming leads to the inability of this method to obtain the recognizability distribution of traffic signs.

Point clouds-based methods. Katz et al. [21] proposed a Hidden Point Removal (HPR) operator to obtain visible points from a given viewpoint, applied it to improving the visual comprehension of point sets [22] and studied the transformation function of the HPR operator [23]. Huang et al. [24], based on the HPR operator, studied traffic sign occlusion detection from a point cloud. They considered occluded distribution and an occlusion gradient.

However, other important factors in the visibility and recognizability of traffic signs, including proportion of the occluded area, influence of vehicle speed, road curvature, number of lanes, and so forth, have not been considered. Besides, the HPR operator cannot detect all the occluding points when the occluding point clouds are composed of multiple objects.

2.2. Traffic Sign Detection and Classification

Most existing traffic sign detection and classification methods are based on color and shape information within images or videos. These methods use color information to segment the sign candidate area, and then, to extract the traffic sign, use shape or edge features, including shape matching [25,26], Hough transform [27], HOG feature and SVM classification [28], and deep learning [29,30], and so forth. Lighting conditions and viewpoint position heavily effect the detection performance of image or video based methods.

With the rapid development of Light Detection And Ranging (LiDAR) technology, especially the application of MLS systems that can collect accurate and reliable 3D point clouds, it is now feasible to survey an urban or roadway environment. Wen et al. [31] stated the attributes of the MLS system and its data acquisition process. To extract urban objects from MLS point clouds, Yang et al. [32] proposed a method based on supervoxels and semantic knowledge. Lehtomäki et al. [33] took spin images and LDHs into account and used a machine learning method to recognize objects. Some researchers segment the objects according to their topology in point clouds first and then classify them by the constructed shape related features [34–37]. Photogrammetry allows, not only the obtaining of geometric data, but also the recovery of some data related to texture and semantics [38]. Other researchers also proposed methods to combine 3D point clouds and 2D images [39–42]. Those methods use the geometric attribute to detect traffic signs in point clouds first and then recognize them or analysis their retroreflectivity condition on their corresponding 2D images.

2.3. Road Marking Detection and Classification

Although we can obtain the position of the viewpoint by the trajectory of the MLS system, the lane change of the vehicle and the unknown road width will make the position of the viewpoint we calculated unreasonable, such as viewpoints not on the road, the horizontal viewpoints in a direction perpendicular to the trajectory line on an inclined road surface, so we obtain the viewpoints based on the road marking detection and classification. This has another advantage that it makes us can calculate the visibility and recognizability of the traffic sign in each lane separately.

Most image-based road marking detections use a deep learning method [43,44]. Because we need the 3D spatial information of the viewpoints on the roads, 2D image based road marking detection is not suitable in this paper.

Point clouds-based detection has increased in recent years. Guan et al. [45] extract road surface by curbs-based methods and generate a geo-referenced intensity image of the surface using an extended inverse-distance-weighted (IDW) approach first and then segment the geo-referenced intensity image into road-marking candidates with multiple thresholds. They further proposed 2D multi-scale tensor voting (MSTV) framework and clustering analysis in Reference [46] based on Reference [45]. Soilán et al. [47] also project the point clouds into intensity image first and then extract and classify the most common road markings, namely pedestrian crossings and arrows. The disadvantage of using geo-referenced intensity image as a medium to extract road marking is that it may lead to incomplete and incorrect in the process of feature extraction. In order to avoid this disadvantage, Yu et al. [48] extract road markings directly from the 3D point clouds via multi-segment thresholding

and Otsu binarization, classify them by Deep Boltzmann Machines (DBM) into seven classes, including boundary line, center line, arrows, pedestrian warning, pedestrian crossings, stop lines and others.

In this paper, when the boundary lines and center lines of road marking can be detected [48], we generate viewpoints in each lane by the extracted boundary lines and center lines, otherwise, we generate viewpoints on the road by trajectory.

3. Definition of Models

Our model, as well as other models [8–10], also relies on human vision. The advantage of our method is it calculates, from the 3D point cloud, an image area of the traffic sign on the retina. Thus, we obtain a continuous, quantitative representation of the visibility and recognizability of traffic signs within Sight Distance (SD) over a real road. In the context of traffic signs design, SD is defined as the length of roadway ahead visible to the driver [49].

In the following section, in addition to defining the models, we also introduce the concepts of visibility and recognizability fields, which are similar to magnetic fields. To more easily understand the meaning of a symbol, except for natural exponential function, the superscript of the symbol is an abbreviation of its meaning. The symbols used to define the model are listed in nomenclature.

3.1. Vem Model

We label a viewpoint at position (w, l) as $p(w, l)$. w is the width from viewpoint to roadside in vertical road direction. l is the length from viewpoint to traffic sign along road direction. The visibility of a traffic sign from a viewpoint $p(w, l)$, except for the occlusion factor, we express all the other factors included in the geometric factors (mentioned in Section 1), as $E_{w,l}^{geo}$. Geometric factor $E_{w,l}^{geo}$, occlusion factor $E_{w,l}^{occ}$ and sight line deviation factor $E_{w,l}^{sight}$ are used to construct the VEM model. The viewpoint visibility of a traffic sign, $E_{w,l}^{visibility}$, is defined as follows:

$$E_{w,l}^{visibility} = E_{w,l}^{geo} * E_{w,l}^{occ} * E_{w,l}^{sight} \quad (1)$$

The correlation between the geometric factors and sight line deviation is as follows: The evaluation of geometric factors depends on the viewpoint and the geometric attributes of traffic sign itself. The geometric attributes of traffic sign include orientation, height, position and so forth. The evaluation of geometric factors is independent of the sight line of driver. Sight line deviation factor reflects the different imaging positions on the retina that may lead to different visibilities and recognizabilities. The angle between sight line and line from viewpoint to the center of a traffic sign affects the sight line deviation factor. When the viewpoint and the geometric attributes of traffic sign is determined, as the angle of the sight line deviation increases, the visibility and recognizability of the traffic sign decrease.

How to evaluate the geometric, occlusion and sight line deviation factors is described below.

- Geometric factor evaluation

We use the principle of retinal imaging to consider the impact of the geometric factor. The evaluation of the geometric factor, $E_{w,l}^{geo}$, is given as follows:

$$E_{w,l}^{geo} = A_{w,l}^{view} / A_{type}^{unit} \quad (2)$$

where $A_{w,l}^{view}$ is the imaging area of retinal of a sign from viewpoint $p(w, l)$; A_{type}^{unit} is the retinal imaging area of a standard traffic sign viewed from the “unit viewpoint”.

The reason for introducing A_{type}^{unit} is to make the value of $E_{w,l}^{geo}$ fall in intervals $[0, 1]$. The “unit viewpoint” is a viewpoint that has a unit distance d^{unit} to the panel and in the normal line passing through the center of the panel. The d^{unit} are illustrated in Figure 2a. To ensure $E_{w,l}^{geo} \leq 1$, we set the unit distance, d^{unit} , at less than 3 m. It is unnecessary to compute the visibility of a traffic

sign when the observation distance is less than three meters. Due to the vehicle almost passing through the traffic sign, it is highly impractical (almost impossible) for the driver to turn around 90° to observe the traffic sign.

$E_{w,l}^{geo}$ is inversely proportional to the following: (1) the angle between the line connecting the viewpoint to the center of the traffic sign panel and the normal passing through the center (orientation factor); (2) observation distance; and (3) the damage degree of the traffic sign panel.

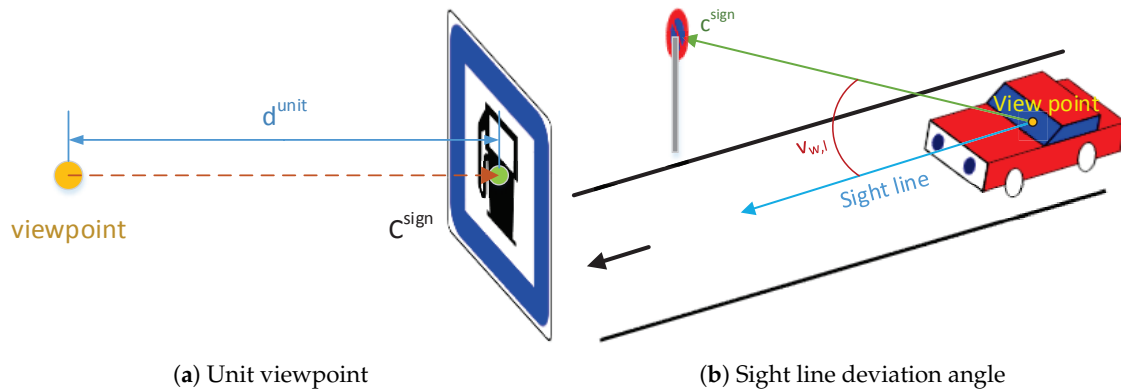


Figure 2. Illustration of the unit viewpoint and the angle of sight line deviation.

- Occlusion factor evaluation

As K. Doman [12] did, we consider the ratio of occlusion into the model. The occlusion factor $E_{w,l}^{occ}$ is given as follows:

$$E_{w,l}^{occ} = e^{-\lambda * A_{w,l}^{occ} / A_{w,l}^{view}} \quad (3)$$

where $A_{w,l}^{occ}$ is the occluded imaging area of retinal of a sign from viewpoint $p(w,l)$; λ is the penalization weight.

when the ratio of occlusion, $A_{w,l}^{occ} / A_{w,l}^{view}$, is constant, $E_{w,l}^{occ}$ decreases as λ increases; when λ is constant, $E_{w,l}^{occ}$ decreases as the ratio of occlusion increases. To ensure $E_{w,l}^{occ}$ is nearly zero when ratio of occlusion is nearly one and the traffic sign cannot be recognized under the situation of half occlusion, λ must satisfy the condition: $\lambda \geq 6$.

- Sight line deviation factor evaluation

Given $E_{w,l}^{geo}$ and $E_{w,l}^{occ}$, the sight line deviation factor reflects the different imaging positions on the retina that may lead to different visibilities. An object viewed from the line-of-sight in front of a driver's eye is seen more clearly than an object viewed from the periphery. When a traffic sign is fall in GFOV, it can be recognized; otherwise, it is unrecognizable. The evaluation of the driver's sight line deviation factor, $E_{w,l}^{sight}$, is established as follows:

$$E_{w,l}^{sight} = \begin{cases} 1, & v_{w,l} \leq G/2 \text{ and } v_{w,l} < V \\ e^{-\eta * \frac{v_{w,l} - G/2}{V - G/2}}, & G/2 < v_{w,l} < V \\ 0, & v_{w,l} \geq V \end{cases} \quad (4)$$

where $v_{w,l}$ is the angle between sight line and line from viewpoint to the center of traffic sign; G is the GFOV angle; V is the maximum angle between sight line and the line from viewpoint to the "A" pillar of the vehicle far from the driver; η is the punishment weight. The $v_{w,l}$ are illustrated in Figure 2b.

The GFOV decreases progressively with increasing vehicle speed [4,6]. G depends on the actual 85th percentile driving speed [50,51]. The actual 85th percentile driving speed on a road can be obtained by traffic big data.

The line-of-sight, the middle line of GFOV, is along the driving direction. If $G/2 \leq v_{w,l} \leq V$, a driver must turn his head to see traffic signs, thereby reducing the visibility of the traffic signs. If $v_{w,l} > V$, then $E_{w,l}^{sight} = 0$, because V is the maximum viewing angle for a driver.

In summary, for a given traffic sign, the visibility, $E_{w,l}^{visibility}$, of viewpoint $p(w, l)$ equals the following:

$$E_{w,l}^{visibility} = \frac{A_{w,l}^{view}}{A_{type}^{unit}} * e^{-\lambda * A_{w,l}^{occ} / A_{w,l}^{view}} * E_{w,l}^{sight} \quad (5)$$

To better understand the TSVREM model, we propose a visibility field concept as follows:

Visibility field: For a given surrounding around a target object, the visibility distribution of viewpoints in a 3D space constitutes a visibility field.

The visibility field reflects the visible distribution around a target object in 3D space. Take a traffic sign as an example. The visibility field of its hemispherical space is shown in Figure 3. The traffic sign (yellow) and distribution of its viewpoints (white points) are shown in Figure 3a. In Figure 3b,c, the color of the pixels changing from black to white, means that, the value of viewpoint visibility change from small to big. Figure 3b shows the visibility field of the traffic sign with occlusion. Figure 3c shows the visibility field of the traffic sign without occlusion.

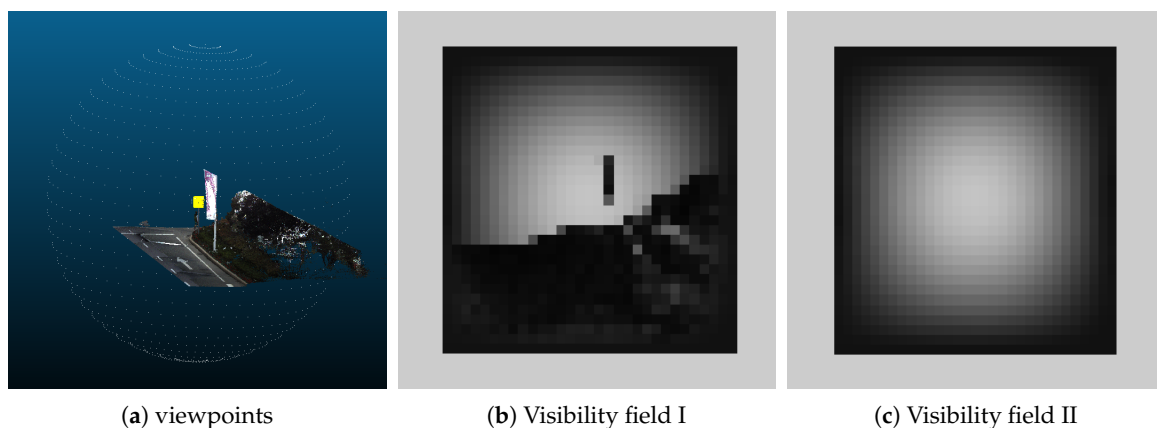


Figure 3. Hemispherical visibility field of a traffic sign.

3.2. Tsvrem Model

In this section, we first introduce how to evaluate the recognizability of a traffic sign from a viewpoint (viewpoint recognizability) and propose the definition of a visual recognizability field. Then, we present how to evaluate the recognizability of a traffic sign (traffic sign recognizability).

3.2.1. Viewpoint Recognizability and Definition of Visual Recognizability Field

Due to differences in language, symbols, habits and so forth, different countries have different design standards [50,52,53] for the location, size and content of signs. This leads to that two viewpoints having the same visibility; the recognition of their representatives in different countries may be different. However, the traffic standards in all countries have a common point at which a human must recognize a traffic sign within the SD distance. Based on this common point, we introduce a concept of standard visibility of a viewpoint to fill the gap between the visibility and recognizability of a sign and use the normalized visibility result as the recognizability. The advantage of introducing the

standard visibility of a viewpoint is that such an introduction, not only normalizes the visibility to interval $[0, 1]$, but also evaluates the degree of difference between actual traffic sign installation and installation specifications requirements.

Standard traffic surrounding: the surrounding in which a standard traffic sign is installed along a straight road, according to traffic sign specifications and free of occlusion at any viewpoint.

In actual traffic surrounding, the definition of viewpoint visibility includes three parts: geometric factor, occlusion factor, sight line deviation factor. In a standard traffic surrounding, there is no occlusion, therefore, the viewpoint visibility degenerates into the product of the standard geometric factors, $E_{w,l}^{geoS}$ and standard sight line deviation factors, $E_{w,l}^{sightS}$. The standard visibility of a viewpoint $E_{w,l}^{visibilityS}$ is defined as follows:

$$E_{w,l}^{visibilityS} = E_{w,l}^{geoS} * E_{w,l}^{sightS} \quad (6)$$

To define the viewpoint recognizability, for a viewpoint in an actual surrounding, we need to find out its corresponding viewpoint in the standard traffic surrounding. The corresponding viewpoint in a standard traffic surrounding is the viewpoint that has the same length along the road from the viewpoint to the traffic sign and the same width in the vertical road direction from the viewpoint to the roadside. Considering the other factors, E^{other} , mentioned in Section I, the visual recognizability, $E_{w,l}^{recognizability}$, of a viewpoint $p(w, l)$ is given as follows:

$$E_{w,l}^{recognizability} = \gamma * \frac{E_{w,l}^{visibility}}{E_{w,l}^{visibilityS}} + (1 - \gamma) * E^{other} \quad (7)$$

where $\gamma = 1$ because we don't consider E^{other} in this paper. Under normal circumstances, $E_{w,l}^{recognizability} \leq 1$, but due to road curvature, the incorrect installation of the traffic sign may cause, at some viewpoints, $E_{w,l}^{recognizability} > 1$. The traffic sign must be recognized by the driver when the actual visibility is greater than its standard visibility. Thus, in this case, we set $E_{w,l}^{recognizability} = 1$.

Visual recognizability field: for a given surrounding around an object, the visual recognizability distribution of viewpoints in a 3D space constitutes a visual recognizability field.

3.2.2. Traffic Sign Visual Recognizability

The visual recognizability of a traffic sign is equal to the evaluation value of its visual recognizability field. The visual recognizability field of a traffic sign is a surface that is parallel and higher than the road surface with a observation height and included in the driving area in front of the traffic sign within the sight distance. We use the integral average value to evaluate the visual recognizability field of a traffic sign. The visual recognizability of a traffic sign E^{sign} is given as follows:

$$E^{sign} = \frac{1}{A^D} \iint_D E_{w,l}^{recognizability} dw dl \quad (8)$$

where the integral area D is expressed by inequality as follows: $D = \{(w, l) \mid 0 \leq w \leq w^{driving}, 0 \leq l \leq l^{SD}\}$. A^D is the area of D . $w^{driving}$ is the width of driving area in vertical road direction. l^{SD} is the length of sight distance.

4. Tsvrem Model Implementation

The inputs of the TSVREM Model are the road point clouds and the trajectory of MLS. The outputs are visibility field, recognizability field and traffic sign visual recognizability. First, we detect and classify the traffic signs and road markings using References [31,48], respectively. The thresholds or parameters used to extract traffic signs and road markings are consistent with the original paper. For example, when extracting traffic signs, the threshold of reflectance intensity is greater than 60,000 and the eigenvalues satisfy the following condition: $\lambda_1 \gg \lambda_2 \approx \lambda_3$ and $\lambda_1/\lambda_2 > 10$. Second,

we segment the traffic sign surrounding point clouds and remove the outliers points [54]. The traffic sign surrounding point clouds are the point clouds on the right (or left) side above the roadway in front of a traffic sign within the SD distance. Then, we select viewpoints from road marking point clouds. Finally, using the traffic sign panel point cloud, the traffic sign surrounding point clouds and the viewpoint together, we compute viewpoint visibility, viewpoint recognizability, traffic sign recognizability and output occluding points. The algorithm is summarized in Algorithm 1.

Algorithm 1 TSVREM Model Implementation

Input: MLS point clouds, trajectory

Output: visibility field, visual recognizability field,
traffic sign visual recognizability, occluding points

```

1: Detect traffic signs [31], road markings [48]
2: Select viewpoints from road markings Section 4.1.1
3: Segment traffic sign surrounding point clouds Section 4.1.2
4: for each viewpoint  $p(w, l)$ 
5:   Translate and rotate [55] a group input data to human view
6:   Compute traffic sign retina imaging area Section 4.2.1
7:   Compute occlusion point clouds retina imaging area Section 4.2.2
8:   Compute sight line deviation degree Section 4.2.3
9:   Compute visibility in actual traffic surrounding Equation (5)
10:  Compute visibility in standard traffic surrounding Section 4.2.4 Equation (6)
11:  Compute recognizability Equation (7)
12: end for
13: Compute traffic sign visual recognizability according to Equation (8)
  
```

4.1. Viewpoints Selection and Segment Traffic Sign Surrounding Point Clouds

4.1.1. Viewpoints Selection

After detecting road markings, we obtain the left and right solid road marking lines in driving area for a traffic sign. In paper Yu et al. [48], the correctness of detection road marking reaches 92%. If a solid line is not continuous, or is partially missing because of low reflectivity, to complete it, we use its attribute that it approximately parallels the trajectory line. If road markings are totally mis-detected, we first subtract the height of the MLS device from the trajectory to obtain a trajectory on the road. Then, the right solid road marking line, is obtained, by moving the trajectory on the road to the right edge of the road according to the traffic sign position in standard road design; the left solid road marking line, is obtained, by moving the trajectory on the road to the left according the default number of lanes.

We segment the left and right solid lines along the trajectory at appropriate intervals [45]. store the intersections in left edge and right edge, respectively. As shown in Figure 4 left, the rectangle $(m1, m2, m3, m4)$ is a segmented area. The center of the segmented area in left solid road marking cloud is stored into left edge. The center of the segmented area in left solid road marking point cloud (roadside) is stored into right edge.

For two adjacent points in each edge, we use interpolation or sampling method to insert or sample any number of points at will, meanwhile, ensure that the number of points in left and right edges is equal for the convenience fo calculation. Observation height is added to each point in the two edges. For each pair points, coming from the left and edges separately and have the same index in each edge, we use interpolation method to insert any number of points at will. The observation height is usually set at 1.2 m above the road surface [56]. Shown in Figure 5 are the results of the calculated viewpoints.

4.1.2. Segment Traffic Sign Surrounding Point Clouds

We segment the traffic sign surrounding point clouds along the trajectory using octree searching method. As shown in Figure 4 left, the rectangle $(r1, r2, r3, r4)$ is a slice of segmented area. The points that are impossible to occlude traffic sign are discarded. Those discarded points include the points too

far from the roadside (buildings along the road in the distance), too close to the ground (plants used to green roads under traffic signs), too high from the ground (tree crown above traffic signs).

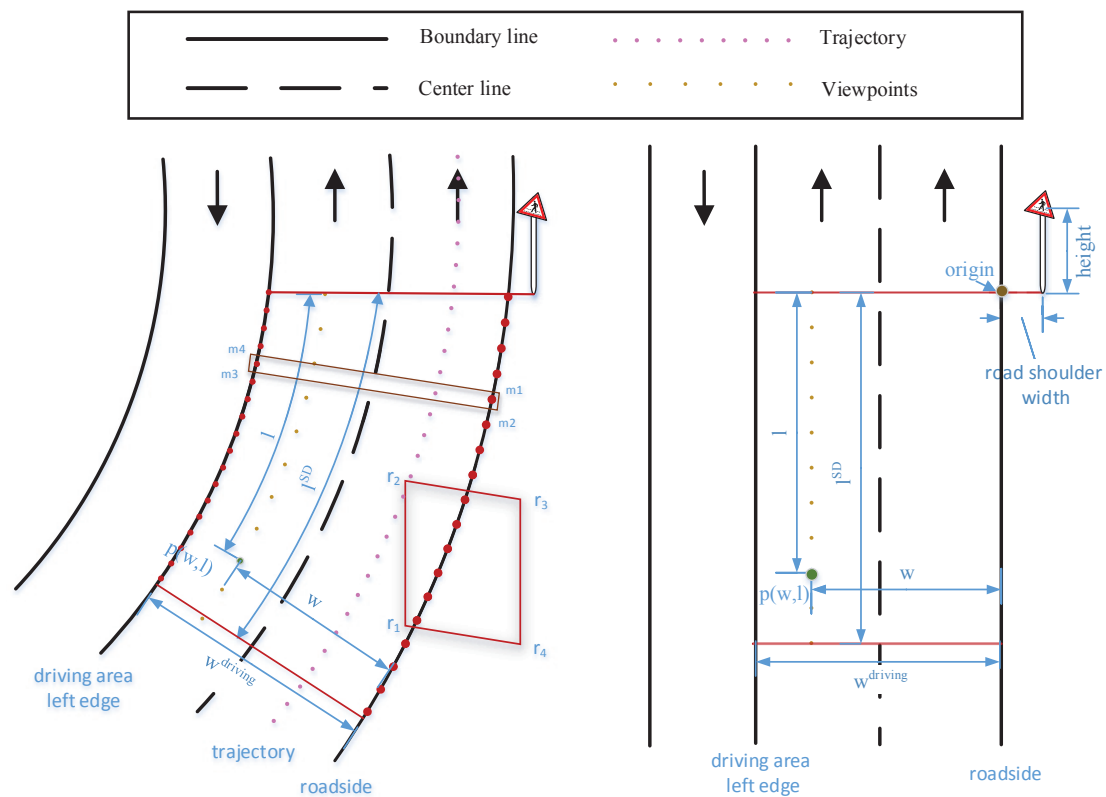


Figure 4. Illustration of the symbols in the TSVREM model. The **left** figure illustrate the actual traffic surrounding. The **right** figure illustrate the standard traffic surrounding

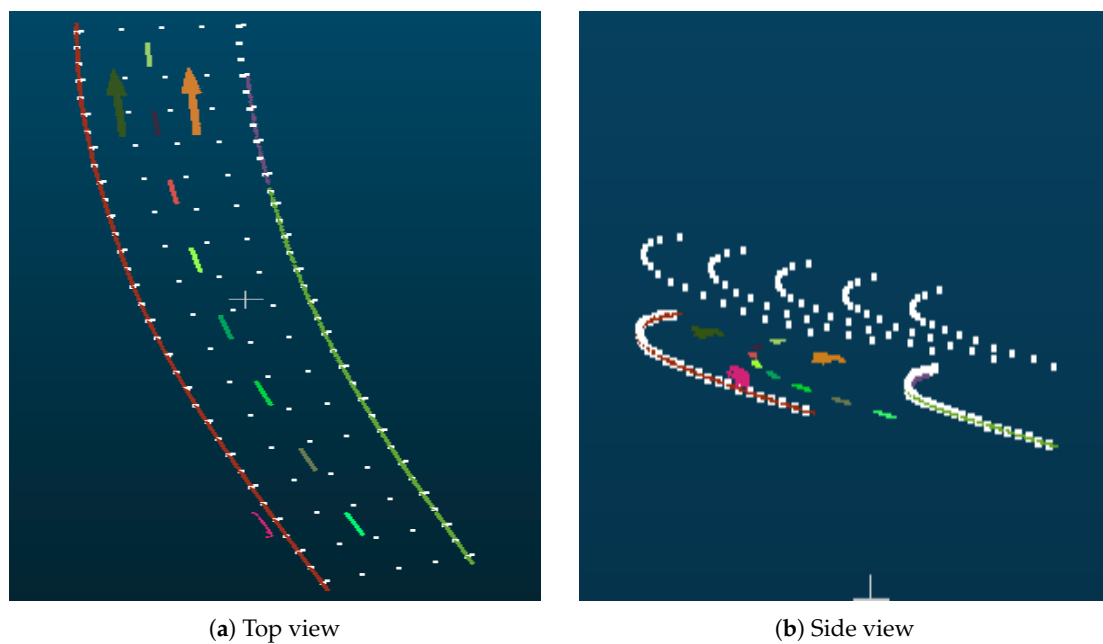


Figure 5. Viewpoints computation result.

4.2. Traffic Sign Visual Recognizability Computing

This Section includes the following procedures: computing the retinal imaging area of a point cloud from a given viewpoint; obtaining the occluding point cloud; obtaining the sight line deviation; and establishing the standard traffic environment.

4.2.1. Traffic Sign Retina Imaging Area Computing

The first step in computing the retinal imaging area is to rotate the coordinates of the input point clouds and viewpoints to the human view. For a group of input data: traffic sign panel point cloud, traffic sign surrounding point clouds, viewpoints, we first translate the origin of coordinate system to traffic sign panel center, by subtracting the coordinate of the panel center; and then rotating their z-axis coordinates to the translated line from the traffic sign panel point cloud center to the viewpoint using the quaternions rotation method [55]. The rotary axis is the vector of the cross product result between the z-axis and the translated line from the traffic sign panel point cloud center to the viewpoint. The rotary angle is the angle between the z-axis and translated the translated line from the traffic sign panel point cloud center to the viewpoint. The following operating about the a group input data in the paper are based on the coordinate-transformed point clouds.

Traffic sign panel point cloud is projected onto XOY plane, then, we get the projected traffic sign point cloud. The outer boundary of the projected traffic sign point cloud is computed by the alpha shape algorithm [57]. The alpha parameter in Reference [57] is set to about twice the interval between points. We use the polygon area formula to compute the area of the projected traffic sign point cloud. Finally, the area is mapped to the retinal imaging area using the human retinal imaging principle [58]. The distance from the center of the entrance pupil of an eye to the retina is set at seventeen millimeters.

4.2.2. Occlusion Point Clouds Retina Imaging Area Computing

From the vertexes of boundary of the projected traffic sign panel point cloud, we select a vertex which has the maximum distance from its center (origin) to the vertexes. Rotating the line segment from viewpoint to the selected vertex around the axis form viewpoint to origin, we get a vertebral body. The points which possible occlude the traffic sign are included in the vertebral body. We segment points in the vertebral body from traffic sign surrounding point clouds by two conditions: (1) the angle between the vector from viewpoint to the point and the vector form viewpoint to origin, is less than, the angle between the vector from viewpoint to the selected vertex and the vector from viewpoint to origin; (2) the distance from viewpoint to the point, is less than, the distance from viewpoint to the selected vertex.

For each point inside in the vertebral body, there is a ray from viewpoint to the point. We compute the intersection of ray to the XOY plane. If the intersection is inside the boundary of the projected traffic sign panel point cloud, then, the point is the occlusion point. All the occlusion points form occlusion point clouds. All intersections of the occluding points form block point clouds. The occlusion point clouds and block point clouds have the same retina imaging area. The retina imaging area of block point clouds is computed by the alpha shape algorithm and human retinal imaging principle. We call the method which obtain occlusion point clouds described above the “Ray Method”. Figure 6 shows an example of the “Ray Method” for computing occluded area of a traffic sign. In Figure 6a,b, the blue line segments are used to illustrate the spatial relationship among viewpoint, occluding object, traffic sign. The blue lines segments start at the viewpoint, pass through the occluding point and intersect the traffic sign (yellow) at the occluded point (red). In Figure 6c, the closed edges of the traffic sign (yellow points) is the obtained polygon for computing the area of the traffic sign. The close edges of the red points are the obtained polygon used for computing occluded area.

4.2.3. Sight Line Deviation Computing

The sight line of a driver is parallel with road direction (Figure 2b). It is a vector and equals the result of the coordinate of viewpoint $p(w, l - 1)$ minus the coordinate of viewpoint $p(w, l)$. The angle of sight line deviation equals the angle between sight line and the line from viewpoint to the center of traffic sign panel.

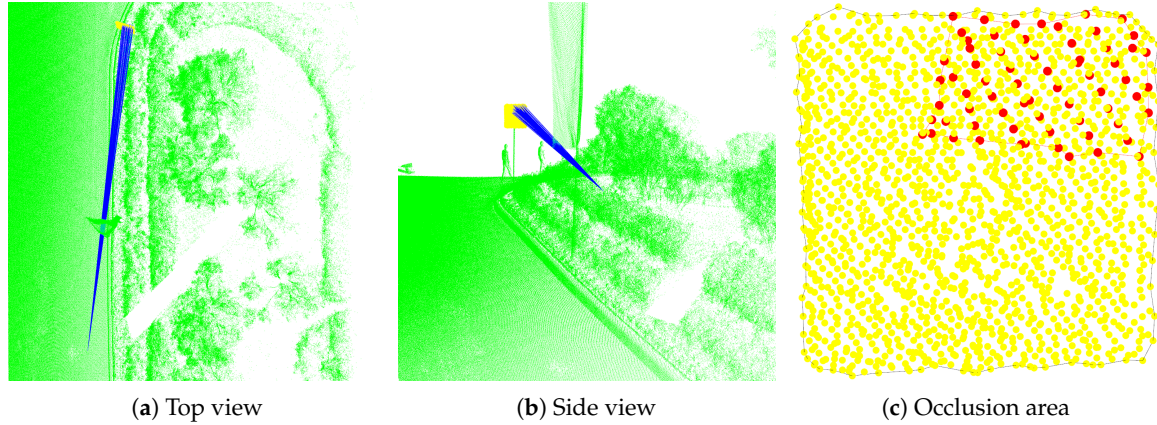


Figure 6. Traffic sign occluded point cloud computing result.

4.2.4. Standard Traffic Surrounding Setting

For all types of traffic signs, we built a traffic sign panel point cloud library, which contains one of each class of traffic sign panel point clouds. Every traffic sign panel point cloud in the library meets the conditions that their normal vector parallels the y-axis and the center is the origin.

According to a traffic sign design manual, such as the *Manual on Uniform Traffic Control Devices* (MUTCD in the United States) [50] and *Road Traffic Signs and Markings* (GB 5768-1999 in China) [53], the height, depression angle, road shoulder width, and the orientation angle are specified. Using the Y axis as roadside and origin as the traffic sign at roadside location. We first translated the traffic sign panel point cloud based on height and road shoulder width. Then we rotated the traffic sign panel point cloud [55] based on depression angle and orientation angle. For a viewpoint $p(w, l)$ in actual traffic surroundings, the coordinate of its corresponding viewpoint in standard traffic surrounding is $(-w, l, 0)$. The setting of standard traffic surroundings is shown in Figure 4 on the right.

5. Experiments And Discussion

5.1. Parameter Sensitivity Analysis

Parameter in $E_{w,l}^{geo}$. Given a viewpoint, a traffic sign, $E_{w,l}^{occ}$ and $E_{w,l}^{sight}$ are ascertained. $E_{w,l}^{visibility}$ is inversely proportional to the size of the parameter $d_{w,l}^{unit}$. $E_{w,l}^{geo}$ equals “1” when a driver observes a complete traffic sign at the unit viewpoint; equals “0” when the normal of traffic sign plane is perpendicular to the line from viewpoint to the center of traffic sign panel.

Parameters in $E_{w,l}^{occ}$. $E_{w,l}^{occ}$ must meet the real situation where $E_{w,l}^{occ} = 1$ under the circumstances of no occlusion and $E_{w,l}^{occ}$ is nearly equal to zero under the circumstances of half occlusion. We depict the value of $E_{w,l}^{occ}$ in Figure 7a under different occlusion ratios and λ . From the figure, we can see $E_{w,l}^{occ}$ meet the real situation when $\lambda \geq 6$. When traffic sign is not be occluded, $A_{w,l}^{occ} = 0$, $E_{w,l}^{occ} = e^0 = 1$. When traffic sign is completely be occluded, $A_{w,l}^{occ} = A_{w,l}^{view}$, $E_{w,l}^{occ} = e^{-6} = 0.0025 \approx 0$.

Parameters in $E_{w,l}^{sight}$. Some research [4,51,59] shows that most traffic signs (fall in the GFOV equal to 100°) are recognized accurately at a speed of 60 km/h; and most of traffic signs (falling in the GFOV

equal to 22°) are recognized accurately at a speed of 120 km/h. Linear interpolation is used to calculate the G at different speeds. It is easily proven that Equation (4) is continuous at $G/2$.

Usually for a car about two meters wide, the angle between the driver's viewpoint to the "A" pillar of vehicle and the sight line is generally less than 80 degrees. In this paper, we set V equals 80 degrees.

$E_{w,l}^{sight}$ is shown in Figure 7b under different $v_{w,l}$, V and η in Equation (4). This figure shows that $E_{w,l}^{sight}$ equals 1 when traffic signs fall in GFOV ($v_{w,l} \leq G/2$ and $v_{w,l} < V$) and decreases when viewing angle beyond GFOV ($G/2 < v_{w,l} < V$). $E_{w,l}^{sight}$ equals 0 When traffic signs fall behind the line from viewpoint to the "A" pillar of vehicle ($v_{w,l} \geq V$). η as the punishment weight should meet the above demand, like λ in Equation (3), $\eta \geq 6$.

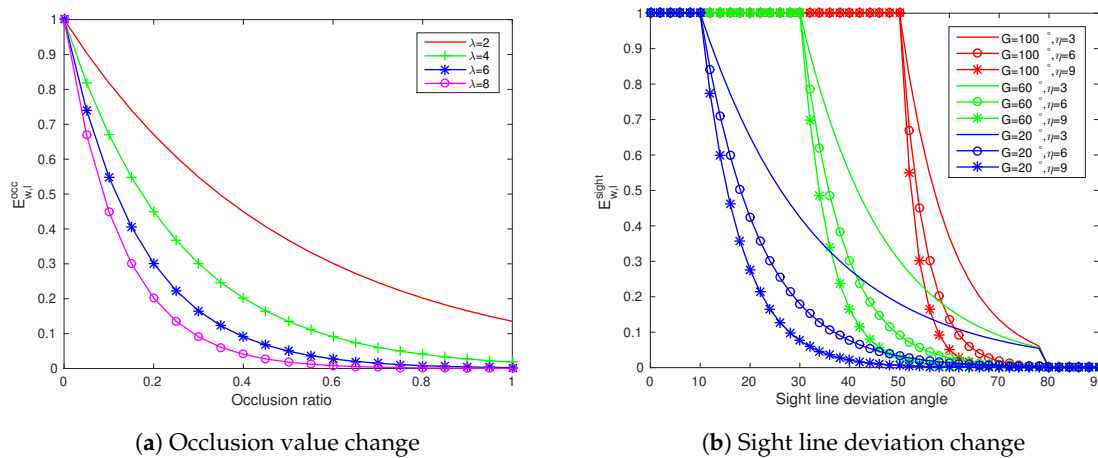


Figure 7. Occlusion change and sight line deviation change under different parameters.

5.2. Datasets Acquisition

To prove the practicality of our models and algorithm in urban roads, mountain roads and highways, using a RIEGL VMX-450 MLS system (Figure 8a), we scanned the following three datasets as shown in Figure 9: Huizhanchongxin Block (HB), Longhu Mountain Road (LMR), Shenyang-Haikou Highway (SHH) in Xiamen Island, China. Once projected on the horizontal plane, the density of the point clouds among the three datasets is about 4000~8000 points per square meter. The density decreases with increasing vehicle speed. The box-shaped legend is the area of a grid in the figure.

Using devices (including Leica Viva GNSS CS15 receiver, Huawei Honor V8 phone, bluetooth camera remote controller and a car), we created another dataset about photos and their GPS positions in street scenes. To improve the accuracy of the GPS positions, before taking a photo and recording its GPS position by CS15, we stopped the car and waited for the GPS signal to stabilize. GPS position accuracy, based on Realtime kinematic (RTK) technology, reached the centimeter level in open areas. Finally, 144 photos, with each photo (3968×2976 pixels) containing a GPS position, were obtained along the Longhu Mountain Road. Photo validation devices installed on the car are shown in Figure 8b. In this figure, the yellow points are traffic signs and the red points are the points occluding the view of the traffic sign. The contrast between the phone (upper left corner) and its corresponding image (upper right corner) shows they contain the same scenario. The traffic sign in the photo cannot be recognized by a human. After computation by our model, the occlusion ratio of the traffic sign is 87.37%; its recognizability is 0.59%; also, it is non-recognizable. Figure 8c shows the consistency between the trajectory collected by the VMX-450 and the photo positions collected by the GNSS CS15 receiver. The red points in Figure 8c are photo positions. In this experiment, the cellphone is fixed on the top of the car together with Leica Viva GNSS CS15 receiver. Different from the actual visibility or recognizability computing on a real road, the height of the viewpoints are the cellphone GPS positions,

not 1.2 m above the road. This is to ensure that the recognizability computed in point clouds is consistent with the image.

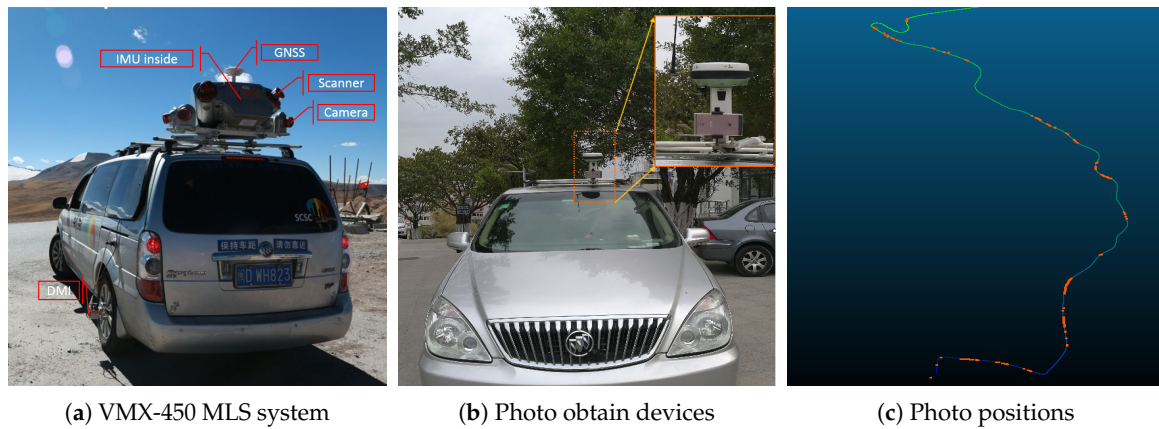


Figure 8. Datasets acquisition device and result.

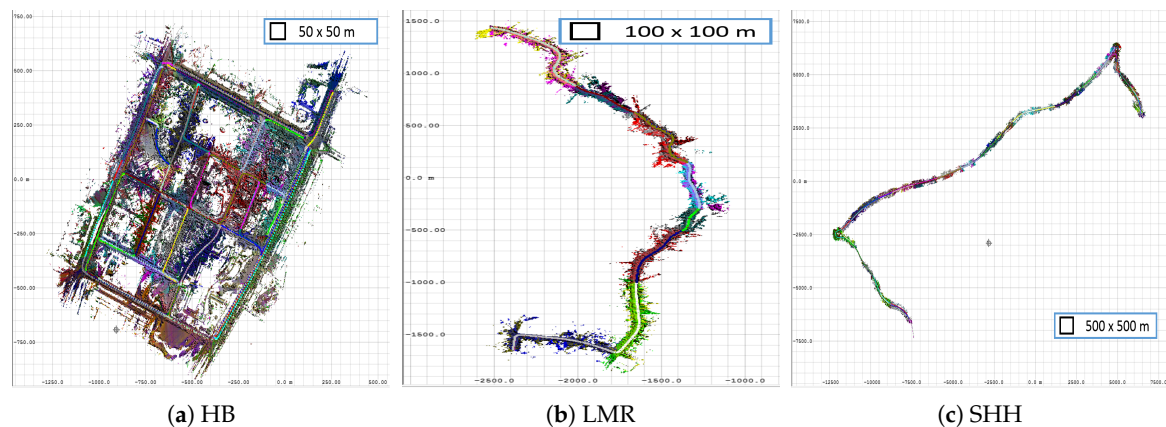


Figure 9. The three datasets used in experiments.

5.3. Verification Experiment and Discussion

We designed a verification experiment to verify that the recognizable result of our algorithm is in accordance with a real street scene. In this experiment, we calibrated the V8 phone to obtain the intrinsic parameters matrix. The extrinsic parameters matrix of V8 phone is computed by its position and drive direction in the point clouds. Then, we obtained an image of the camera in the point cloud scenes from the position of photo. This image contain the same scenario with the photo at same position. One photo has a corresponding image. The comparative result of them is shown in Figure 10. This figure shows that the recognizable result of our model consistent with the image. Additionally, in this example, we see the advantage of computing visibility and recognizability based on point clouds. Because the traffic sign is almost completely occluded, it is impossible to detect a traffic sign using image based methods. More results are shown in Figure 11. In this figure, The first row: photos taken at each viewpoint. The second row: the point cloud at the same viewing angle with each photo. It includes traffic sign (yellow) and traffic sign surrounding point cloud (green). The third row: occlusion detection result. the red points are occlusion points. O means occlusion area ratio and R means recognizability. From this figure, we can see the visual recognizability of traffic sign is inversely proportional to the occlusion ratio. This value is basically in line with human cognition.



Figure 10. Comparison of results between camera photo and point clouds image and between our method and the HPR method at the same viewpoint with camera. (a) Camera photo; (b) Point clouds image; (c) Occlusion detection by our method; (d) Occlusion detection by the HPR method [24].

We compared our method with the recent work of Huang et al. [24] on occlusion detection in this example. They use the HPR method to detect occlusion of traffic signs. Their last step is to obtain the visible point cloud by computing the convex hull of point clouds. Due to using this step, some occluded points are lost. The results of comparing our method with theirs are shown in Figure 10c,d. Our results extract all the occluding points. Thus the HPR method lost most of the occluding points.

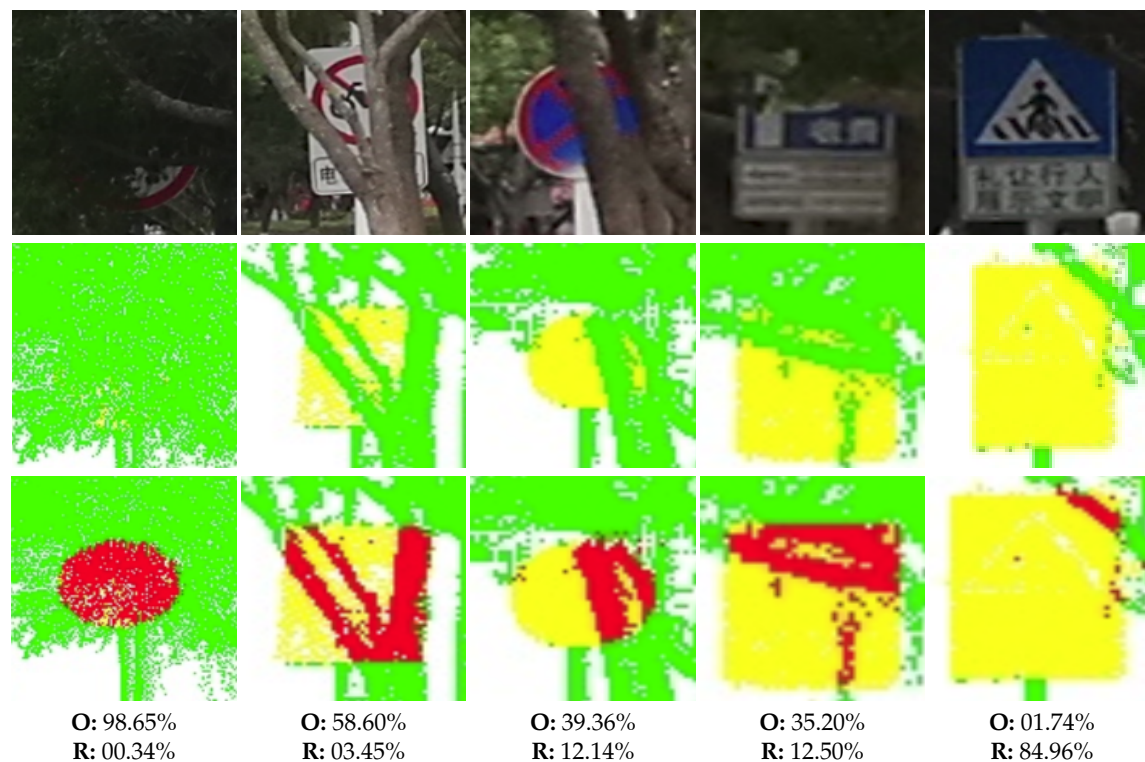


Figure 11. The visual recognizability results of traffic signs under occlusion.

5.4. Accuracy Analysis and Reliability Analysis

Does different point density and noise of point clouds impact the accuracy of the retinal imaging area in our model? Do challenging data, such as crossroads, street corners, roundabouts and mountain road impact the reliability of our algorithm? We analysis accuracy and reliability in the following two Sections:

5.4.1. Accuracy Analysis

In LiDAR, some studies [60–62] rely on accuracy problem, because scan patterns produce different footprints, which, depending on the area, leads to different accuracies. MLS has a small footprint, which provides high point cloud density on the road [63]. The accuracy analysis consists three parts. The first part is occlusion accuracy analysis based on image verification experiment. The second part is a experiment which compare our method with Huang et al. [24]. The last part is alpha shape algorithm experiment to analyze the influence by its parameter on accuracy.

Two metrics are used in accuracy analysis: accuracy and occlusion detection rate. The accuracy equals to the occlusion area ratio of a traffic sign obtained by the algorithm divided by its ground truth occlusion area ratio. It is closely related to the accuracy of visibility and recognizability calculations. The occlusion detection rate is the number of detected occlusion points in surrounding point clouds divided by the ground truth number of those points. It reflects the completeness of the detected occlusion points. Complete detection of the occlusion points helps to manual maintain traffic signs.

The accuracy analysis based on image verification experiment. We use the ratio of occlusion area in the image as the ground truth to evaluate the accuracy of the corresponding ratio of occlusion area computed in point clouds. Based on image verification data, we manually segment the traffic signs and occluded areas of the sign in the photo and count their number of pixels first. Then the ratio of occlusion area of traffic signs are computed by their number of pixels as the ground truth of the accuracy metric. An example of our method to evaluate accuracy is shown in Figure 12. In this figure, there are some images inside in some boxes. The image in middle box is the side view of a point clouds environment. The images inside the box in the upper left corner are the front and top views of

the point clouds environment. The images inside the boxes in other three corners are the observed traffic sign images at three viewpoints. The observed traffic sign images include the original image, the segmented image, the image generated by the point clouds and the image of boundaries calculated in point clouds. From Figure 12, we can see the computed boundaries of the sign and occlusion area are basically in line with the boundaries in original image. The detailed computed results of the three viewpoints are shown in Table 1. The sign area and occlusion area computed in point clouds are the retina imaging area. Their units are 10^{-4} square meters. From this table, we can see that our algorithm can achieve about 95% accuracy in computing occlusion ratio. For all 144 photos in the verification experiment, the accuracy of our algorithm is 94.89%. This experiment proves that the accuracy of our algorithm is in line with the actual situation.

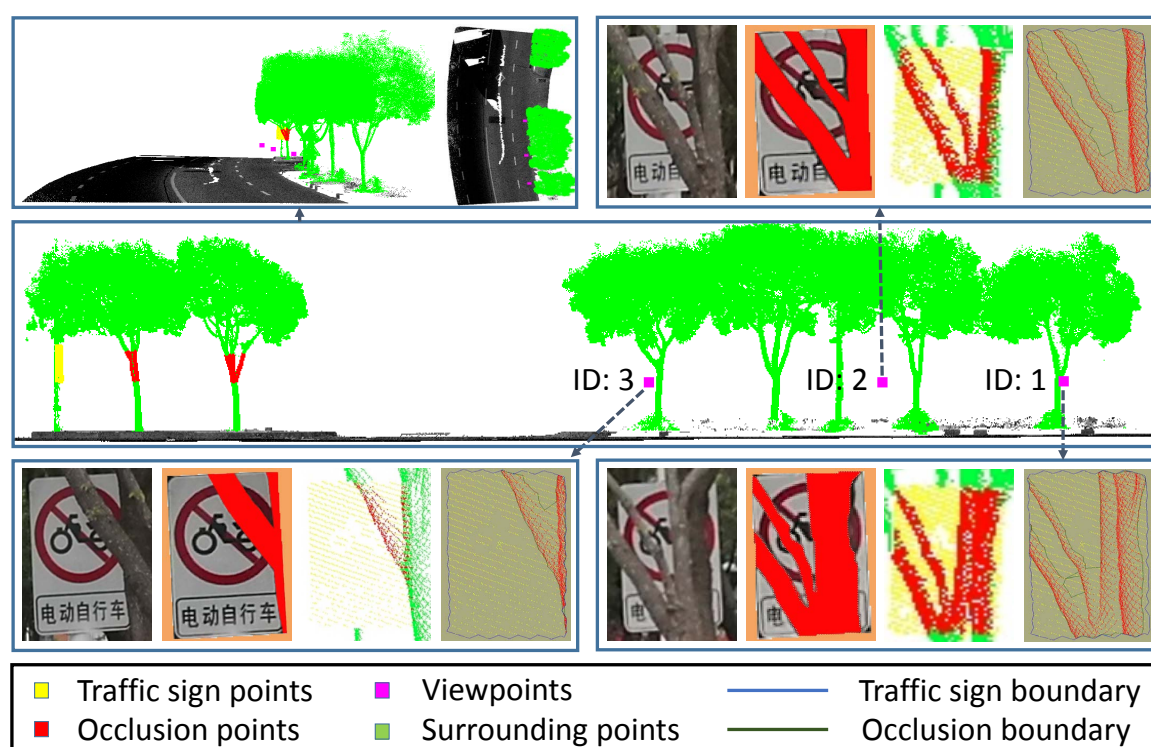


Figure 12. The example of evaluating the occlusion ratio accuracy in a real environment.

Table 1. The occlusion ratio accuracy calculation result of viewpoints in Figure 12.

Viewpoint ID	Image			Point Clouds			Accuracy (%)
	Sign Pixels	Occlusion Pixels	Ratio (%)	Sign Area	Occlusion Area	Ratio (%)	
1	6917	3971	57.41	2.76	1.50	54.48	94.89
2	9738	4285	44.00	3.44	1.42	41.39	94.06
3	18,477	3170	17.16	7.70	1.27	16.48	96.04

The comparison experiment. A comparison of the accuracy of computing a retinal imaging area with our method and the HPR method used in Reference [24] is shown in Figure 13. The point clouds used in the experiment were generated by uniform distribution according to Figure 13a. Shown in this figure are side views of the traffic sign (yellow) and occluding objects (red). The traffic sign center, occluding object centers and the viewpoint, are collinear. The traffic sign and three occlusion objects are square and parallel with each other. We generated experimental data at different densities changed from 101×101 to 11×11 points per square meter. According to normal distribution, noise points were added to the bounding box of the traffic sign point cloud and viewpoint. Before calculating, radius

filtering was applied to all point clouds. In Figure 13b,c the HPR parameter is set to five. The density of the point clouds in Figure 13c,d is set to 51×51 points per square meter.

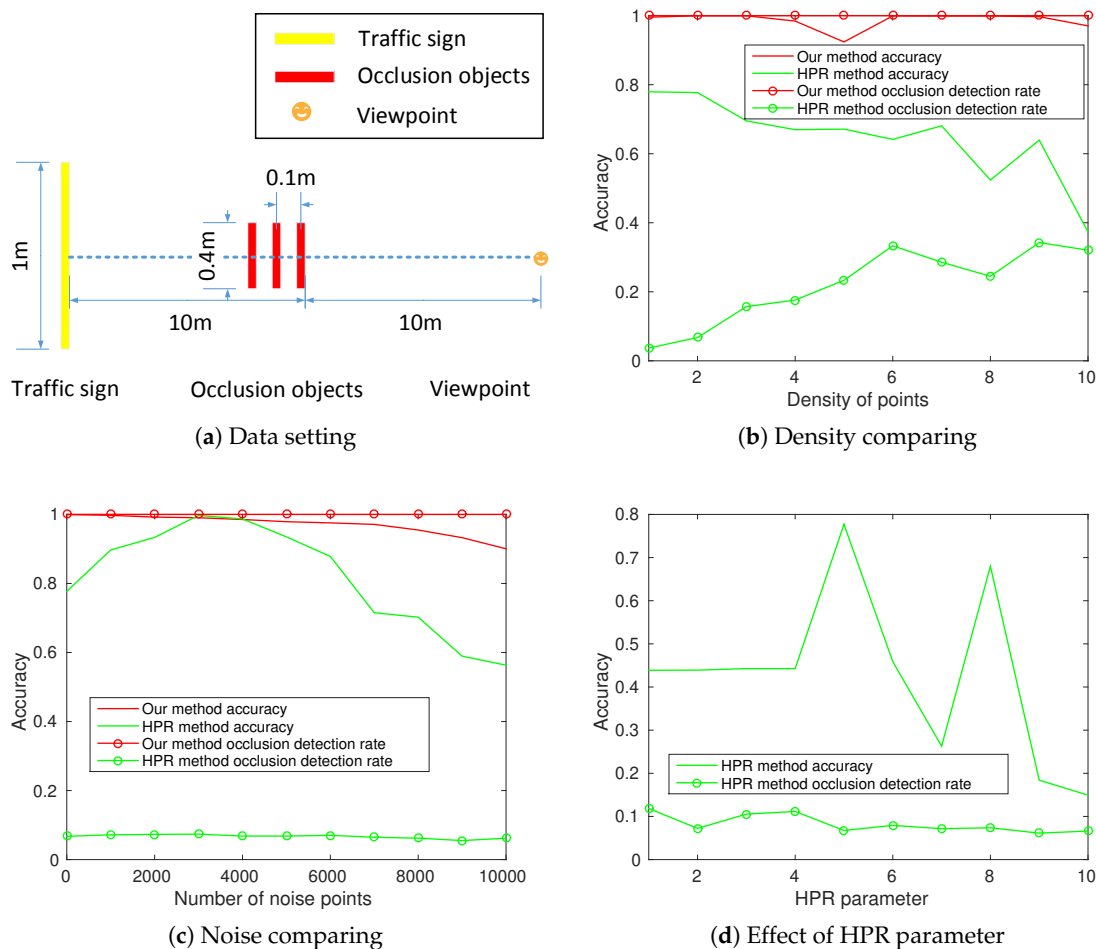


Figure 13. Comparison of data setting and accuracy results for different conditions for both our method and the HPR method. In (b), the density of points is expressed by the distance (cm) of the neighbor points.

Figure 13b shows our method is stable to the density of point clouds. Our method achieves an accuracy of at least 92.4% (average 98.68%) when density changes from 101×101 (points' interval 1 cm) to 11×11 (points' interval 10 cm) points per square meter. When the density changes, the fluctuation in the accuracy is not large. Thus, the accuracy of the HPR methods decreases with decreasing density, is unstable. For occlusion detection rate, our method detects all the occluded points and is not affected by the density; the HPR method lost most occluded points and is affected by the density. Figure 13c shows that the effect of noise on the two methods when noise points less than ten thousand (500 noise points per cubic meter). In the data preprocessing of our method, a filtering operation is included. The filtering condition is there must be six points within the sphere with a radius of ten cm. To be fair, we used the same filtering operation before using HPR calculations. From the figure, we can see, when the number of noise points increases, our method is stable and has high accuracy and occlusion detection rate. HPR method is unstable. Therefore the discrete points floating in the air will affect the evaluation results in the practical application. Figure 13d shows that the retinal imaging area changes as the HPR parameter changes. The HPR parameter seriously affects the result. Thus, the HPR method is unsuitable for calculating the retinal imaging area.

The alpha shape algorithm parameter experiment. In the area calculation of a traffic sign or the occlusion region, the outer boundary is firstly derived using the alpha shape algorithm. However, the different choices of the parameters of the algorithm will have a certain impact on the calculation of the boundary, especially when the object is of non-convex characteristic. Therefore, we designed another experiment for evaluating its influence on accuracy metric. In this experiment, we replace the occlusion objects in Figure 13a with square concave point clouds and half circular concave point clouds. The other setting includes: the side of the square is 0.2 m long, the radius of the circle is 0.15 m and the interval between points is 1 cm. The result of this experiment is shown in Figure 14. From this figure, we can see the accuracy decreases as the parameter of alpha shape becomes larger. Both of concave point clouds achieved almost best accuracy, 99.67%, when the value of the alpha shape parameter equals to 2, that is, twice of point density. The parameter of alpha shape is set to two times the density of point clouds in our practical applications, too. Images at the upper row in Figure 14b are shown the boundary of square concave point clouds when alpha shape parameters from one to ten. Images at the down row in Figure 14b are shown the boundary of half circular concave point clouds when alpha shape parameters equal to two and eight, respectively.

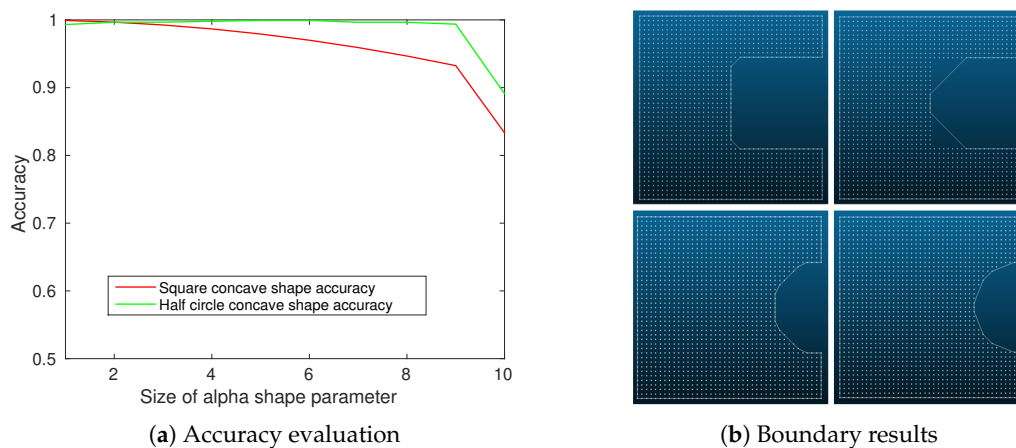


Figure 14. The accuracy evaluation under different alpha shape parameters.

5.4.2. Reliability Analysis

In crossroads, street corners, roundabouts, and mountain roads, the traffic environments are complicated. Road curvature change is large and no road markings in the middle of crossroads and corners, thus increasing the difficulty of viewpoint selection. Besides, severe occlusion are most likely to occur in these areas.

Figure 15 shows that our algorithm is still reliable in challenging data of crossroads, street corners, roundabouts. In this figure, the color is expressed as follows: the detected traffic sign (yellow), occluding point clouds (red), visibility field and recognizability field results (mesh planes). The box marked with cross ("X") means that the traffic sign is occluded in that area. The color of the mesh planes, changing from green to red, means that the values of visibility or recognizability change from big to small. The visual recognizability of this three traffic sign are 0.977, 0.547, 0.906, respectively. From this figure, we can intuitively observe the visibility field and visual recognizability of traffic sign. In visibility field, the visibility equals one only at "unit viewpoint", the farther the viewpoint is from the traffic sign, the smaller the visibility value. So the visibilities are shown mostly by red. In the visual recognizability field, only positions with occlusion and positions with large sight line deviation are shown by red. In addition, we have actual GPS coordinates in areas with low visual recognizability. It is convenient for road management departments to discover and optimize traffic signs with problems (with low visual recognizability).

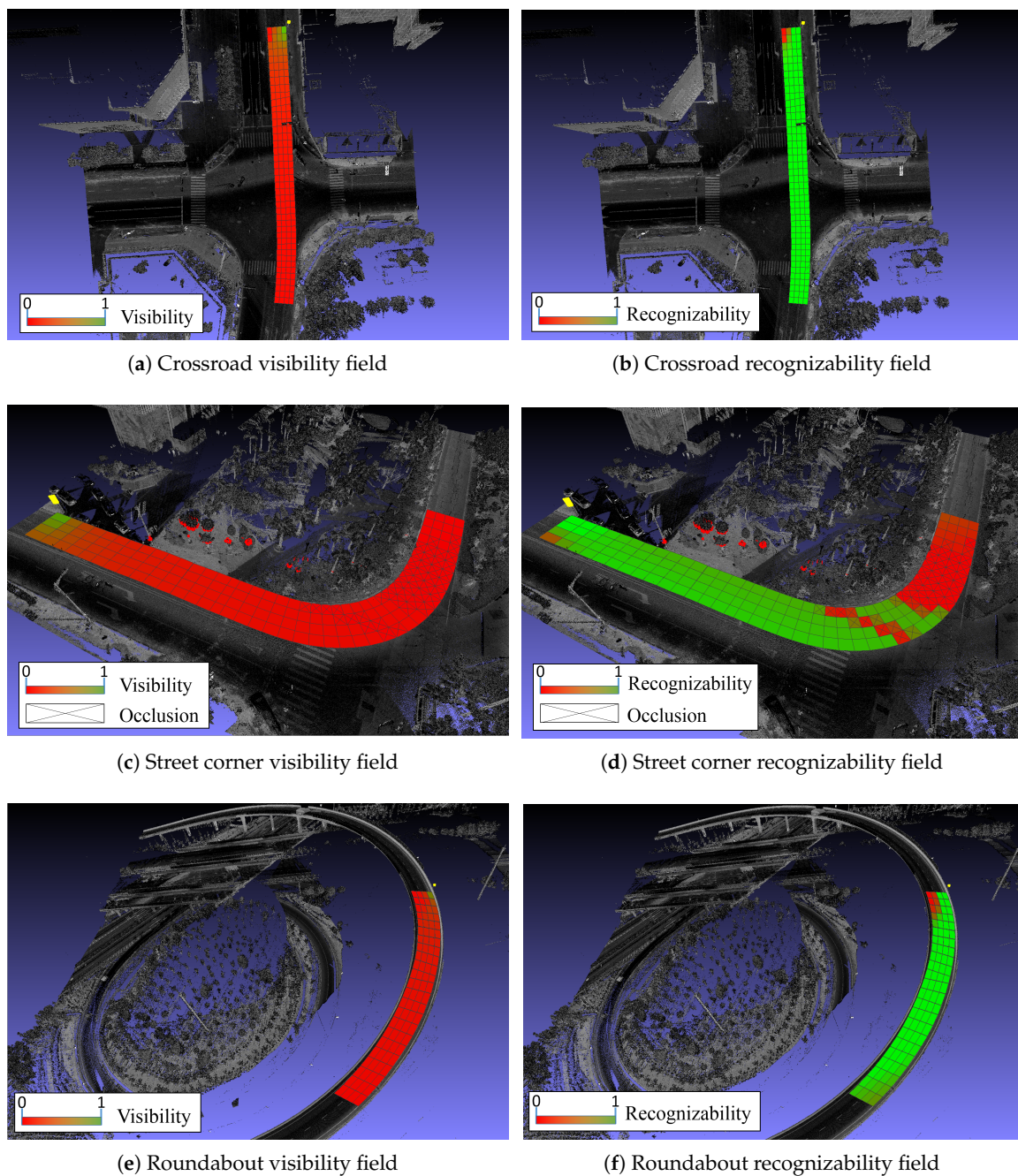


Figure 15. The TSVREM model applies on challenging data.

5.5. Large-Scale Application Experiment and Discussion

To verify its usability, we applied the TSVREM model to a large-scale traffic environment. The proposed traffic sign visual recognizability model was implemented using C++ running on an Intel (R) Core (TM) i5-4460 computer. The Figure 16 shows the large-scale application of our algorithm on mountain road. The surrounding is shown by color point clouds. The meaning of color representation of the traffic sign visual recognizability fields is the same with Figure 15. The computing times for each processing step in the three datasets are given in Table 2, which shows that the computed speed is fast enough to meet the engineering application demand of off line processing. The biggest time cost is the viewpoints selection progress.

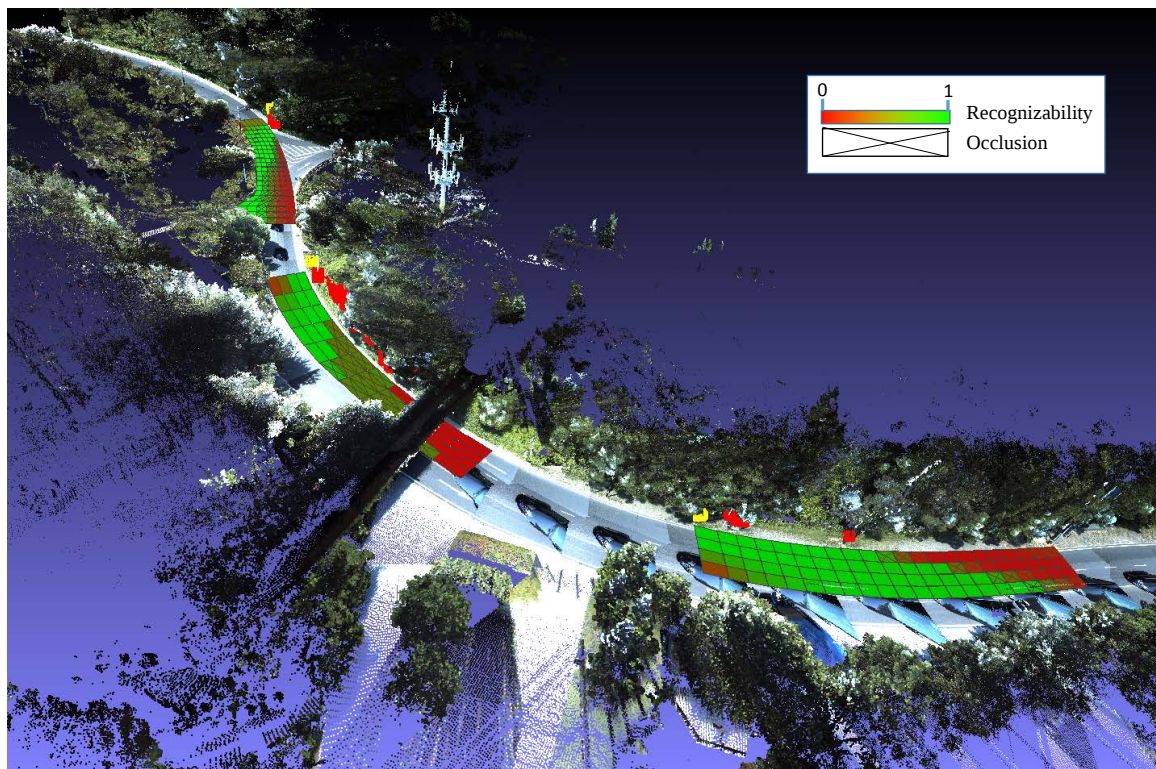


Figure 16. The large-scale application of TSVERM model on mountain road.

Table 2. The time calculation cost of datasets.

Dataset	Length (km)	Speed Limit (km/h)	Steps (min)			Average Cost (m/min)
			Sign Detection	Viewpoints Selection	Recognizability Evaluation	
HB	24.12	30	188.67	477.62	8.63	35.74
LMR	9.49	40	96.25	133.63	5.73	40.29
SHH	62.48	120	177.18	215.00	19.75	151.68

The statistical results of TSVREM model implemented in three datasets are shown in Table 3. In Table 3, “m/min” is the average length (meters) of the road calculated per minute. This table shows that about three quarters of the traffic signs have occlusion in three datasets. The average values of occlusion area ratio and recognizability about three datasets are about fifteen percent and seventy-five percent, respectively. Xiamen city is located in the south of China, where plants perennially grow densely. So far, there is no effective detection method that can accurately detect occlusion covering all road surfaces within a sight distance range, leading to an increasingly serious occlusion of traffic signs, especially the mountain road. This shows that our algorithm is of great significance in the visibility and recognizability maintenance of traffic signs.

Table 3. The TSVREM model calculation result in three datasets.

Dataset	Traffic Signs' Number	Occlusion		Average (%)	
		Number	Ratio (%)	Occlusion Area Ratio	Recognizability
HB	135	101	74.81	10.97	78.43
LMR	73	61	83.56	20.73	62.61
SHH	127	90	70.87	12.71	80.51
Total	335	252	75.22	14.80	73.85

6. Conclusions

This paper, based on human visual recognition theory and point clouds collected by an MLS system, presented a new way to evaluate traffic sign visual recognizability in each lane. The proposed model not only quantitatively expresses the visibility and recognizability of a traffic sign from a viewpoint, but also continuously expresses visibility and recognizability, within sight distance, over the entire road surface. Unlike the existing methods for studying visibility and recognizability limited by position of viewpoint in 2D space or cannot be applied in the real road environment, we proposed a new way to evaluate visibility and recognizability in 3D space conquered those problem. Based on traffic sign detection method [31], our algorithm can automatically process more than 90% (92.61% in [31]) traffic signs. The rest traffic signs can be manually detected and processed by our algorithm. Our methods also can be used to detect occlusion and inspect spatial installation information of traffic signs for inventory purposes. Moreover, our model, because it has a process similar to traffic signs, can be easily expanded to other traffic devices, such as traffic lights.

Author Contributions: Conceptualization, S.Z. and C.W. (Cheng Wang); Methodology, S.Z.; Software, S.Z.; Validation, S.Z., C.W. (Chenglu Wen) and L.L.; Formal Analysis, Z.Z.; Investigation, C.Y.; Resources, S.Z.; Data Curation, L.L.; Writing-Original Draft Preparation, S.Z.; Writing-Review & Editing, S.Z.; Visualization, C.W. (Chenglu Wen); Supervision, J.L.; Project Administration, J.L.; Funding Acquisition, C.W. (Cheng Wang).

Funding: This research was funded by National Natural Science Foundation of China, grant number U1605254 and 61728206.

Acknowledgments: We thank the reviewers for their careful reading and valuable comments, which helped us to improve the manuscript. This work is partially supported by the Collaborative Innovation Center of Haixi Government Affairs Big Data Sharing and Cloud Services.

Conflicts of Interest: The authors declare no conflict of interest.

Abbreviations

η	The punishment weight of sight line deviation
γ	The weight of other factors
λ	The penalization weight of occlusion
A^D	The area of driving area within sight distance of a traffic sign
A_{type}^{unit}	The retinal imaging area of a standard traffic sign viewed from the “unit viewpoint”
$A_{w,l}^{occ}$	The occluded imaging area of retinal of a sign from viewpoint $p(w, l)$
$A_{w,l}^{view}$	The imaging area of retinal of a sign from viewpoint $p(w, l)$
D	The driving area within sight distance of a traffic sign
E^{sign}	The visual recognizability of a traffic sign
$E_{w,l}^{geoS}$	The estimated value of standard geometric factors
$E_{w,l}^{geo}$	The estimated value of geometric factors
$E_{w,l}^{occ}$	The estimated value of occlusion factors
$E_{w,l}^{sightS}$	The estimated value of standard sight line deviation factors
$E_{w,l}^{sight}$	The estimated value of sight line deviation factor
$E_{w,l}^{visibility}$	The estimated value of viewpoint visibility
$E_{w,l}^{visibilityS}$	The estimated value of standard viewpoint visibility
G	The value of GFOV
l	The length from viewpoint to traffic sign along road direction
$p(w, l)$	A viewpoint
V	The maximum angle between sight line and the line from viewpoint to the “A” pillar of the vehicle
$v_{w,l}$	The angle between sight line and line from viewpoint to the center of traffic sign
w	The width from viewpoint to roadside in vertical road direction
$w^{driving}$	The driving area width of the road
GFOV	Geometric Field Of View
HPR	Hidden Point Removal
MLS	Mobile Laser Scanning]
SD	Sight Distance
TSVREM	Traffic Sign Visual Recognizability Evaluation Model
VEM	Visibility Evaluation Model

References

1. Liu, B.; Wang, Z.; Song, G.; Wu, G. Cognitive processing of traffic signs in immersive virtual reality environment: An ERP study. *Neurosci. Lett.* **2010**, *485*, 43–48. [[CrossRef](#)] [[PubMed](#)]
2. Kirmizioglu, E.; Tuydes-Yaman, H. Comprehensibility of traffic signs among urban drivers in Turkey. *Accid. Anal. Prev.* **2012**, *45*, 131–141. [[CrossRef](#)] [[PubMed](#)]
3. Ben-Bassat, T.; Shinar, D. The effect of context and drivers' age on highway traffic signs comprehension. *Transp. Res. Part Traffic Psychol. Behav.* **2015**, *33*, 117–127. [[CrossRef](#)]
4. Mourant, R.R.; Ahmad, N.; Jaeger, B.K.; Lin, Y. Optic flow and geometric field of view in a driving simulator display. *Displays* **2007**, *28*, 145–149. [[CrossRef](#)]
5. Belaroussi, R.; Gruyer, D. Impact of reduced visibility from fog on traffic sign detection. In Proceedings of the 2014 IEEE Intelligent Vehicles Symposium Proceedings, Dearborn, MI, USA, 8–11 June 2014; pp. 1302–1306.
6. Rogé, J.; Pébayle, T.; Lambilliotte, E.; Spitzenstetter, F.; Giselbrecht, D.; Muzet, A. Influence of age, speed and duration of monotonous driving task in traffic on the driver's useful visual field. *Vis. Res.* **2004**, *44*, 2737–2744. [[CrossRef](#)] [[PubMed](#)]
7. Costa, M.; Simone, A.; Vignali, V.; Lantieri, C.; Bucchi, A.; Dondi, G. Looking behavior for vertical road signs. *Transp. Res. Part Traffic Psychol. Behav.* **2014**, *23*, 147–155. [[CrossRef](#)]
8. Lyu, N.; Xie, L.; Wu, C.; Fu, Q.; Deng, C. Driver's cognitive workload and driving performance under traffic sign information exposure in complex environments: A case study of the highways in China. *Int. J. Environ. Res. Public Health* **2017**, *14*, 203. [[CrossRef](#)]
9. Motamedi, A.; Wang, Z.; Yabuki, N.; Fukuda, T.; Michikawa, T. Signage visibility analysis and optimization system using BIM-enabled virtual reality (VR) environments. *Adv. Eng. Inform.* **2017**, *32*, 248–262. [[CrossRef](#)]
10. Li, L.; Zhang, Q. Research on Visual Cognition About Sharp Turn Sign Based on Driver's Eye Movement Characteristic. *Int. J. Pattern Recognit. Artif. Intell.* **2017**, *31*, 1759012. [[CrossRef](#)]
11. González, Á.; Garrido, M.Á.; Llorca, D.F.; Gavilán, M.; Fernández, J.P.; Alcantarilla, P.F.; Parra, I.; Herranz, F.; Bergasa, L.M.; Sotelo, M.Á.; et al. Automatic traffic signs and panels inspection system using computer vision. *IEEE Trans. Intell. Transp. Syst.* **2011**, *12*, 485–499. [[CrossRef](#)]
12. Doman, K.; Deguchi, D.; Takahashi, T.; Mekada, Y.; Ide, I.; Murase, H.; Sakai, U. Estimation of traffic sign visibility considering local and global features in a driving environment. In Proceedings of the 2014 IEEE Intelligent Vehicles Symposium Proceedings, Dearborn, MI, USA, 8–11 June 2014; pp. 202–207.
13. Khalilikhah, M.; Heaslip, K. Analysis of factors temporarily impacting traffic sign readability. *Int. J. Transp. Sci. Technol.* **2016**, *5*, 60–67. [[CrossRef](#)]
14. Balsa-Barreiro, J.; Valero-Mora, P.M.; Berné-Valero, J.L.; Varela-García, F.A. GIS Mapping of Driving Behavior Based on Naturalistic Driving Data. *Isprs Int. J. Geo-Inf.* **2019**, *8*, 226. [[CrossRef](#)]
15. Balsa-Barreiro, J.; Valero-Mora, P.M.; Montoro, I.P.; García, M.S. Geo-referencing naturalistic driving data using a novel method based on vehicle speed. *IET Intell. Transp. Syst.* **2013**, *7*, 190–197. [[CrossRef](#)]
16. Sun, L.; Yao, L.; Rong, J.; Lu, J.; Liu, B.; Wang, S. Simulation analysis on driving behavior during traffic sign recognition. *Int. J. Comput. Intell. Syst.* **2011**, *4*, 353–360. [[CrossRef](#)]
17. Li, N.; Busso, C. Predicting perceived visual and cognitive distractions of drivers with multimodal features. *IEEE Trans. Intell. Transp. Syst.* **2015**, *16*, 51–65. [[CrossRef](#)]
18. Bohua, L.; Lishan, S.; Jian, R. Driver's visual cognition behaviors of traffic signs based on eye movement parameters. *J. Transp. Syst. Eng. Inf. Technol.* **2011**, *11*, 22–27.
19. Doman, K.; Deguchi, D.; Takahashi, T.; Mekada, Y.; Ide, I.; Murase, H.; Tamatsu, Y. Estimation of traffic sign visibility toward smart driver assistance. In Proceedings of the 2010 IEEE Intelligent Vehicles Symposium (IV), San Diego, CA, USA, 21–24 June 2010; pp. 45–50.
20. Doman, K.; Deguchi, D.; Takahashi, T.; Mekada, Y.; Ide, I.; Murase, H.; Tamatsu, Y. Estimation of traffic sign visibility considering temporal environmental changes for smart driver assistance. In Proceedings of the 2010 IEEE Intelligent Vehicles Symposium (IV), Baden-Baden, Germany, 5–9 June 2011; pp. 667–672.
21. Katz, S.; Tal, A.; Basri, R. Direct visibility of point sets. *Acm Trans. Graph. (TOG) ACM* **2007**, *26*, 24. [[CrossRef](#)]
22. Katz, S.; Tal, A. Improving the visual comprehension of point sets. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Portland, OR, USA, 23–28 June 2013; pp. 121–128.
23. Katz, S.; Tal, A. On the Visibility of Point Clouds. In Proceedings of the IEEE International Conference on Computer Vision, Santiago, Chile, 7–13 December 2015; pp. 1350–1358.

24. Huang, P.; Cheng, M.; Chen, Y.; Luo, H.; Wang, C.; Li, J. Traffic sign occlusion detection using mobile laser scanning point clouds. *IEEE Trans. Intell. Transp. Syst.* **2017**, *18*, 2364–2376. [[CrossRef](#)]
25. Lillo-Castellano, J.; Mora-Jiménez, I.; Figuera-Pozuelo, C.; Rojo-Álvarez, J.L. Traffic sign segmentation and classification using statistical learning methods. *Neurocomputing* **2015**, *153*, 286–299. [[CrossRef](#)]
26. Li, H.; Sun, F.; Liu, L.; Wang, L. A novel traffic sign detection method via color segmentation and robust shape matching. *Neurocomputing* **2015**, *169*, 77–88. [[CrossRef](#)]
27. Qin, K.h.; Wang, H.Y.; Zheng, J.T. A unified approach based on hough transform for quick detection of circles and rectangles. *J. Image Graph.* **2010**, *1*, 109–115.
28. Greenhalgh, J.; Mirmehdi, M. Real-time detection and recognition of road traffic signs. *IEEE Trans. Intell. Transp. Syst.* **2012**, *13*, 1498–1506. [[CrossRef](#)]
29. Yuan, Y.; Xiong, Z.; Wang, Q. An incremental framework for video-based traffic sign detection, tracking, and recognition. *IEEE Trans. Intell. Transp. Syst.* **2017**, *18*, 1918–1929. [[CrossRef](#)]
30. Zeng, Y.; Xu, X.; Shen, D.; Fang, Y.; Xiao, Z. Traffic sign recognition using kernel extreme learning machines with deep perceptual features. *IEEE Trans. Intell. Transp. Syst.* **2017**, *18*, 1647–1653. [[CrossRef](#)]
31. Wen, C.; Li, J.; Luo, H.; Yu, Y.; Cai, Z.; Wang, H.; Wang, C. Spatial-related traffic sign inspection for inventory purposes using mobile laser scanning data. *IEEE Trans. Intell. Transp. Syst.* **2016**, *17*, 27–37. [[CrossRef](#)]
32. Yang, B.; Dong, Z.; Zhao, G.; Dai, W. Hierarchical extraction of urban objects from mobile laser scanning data. *Isprs J. Photogramm. Remote. Sens.* **2015**, *99*, 45–57. [[CrossRef](#)]
33. Lehtomäki, M.; Jaakkola, A.; Hyypä, J.; Lampinen, J.; Kaartinen, H.; Kukko, A.; Puttonen, E.; Hyypä, H. Object classification and recognition from mobile laser scanning point clouds in a road environment. *IEEE Trans. Geosci. Remote. Sens.* **2016**, *54*, 1226–1239. [[CrossRef](#)]
34. Wang, J.; Lindenbergh, R.; Menenti, M. SigVox—A 3D feature matching algorithm for automatic street object recognition in mobile laser scanning point clouds. *Isprs J. Photogramm. Remote. Sens.* **2017**, *128*, 111–129. [[CrossRef](#)]
35. Huang, J.; You, S. Pole-like object detection and classification from urban point clouds. In Proceedings of the 2015 IEEE International Conference on Robotics and Automation (ICRA), Seattle, WA, USA, 26–30 May 2015; pp. 3032–3038.
36. Golovinskiy, A.; Kim, V.G.; Funkhouser, T. Shape-based recognition of 3D point clouds in urban environments. In Proceedings of the 2009 IEEE 12th International Conference on Computer Vision, Kyoto, Japan, 29 September–2 October 2009; pp. 2154–2161.
37. Li, F.; Elberink, S.O.; Vosselman, G. Semantic labelling of road furniture in mobile laser scanning data. The International Archives of the Photogrammetry, Remote Sensing and Spatial Information Sciences, Wuhan, China, 18–22 September 2017.
38. Balsa-Barreiro, J.; Fritsch, D. Generation of visually aesthetic and detailed 3D models of historical cities by using laser scanning and digital photogrammetry. *Digit. Appl. Archaeol. Cult. Herit.* **2018**, *8*, 57–64. [[CrossRef](#)]
39. Yu, Y.; Li, J.; Wen, C.; Guan, H.; Luo, H.; Wang, C. Bag-of-visual-phrases and hierarchical deep models for traffic sign detection and recognition in mobile laser scanning data. *ISPRS J. Photogramm. Remote. Sens.* **2016**, *113*, 106–123. [[CrossRef](#)]
40. Soilán, M.; Riveiro, B.; Martínez-Sánchez, J.; Arias, P. Traffic sign detection in MLS acquired point clouds for geometric and image-based semantic inventory. *ISPRS J. Photogramm. Remote. Sens.* **2016**, *114*, 92–101. [[CrossRef](#)]
41. Tan, M.; Wang, B.; Wu, Z.; Wang, J.; Pan, G. Weakly supervised metric learning for traffic sign recognition in a LIDAR-equipped vehicle. *IEEE Trans. Intell. Transp. Syst.* **2016**, *17*, 1415–1427. [[CrossRef](#)]
42. Ai, C.; Tsai, Y.J. An automated sign retroreflectivity condition evaluation methodology using mobile LIDAR and computer vision. *Transp. Res. Part Emerg. Technol.* **2016**, *63*, 96–113. [[CrossRef](#)]
43. Lee, S.; Kweon, I.S.; Kim, J.; Yoon, J.S.; Shin, S.; Bailo, O.; Kim, N.; Lee, T.H.; Hong, H.S.; Han, S.H. VPGNet: Vanishing Point Guided Network for Lane and Road Marking Detection and Recognition. In Proceedings of the IEEE International Conference on Computer Vision (ICCV), Venice, Italy, 22 October–29 October 2017; pp. 1965–1973.
44. Ahmad, T.; Ilstrup, D.; Emami, E.; Bebis, G. Symbolic road marking recognition using convolutional neural networks. In Proceedings of the 2017 IEEE Intelligent Vehicles Symposium (IV), Los Angeles, CA, USA, 11–14 June 2017; pp. 1428–1433.

45. Guan, H.; Li, J.; Yu, Y.; Wang, C.; Chapman, M.; Yang, B. Using mobile laser scanning data for automated extraction of road markings. *ISPRS J. Photogramm. Remote. Sens.* **2014**, *87*, 93–107. [CrossRef]
46. Guan, H.; Li, J.; Yu, Y.; Ji, Z.; Wang, C. Using mobile LiDAR data for rapidly updating road markings. *IEEE Trans. Intell. Transp. Syst.* **2015**, *16*, 2457–2466. [CrossRef]
47. Soilán, M.; Riveiro, B.; Martínez-Sánchez, J.; Arias, P. Segmentation and classification of road markings using MLS data. *ISPRS J. Photogramm. Remote. Sens.* **2017**, *123*, 94–103. [CrossRef]
48. Yu, Y.; Li, J.; Guan, H.; Jia, F.; Wang, C. Learning hierarchical features for automated extraction of road markings from 3-D mobile LiDAR point clouds. *IEEE J. Sel. Top. Appl. Earth Obs. Remote. Sens.* **2015**, *8*, 709–726. [CrossRef]
49. Garber, N.J.; Hoel, L.A. *Traffic and Highway Engineering*; Cengage Learning: Boston, MA, USA, 2014.
50. Administration, F.H. Manual on Uniform Traffic Control Devices, 2009. Available online: https://mutcd.fhwa.dot.gov/pdfs/2009/pdf_index.htm (accessed on 1 May 2019).
51. Diels, C.; Parkes, A.M. Geometric field of view manipulations affect perceived speed in driving simulators. *Adv. Transp. Stud.* **2010**, *22*, 53–64.
52. for Transport, D. The Traffic Signs Regulations and General Directions 2016. 2016. Available Online: https://assets.publishing.service.gov.uk/government/uploads/system/uploads/attachment_data/file/523916/DfT-circular-01-2016.pdf (accessed on 1 May 2019).
53. Yang, J.; Liu, H. GB 5768-1999, Road Traffic Signs and Markings, 1999. Available Online: <http://www.gb688.cn/bzgk/gb/newGbInfo?hcno=A009EE301906F810B586264BDA029FD3> (accessed on 1 May 2019).
54. Byers, S.; Raftery, A.E. Nearest-neighbor clutter removal for estimating features in spatial point processes. *J. Am. Stat. Assoc.* **1998**, *93*, 577–584. [CrossRef]
55. Kuipers, J.B.; others. *Quaternions and Rotation Sequences*; Princeton University Press: Princeton, NJ, USA, 1999; Volume 66.
56. Banks, J.H. *Introduction to Transportation Engineering*; McGraw-Hill: New York, NY, USA, 2002; Volume 21.
57. Edelsbrunner, H.; Kirkpatrick, D.; Seidel, R. On the shape of a set of points in the plane. *IEEE Trans. Inf. Theory* **1983**, *29*, 551–559. [CrossRef]
58. Kaiser, P.K. *The Joy of Visual Perception*; York University: York, UK, 2009.
59. Binghong Pan, Y.Z.; Liang, X. Application of dynamic vision theory in highway alignment design. *J. Chang. Univ. Nat. Sci. Ed.* **2004**, *24*, 20–24.
60. Ullrich, A.; Pfennigbauer, M. Noisy lidar point clouds: Impact on information extraction in high-precision lidar surveying. *Laser Radar Technology and Applications XXIII. Int. Soc. Opt. Photonics*, **2018**, 10636, 106360M.
61. Gargoum, S.; El-Basyouny, K. Effects of LiDAR Point Density on Extraction of Traffic Signs: A Sensitivity Study. *Transp. Res. Rec.* **2019**. [CrossRef]
62. Järemo Lawin, F.; Danelljan, M.; Shahbaz Khan, F.; Forssén, P.E.; Felsberg, M. Density Adaptive Point Set Registration. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Salt Lake City, UT, USA, 18 June 2018; pp. 3829–3837.
63. Kukko, A.; Kaartinen, H.; Hyypä, J.; Chen, Y. Multiplatform mobile laser scanning: Usability and performance. *Sensors* **2012**, *12*, 11712–11733. [CrossRef]



© 2019 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<http://creativecommons.org/licenses/by/4.0/>).