

NaviView: Virtual Mirrors for Visual Assistance at Blind Intersection

Fumihiko TAYA^{*1}, Kazuhiro KOJIMA^{*2}, Akihiko SATO^{*3},
Yoshinari KAMEDA^{*4} and Yuichi OHTA^{*5}

*Graduate School of Science and Engineering, University of Tsukuba^{*1,*2}*

*Graduate School of Systems and Information Engineering, University of Tsukuba^{*3}*

*Department of Intelligent Interaction Technologies, University of Tsukuba^{*4,*5}*

*Center for Computational Sciences, University of Tsukuba^{*4,*5}*

*{taya, kojima, sato}@image.esys.tsukuba.ac.jp^{*1,*2,*3}*

*{kameda, ohta}@iit.tsukuba.ac.jp^{*1,*2,*3}*

“NaviView” represents a concept of our visual assistance system that we have been working on for driving safety. We utilize surveillance cameras on roadsides in order to provide views of hazardous areas that drivers cannot see directly. In this paper, we propose a visualization method to show images of the surveillance cameras in a shape of a virtual traffic mirror. We examined three types of the virtual mirrors and we conducted experiments on a driving simulator to verify the effectiveness of the virtual mirrors. The results show that the virtual mirrors can shorten the time to recognize the other vehicle running from the blind area of an intersection.

Keywords: AHS, NaviView, Visual assistance, Video warping, Mixed reality, Blind intersection

1. Introduction

Traffic accidents at blind intersections are serious social problems. In the field of the ITS, various researches and developments are proceeding in order to avoid the traffic accidents. Especially, the AHS, which is the system to assist drivers for safety driving by the interaction between road infrastructures and vehicles, is putting the developed technology to practical use.

On the AHS, a large number of various sensors on roadsides are installed and used mainly for total traffic management. In addition to the traffic management, some of the sensors are useful for individual support. We focus on surveillance cameras (roadside cameras) because we think they are the most useful ones for that purpose.

As for researches using image sensors, object recognition by computer vision is one of the most popular technique in ITS literatures. Some researches reported that they use pattern recognition techniques and they achieved high recognition rates over 90 percents [1] [2]. However, in the ITS, just one failure on the system may cause a serious traffic accident.

Alarming for a failure may be a feasible solution for false-positive cases, but it means that drivers have to always be aware of alarms and it might increase the chance of accidents because alarms will degrade driver's concentration.

Therefore, we consider the social consensus currently stays within the opinion that object recognition in

images should be responsible for drivers and it should be done by themselves, not by computers.

As a solution of this problem, we expect that visual assistance systems which include no recognition process will be widely accepted. We have been proposing a series of “NaviView” concept that embodies visual assistance for drivers [3] [4] [5] [6] [7].

The main feature of NaviView is to utilize images of roadside surveillance cameras so as to enhance driver's view.

In the NaviView system, the images of blind areas are captured by roadside surveillance cameras, and an image warping technique that utilizes geometric transformation visualizes the blind areas. An in-vehicle device (e.g., a small monitor) presents the warped and processed images to a driver.

Since it is not easy for drivers to recognize what is shown in the images of the roadside cameras by just watching the images in a monitor within a short time, the images should be warped and visualized to fit the driver's view by aligning the images to the real world.

Therefore, it is important to carefully design the way of visualization of the roadside camera images. In mixed reality literatures, many researches have been proposed for integrating video camera images with real user's view (refer [8] [9] [10] to see some advanced research results). However, as they mainly intended to support collaborative works, virtual reality applications, and/or

telecommunications, their proposed methods are not useful for the driver support in ITS.

In order to let drivers recognize objects in blind areas by showing the images of those areas, it is necessary to visualize the blind areas with less psychological burden for driver's vision system.

To satisfy this constraint, we focus attention on traffic mirrors, and we propose to show images of roadside cameras in a shape of a virtual traffic mirror in this paper. The traffic mirror is the most common infrastructure for safety driving assistance at an intersection that has blind areas and people get used to seeing them. We think that the visualizing method that imitates traffic mirrors reduces the psychological burden of recognizing what is shown in the images. Drivers can naturally recognize objects in the blind areas through the virtual traffic mirror.

In the rest of this paper, we first discuss the properties of traffic mirror and implementation plans of virtual mirrors that exceed the function of normal traffic mirrors. Then we describe the method to warp and display images. After that, experimental results of our method are shown. They proved that the virtual mirrors are useful for individual vehicle support. In the last section, we conclude the proposed method.

2. Visual assistance by mirrors

2.1. Traffic mirror

Traffic mirrors are generally installed at blind intersections. Through the traffic mirrors, drivers can see the areas that they cannot see directly. The mirrors help the drivers to see the blind areas hidden by walls or buildings before they go into the intersection, and the delay of recognizing the objects in the blind areas is shortened. Since the recognition delay of objects in blind areas is a main factor of traffic accidents, traffic mirrors are considered to be useful and are widely used in the traffic environment. However, traffic mirrors still have dead zones, and the dead zones may cause traffic accidents. Therefore, the dead zones should be eliminated as much as possible on designing virtual mirrors.

2.2. Geometric transformation

To avoid traffic accidents, it is necessary to let drivers recognize objects in blind areas. As a method to speed up the recognition, we propose a "virtual mirror," a new visual assistance system, which warps images of surveillance cameras in a shape of a traffic mirror and shows them to drivers.

The images to be presented to drivers are reshaped by geometric transformation. Although the basic idea is to reconstruct an image of real traffic mirror from the images taken by the surveillance cameras, we do not aim to realize a complete simulation of a real traffic mirror. We rather propose new visual assistance systems that look like mirrors, which technically are not mirrors. There are several reasons to do that. The first reason is that a normal traffic mirror is not a perfect device as mentioned in the previous section. The second is that surveillance cameras are sometimes placed where it is almost impossible to provide sufficient video images in order to reconstruct an image of real traffic mirror due to optical geometry constraints. If we want to utilize visual information of blind areas from surveillance cameras efficiently, it is better not to follow the optical limitations that derive from the strict simulation of real traffic mirrors.

We list up possible forms of geometric transformation and discuss their merits and demerits. The form of geometric transformation can be classified into three forms.

- (1) No transformation (direct use of original images).
- (2) Transformation based on driver's viewpoint.
- (3) Transformation based on fixed viewpoint.

Form (1) directly uses the original image obtained by surveillance cameras and shows them on a square region in driver's view. This is the simplest method, and whole area in the image can be used. In other words, there is no loss of visual information. However, as the locations of surveillance cameras vary in intersections, sometimes it is difficult for drivers to understand the spatial relationship between the location of the camera and their current location in the intersection. It may cause the recognition delay.

Form (2) is a method to display images that seem to be taken from the viewpoint that moves as the driver's viewpoint move. An example is a virtual camera that is located several meters ahead of the vehicle and is directed right or left. This virtual camera is really useful to check crossing streets at blind intersections and it can shorten the recognition delay. However, it is usually difficult to find surveillance cameras that can provide sufficient image materials to synthesize the images of such virtual camera because the virtual camera moves with the vehicle's motion.

Form (3) is a method to display images as if they were taken at a certain fixed viewpoint at intersections, regardless of the actual location of surveillance cameras. In this method, it is not necessary to measure the viewpoint of the driver because the virtual viewpoint is fixed against intersections. It is important to determine

the appropriate location of the virtual viewpoint in intersections so as to shorten the recognition delay. For example, the virtual viewpoint should provide a wide view of blind roads in the intersection. Providing the image materials to synthesize the images of the virtual viewpoint in form (3) is relatively easy because the virtual viewpoint and the actual surveillance cameras are all fixed.

2.3. Display device

In virtual mirror systems, a display device should satisfy the following constraints.

- (a) Less eye movement of drivers.
- (b) Less burden of using/wearing the device.
- (c) Wider view angle so that drivers can see any directions.

A Windshield Display (WSD) is a device that can show images on a windshield. It meets the demands of (a), (b), and (c).

In order to display a virtual mirror at a designated location in the real world for a driver's viewpoint, the pose of driver's eyes in the car should be estimated precisely. In this paper, we assume the driver's viewpoint is estimated precisely by some other techniques.

3. Virtual mirrors

We propose three types of virtual mirrors; VM1, VM2, and VM3. The first one corresponds to form (1) while VM2 and its improved version VM3 correspond to the idea of form (3). Since the idea of form (2) requires a number of image sources as the user's vehicle moves, we do not think it is practical and we do not deal with it.

In order to explain the definitions of virtual mirrors in the following sections, we first prepare an intersection as an example. The layout of the intersection is shown in Fig. 1. A driver who receives the visual support of virtual mirrors is in the vehicle circled in the figure. The narrow road on which the vehicle stops is 2.8-meter width. The crossing street is a double lane and the width of road is 5.6 meters. The origin of the world coordinates is set on the road surface of the center of intersection O , and Y-axis is vertical to the road surface. There are ten real and virtual candidates of camera mounting point, and their positions and orientations are shown in Table 1. Six (① - ⑥) out of the ten points are exploited as real camera mounting points, while the rest four (% , ③ , # , \$) are for virtual cameras. When a camera is set at ① directing along X axis, that means the camera is looking right in Fig.1, it is denoted as "①-right" in this paper.

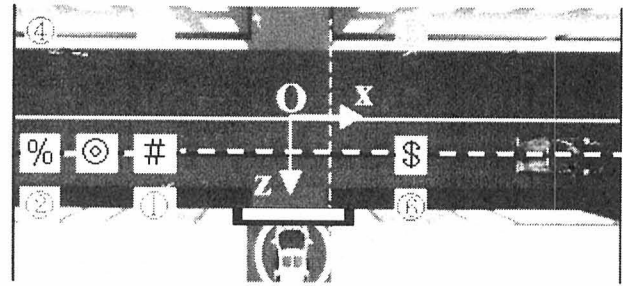


Fig.1. Locations of surveillance cameras

Table.1 Candidates of camera mounting point

Cam	Position (X,Y,Z)	Orientation Looking right/left
Real cameras		
①	(-5.0, 2.5, 2.8)	(0.250, -0.067, -0.966) / (-0.250, -0.067, -0.966)
②	(-10.0, 2.5, 2.8)	(0.250, -0.067, -0.966) / (-0.250, -0.067, -0.966)
③	(-5.0, 2.5, -2.8)	(0.250, -0.067, 0.966) / (-0.250, -0.067, 0.966)
④	(-10.0, 2.5, -2.8)	(0.250, -0.067, 0.966) / (-0.250, -0.067, 0.966)
⑤	(5.0, 2.5, -2.8)	(-0.250, -0.067, 0.966) / (0.250, -0.067, 0.966)
⑥	(5.0, 2.5, 2.8)	(-0.250, -0.067, -0.966) / (0.250, -0.067, -0.966)
Virtual cameras		
%	(-10.0, 2.5, 1.4)	(0.966, -0.259, 0.0) / (-0.966, -0.259, 0.0)
③	(-7.5, 2.5, 1.4)	(0.966, -0.259, 0.0) / (-0.966, -0.259, 0.0)
#	(-5.0, 2.5, 1.4)	(0.966, -0.259, 0.0) / (-0.966, -0.259, 0.0)
\$	(5.0, 2.5, 1.4)	(-0.966, -0.259, 0.0) / (0.966, -0.259, 0.0)

3.1. VM1: Monitor mirror

This method is based on the form (1). Original images obtained by surveillance cameras are directly displayed on the mirror area in a virtual mirror. This is technically the simplest method. Drivers will see the mirrors as if they see surveillance camera monitors on roads. To realize this, we set several assumptions shown below.

- (I) When a 3D point X of a virtual object is mapped to x on the image plane of the display device by projection P_D , drivers see the virtual object as if it is in the real world.

(II) Surveillance camera i takes images of an intersection by a perspective projection P_i .

When the assumptions (I) and (II) are satisfied, a point X is mapped to x_i on the image plane of a surveillance camera i by:

$$x_i = P_i X$$

Since we assume the point X exists on a certain surface, the relation between x and x_i can be expressed by a homography H_{VM1} .

$$x = H_{VM1} x_i = H_{VM1} P_i X$$

H_{VM1} is defined by P_D and the location of the virtual mirror in the real world. P_i is uniquely determined by the location of the surveillance camera i in the intersection, and P_D by the location of driving vehicle. In order to determine H_{VM1} , four corners of the mirror on the image plane of display device should be specified when the viewpoint of the driver is given.

Fig.2 is an example of VM1. In Fig.2, the image in the right mirror is an image obtained by a surveillance camera located at “#-right” in Fig.1, and the image in the left mirror is obtained by a surveillance camera located at “\$-left” in Fig.1. This figure shows a good realization of VM1 because it is usually rare to find cameras at a good location like # and \$. If there exist only roadside cameras like ① - ⑥, the virtual mirror VM1 might be less recognizable.

The advantage of VM1 is that the whole area of the original image is used and displayed. Dead zone of VM1 depends on the location and direction of the surveillance camera against the intersection.

3.2. VM2: Virtual mirror based on the virtually mounted camera

This is a method based on the form (3). This method synthesizes images whose viewpoint does not depend on the real locations of surveillance cameras. The virtual camera is placed at a certain point in the intersection. As

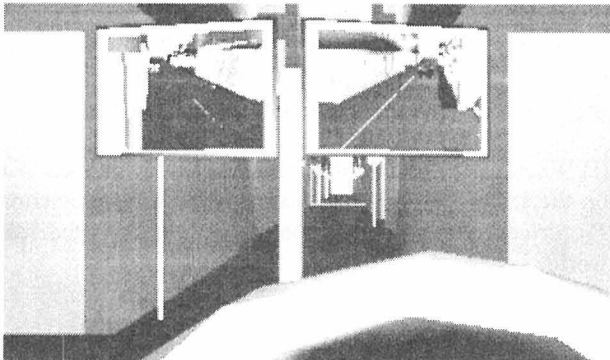


Fig.2. Virtual mirror (VM1)

a result, drivers can feel as if the virtual cameras are always mounted at the same location even when they proceed into different intersections. This helps drivers easily recognize the spatial relationship between the virtual camera and the driver himself.

To realize this method, we set assumptions (III) and (IV) in addition to (I) and (II).

(III) Objects in the real world exist on a flat road surface and they should be planar and infinitely thin.

(IV) Virtual surveillance camera V takes an image of the intersection by perspective projection P_v .

When all the assumptions are satisfied, a point x on the display device is computed from X by;

$$x = H_{VM2} x_i = H_{VM2} P_i X$$

H_{VM2} is a homography matrix defined by P_D , P_v , and the location of the virtual mirror. P_v is determined uniquely once after the location of a surveillance camera is given. H_{VM2} is determined by locating four points, which are on a same plane (e.g. corners of a rectangle on road surface) in the real world, on the image plane of the display device when driver's viewpoint is given.

One problem of this method is that the visible area may be smaller than that of VM1 depending on the relationship between P_i and P_v . The other problem is that target X needs to be on the specified flat surface. In other words, if there are 3D objects, e.g. vehicles, road signs, buildings, will be distorted in unrealistic way through the transformation.

Fig.3 shows an example of VM2. A virtual surveillance camera is virtually mounted above the center of the lane at ⑥ in Fig.2. The virtual camera is directed to take images of the intersection so that drivers can see other coming vehicles from blind directions. Note that the vanishing points of the lanes will always be found at a same position in the mirrors of VM2.

As mentioned before, the advantage of VM2 is that it can provide fixed unique view of intersections because the image is synthesized as if it is taken at a virtual fixed viewpoint. Drivers do not need to imagine the actual locations of real surveillance cameras. On the contrary, VM2 has a disadvantage that it may include black-out regions due to the possible lack of the image sources for the regions.

Fig.3 is the case that the image of a surveillance camera located at “①-right” in Fig.1 is reshaped to the image in the right mirror and the image of a surveillance

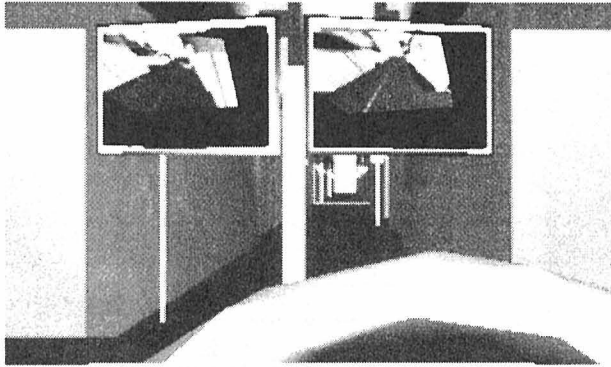


Fig.3. Virtual mirror

camera located at “⑥-left” in Fig.1 is reshaped to the left one. It means that a virtual surveillance camera of the right mirror is located at “②-right” in Fig.1. The black area in the mirror is the black-out region caused by missing image source.

3.3. VM3: Virtual mirror based on the virtually relocated camera

VM2 has an advantage of realizing a unique fixed view against intersections. However, it is inevitable to eliminate black-out regions in VM2 images because the viewing angle of real surveillance cameras is not wide and the viewpoint of the real camera is different from that of the fixed viewpoint of the virtual camera. The size of the black-out region of synthesized image depends on the difference between the poses of a real camera and a virtual camera. If the difference becomes larger, the black-out regions grow bigger.

Therefore, if there is a new virtual viewpoint that includes less black-out regions and yet it is easily located in drivers' mind, it will be better. Because the direction of the ideal virtual viewpoint should not be changed to keep its understandability, we take up only horizontal translation of the virtual camera so that the displacement between the virtual viewpoint and the real viewpoint becomes smaller.

VM3 is an improved version of VM2. The improved virtual mirror exploits a new position that is closer than that of VM1 to the real camera. This method sets a virtual surveillance camera at a location that is above the center of the lane and is the nearest location to a surveillance camera on the center of the lane (the white dashed line in Fig.1). Fig.4 is an example of VM3. A virtual surveillance camera corresponding to the real camera at “②-right” is placed at “②-right”, and its image is displayed in the right mirror in Fig.4. In the same manner, the virtual camera for “⑥-left” is placed at “⑥-left”.

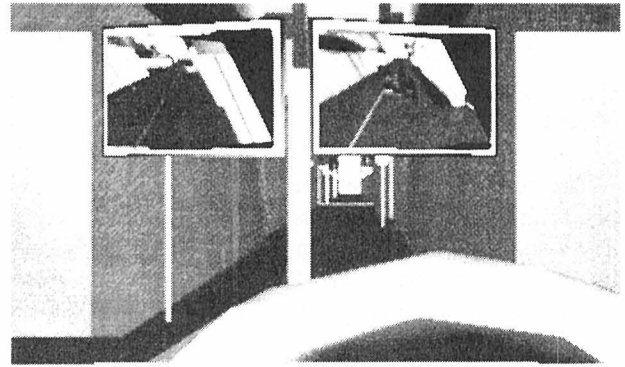


Fig.4. Virtual mirror (VM3)

As shown in Fig.4, VM3 can provide wider visible regions and less black-out regions in the virtual mirror because the virtual viewpoint is relatively close to the corresponding real surveillance camera. On the contrary, it might take extra time to recognize the spatial relationship of the images in the virtual mirror since the virtual viewpoint is not rigidly fixed against intersections. However, we think this disadvantage is not critical because the direction of the virtual viewpoint can be set constant and parallel to that of VM2. So far as the direction of the virtual cameras of VM3 is parallel and the mounting point moves on a line which is parallel to the direction, the synthesized images look similar.

4. Experiments

We conducted experiments on a driving simulator to examine the validity of our method. We use the driving simulator because we have to set all the subjects in equal condition.

In the experiments, we prepare a situation of a blind intersection where a vehicle is coming from the right without noticing the vehicle of subjects.

We implemented the simulator shown in Fig.5 and Fig.6. The system is composed of a simulator part and a windshield display part. The simulator part renders the blind intersection seen from driver's viewpoint by computer graphics. The rendered scene is projected on a screen by a projector. The windshield display part synthesizes a virtual mirror on a plasma display. A virtual mirror can be found on the plasma display device in Fig.5, which is marked with a circle. The mirror is reflected at the windshield because the windshield works as a half mirror. Subjects see the virtual mirror and the CG scene of the blind intersection simultaneously. As a virtual mirror is rendered as if it were in the fixed position during the experiment, the appearance of the virtual mirror is changed by the vehicle position. During the experiments, the head of subjects are fixed to a certain position.

As the intersection is constructed in computer graphics, output images of surveillance cameras are also synthesized according to the preset location and direction of the cameras.

The resolution of the screen is 800 pixels by 600 pixels, and its physical size is 2.6 meters by 2.0 meters. The resolution we use in the plasma display is 800 pixels by 600 pixels, and it has 42 inch display area. The virtual mirror is designed as 150cm wide and 100cm high.

The vehicle that subjects drive moves automatically to a stop line and stops there. The task of subjects is to examine the safety before they proceed into an intersection by confirming the existence of coming vehicle by their eyes as early as possible after the subject's vehicle restarts from the stop line. Subjects can move their vehicle forward by pressing a key on a keyboard. This keyboard operation corresponds to a driving maneuver. They can also look around the intersection by pressing another key. This operation corresponds to head rotation action in actual driving.

We measured the response time (T_r) for estimation. The response time is counted from the time when the subject's vehicle stops at the stop line by the time the subject presses a key on a keyboard when the subject finds the coming vehicle. When the subject's vehicle is at the stop line, a virtual mirror, that is placed on the other side of the road and is 5.6 meters away, will be rendered by 84 pixels in width on our driving simulator.

As T_r indicates the total time from the start at the stop line till the confirmation of the coming vehicle, we think it clearly shows what we can expect if the virtual mirrors are installed in ordinary traffic situations.

The experiments are conducted by subjects whose age is between 22 and 26. All of the subjects have a driver license. At first, they are told the concept and procedure of the experiment while they are receiving training. Then, the data collection phase is started. In

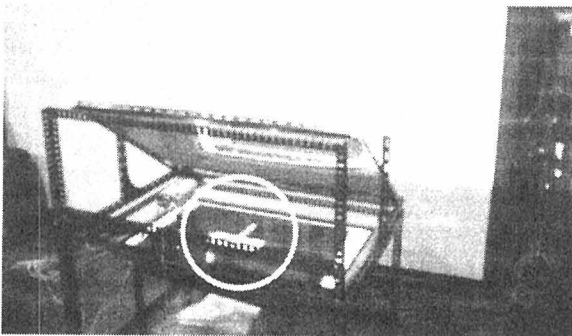


Fig.5. Snapshot of a driving simulator

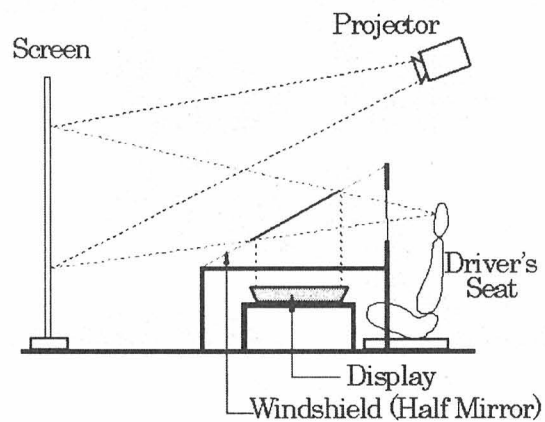


Fig.6. Driving simulator

this phase, we measure the response time for ten subjects for each situation. Then, a subject crosses six intersections those are in different situations each other. The distance between the intersections is 30 meter.

4.1 VM1 against no mirrors and normal mirrors

We measured the response time on using virtual mirrors, and collected the data by changing the time T_i that is the time until the front end of the body of coming vehicle reaches the intersection (white thick line in Fig.1). In the training phase, the time was set randomly from 0.0sec to 5.5sec. The time 0.0 sec means that the coming vehicle just arrives at the intersection when the subject's vehicle stops. The training phase ends when the subject claims he is well trained. T_i was pre-determined for each intersection. However, the subjects were told that the time was set randomly.

In this experiment, we collected the data for "No Mirror," "CG Traffic Mirror," and "VM1". "No Mirror" means there is no mirror in the intersection. "CG Traffic Mirror" means there are normal traffic mirrors shown in Fig.7. They optically simulate real traffic mirrors in CG. "VM1" means that there are virtual mirrors of VM1.

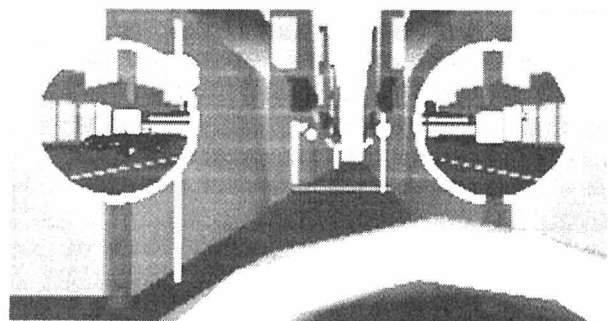


Fig.7. CG traffic mirrors

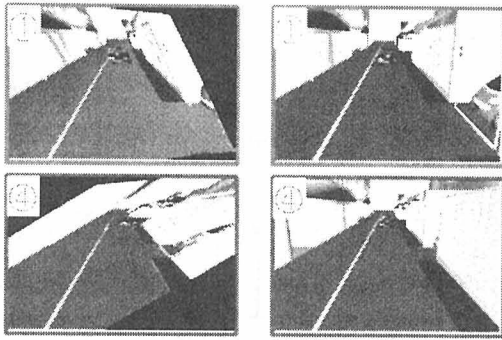


Fig.10. VM3(left) and VM3'(right)

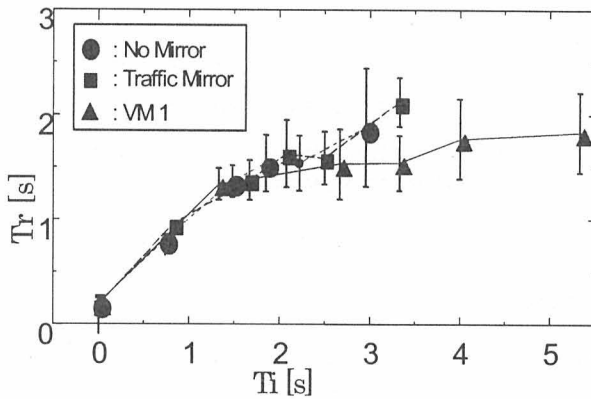


Fig.8. Response time by virtual mirrors

The procedure of the experiment is as follows.

1. Collect the data after training with no visual assistance.
2. Collect the data after training with CG traffic mirrors.
3. Collect the data after training with VM1.

Fig.8 shows the average of all subject's response time T_r for "No mirrors", "Traffic mirror", and "VM1". The horizontal axis indicates T_i , and the vertical axis indicates T_r . The bar on each element stands for its standard deviation.

T_r of VM1 does not become longer as T_i increases. This is because subjects could recognize the coming vehicle at far position by VM1. This means that the virtual mirror is useful to find vehicles far away from the intersection. On the other hand, there are little difference while T_i is small because the coming vehicle is near the intersection and it can be easily found without mirrors if subjects turn their view to their right slightly.

4.2. Location of surveillance cameras

In real traffic environment, the locations of surveillance cameras are different at intersections. If the original images taken by the cameras are directly shown

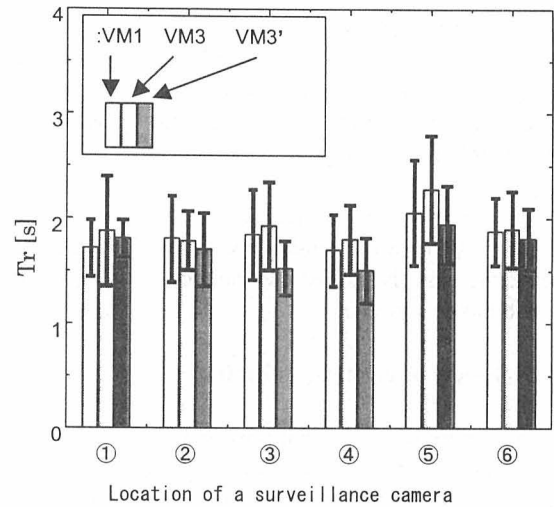


Fig.9. Response time by location of surveillance cameras

to drivers as a virtual mirror (VM1), it is sometimes difficult for drivers to understand the spatial relationship.

In order to evaluate the influence of location changes of the surveillance cameras, we installed surveillance cameras at different locations in the intersection and measured the response time. In training phase, the locations of surveillance cameras were set randomly among ① to ⑥ in Fig.2. The mark "%" indicates a location of virtual surveillance camera corresponding to ②④. The mark "\$" corresponds to ①③, and the mark "\$" corresponds to ⑤⑥. In the data collection phase, the location was changed in a pre-determined order. The subjects were told that the locations were randomly selected.

In this experiment, we collected the data of VM1, VM3, and VM3'. As for VM3', a real surveillance camera is set at the location where the virtual surveillance camera would be set in VM3. It means VM3' has no black-out region in its image.

The procedure of the experiment is as follows.

1. Collect the data after training with VM1.
2. Collect the data after training with VM3.
3. Collect the data after training with VM3'.

Fig.9 shows the average response time T_r of the three methods and the bar on each element stands for its standard deviation. From the result, we can say:

<1> The average response time of VM3' is always shorter than that of VM3.

<2> The average response time of VM1, VM3, and VM3' do not highly depend on the location of surveillance cameras.

<1> is thought to have two reasons. One is that VM3 has black-out regions in the mirror while VM3' does not. Fig.10 shows an example. The other is that since the real world does not meet the assumption (III), distortion makes it hard to recognize the objects in VM3.

As for <2>, the difference of response times was smaller than the standard deviation. It means that response time is almost independent of the location of surveillance cameras.

4.3. Speed of coming vehicle

We examined the relationship between the response time and the speed of the coming vehicle. We collected the data for different speed of the coming vehicle of 40km/h, 50km/h, and 60km/h. A virtual mirror of VM1 was set in the center of the road in front of the driver, and it is always displayed. Fig.11 shows the response time for each speed of the coming vehicle. Though a little difference can be seen, no salient difference exists. Thus, we can say that the speed of coming vehicles makes little influence on the response time.

Fig.12 shows the relationship between the distance of the coming vehicle and the response time. The horizontal axis indicates the distance of the coming vehicle at the time when the subject's vehicle stops at the stop line. The bar on each element stands for its standard deviation. When the distance is longer than a certain distance, the response time becomes constant and they are within 1.5 seconds.

4.4. Timing to display virtual mirror

Virtual mirrors can be displayed at arbitrary timing. If virtual mirrors show up too early, they may draw driver's attention too much and disturb their safe driving. For

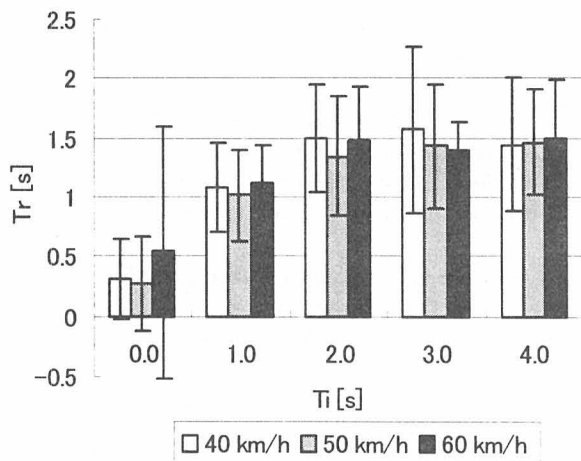


Fig.11. Response time by changing the speed of coming vehicle

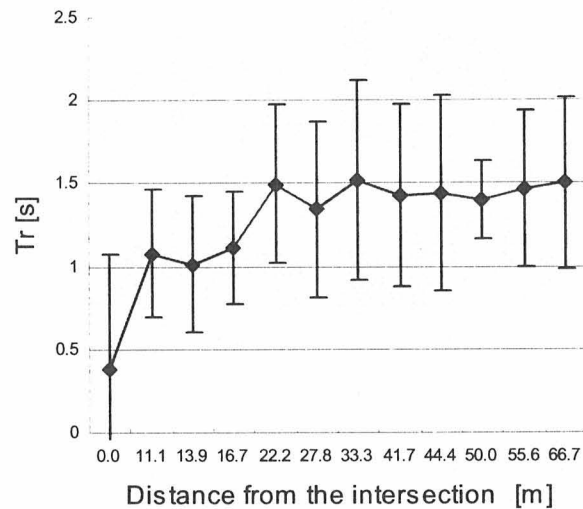


Fig.12. Response time against distance of the coming vehicle

these reasons, the timing of showing virtual mirrors should be examined as the user's vehicle comes close to the intersection.

In order to examine the best timing, we measured the response time by changing the timing to display virtual mirrors. We collected the data for four timings:

- When the subject's vehicle stops (0.0 second).
- 1.2 seconds before the vehicle stops (10 meters before the intersection).
- 2.4 seconds before the vehicle stops (20 meters before the intersection).
- Always.

We used VM1 for this experiment.

Fig.13 shows the response time of four timings. Compared to the timing to display virtual mirrors when the vehicle stops, other timings show shorter response time. It is probably because drivers can recognize the virtual mirrors before the vehicle arrives at the intersection in case virtual mirrors are displayed before the vehicle stops.

We can say that virtual mirrors should be displayed as early as possible to shorten the recognition delay. We think this is because the subjects were asked to concentrate on finding the coming vehicle and they did not need to drive their vehicle as it is automatically driven to the stop line in this experiment. In certain traffic environment, this consequence is valid when there are no obstructive objects around a driver's vehicle because the driver does not need to pay much attention on maneuvering the vehicle in that case. However, it might not be true in other real driving situations.

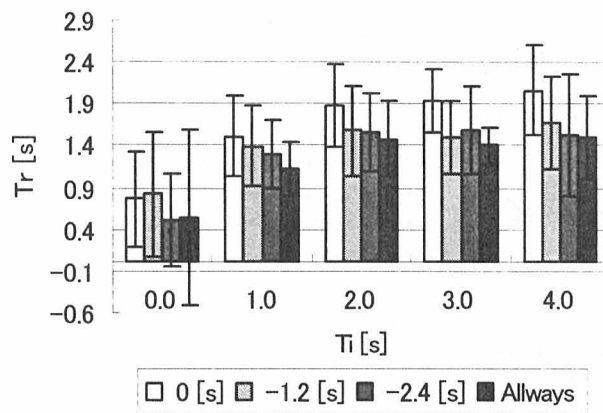


Fig.13. Response time by changing timing to display

4.5. Location of virtual mirror

Virtual mirrors can be placed at any locations. Therefore, a virtual mirror can be placed at a location where a real traffic mirror cannot be installed physically. We examined the response time by changing the location of virtual mirrors. We set virtual mirror of VM1 at four locations of (A), (B), (C), and (D) shown in Fig.14. The

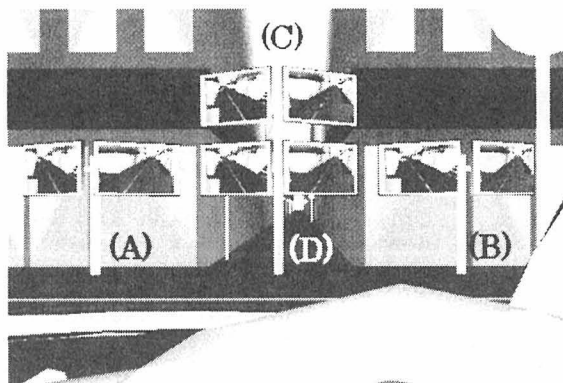


Fig.14. Location of virtual mirrors

images in the mirrors are same regardless of the locations of the virtual mirrors.

Fig.15 shows the result. Elements of the graph from the left correspond to (A), (B), (C) and (D) respectively in Fig.14.

The response time of (D) marks longer compared to the others. It is because (D) is an impossible position to install a mirror in the real world, while (A) and (B) are possible. (C) is also an impractical location, but it is still acceptable because it does not disturb drivers to see the direction in front of their vehicle. Thus, (C) is thought to mark a shorter response time compared to (D).

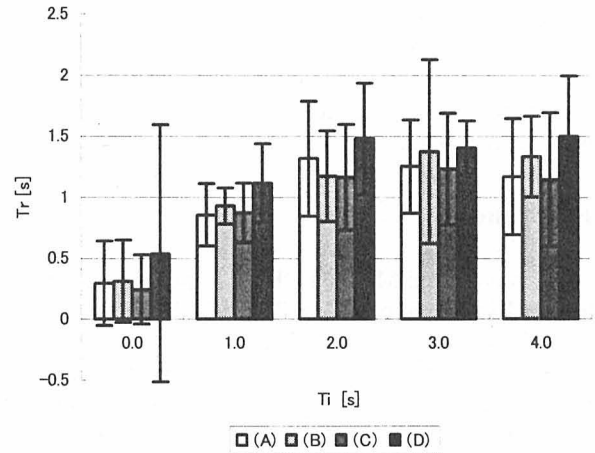


Fig.15. Response time by changing location of virtual mirrors

5. Conclusion

In this paper, we proposed a novel visual assistance system in which a virtual mirror helps drivers to see blind intersections. The system utilizes surveillance cameras. To verify the proposed method, we implemented a driving simulator, and conducted various experiments. The results revealed that drivers can shorten the delay for recognition of coming vehicles in blind intersections by virtual mirrors. As a result, we can conclude that they will be useful for reducing traffic accidents at blind intersections.

We conducted experiments on a driving simulator by using a simple road model made in CG. Hence further investigation with more realistic human interface and finer graphics of the driving simulator will be needed to confirm the results.

There are differences between the tasks for a subject in the experiments and the tasks in a real driving situation. Further evaluation of the response time T_r is needed by improving the interface of our driving simulator from keyboard operation to driver's wheel and head tracking. We have to examine the conceivable problems in the real driving situation, such as the resolution of the surveillance cameras, the delay and throughput for transmission of the surveillance camera images to the vehicle, consistency of the virtual mirrors, and so on.

When we think about operation of virtual mirrors in real traffic environment, we need to care about fail-safe functionality. For example, in case of unstable video transmission and/or malfunction of the system and network, we think it is better to alarm drivers about the

malfunction and to stop showing virtual mirrors, as virtual mirrors are merely a visual support. In addition, we need to improve photometric consistency of virtual mirrors against real scene by utilizing recent techniques of computer vision.

Reference

- [1] T. Nakata and M. Aoki, "Field Experiment of Image Sensor Use in a Snowy Area", CD-ROM Proceedings of 11th World Congress on ITS, 2004.
- [2] Y. Ishizuka and Y. Hirai, "Segmentation of Road Sign Symbols Using Opponent-Color Filters", CD-ROM Proceedings of 11th World Congress on ITS, 2004.
- [3] E. Ichihara, H. Takao, and Y. Ohta, "NaviView: Visual Assistance of Drivers by Using Roadside Cameras -Visualization of Occluded Cars at an Intersection-", World Multiconference on Systems, Cybernetics and Informatics, XIII, part II, pp.175-180, 2001.
- [4] E. Ichihara and Y. Ohta, "NaviView: Visual Assistance Using Roadside Cameras -Evaluation of Virtual Views-", Proceedings of IEEE Intelligent Transportation Systems, pp.322-327, 2000.
- [5] E. Ichihara, H. Takao, and Y. Ohta, "Visual Assistance for Drivers Using Roadside Cameras", Proceedings of IEEE /IEEJ/JSAL International Conference on Intelligent Transportation Systems, pp.170-175, 1999.
- [6] E. Ichihara, H. Takao, and Y. Ohta, "NaviView: Bird's-Eye View for Highway Drivers Using Roadside Cameras", IEEE International Conference on Multimedia Computing and Systems, vol.2, pp.559-565, 1999.
- [7] E. Ichihara, H. Takao, and Y. Ohta, "Bird's-Eye View for Highway Drivers Using Roadside Cameras", Proceedings of the 1998 IEEE International Conference on Intelligent Vehicles, vol.2, pp.640-645, 1998.
- [8] Y. Kameda, T. Takemasa, and Y. Ohta, "Outdoor See-Through Vision Utilizing Surveillance Cameras", IEEE and ACM International Symposium on Mixed and Augmented Reality, pp.151-160, 2004.
- [9] R. Bane and T. Hollerer, "Interactive Tools for Virtual X-Ray Vision in Mobile Augmented Reality", IEEE and ACM International Symposium on Mixed and Augmented Reality, pp.231-239, 2004.
- [10] S. Morita, K. Yamazawa, and N. Yokoya, "Networked Video Surveillance Using Multiple Omnidirectional Cameras", IEEE International Symposium on Computational Intelligence in Robotics and Automation, pp.1245-1250, 2003.



Fumihito Taya received the B.E. degree from the College of Engineering Systems, University of Tsukuba, Ibaraki, Japan, in 2004. He is currently a student in University of Tsukuba. His research interests include visual assistance for driver and image processing.



Kazuhiro Kojima received the B.E. degree from Himeji Institute of Technology, Hyogo, Japan, in 2002 and received the M.E. degree from the Sciences and Engineering in University of Tsukuba, Ibaraki, Japan, in 2004. In University of Tsukuba, his research interest was in visual assistance for drivers.



Akihiko Sato received the B.E. degree from the College of Engineering Systems, University of Tsukuba in 2005. He is currently a student in University of Tsukuba. His research interests include visual assistance for driver and image processing.



Yoshinari Kameda received the B.E., M.E., and Ph.D. from Kyoto University in 1991, 1993, and 1999 respectively. He worked at Kyoto University from 1996 to 2003 as a research associate. He moved to University of Tsukuba in 2003. He is now an associate professor at Department of Intelligent Interaction Technologies, Graduate School of Systems and Information Engineering, University of Tsukuba. His research interests include ITS, video media processing, massive sensing, sensor fusion, mixed reality, and automatic video filming.



Yuichi Ohta received the B.E. degree in electronics and the M.Sc. and Ph.D. degrees in information science, all from Kyoto University, Kyoto, Japan, in 1972, 1974, and 1980. From 1978 to 1981 he was on the faculty at the Department of Information Science, Kyoto University. In 1981, he joined the faculty at University of Tsukuba, where he is currently a Professor of the Department of Intelligent Interaction Technologies. He was Visiting Scientist at the Computer Science Department, Carnegie Mellon University, Pittsburgh, PA, from 1982 to 1983. His research interests are in the fields of computer vision, artificial intelligence, and intelligent visual interface. He is a fellow of IAPR and IEICE.

Received: 25 April 2005

Received in the first revised form: 22 July 2005

Received in the second revised form: 28 September 2005

-38- Accepted: 4 October 2005

Editor: Shinji Tanaka