

Journal of Experimental Psychology: Human Perception and Performance

Eye Movements Reveal That Drivers Can Predict the Location of Hazards in Dynamic Road Scenes but Gaze and Awareness Are Dissociable

Jiali Song, Ido Zivli, and Benjamin Wolfe

Online First Publication, June 8, 2026. <https://dx.doi.org/10.1037/xhp0001437>

CITATION

Song, J., Zivli, I., & Wolfe, B. (2026). Eye movements reveal that drivers can predict the location of hazards in dynamic road scenes but gaze and awareness are dissociable. *Journal of Experimental Psychology: Human Perception and Performance*. Advance online publication. <https://dx.doi.org/10.1037/xhp0001437>

Eye Movements Reveal That Drivers Can Predict the Location of Hazards in Dynamic Road Scenes but Gaze and Awareness Are Dissociable

Jiali Song, Ido Zivli, and Benjamin Wolfe

Department of Psychological and Brain Sciences, University of Toronto Mississauga

Parsing complex dynamic scenes is critical for navigating our visual world, and driving safely is a daily task that requires responding quickly to hazardous events. Theories of driver awareness suggest that drivers need to look at hazards to detect them, particularly in anticipation of hazard onset. Indeed, scene context may guide eye movements to likely hazard locations. However, given the complexity of real road scenes, drivers may rely instead on explicit cues of an impending collision before identifying the correct location. In 2024, we recorded eye position while 30 licensed drivers localized hazards in annotated dashcam footage. On correct trials, drivers started to look at where the hazard will be 2 s before onset, highlighting the importance of anticipatory processes in dynamic scene perception and safe driving. However, hazard-directed looking is not indicative of awareness, as 40% of missed hazards were foveated before response. Our results demonstrate early gaze guidance by scene context, which can help drivers anticipate hazards, an idea consistent with theories of driver awareness and with theories of visual search. However, looking directly at the hazard alone is not necessary or sufficient for hazard awareness and should not be used to index awareness in driver monitoring systems.

Public Significance Statement

Safe driving requires quickly extracting meaningful information from a complex world over time. We might think drivers always need to look at hazards on the road, and that if they look, they are aware of hazards. We found that drivers anticipate hazards, looking where it will be several seconds before it unfolds, and that looking earlier is associated with responding faster. However, simply looking at the hazard or where it will be is a poor indicator of hazard awareness because drivers can look at the hazard and still miss it. Our results reveal the gap between assumptions we make about where and when drivers look, what they see, and what they think is happening around them. Our findings have implications beyond merely what drivers see and do not see, and point to the perceptual and attentional mechanisms that support gathering task-relevant information in complex, dynamic natural scenes.

Keywords: eye movements, dynamic natural scenes, driving, scene perception

Supplemental materials: <https://doi.org/10.1037/xhp0001437.supp>

Our visual world is full of complex, dynamic scenes. How does our visual system make sense of them? From a single, brief glance, we can already access rich gist-level information from a scene, such as global layout and structure, some semantic information, and a few objects (Oliva, 2005; Sampanes et al., 2008). However,

not all objects can be processed to the same extent in a single glance, so we must make eye movements to explore the scene over time. Prior work on static scenes suggest that eye movements are goal driven and reflect cognitive processes (Andrews & Coppola, 1999; Castelano et al., 2009; Yarbus, 1967). Moreover,

Branka Spehar served as action editor.

Jiali Song  <https://orcid.org/0000-0002-5309-6398>

We thank Anna Kosovicheva for her advice and support during the design and analysis of this study. This work was funded by a Transport Canada Road Safety Transfer Payment Program Award to Birsan Donmez and Benjamin Wolfe.

Jiali Song served as lead for formal analysis, investigation, methodology, software, validation, visualization, writing—original draft, and writing—review and editing, contributed equally to conceptualization, project administration, and supervision, and served in a supporting role

for funding acquisition. Ido Zivli served in a supporting role for formal analysis, investigation, software, visualization, and writing—review and editing. Benjamin Wolfe served as lead for funding acquisition and resources; contributed equally to conceptualization, project administration, and supervision; and served in a supporting role for formal analysis, investigation, methodology, validation, visualization, and writing—review and editing.

Correspondence concerning this article should be addressed to Jiali Song, Department of Psychological and Brain Sciences, University of Toronto Mississauga, 3359 Mississauga Road, Mississauga, Ontario, L5L 1C6, Canada. Email: jiali.song@utoronto.ca

although eye movements can be highly idiosyncratic between individuals during free-viewing, an explicit search task reduces variance between participants (Andrews & Coppola, 1999; Eckstein et al., 2006; Ehinger et al., 2009; Torralba et al., 2006), suggesting that our eye movements tend to prioritize task-relevant regions of the scene. Such prioritized regions are based on a combination of salience, semantic meaning, and target features (Wolfe, 2021). Given that eye movements are used to gather task-relevant information, does where we look and when we look at it dictate what we perceive?

Although much has been learned from studying static scenes, daily life requires parsing scenes that change over time due to self- or object-motion. The temporal unfolding of events in these dynamic scenes adds a layer of complexity to interpreting eye movements because task-relevant information may be present at some times, but not other times. Choosing to look in the wrong location at an unfortunate time may lead an observer to miss critical information as the information may not be present at a later point. For this reason, not only where, but *when* we look is paramount for understanding how we parse dynamic scenes.

Driving and road scenes are a particularly informative space in which to study questions about how we gather information to understand dynamic scenes. A crucial part of navigating the road safely is locating and responding to hazards quickly and accurately. Given the high speeds at which motor vehicles travel, even a fraction of a second of delay can be the difference between life or death, yet drivers meet these temporal demands most of the time. Understanding how drivers accomplish this task is required for a comprehensive picture of how humans process dynamic scenes under time pressure in highly consequential situations.

Situation Awareness and Driver Eye Movements

Situation awareness (SA) is a framework for characterizing and evaluating how well human operators understand and is able to respond to their environment (Endsley, 1995a). This framework was initially designed in an aviation context but have also been used to characterize a driver's awareness of road hazards as the operator of the vehicle they drive (Endsley, 2020). SA is thought to have three levels. Level 1 is perception of the elements in the environment, such that the operator is aware of task-relevant stimuli in the environment such as a vehicle in the road. Level 2 is comprehension of the current situation, such that an operator can correctly interpret and evaluate the current state of the environment, such as recognizing a hazard that requires an evasive maneuver. Level 3 is projection of the future, such that an operator can predict what will happen next in their environment, such as predicting that a hazard may require an evasive maneuver in the future and hovering the foot over the brake pedal. It is thought that each level of SA is predicated on the previous levels, however, the relationship among the three levels are not linear, such that the awareness of (Level 1) and judgments about (Level 2) stimuli in the environment (Level 1) may be informed by prior predictions made by the operator (Level 3; Endsley, 1995a). One common index of SA is the Situation Awareness Global Assessment Technique (SAGAT), whereby drivers are asked to verbally describe the scene while the simulated situation is paused and visual information is occluded (Endsley, 1988, 1995b). However, this test is unfeasible to implement on the road because questions must be tailored to

the specific scenarios and tasks tested, and it requires pausing the driving task, which would disrupt real on-road driving.

Gaze statistics have been proposed as an attractive alternative tool to measure SA because they are less disruptive than verbal report measures such as the SAGAT and can be measured continuously (Arias-Portela et al., 2024; Zhang et al., 2023). Prior research examining the relation between gaze and SA focused on aggregate eye movement statistics, such as duration of looking at the front window and saccade amplitude, and found them to differ between high and low SA situations in simulated driver take-over situations (Kim & Yoon, 2025; Liang et al., 2021; Tan & Zhang, 2022), suggesting that eye movements contribute to developing SA.

One issue for these gaze metrics is that they are often computed without considering the objects drivers looked at, why they looked there, and what information they needed (Ahlström et al., 2021). That is, regions of interest are often defined arbitrarily (Arutyunova et al., 2025), or by parts of the ego-vehicle, such as front windshield, and side or rear mirrors (Ding et al., 2024; Kim & Yoon, 2025; Lai, 2025), rather than by scene content, such as objects in the scene or other road users (but see Hofbauer et al., 2020). For this method to be generalizable across studies and across different on-road scenarios, drivers under similar task load with similar SA should have similar eye movement statistics regardless of the specific contents of the scene at hand. That is, this method implicitly assumes that drivers search for hazards based on a general understanding of where hazards are likely to be on roads, rather than the specific road ahead. Yet it is unlikely that this is an accurate characterization of driver gaze behavior. Gaze measures that account for objects in the scene are better correlated with SA than scene unaware measures (Biswas et al., 2024; Jiang et al., 2024), suggesting that driver gaze is driven at least partially by the content of the scene. However, it is unclear whether this improvement is due to successful localization of the hazard, in which case scene-driven gaze should be observed only after hazard onset.

In contrast, according to Information Acquisition (Wolfe, Sawyer, & Rosenholtz, 2022), a theory of how the visual system gathers and interprets information in a driving context, drivers use eye movements strategically to search a scene. Because gist-level information can be accessed within a single glance in peripheral vision, some awareness of the environment should be achieved early on. Furthermore, drivers may use prior knowledge of the structure of the road scene and its semantic content to guide eye movements, an idea supported by eye movement work in static scene search (Henderson et al., 1999; Neider & Zelinsky, 2006; Vö et al., 2019). This prior knowledge is gained through daily lived experience and can include the spatial-relationships between objects (scene grammar; Vö et al., 2019), and the strength of the spatial-association between objects and their location in a scene (object spatial certainty; Krzyś et al., 2025). If this is the case, then eye movements should be strategically deployed and driven by the specific scenario at hand well before hazard onset. Examining when drivers exhibit gaze driven by content of the specific scenario, rather than based solely on a general idea of where road hazards are likely to be overall, will reveal the extent to which these assumptions are true.

Hazard Anticipation

Given that road scenes are dynamic, and the visual system takes time to process inputs, prediction is a fundamental aspect

of navigating roads. Hazard perception (HP) is a theory of hazard and collision avoidance that postulates that successfully avoiding collisions requires anticipating and attending to hazards before they require immediate evasive action (Cao et al., 2022; Horswill, 2016). Typically, HP is tested with a computerized task that asks drivers to identify and describe objects in the scene that are likely to become hazards requiring evasive maneuvers in the future (latent hazards; Vlakoveld et al., 2011) in real road videos. Considerable attention has been given to HP, as performance in these prediction tasks have been found to correlate with collision risk on the road (Boufous et al., 2011; Fisher et al., 2006; Horswill et al., 2015; Wells et al., 2008; Wetton et al., 2010). Under this framework, drivers are theorized to draw on their prior knowledge and expectations to direct attention to locations with a high likelihood of producing hazards, even when no visible hazard is present in the scene.

This idea is consistent with the literature on visual search in static natural scenes that find that participants tend to look where the target is likely to appear (Torralba et al., 2006). For example, fixations during search for a toilet are more likely to be constrained along the ground in bathrooms, whereas fixations are more distributed during search for a cat that may appear anywhere in a scene. Such attentional guidance by scene context can speed search for objects with strong or reliable spatial associations (Krzyś et al., 2025). Together, these studies suggest that statistical regularities of scenes should guide the eyes toward plausible locations of hazards and speed detection.

For drivers, it is thought that knowledge of these statistical regularities (i.e., locations wherein hazards are likely to appear) is acquired through experience, because eye movement patterns in novice (less than 1 year of driving experience) and experienced drivers differ when scanning road scenes for hazards (e.g., Konstantopoulos et al., 2010; Mourant & Rockwell, 1972; Underwood et al., 2003). Experienced drivers look at locations where there are potential hazards more often than novice drivers (Fisher et al., 2006; Pollatsek et al., 2006; Stahl et al., 2019), and hazards that are hidden by another object are particularly challenging for novice drivers (Crundall et al., 2012; Underwood et al., 2002; Ventsislavova et al., 2016; Vlakoveld et al., 2011). Even in novice drivers, fixating on a pedestrian who is about to enter the road made drivers more likely to detect them when they later crossed the street (Åbele et al., 2018). These findings suggest that looking at a task-relevant location at least facilitates hazard detection and anticipation, particularly when drivers look early prior to the onset of an event that requires intervention, an idea that is in line with HP. If anticipation of hazards is indeed required for appropriate response, then drivers should look at the location of the upcoming hazard prior to hazard onset when they correctly localize a hazard, whereas this early looking should not occur when they fail to localize the hazard. In other words, HP would predict that drivers miss hazards due to failure to look at it early enough.

Is Looking at a Hazard Required for Awareness?

Much of the eye-movement research on SA and HP focusses on where drivers look (Ahlström et al., 2021). The goal of such focus is to find eye-movement measures that can serve as a proxy for hazard awareness, which would be extremely useful in driver-monitoring applications. However, a common assumption of this approach is that looking at hazards is required for detecting them

(Åbele et al., 2018; Hofbauer et al., 2020). A failure to look at a hazard has even been suggested as an indication of lack of awareness (Åbele et al., 2018; Hofbauer et al., 2020). However, other empirical evidence suggests that peripheral vision plays an important role in detecting hazards and guiding the eyes to them (Vater et al., 2022; Wolfe et al., 2017). Even in the findings of Åbele et al. (2018), fixating the pedestrian did not guarantee that the driver would brake in response. Moreover, in some circumstances, drivers responded without ever fixating the pedestrian, suggesting that simple accounts of gaze and awareness do not capture the reality of driver awareness.

Indeed, successful detection of highly dangerous hazards can occur in the absence of foveating the hazard without statistically significant costs to accuracy or speed (Huestegge & Böckler, 2016), indicating that peripheral vision is sufficient for hazard detection. However, this study was done using static images of hazards that were shown briefly to prevent eye movements, and scene motion was absent. It may be the case that in dynamic scenes where motion is present, drivers still prefer relying on foveal vision when given the choice to resolve ambiguity. If so, then drivers should always look at a hazard before correctly localizing it, in line with the assumptions of prior work on HP.

Another reason to look beyond simply whether the hazard is foveated is the phenomenon of looked-but-failed-to-see (LBFTS) errors. LBFTS errors occur when a participant fails to report the presence of a target despite having previously fixated on it (Ahlström et al., 2021; Herslund & Jørgensen, 2003; Hills, 1980; Wolfe, Kosovicheva, & Wolfe, 2022). LBFTS errors were first observed in police reports, as a self-reported phenomenon by at-fault drivers absent other reasons for the collision such as distraction (Hills, 1980). This idea is contrary to predictions of HP, as HP would not predict that a driver would miss a hazard despite looking right at it. One possible explanation is that self-reports are unreliable because awareness of eye movement history is limited (Clarke et al., 2017; Foulsham & Kingstone, 2013; Mahon et al., 2018; Marti et al., 2015). This phenomenon has been observed previously in hazard detection contexts, particularly in young-novice drivers (Kahana-Levy et al., 2019). If LBFTS errors occur as reported by drivers, we should observe cases where drivers did precisely as they described—looking directly at a hazard, but failing to report it, indicative of a failure of awareness. However, if LBFTS errors were false reports due to poor awareness of eye movements, and HP's prediction that looking at the hazard location is sufficient for awareness, then we should not observe these events.

When Do Drivers Look at the Hazard?

HP suggests that successful collision evasion requires anticipating the hazard, but provide no mechanistic explanation about how anticipation occurs. Under the strongest interpretation of this idea, drivers would need to look at a hazard-relevant location before the first visible sign that it will become a hazard. In the driving literature, such locations are typically determined by experimenters or expert drivers with knowledge of what will occur later in the video, whereas drivers on the road cannot know exactly which object in the scene will materialize into an immediate hazard. One possibility is that a collision likelihood is assigned for all possible locations in the scene. However, in an otherwise mundane road scene, nearly anything could be a potential hazard in the future,

such as a tree falling or a cyclist appearing from behind a tall bush. Given the time pressure inherent to road situations, it is impossible to process every single object in the scene in detail. How and why does the visual system prioritize some locations over others?

Another possibility given by the information acquisition theory is that search for road hazards operate under similar processes as visual search in other contexts, and that the locations where drivers search for a hazard are guided by the visual characteristics of the scene and its constituent objects (Wolfe, Sawyer, & Rosenholtz, 2022). As such, drivers should attend to a particular location when there is a visible sign that the situation deviates from a safe trajectory. For example, a driver may ignore a tree in the scene until it tilts, indicating that it will fall into the road. If this is true, then drivers should only look at the tree when it tilts. However, if drivers can anticipate this event to some extent using other contextual factors, such as heavy winds or the presence of other fallen trees, then drivers should make anticipatory glances at the tree before it starts to fall.

In addition, drivers should also respond earlier when they successfully anticipate a hazard, and later when they could not. Therefore, we expect a relation between when drivers look and their response time. Furthermore, prior research using static images of road scenes suggests that faster hazard responses may be due to increased efficiency in guiding the eyes to the hazard from interpreting peripheral visual input (Huestegge et al., 2010), while processes that occur after foveation, such as recognition and decision making stay relatively stable. This idea is unaccounted for by SA or HP, as neither theories offer mechanistic explanations of eye movement guidance. If guidance by peripheral input determines the speed of response, then hazards that are detected earlier should also be looked at earlier, whereas the duration between looking at the hazard and response should be relatively stable across response times. However, scene context can also facilitate object recognition (Peelen et al., 2024), which may speed hazard recognition after foveation. If this is true, then the when drivers look at the hazard relative to response should also vary with RT.

The Current Study

This study investigates whether and when drivers look at the hazard when they search for hazards in real dynamic road videos. We recorded eye position while drivers reported the location of hazards that required immediate evasive maneuver (left or right) using a standard computer keyboard. We then analyzed gaze position over time as a function of localization accuracy and response time to determine when gaze strategy became nonrandom, whether looking at the hazard is required for hazard detection, and how looking is related to response time.

Method

Sample

As per our preregistration, we recruited a sample of 30 licensed drivers in 2024 from the University of Toronto Mississauga undergraduate research participant pool. They were aged 18–35 ($M = 18.83$ years, $SD = 0.98$ years, 25 female, 5 male, 19 G2 restricted license or equivalent, 11 full G license or equivalent), had at least 1 year of driving experience ($M = 3.29$ years, $SD = 1.63$ years), and had normal or corrected-to-normal vision with near acuities of

20/25 or better confirmed with near Early Treatment Diabetic Retinopathy Study (ETDRS) test (Bühren, 2014) at 40 cm viewing distance. One participant included in the data set is missing seven out of 262 trials due to a software error. Based on a sensitivity analysis with an alpha of .05 and power of .95, this sample size can detect a minimum t value of 2.0 and a Cohen's d of 0.68 in a paired-sample comparison. This sensitivity analysis was conducted based on a paired-sample comparison between correct and incorrect trials, as this difference is the main effect of interest.

Only participants who maintained overall localization accuracy of 80% or above were included in the study. All participants met this criterion. Participants with fewer than 20% of missing samples in the eye tracking data were included in the final sample ($M = 4.09\%$, $SD = 4.57\%$, range = [2.06%, 18.44%]), and one participant in the sample was replaced for failing to meet this criterion such that 31 participants were recruited for a final sample size of 30.

Stimuli and Equipment

Video stimuli were forward-facing dashcam footage of real road scenes from the Road Hazard Stimuli data set (Song et al., 2024; Wolfe et al., 2020). Each video was between 1.3–8 s long (mean duration = 3.18 s) with a framerate of 30 Hz at 720p. All videos contained a hazard located to the left or right half of the video (138 left and 123 right). We defined hazards as events on the road that require immediate intervention (e.g., braking or swerving) to avoid a collision or near-collision. For each hazard, the timing and location were annotated by two independent annotators and checked by a third for agreement (Song et al., 2024).

Hazard onset timing was defined as the first instance that the road situation deviates from a safe trajectory. This definition was used as a ground truth baseline to evaluate the timing and accuracy of hazard anticipation, as not all latent hazards materialize into hazards. However, there were also hazards whose onset is the moment of the object's appearance due to the hazard already unfolding while it was out of view (e.g., not visible in frame).

Hazard location was previously annotated with a rectangle fully surrounding the hazardous object on each frame when the hazardous situation was unfolding (i.e., between hazard onset time and when the driver in the video visibly responds to the hazard). These annotations were only used in stimulus selection and data analyses and were not shown to the participants in the study. We included only videos in which the annotated hazard location did not cross the midline of the video to facilitate our definition of localization accuracy (left/right judgment). We chose a highly variable set of 261 hazards to mimic the uncertainty of searching for hazards on the road and to avoid setting up artificial expectations about when and where the hazard will appear. Hazards were carefully balanced for their left versus right location, and the rest of the features of the videos were allowed to vary. Videos included day and night scenes, urban and high-way driving, left and right hand traffic, and a variety of hazards, including but not limited to vehicles, pedestrians, animals, and objects on the road. These labels are available in the Road Hazard Stimuli data set.

All stimuli were presented using PsychToolbox-3 (Brainard, 1997; Kleiner et al., 2007; Pelli, 1997) in MATLAB Version 2019 b using a Mac Pro computer running MacOS Mojave. The size of the maximum display area was 59.6 cm by 33.6 cm, corresponding to 55.2 degrees visual angle (DVA) wide and 32.8 DVA tall at a

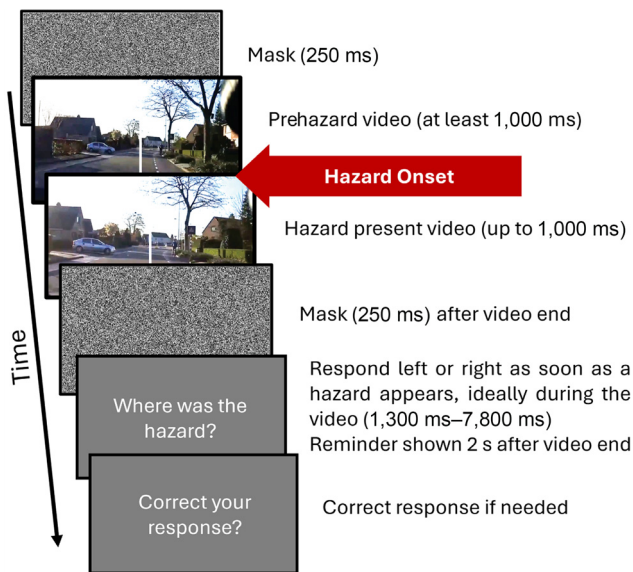
viewing distance of 57 cm, which was maintained with a chinrest. The refresh rate of the monitor was 60 Hz. Movements of the right eye were captured using a tower-mounted Eyelink 1000 at 1000 Hz and recorded using the Eyelink Toolbox in MATLAB (Cornelissen et al., 2002).

Procedure

Written informed consent was obtained from all participants before any data was collected. The experimental procedure was approved by the University of Toronto Social Sciences, Humanities and Education Research Ethics Board. At the beginning of the experiment, the eye-tracker was calibrated to each participant using a 9-point calibration procedure. The maximum validation error accepted was 1.5 DVA.

Figure 1 shows a schematic of the trial procedure. At the beginning of each trial, participants were asked to fixate on a green square fixation point shown at the center of the screen. After a randomly generated white-noise luminance mask for 250 ms, participants

Figure 1
Schematic of a Single Trial



Note. On each trial, participants fixated a green square fixation point at the center of frame (not shown in the figure). Then a mask appeared for 250 ms, followed by a video lasting between 1,300 and 7,800 ms. In each video, a hazard that required immediate response to avoid a collision occurred at least 1 s after the beginning of the video and at least 300 ms before the end of the video. In the example stimulus illustrated in the figure, the hazard (silver sedan) was stationary during the prehazard video and was not on a collision course with the dashcam-equipped vehicle. At the moment of hazard onset, the silver sedan started moving into the road, creating a situation that requires braking or swerving to avoid a collision. Participants were asked to indicate the location of the hazard relative to the white vertical line that is always visible during the video (width increased for illustration purposes, the actual line was three pixels wide). Participants were asked to respond as soon as they saw a hazard, using the left and right arrow keys. If they did not respond after 2 s after the end of the video, a text prompt was displayed on screen until response. At the end of each trial, they were given an opportunity to correct their response. See the online article for the color version of this figure.

viewed a short road video. Although the video was on screen, a visible white vertical line bisected the video into a left and right half. At the end of the video, a second noise mask was shown for 250 ms. Both masks and video stimuli subtended 40 DVA horizontally and 22 DVA vertically.

Given that every video contained a materialized hazard (i.e., an on-road situation that requires immediate driver response to avoid a collision), participants were asked to localize the hazard to the left or right of the white vertical line using the left and right arrow keys of a standard computer keyboard. The hazard occurred at an unpredictable time point within each video. Hazard onset occurred at least 1 s after the beginning of the video, and at least 300 ms before the end of the video. Participants were instructed to respond as soon as they saw a hazard. The video ended when the driver in the video responded to the hazard in the video regardless of whether the participant has responded. If participants responded while the video was playing, auditory feedback in the form of a beep at approximately 1.65 kHz for 300 ms duration signaled the registration of their response. If the video ended before participants responded, on-screen text prompted them to report the location of the hazard. The trial continued once a response had been made.

Following this localization response, participants were given a chance to correct their response, prompted by on-screen text. Once satisfied with their response, participants pressed the space bar to continue to the next trial. The response correction prompt was provided at the end of every trial to allow participants to correct motor errors.

All participants viewed all 262 videos in a randomized order. The first 10 trials were presented as practice trials so that participants could become familiar with the task, although nothing about them differs from the rest of the trials. Furthermore, participants were given a chance to ask questions at the end of the practice trials. Participants were given an opportunity to take a break halfway through the experiment. After completing the computerized task, participants were asked to complete a driving experience survey, the question and responses can be found in the Open Science Framework repository for this study (<https://osf.io/fv6e4/>). The entire experimental procedure took 1 hr, and participants were compensated \$15 CAD for their time.

Data Analysis

We used the recorded gaze positions to determine whether participants looked directly at the hazard. We defined direct looking as gaze points that landed within the annotated region of the hazard. Then, for each time point, we calculated the proportion of trials on which participants looked directly at the hazard out of all possible trials in the relevant condition. For times before hazard onset, the annotated hazard location at the time of hazard onset was used to look for evidence of anticipatory looking. Hazard localization accuracy and response time were used to categorize trials to compare gaze behavior measures, which will be described later in the text. Where applicable, effect sizes for paired *t* tests are adjusted by within-subject variance using the *rm* method in *effectsize* package (Ben-Shachar et al., 2020) for R, and are notated as d_{rm} .

When Does Gaze Position Differ From Chance?

We investigated whether gaze location differed from chance at each time point by testing the observed data against the null

hypothesis that all videos are interchangeable, such that participants fail to use video specific scene context to guide where they look. To estimate chance performance under such a null hypothesis, we conducted permutation tests (Chihara & Hesterberg, 2018; Collingridge, 2013; Hesterberg et al., 2003) at each time point by shuffling the annotated hazard location among trials within each participant. Then, we calculated a permuted null distribution for the proportion of direct gaze points at each available millisecond. To estimate the null distribution, we did this 1,000 times for each participant, which allows estimation of the p value to .001 precision. An observed proportion was considered significantly above chance if it fell beyond the central 95% of the permuted null distribution for that participant, which corresponds to a two-tailed test using a critical α of .05. Family-wise error rate was controlled using a false discovery rate (FDR) of 5% for each participant. This method accounts for individual differences in gaze behavior observed across videos, and was preregistered.

Then, to further examine when participant's gaze behavior became different from chance, we extracted the *time threshold* of scene-guided looking for each participant and condition. Time threshold is defined as the earliest time that the observed data differed from the null distribution and that was continuous with the annotated hazard onset time. This time threshold is a measure of the earliest time when participant's gaze position starts to differ from chance. This measure assumes that once participants start searching for a road hazard based on the scene at hand rather than simply orienting to the experimental situation on each trial, eye movements should remain strategic for the rest of the trial.

We repeated the same procedure relative to response time. Given that looking behavior may have differed between correct and incorrect trials, we extracted the time threshold for correct and incorrect trials separately for each participant following the procedure described earlier. To quantitatively examine the difference between correct and incorrect trials, we conducted a two-tailed t test to determine whether the time-thresholds on correct and incorrect trials differed from each other. This analysis was not preregistered.

Comparing Gaze Behavior for Correct and Incorrect Trials

To further examine how eye movements varied over time for correct and incorrect trials, we calculated the mean proportion of trials with direct gaze at each time point for correct and incorrect trials separately. We calculated an average proportion per time point, participant, and condition, aligning the trials by hazard onset and by participant response separately. To quantify whether the likelihood of looking at the hazard differed for between correct and incorrect trials, we submitted them to Welch's two-sample t tests (two-tailed) to account for the fact that not all participants had available data at every single time point we tested for incorrect trials, because there were few incorrect trials overall.

We examined gaze distance from hazard center using the same analyses methods described earlier. However, the results were less consistent and more difficult to interpret (see [online supplemental materials](#) for more details). For brevity, we focus our current reporting on direct looking as the main dependent variable.

To account for the possibility of looking at the hazard early and then looking away, we examined the proportion of trials on which participants looked at the hazard at any time within 1 s before responding. This analysis was done to examine the incidence of

LBFTS errors, where participants missed the hazard despite having looked directly at the hazard, as well as incidence of successful hazard localization using peripheral vision only, without directly looking at the hazard. We submitted the proportions to a two-tailed paired t test to determine whether correct and incorrect trials differed. This analysis was not preregistered.

Hazard Localization RT and Gaze Behavior

To examine the relation between gaze behavior and RT, we categorized each correct trial according to localization response time relative to the annotated hazard onset. On trials on which participants corrected their response, the initial response was used for RT categorization. On average, trials corrected per participant is 12 trials ($SD = 6.83$), and corrected trials represent 4.5% of all trials. For the rest of the article, we will refer to this categorization as RT Category. Trials were labeled early when the localization response occurred before the annotated hazard onset, on-time when the localization response occurred after the annotated hazard onset but before the end of the video (i.e., the annotated driver response in the original video), and late when localization response occurred after the end of the video. A late response means that the participant would have failed to avoid a collision if they were the driver in the video. Both accuracy and response time category were then used to examine whether and how gaze measures vary among these categories. Two participants with missing data in any RT Category were excluded from statistical tests involving RT Category.

Due to the variability in RT, participants had a variable number of trials in each RT Category. To account for such variability, proportion of direct-looking were calculated according to each participant's total number of possible trials within each RT category. Then, we conducted pairwise permutation tests on the proportion of direct-looking to determine whether each category significantly differed from each other. For each pairwise comparison, and each participant, we took the relevant gaze position time series for each relevant trial, and randomly shuffled the RT category labels of those trials. Then, we calculated the proportion of direct-looking for that participant at each time point, keeping the total number of trials in each RT category equal as the observed, unshuffled data. After applying this process for all participants, we averaged each participant's average time series together in each condition and calculated a difference score between the two conditions being compared.

To generate the null distribution against which to test the observed data, we repeated this process 1,000 times, which allows estimating p values to .001 precision, and the significance threshold is based on the limits of the center 95% of the null distribution. For each time point, an observed difference between two conditions is deemed statistically significant if it was more extreme than the central 95% of data in the permuted null distribution (corresponding to two-tailed critical alpha of .05). We then used an FDR of 5% to account for family-wise error rate. This permutation test method controls for individual differences, varying number of trials in each RT category, and the effects of videos in the statistical testing, when determining at which time points each condition differed from another. This procedure was conducted twice, once based on time relative to hazard onset, and another based on time relative to participant response. The analyses involving RT were exploratory and not preregistered.

Transparency and Openness

We report how sample size is determined, all data exclusions, study manipulations, and dependent measures. Stimulus presentation code, data, and analysis materials are publicly available on OSF repository at <https://osf.io/fv6e4/>, and the stimuli are publicly available on OSF at <https://osf.io/tgzb7/>. Data were analyzed using R, Version 4.1.3 (R Core Team, 2017) and the packages *tidyverse* (Wickham et al., 2019) and *eyelinkReader* (Pastukhov, 2024), and in MATLAB, Version 2023 b (The MathWorks Inc, 2023). The study's design and some analyses were preregistered, available at <https://osf.io/hv9ye/>. However, whereas the preregistration stated that the main variable of interest was gaze distance from hazard and the proportion of gaze landing in hazard was complementary, the reported data focuses on the proportion measures as they are more interpretable. Distance from hazard results can be found in the [online supplemental materials](#). Furthermore, although we registered the permutation tests on whether the observed data differed from chance, we did not preregister the time-threshold analysis. We indicate other nonpreregistered analyses in text where applicable, including testing time-threshold between correct and incorrect trials, RT category analyses, and proportion of trials with direct looking within 1 s to response.

Results

Localization Performance

All participants in the final sample performed above 80% in the localization task ($M = 97.1\%$, $SD = 1.8\%$). One video was excluded from all subsequent analyses because no participants

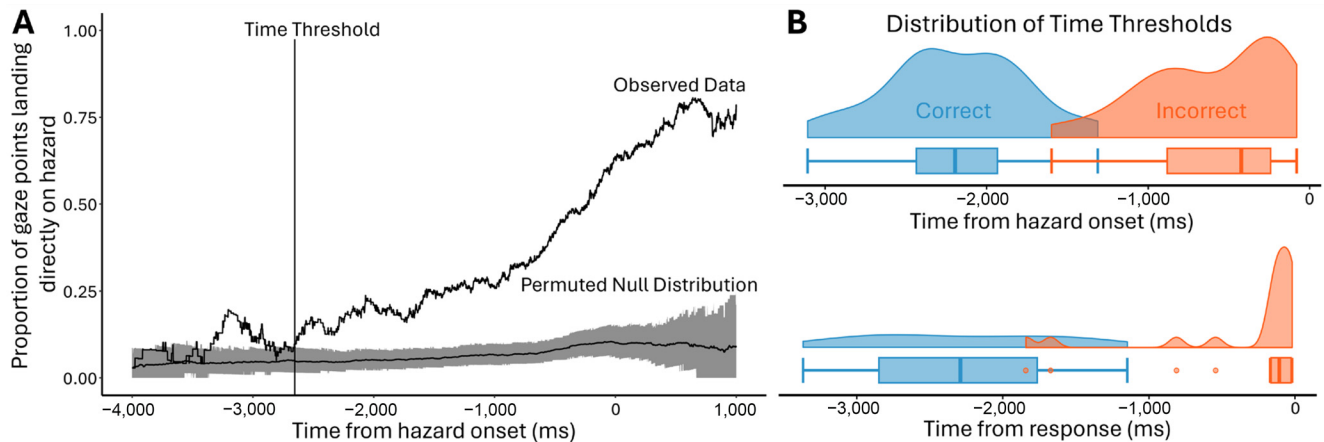
correctly localized the hazard, resulting in a total of 261 hazards. After exclusion, the average accuracy was 97.9% ($SD = 3.8\%$), 98.7% ($SD = 1.2\%$), and 93.7% ($SD = 4.6\%$) for early, on-time, and late trials, respectively, resulting in, on average, 66.8 ($SD = 38.1$) early, 102 ($SD = 26.8$) on-time, and 82.6 ($SD = 51.2$) late correct trials per participant.

When Does Gaze-Behavior Become Different From Chance?

The permutation analysis revealed that direct looking differed from chance well before hazard onset. Figure 2 shows the distribution of time thresholds of nonrandom looking for correct and incorrect trials, relative to hazard onset (Figure 2A), and to response (Figure 2B). Using our procedure, we found a valid time threshold for correct trials for all participants. However, due to the low number of incorrect trials, no time threshold for incorrect trials could be found for seven participants relative to hazard onset, and for 10 participants relative to response time. We subsequently excluded the participants with missing data from all subsequent statistical tests involving time thresholds. In the rest of the sample, time-thresholds for different-from-chance looking were significantly earlier for correct than incorrect trials both relative to hazard onset, $t(22) = 14.54$, $p < .001$, $d_{rm} = -3.78$, mean difference = $-1,655.78$, 95% confidence interval [CI; $-1,892.01$, $-1,419.56$], and relative to response, $t(19) = -11.53$, $p < .001$, $d_{rm} = -3.25$, mean difference = $-1,924.45$, 95% CI [$-2,273.882$, $-1,575.018$]. These results indicate that gaze behavior depended on the specific scene at hand well before hazard onset and continued until response, particularly for correct trials.

Figure 2

Time When Participant Gaze Position Differs From Permuted Null Distribution



Note. This figure shows an example of observed data and the permuted null distribution in a single participant, for all correct trials. The permuted null distribution represents the null hypothesis that all videos were interchangeable and no video-specific information was used by participants. Time threshold was defined as the earliest time at which the participant's looking differed from the null distribution, which was also continuous with time point 0. In the example in panel A, time zero denotes hazard onset. We also conducted an analysis yoked to response onset, such that time 0 ms represents the onset of the participant's response. Panel B shows the estimated time thresholds of correct and incorrect trials yoked to hazard onset at time 0 (top) and yoked to response onset at time 0 (bottom). Time thresholds were much earlier for correct trials compared with incorrect trials relative to hazard onset and participant response, suggesting that nonrandom gaze behavior is related to gathering information about the hazard. See the online article for the color version of this figure.

Comparing Direct Looking Behavior on Correct and Incorrect Trials

Next, we examined the hypothesis that direct looking would vary with localization accuracy, by comparing gaze behavior on correct and incorrect trials. Figure 3A shows the proportion of trials with direct gaze at the hazard at each time point relative to hazard onset. Visual inspection of Figure 3A shows that the proportion of direct gaze increase over time, but the increase is much steeper for correct trials (blue), than for incorrect trials (orange), particularly within 1 s before annotated hazard onset. Furthermore, there was more direct looking at the hazard for correct trials (peaking at .63 at 852 ms) than incorrect trials (peaking at .36 at 1,013 ms) surrounding hazard onset. This observation is corroborated by the fact that correct and incorrect trials significantly differ from each other between -688 ms to 749 ms ($t > 2.36$, $p < .24$ after correction for family-wise error rate according to an FDR of 5%), as indicated by pairwise t tests conducted at each time point. Furthermore, there was an earlier cluster of timepoints between approximately -2,308 ms to -1,748 ms, when participants looked directly at the hazard more often during correct trials than incorrect trials ($t > 2.34$, $p < .025$).

A similar pattern of results can be observed in Figure 3B, showing the proportion of direct looking as a function of time from response. The proportion of direct looking is higher for correct (peaking at .68 at -80 ms) than incorrect trials (peaking at .35 at 1,564 ms) surrounding the time of response from -888 ms to 749 ms ($|t| > 2.21$, $p < .03$).

These same analyses were repeated on the subset of videos for which there were both correct and incorrect responses. The pattern of results was very similar albeit with smaller effects due to the reduced number of trials (see online supplemental materials for details). These results indicate that the increased looking on correct trials was not due to stimulus effects.

Figure 4 shows the incidence of trials on which participant directly looked at the hazard immediately before responding for incorrect and correct trials. In 40% of missed trials, participant gaze position landed on the hazard, indicating LBFTS errors. For correct trials, this percentage was much higher at 83%. However, there remained 17% of trials in which participants failed to look directly at the hazard and still correctly reported hazard location. The paired t test indicates that the proportion of trials with hazard-directed looking is significantly higher for correct trials than incorrect trials, $t(29) = 9.3227$, $p < .001$, $d_m = 1.84$, mean difference = 0.42.

Together, these results indicate that looking directly at the hazard is associated with correct hazard localization. However, direct looking is not perfectly associated with correct localization. Moreover, in the time series analyses, given that the proportion of direct looking never reaches near ceiling in any condition, there is no single timepoint that is diagnostic of hazard awareness. In sum, although looking at the hazard is helpful for correct localization, it is not always necessary in a substantial number of cases.

Does Looking Behavior Vary With Speed of Hazard Localization?

To examine whether direct gaze is correlated with hazard localization speed, we examined the correct trials divided into three RT categories, early, on-time, and late. Comparing how the proportion of direct looking over time varies with RT category can reveal the

temporal aspects of gaze behavior most likely to determine subsequent RT. Visual inspection of Figure 5A suggests that participants look earlier when they respond earlier. The overall proportion of direct looking increases over time and plateaus at a maximum around the time of hazard onset, and the timing of the maximum is earlier for earlier responses (~ -200 ms for early, ~ 100 ms for on time, and ~ 200 ms for late trials). The increase in the proportions also begin in the order of increasing RT. The maximum proportion is 0.68 for early trials, 0.70 for on-time trials, and 0.51 for late trials at 714 ms, 653 ms, and 653 ms, respectively. These results are in line with the idea that earlier responses are associated with earlier looking.

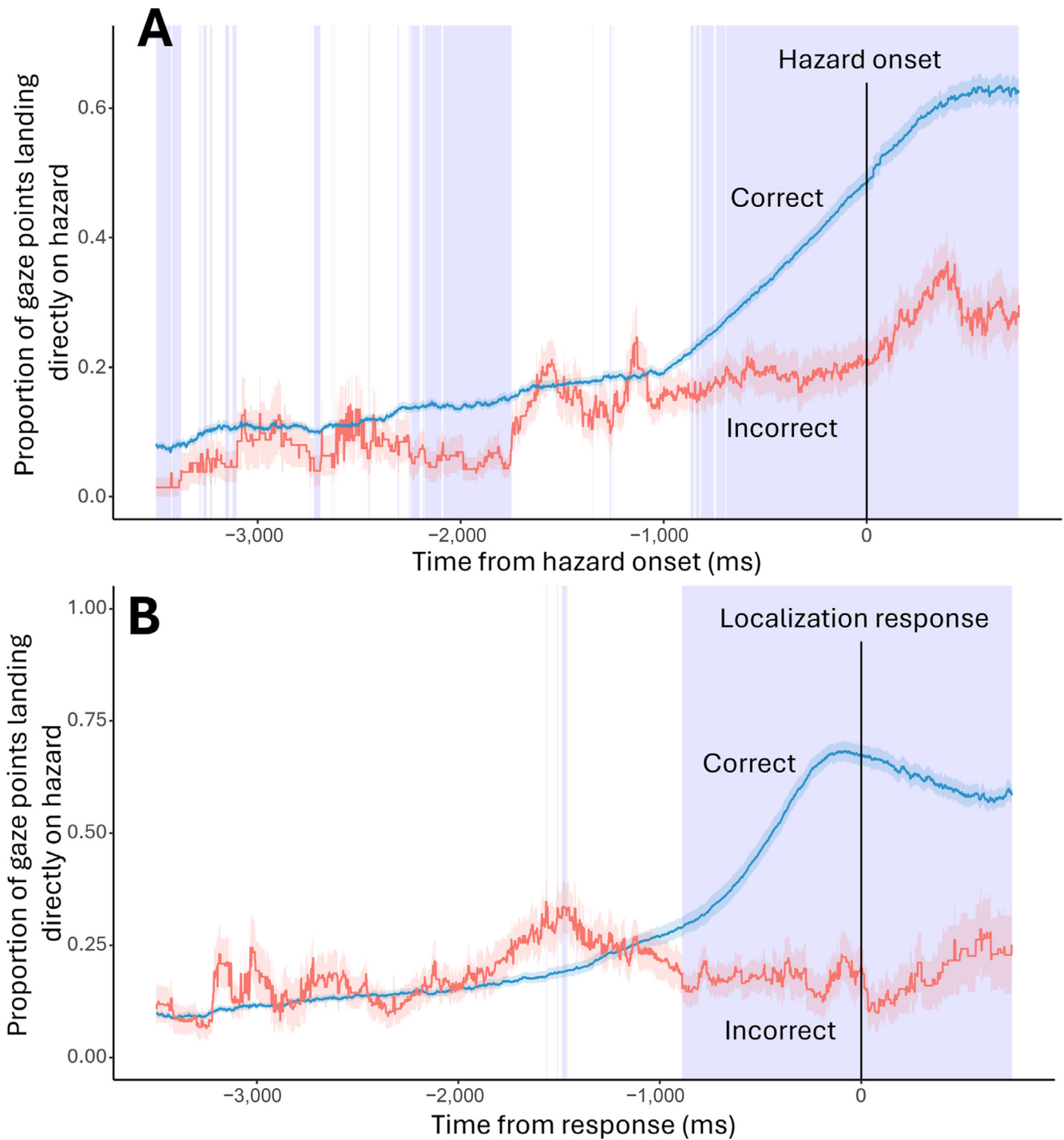
Pairwise comparisons at each time point revealed when gaze behavior differentiated between response speeds. Early, on-time, and late trials all differed significantly from each other for 796 ms just before hazard onset at -884 ms to 89 ms. However, early trials significantly differed from late and on-time trials much earlier than late and on-time trials differed from each other, suggesting that there was anticipatory looking when there was anticipatory responding. Early trials differed significantly from late trials as early as -2,837 ms to -2,342 ms, although the largest segment of significant times runs from -2,255 ms to 1,011 ms. Early trials differed significantly from on-time trials from -2,005 ms to 89 ms. These results indicate that on trials where participants anticipated the annotated hazard, they were more likely to show anticipatory looking before hazard onset. Whereas on-time and late trials only differed from each other starting from -884 ms to 1,005 ms, suggesting that there is more anticipatory looking for on-time trials than late trials, but it occurred later than early trials. Altogether, these results suggest that anticipatory looking occurred before hazard onset, and that earlier localization responses are associated with earlier foveation of the hazard.

The results of direct gaze proportion as a function of time from response show different results with regards to RT Category. Visual inspection of Figure 5B indicates that proportion of gaze rises over time and plateaus just before participant response at time zero. The maximum values before the time of response are 0.61 at -117 ms for early trials, 0.73 at -78 ms for on-time trials, and 0.73 at -78 ms for late trials. Notably, all peaks occur just before the response, indicating that the relationship between looking and response cannot distinguish between RT categories. These qualitative observations are corroborated by quantitative tests, as only late trials differ significantly from other conditions (vs. early trials: -2,340 ms to -463 ms, and vs. on-time trials: -300 ms to -713 ms). Early and on-time trials did not significantly differ from each other at any time point. These data suggest that processing time after foveation did not systematically differ among RT Categories.

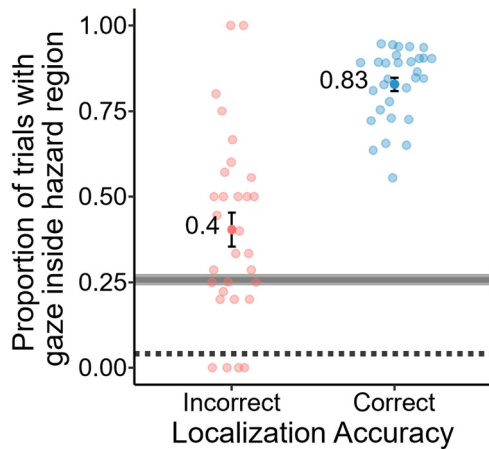
Discussion

In this study, we investigated how the visual system supports the detection and appropriate response to road hazards. We evaluated several hypotheses motivated by prevailing theories of driver perception.

SA is a framework with which driver awareness is described but provides no mechanistic explanation for how it is achieved. Prior work investigating gaze correlates of driver SA without considering the content of the scene assumes that gaze behavior among different scenes are essentially interchangeable until hazard onset. However, we found that gaze was consistently guided by the

Figure 3*Proportion of Hazard-Directed Looking Over Time Grouped by Hazard Localization Accuracy*

Note. This figure shows the proportion of trials with the proportion of direct looking at each time point. Panel A shows the proportion as a function of time from hazard onset such that time 0 represents the moment of hazard onset, and panel B shows the proportion as a function of time from response such that time 0 represents the onset of the localization response. The blue and orange lines represent correct and incorrect trials, respectively. Solid lines and shaded regions around them indicate mean and standard error of the mean, respectively. Regions highlighted in light purple indicate the time points at which incorrect and correct trials differ significantly according to two-tailed t tests. See the online article for the color version of this figure.

Figure 4*Proportion of Trials With Direct Looking Within 1 s to Response*

Note. Mean proportion of trials with gaze points landing on the hazard any time within 1 s before response. Orange and blue solid dots represent incorrect and correct hazard localization trials, respectively, and transparent dots represent individual participants. Error bars indicate standard error of the mean. The dark gray horizontal solid line indicates the mean proportion of the permuted null distribution (1,000 iterations), and the surrounding lighter gray region indicates the upper and lower limits of the permuted 95% confidence interval. The black dotted line indicates the average proportion of display area taken up by the annotated hazard region for the frames included in this analysis. The incorrect trials with gaze landing within the hazard region represent looked-but-failed-to-see errors. See the online article for the color version of this figure.

scene context at hand as early as 3,000 ms before hazard onset. Furthermore, such scene guidance was observed earlier for correct trials compared with incorrect trials, suggesting that participants used the structure of specific road scenes on each trial to guide their gaze strategy and aligns with findings using static natural scenes (Henderson et al., 1999; Neider & Zelinsky, 2006; Torralba et al., 2006; Vö et al., 2019). These results also speak to the non-linear nature of achieving SA, as perceptual and predictive processes iterate over time (Endsley, 1995a), and are incompatible with the idea that gaze is essentially similarly distributed across different scenes. These results suggest that gaze measures should be analyzed and interpreted in the context of the contents of the scene, because such scene context affects where people look.

These results are consistent with an Information Acquisition (IA) account of hazard detection, such that scene context is captured quickly using peripheral vision, which is used to strategically guide search for specific hazards (Wolfe et al., 2020). According to IA, information about scene structure is captured quickly by peripheral vision, which is then be used to guide where drivers look. Our finding that gaze differs from the null distribution several seconds before hazard onset is consistent with this idea. Moreover, when examining hazard-directed looking behavior over time, we found that participants started to look reliably at the future location of the hazard before the hazard unfolded, indicating predictive processing within approximately 1,000 ms ahead of hazard onset. These results are consistent with the idea that eye movements are predictive (Henderson, 2017) and fit well within the 500–1,000 ms duration range at which the eyes lead action in other naturalistic

tasks (Hayhoe et al., 2003, 2012; Land et al., 1999). This result suggests that anticipation and predictive processes contribute to hazard detection, and that diagnostic information about where the hazard will be may often be present in the scene prior to hazard onset.

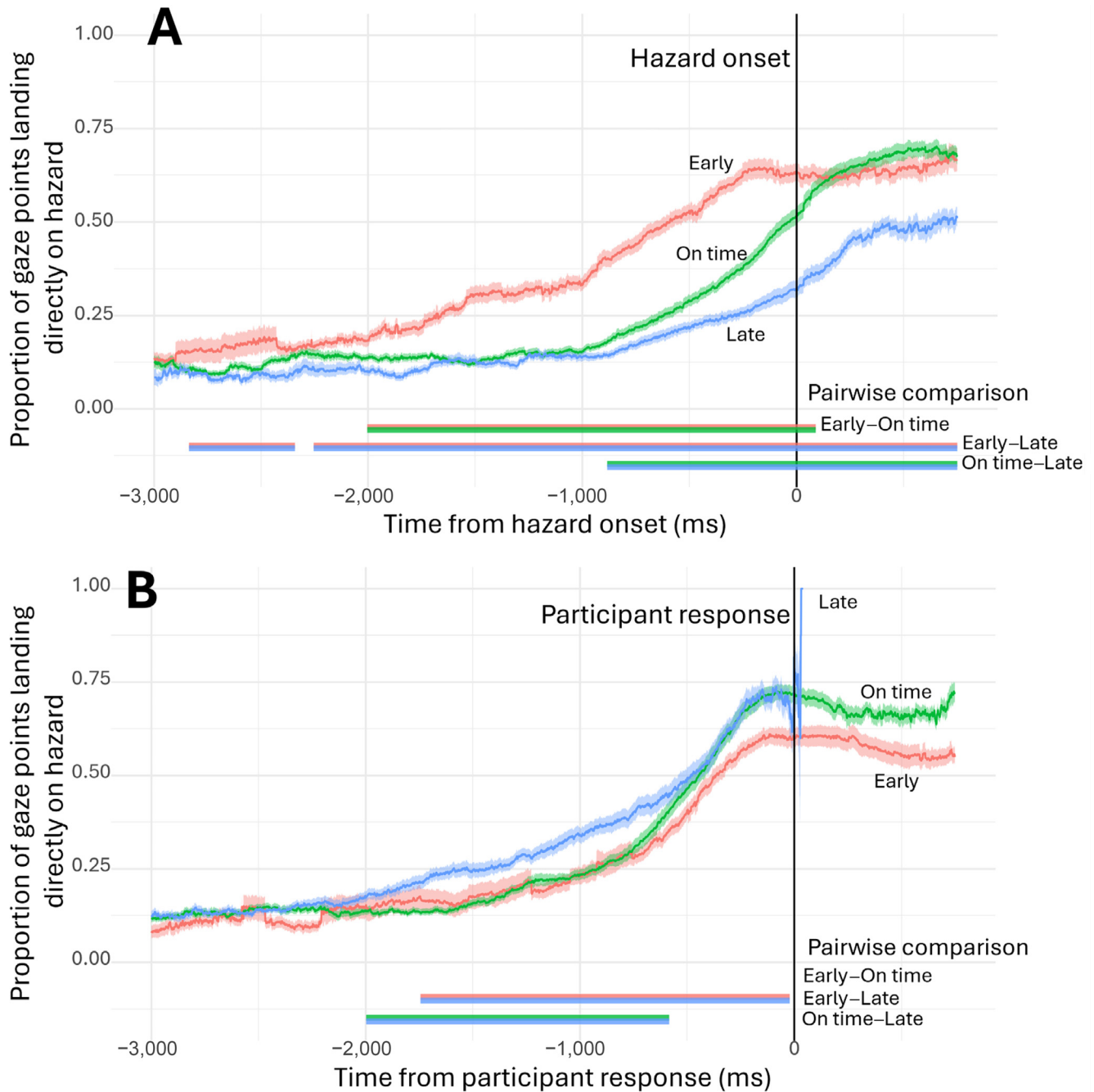
This interpretation of anticipatory looking is congruent with the idea of latent hazards, that is, hazards that are predictable and gradual in their onset and that elicit responses based on anticipated collisions rather than an ongoing collision (Crundall et al., 2012; Horswill, 2016; Vlakveld et al., 2011). The idea of latent hazards is championed by HP, a theory of collision avoidance that suggests that anticipation of hazards is crucial for effectively responding to hazards. However, not all hazards tested here involved anticipatory looking, likely because some hazards are more surprising and less predictable than other hazards. Our results do not support the idea that there is a critical time point at which looking at the hazard is required for its detection, as no single time point reached 100% directly looking for correct trials. However, the difference in gaze over time between correct and incorrect trials suggest that such anticipatory looking is helpful for hazard localization. Whether drivers looked at a location may be a data-driven way to characterize parts of the scene that drivers believe to be latent hazards, but it is not a perfect index as other candidate hazard locations are likely also identified in peripheral vision at each moment. Therefore, examining agreement among drivers about objects judged to be latent hazards may shed light on the processes involved in identifying such candidate hazard locations.

We also found an association between *when* participants looked at the hazard and RT for correct hazard localization. When participants correctly localized hazards faster, they also looked at them earlier, suggesting that anticipatory looking plays an important role in the driver's ability to respond early enough to avoid collisions. However, gaze behavior relative to response did not show such temporal separation among RT categories, suggesting that guiding gaze to the hazard using peripheral vision is a more crucial factor in determining response time, rather than hazard recognition or decisional processes after foveation. These results are consistent with the idea that peripheral vision is used to guide gaze and understanding of the situation overall.

Interestingly, we also found a relation between looking and hazardousness (as rated by an independent sample; see [online supplemental materials](#) for more details). Participants were more likely to look at high-rated hazards compared with low-rated hazards, but only after responding to the hazard. Before response, there was less hazard-directed looking for high-rated hazards compared with low-rated hazards, but the pattern is reversed after hazard onset. This suggests high-rated scenarios had less scene context cues to guide anticipatory gaze, but hazard onsets were better able to capture gaze than low-rated hazards. These results are consistent with a previous study that found that participants make longer saccades toward highly hazardous situations than less hazardous situations (Huestegge & Böckler, 2016). These results suggest that highly rated hazard onsets draw attention more easily than low hazardousness situations and are congruent with the idea that threats capture attention (Schmidt et al., 2017).

Dissociation Between Gaze and Awareness

Although foveating the hazard is associated with correct detection, it is neither necessary nor sufficient for correct hazard

Figure 5*Proportion of Trials With Hazard-Directed Looking Over Time Grouped by Hazard Localization Speed*

Note. Panels A and B show the proportion of trials with gaze position inside the hazard region for each RT category, such that orange, green, and blue represent early, on-time, and late trials, respectively. Solid lines indicate the mean, and the highlighted regions around the solid lines indicate the standard error of the mean. Panel A shows proportion as a function of time from hazard onset, and panel B shows proportion as a function of time from response. Only trials with correct hazard localization were included. For panels A and D, the horizontal bars just above the x-axis indicate times at which a pair-wise permutation test yielded a statistically significant difference. See the online article for the color version of this figure.

detection. Our results suggest that gaze position alone cannot indicate hazard awareness and should not be used as such in driver-monitoring applications. When hazard localization is correct, we found that the hazard was foveated 83% of the time

within 1 s of response. Furthermore, there was no time point at which foveation of the hazard is diagnostic of awareness. Even at the time point with the highest probability of looking, we found that gaze position landed in the hazard region less than

70% of the time. In other words, at any single moment, there was at least 30% chance that a participant is not looking at the hazard even though they will subsequently locate it accurately. These results suggest that peripheral vision can be sufficient for localizing hazards in road videos and are consistent with prior findings with static images of road hazards (Huestegge & Böckler, 2016). Hence, gaze is dissociable from and should not be used as a proxy for any level of SA.

A consideration that tempers this conclusion is that, given that a lateralization response was required on every trial, a correct response may not necessarily indicate awareness of the exact location of the hazard. Participants may have been aware of the general hazardousness of the situation without knowing the exact location of the hazard, or no signal of a hazard was noticed at all, and participants simply guessed after the video ended. In these cases, participants would have a 50% chance of guessing the trial incorrectly. We observed 3% incorrect trials, and from that we can infer that approximately 3% of all trials were guessed correctly, which equals approximately 3% of all correct trials. Assuming that participants never look at the hazard when guessing even by chance, of the 17% of correctly localized hazards without foveation, 14% of the trials remains not guessed, suggesting that participants could use peripheral information to inform their response on these trials.

Another reason why gaze position is a poor indicator of awareness is the presence of LBFTS errors. For missed hazards, we found an average of 40% of trials on which participants looked directly at the hazard within 1 s before response. Furthermore, there is no single diagnostic time at which not looking is perfectly predictive of missing the hazard. The time series analysis of gaze position found that, at any single moment, there was up to 36% chance that the participant looks at the hazard before an incorrect localization response. Although most misses were due to looking too late or not looking at all, there was still a substantial possibility of LBFTS errors. These results suggest that LBFTS errors did happen at least sometimes as drivers reported them (Hills, 1980). As such, failing to look at a hazard is a poor indicator for unawareness of the hazard. Combining gaze with other measures such as information regarding objects in the scene (Biswas et al., 2024; Jiang et al., 2024), and electroencephalography (Yang et al., 2024) may result in more accurate physiological indices of awareness.

The proportion of LBFTS reported here is slightly higher than the range of previously reported LBFTS rates of 34% for complex targets in a laboratory search task (Hout et al., 2015). In many ways these results are not comparable because the tasks and conditions tested are very different. Moreover, given our relatively naïve time-series analysis, likely not all misses with foveated hazards are true LBFTS errors. According to the Normal Blindness framework (Wolfe, Kosovicheva, & Wolfe, 2022), LBFTS errors can arise through two mechanisms. They can occur due to not looking long enough to process the object (selection failure) because another location is deemed more relevant (misguidance), for example, a cyclist was looked at too briefly because a clown appeared on the other side of the sidewalk resulting in a failure to process the cyclist. LBFTS errors can also occur due to a failure in decision making, for example, the cyclist is identified but is deemed not to warrant a response. Our analyses could not distinguish between these two possibilities. Conversely, participants may not have fixated on the hazard at all, as gaze position may have swept over the hazard region during a saccade. In the latter

case, misses would not be true LBFTS errors. We have adopted a liberal definition of what constitutes looked at the hazard to minimize missing any LBFTS errors, and so that trials without looking are defined as precisely as possible.

In some exploratory, nonregistered analyses, we also investigated when participants were likely to blink under the consideration that participants will miss visual information during a blink (full analyses and details in [online supplemental materials](#)). We found that blinks are rare and distributed across the entire trial. Crucially, the likelihood of blinks was consistently nearly zero immediately before the time of response, suggesting that blinking decreased when participants became aware of the hazard. These results correspond well to a prior finding that reduced blink rate are associated with the presence of a hazardous situation (Costela & Castro-Torres, 2020). However, the actual usefulness of such a biomarker for hazard-awareness in on-road applications is predicated on already knowing that a hazard is present and that it is the cause of decreased blink rate. It is difficult to meet these conditions on the road because increased visual processing demand more generally has been shown to reduce blink rate (Cardona et al., 2011; Chen et al., 2023; Drew, 1951; Poulton & Gregory, 1952), and the source of increased demand may be unrelated to the road environment, such as answering a phone call. Therefore, although we found that a decrement in blink rate may predict manual response, its usefulness is limited in more complex settings.

Constraints on Generality

We have demonstrated several aspects of gaze behavior when viewing dynamic road hazards that converge with prior work investigating other everyday tasks. However, not all dynamic scenes are equal. For one, travel speeds during driving are much faster than during other kinds of locomotion such as walking or running. To account for the same amount of time, the distances one must look ahead toward are much higher at increased speeds. For example, looking 2 s ahead requires looking 5.6 m ahead when traveling at 20 km/hr, but requires looking 28 m ahead when traveling 100 km/hr. Many other daily tasks, such as walking in-doors, or making tea, seldom require looking ahead at such far distances. Consequently, the information important for safely navigating the environment during driving would likely differ from other contexts. As a result, one may expect where and when participants look in a scene to differ among these contexts. The limits of such anticipatory looking is of interest for understanding driving safety and scene perception and should be investigated in future work.

Another limitation of this work is that participants from other regions of the world may perform differently than our sample of Canadian drivers. Road environments can vary dramatically in different regions of the world. One such aspect is the types of hazards drivers may encounter on the road. For example, motorcyclists are much more commonly encountered in Thailand than in Canada. Previous research showed that drivers are more likely to miss rare hazards compared with common hazards (Beanland et al., 2014), and more like to miss hazards in general if hazards are rarer overall (Kosovicheva et al., 2023). One may expect Thai drivers and Canadian drivers to respond differently to the same stimuli due to the difference frequencies at which they encounter hazards in their prior experiences. Given that prior knowledge and experience may

affect gaze behavior, we may expect different gaze behavior from different groups of drivers. For this reason, cross-cultural studies may be fruitful for shedding light on the intersection between eye movements, prior experience, and scene perception.

Several aspects of the task design are sufficiently different from on-road driving that limits generalization of hazard localization performance results. Some aspects of the video may have made localization performance worse than what it would have been on the road. For example, the videos tested in the current study only contained forward facing footage, whereas some hazards are visible to a careful driver through the side windows or in mirrors. There may not be enough information contained in the forward facing footage only to detect these hazards, but could otherwise have been detected earlier in an on-road situation. Furthermore, the video stimuli are crowd-sourced from social media and may be biased toward rarer or more unexpected hazards, which may reduce the ability of drivers to detect them in time.

Other aspects would have made localization performance in this study better than on road situations. Participants in the study would be primed to detect hazards given there was a hazard on every trial, whereas it would be extremely unlikely for drivers to encounter over 200 hazards on the road. This may have made participants report hazards more liberally than on-road driving. In addition, the current study tested hazard detection under ideal conditions where participants were engaged in no other task. However, on the road, drivers also need to operate the vehicle simultaneously and may choose to engage in other activities that distract from searching for hazards, such as texting or having a conversation on the phone. Under less optimal conditions such as multitasking paradigms, LBFTS errors may be more common (Wolfe, Kosovicheva, & Wolfe, 2022). Furthermore, it may take longer to achieve SA under distraction and the additional cognitive load of task switching may affect gaze strategy as well. We are currently investigating the effect of distraction on eye movements and LBFTS errors during viewing of dynamic road scenes.

Conclusion

We recorded eye movements while drivers detected immediate hazards from real road videos to investigate several assumptions from theories of models of driver awareness of hazards. We found that several principles of eye movement guidance found in work using static scenes apply to dynamic road scenes, and these principles offer mechanistic theories that are otherwise missing from the prevalent driver awareness models such as SA and HP. Eye movements are guided by scene context and strategically deployed to facilitate the anticipation and detection of hazards, which aligns with the idea that anticipating hazard before they appear facilitates their detection. Whereas eye-movement studies try to find gaze metrics as an indicator for achieving SA or HP, our results indicate that gaze alone is neither necessary nor sufficient for detecting hazards. Drivers can sometimes use only peripheral vision to look at hazards for successful detection, other times, they may look at the hazard directly and still miss it. This evidence indicates that whether the driver has looked at the hazard alone should not be used as an indicator for hazard awareness in driver monitoring systems. Instead, gaze measures are most useful when they are interpreted in the context of what drivers looked at and why, to elucidate the cognitive processes that support hazard detection and

dynamic scene perception (Ahlström et al., 2021; Rayner, 2009). These results align with a body of literature suggesting that examining both the location and timing of gaze can reveal important aspects of scene perception and provide empirical evidence that elaborates on information acquisition framework and some aspects of the HP framework. This work extends our knowledge of how the visual system copes with the demands of complex, daily tasks, and has practical implications for the design of gaze-based driver monitoring systems.

References

- Åbele, L., Hausteine, S., & Möller, M. (2018). *Novice drivers' eye movement patterns in potentially hazardous pedestrian events: Differences between novice drivers with high and low hazard perception skills*. In Proceedings of the 7th Transport Research Arena (TRA 2018), Vienna, Austria, April 16–19, 2018.
- Ahlström, C., Kircher, K., Nyström, M., & Wolfe, B. (2021). Eye tracking in driver attention research—How gaze data interpretations influence what we learn. *Frontiers in Neuroergonomics*, 2, Article 778043. <https://doi.org/10.3389/fnrgo.2021.778043>
- Andrews, T. J., & Coppola, D. M. (1999). Idiosyncratic characteristics of saccadic eye movements when viewing different visual environments. *Vision Research*, 39(17), 2947–2953. [https://doi.org/10.1016/S0042-6989\(99\)00019-X](https://doi.org/10.1016/S0042-6989(99)00019-X)
- Arias-Portela, C. Y., Mora-Vargas, J., & Caro, M. (2024). Situational awareness assessment of drivers boosted by eye-tracking metrics: A literature review. *Applied Sciences*, 14(4), Article 1611. <https://doi.org/10.3390/app14041611>
- Arutyunova, K., Burashnikov, E., Timakin, N., Shishalov, I., Filimonov, A., & Bakhchina, A. (2025). Eye gaze entropy reflects individual experience in the context of driving. *Entropy*, 28(1), Article 8. <https://doi.org/10.3390/e28010008>
- Beanland, V., Lenné, M. G., & Underwood, G. (2014). Safety in numbers: Target prevalence affects the detection of vehicles during simulated driving. *Attention, Perception & Psychophysics*, 76(3), 805–813. <https://doi.org/10.3758/s13414-013-0603-1>
- Ben-Shachar, M., Lüdtke, D., & Makowski, D. (2020). effectsizes: Estimation of effect size indices and standardized parameters. *Journal of Open Source Software*, 5(56), Article 2815. <https://doi.org/10.21105/joss.02815>
- Biswas, A., Gupta, P., Khurana, S., Held, D., & Admoni, H. (2024). *Modeling drivers' situational awareness from eye gaze for driving assistance*. 8th Annual Conference on Robot Learning. <https://openreview.net/forum?id=MwZJ96Ok13>
- Boufous, S., Ivers, R., Senserrick, T., & Stevenson, M. (2011). Attempts at the practical on-road driving test and the hazard perception test and the risk of traffic crashes in young drivers. *Traffic Injury Prevention*, 12(5), 475–482. <https://doi.org/10.1080/15389588.2011.591856>
- Brainard, D. H. (1997). The psychophysics toolbox. *Spatial Vision*, 10(4), 433–436. <https://doi.org/10.1163/156856897X00357>
- Bühren, J. (2014). ETDRS visual acuity chart. In U. Schmidt-Erfurth & T. Kohnen (Eds.), *Encyclopedia of ophthalmology* (p. 1). Springer. https://doi.org/10.1007/978-3-642-35951-4_617-2
- Cao, S., Samuel, S., Murzello, Y., Ding, W., Zhang, X., & Niu, J. (2022). Hazard perception in driving: A systematic literature review. *Transportation Research Record*, 2676(12), 666–690. <https://doi.org/10.1177/03611981221096666>
- Cardona, G., García, C., Serés, C., Vilaseca, M., & Gispets, J. (2011). Blink rate, blink amplitude, and tear film integrity during dynamic visual display terminal tasks. *Current Eye Research*, 36(3), 190–197. <https://doi.org/10.3109/02713683.2010.544442>

- Castelhano, M. S., Mack, M. L., & Henderson, J. M. (2009). Viewing task influences eye movement control during active scene perception. *Journal of Vision*, 9(3), Article 6. <https://doi.org/10.1167/9.3.6>
- Chen, S., Epps, J., & Paas, F. (2023). Pupillometric and blink measures of diverse task loads: Implications for working memory models. *The British Journal of Educational Psychology*, 93(Suppl 2), 318–338. <https://doi.org/10.1111/bjep.12577>
- Chihara, L. M. & Hesterberg, T. (2018). Introduction to hypothesis testing. In L. M. Chihara & T. C. Hesterberg (Eds.), *Mathematical statistics with resampling and R* (2nd ed., pp. 47–73). Wiley. <https://doi.org/10.1002/9781119505969.ch3>
- Clarke, A. D. F., Mahon, A., Irvine, A., & Hunt, A. R. (2017). People are unable to recognize or report on their own eye movements. *Quarterly Journal of Experimental Psychology*, 70(11), 2251–2270. <https://doi.org/10.1080/17470218.2016.1231208>
- Collingridge, D. S. (2013). A primer on quantized data analysis and permutation testing. *Journal of Mixed Methods Research*, 7(1), 81–97. <https://doi.org/10.1177/1558689812454457>
- Cornelissen, F. W., Peters, E. M., & Palmer, J. (2002). The Eyelink toolbox: Eye tracking with MATLAB and the Psychophysics Toolbox. *Behavior Research Methods, Instruments, & Computers*, 34(4), 613–617. <https://doi.org/10.3758/BF03195489>
- Costela, F. M., & Castro-Torres, J. J. (2020). Risk prediction model using eye movements during simulated driving with logistic regressions and neural networks. *Transportation Research Part F: Traffic Psychology and Behaviour*, 74, 511–521. <https://doi.org/10.1016/j.trf.2020.09.003>
- Crundall, D., Chapman, P., Trawley, S., Collins, L., van Loon, E., Andrews, B., & Underwood, G. (2012). Some hazards are more attractive than others: Drivers of varying experience respond differently to different types of hazard. *Accident Analysis and Prevention*, 45, 600–609. <https://doi.org/10.1016/j.aap.2011.09.049>
- Ding, W., Murzello, Y., Cao, S., & Samuel, S. (2024). Exploring the relationship between drivers' stationary gaze entropy and situation awareness in a level 3 automation driving simulation. *Proceedings of the Human Factors and Ergonomics Society Annual Meeting*, 68(1), 879–884. <https://doi.org/10.1177/10711813241275910>
- Drew, G. C. (1951). Variations in reflex blink-rate during visual-motor tasks. *Quarterly Journal of Experimental Psychology*, 3(2), 73–88. <https://doi.org/10.1080/17470215108416776>
- Eckstein, M. P., Drescher, B. A., & Shimozaki, S. S. (2006). Attentional cues in real scenes, saccadic targeting, and Bayesian priors. *Psychological Science*, 17(11), 973–980. <https://doi.org/10.1111/j.1467-9280.2006.01815.x>
- Ehinger, K. A., Hidalgo-Sotelo, B., Torralba, A., & Oliva, A. (2009). Modelling search for people in 900 scenes: A combined source model of eye guidance. *Visual Cognition*, 17(6–7), 945–978. <https://doi.org/10.1080/13506280902834720>
- Endsley, M. R. (1988). Situation awareness global assessment technique (SAGAT). In *Proceedings of the IEEE 1988 National Aerospace and Electronics Conference* (Vol. 3, pp. 789–795). IEEE. <https://doi.org/10.1109/NAECON.1988.195097>
- Endsley, M. R. (1995a). Toward a theory of situation awareness in dynamic systems. *Human Factors*, 37(1), 32–64. <https://doi.org/10.1518/001872095779049543>
- Endsley, M. R. (1995b). Measurement of situation awareness in dynamic systems. *Human Factors*, 37(1), 65–84. <https://doi.org/10.1518/001872095779049499>
- Endsley, M. R. (2020). Situation Awareness in Driving. In J. D. Lee, W. J. Horrey, D. L. Fisher, & M. A. Regan (Eds.), *Handbook of human factors for automated, connected, and intelligent vehicles* (1st ed., pp. 127–152). CRC Press.
- Fisher, D. L., Pollatsek, A. P., & Pradhan, A. (2006). Can novice drivers be trained to scan for information that will reduce their likelihood of a crash? *Injury Prevention*, 12(Suppl 1), i25–i29. <https://doi.org/10.1136/ip.2006.012021>
- Foulsham, T., & Kingstone, A. (2013). Where have eye been? Observers can recognise their own fixations. *Perception*, 42(10), 1085–1089. <https://doi.org/10.1068/p7562>
- Hayhoe, M. M., McKinney, T., Chajka, K., & Pelz, J. B. (2012). Predictive eye movements in natural vision. *Experimental Brain Research*, 217(1), 125–136. <https://doi.org/10.1007/s00221-011-2979-2>
- Hayhoe, M. M., Shrivastava, A., Mruczek, R., & Pelz, J. B. (2003). Visual memory and motor planning in a natural task. *Journal of Vision*, 3(1), 49–63. <https://doi.org/10.1167/3.1.6>
- Henderson, J. M. (2017). Gaze control as prediction. *Trends in Cognitive Sciences*, 21(1), 15–23. <https://doi.org/10.1016/j.tics.2016.11.003>
- Henderson, J. M., Weeks, P. A., Jr., & Hollingworth, A. (1999). The effects of semantic consistency on eye movements during complex scene viewing. *Journal of Experimental Psychology: Human Perception and Performance*, 25(1), 210–228. <https://doi.org/10.1037/0096-1523.25.1.210>
- Herslund, M.-B., & Jørgensen, N. O. (2003). Looked-but-failed-to-see errors in traffic. *Accident Analysis and Prevention*, 35(6), 885–891. [https://doi.org/10.1016/S0001-4575\(02\)00095-7](https://doi.org/10.1016/S0001-4575(02)00095-7)
- Hesterberg, T., Monaghan, S., Moore, D. S., Clipson, A., & Epstein, R. (2003). *Bootstrap methods and permutation tests*. Freeman and Company.
- Hills, B. L. (1980). Vision, visibility, and perception in driving. *Perception*, 9(2), 183–216. <https://doi.org/10.1068/p090183>
- Hofbauer, M., Kuhn, C. B., Püttner, L., Petrovic, G., & Steinbach, E. (2020). Measuring driver situation awareness using region-of-interest prediction and eye tracking. In *2020 IEEE International Symposium on Multimedia (ISM)* (pp. 91–95). IEEE. <https://doi.org/10.1109/ISM.2020.00022>
- Horswill, M. S. (2016). Hazard perception in driving. *Current Directions in Psychological Science*, 25(6), 425–430. <https://doi.org/10.1177/0963721416663186>
- Horswill, M. S., Hill, A., & Wetton, M. (2015). Can a video-based hazard perception test used for driver licensing predict crash involvement? *Accident Analysis and Prevention*, 82, 213–219. <https://doi.org/10.1016/j.aap.2015.05.019>
- Hout, M. C., Walenchok, S. C., Goldinger, S. D., & Wolfe, J. M. (2015). Failures of perception in the low-prevalence effect: Evidence from active and passive visual search. *Journal of Experimental Psychology: Human Perception and Performance*, 41(4), 977–994. <https://doi.org/10.1037/xhp0000053>
- Huestegge, L., & Böckler, A. (2016). Out of the corner of the driver's eye: Peripheral processing of hazards in static traffic scenes. *Journal of Vision*, 16(2), Article 11. <https://doi.org/10.1167/16.2.11>
- Huestegge, L., Skottke, E.-M., Anders, S., Müsseler, J., & Debus, G. (2010). The development of hazard perception: Dissociation of visual orientation and hazard processing. *Transportation Research Part F: Traffic Psychology and Behaviour*, 13(1), 1–8. <https://doi.org/10.1016/j.trf.2009.09.005>
- Jiang, Y., Xu, Q., Zhen, K., & Chen, Y. (2024). *Quantitative evaluation of driver's situation awareness in virtual driving through eye tracking analysis* (arXiv:2404.14817). arXiv. <https://doi.org/10.48550/arXiv.2404.14817>
- Kahana-Levy, N., Shavitzky-Golkin, S., Borowsky, A., & Vakil, E. (2019). The effects of repetitive presentation of specific hazards on eye movements in hazard perception training, of experienced and young-inexperienced drivers. *Accident Analysis and Prevention*, 122, 255–267. <https://doi.org/10.1016/j.aap.2018.09.033>
- Kim, Y. W., & Yoon, S. H. (2025). Understanding drivers' situation awareness in highly automated driving using SAGAT, SART, and eye-tracking data. *Transportation Research Part F: Traffic Psychology and Behaviour*, 109, 1437–1450. <https://doi.org/10.1016/j.trf.2025.02.003>
- Kleiner, M., Brainard, D., & Pelli, D. (2007). What's new in psychtoolbox-3? *Perception*, 36(Suppl 1), 1–16. <https://doi.org/10.1177/030100660703605101>

- Konstantopoulos, P., Chapman, P., & Crundall, D. (2010). Driver's visual attention as a function of driving experience and visibility. Using a driving simulator to explore drivers' eye movements in Day, night and rain driving. *Accident Analysis and Prevention*, 42(3), 827–834. <https://doi.org/10.1016/j.aap.2009.09.022>
- Kosovicheva, A., Wolfe, J. M., & Wolfe, B. (2023). Taking prevalence effects on the road: Rare hazards are often missed. *Psychonomic Bulletin & Review*, 30(1), 212–223. <https://doi.org/10.3758/s13423-022-02159-0>
- Krzyś, K. J., Avitur, C., Williams, C. C., & Castelano, M. S. (2025). Object spatial certainty as a measure of spatial variability and its influence on attention. *Scientific Reports*, 15(1), Article 11263. <https://doi.org/10.1038/s41598-025-93265-1>
- Lai, H.-Y. (2025). Framework of adaptive driving: Linking situation awareness, driving goals, and driving intentions using eye-tracking and vehicle kinetic data. *IEEE Transactions on Intelligent Transportation Systems*, 26(3), 3295–3306. <https://doi.org/10.1109/TITS.2025.3530252>
- Land, M., Mennie, N., & Rusted, J. (1999). The roles of vision and eye movements in the control of activities of daily living. *Perception*, 28(11), 1311–1328. <https://doi.org/10.1068/p2935>
- Liang, N., Yang, J., Yu, D., Prakah-Asante, K. O., Curry, R., Blommer, M., Swaminathan, R., & Pitts, B. J. (2021). Using eye-tracking to investigate the effects of pre-takeover visual engagement on situation awareness during automated driving. *Accident Analysis and Prevention*, 157, Article 106143. <https://doi.org/10.1016/j.aap.2021.106143>
- Mahon, A., Clarke, A. D. F., & Hunt, A. R. (2018). The role of attention in eye-movement awareness. *Attention, Perception & Psychophysics*, 80(7), 1691–1704. <https://doi.org/10.3758/s13414-018-1553-4>
- Marti, S., Bayet, L., & Dehaene, S. (2015). Subjective report of eye fixations during serial search. *Consciousness and Cognition*, 33, 1–15. <https://doi.org/10.1016/j.concog.2014.11.007>
- Mourant, R. R., & Rockwell, T. H. (1972). Strategies of visual search by novice and experienced drivers. *Human Factors*, 14(4), 325–335. <https://doi.org/10.1177/001872087201400405>
- Neider, M. B., & Zelinsky, G. J. (2006). Scene context guides eye movements during visual search. *Vision Research*, 46(5), 614–621. <https://doi.org/10.1016/j.visres.2005.08.025>
- Oliva, A. (2005). Gist of the scene. In O. P. John, R. W. Robins, & L. A. Pervin (Eds.), *Neurobiology of attention* (pp. 251–256). Elsevier.
- Pastukhov, A. (2024). *eyelinkReader: Import gaze data for EyeLink eye tracker*. <https://github.com/alexander-pastukhov/eyelinkReader/>
- Peelen, M. V., Berlot, E., & de Lange, F. P. (2024). Predictive processing of scenes and objects. *Nature Reviews Psychology*, 3(1), 13–26. <https://doi.org/10.1038/s44159-023-00254-0>
- Pelli, D. G. (1997). The VideoToolbox software for visual psychophysics: Transforming numbers into movies. *Spatial Vision*, 10(4), 437–442. <https://doi.org/10.1163/156856897X00366>
- Pollatsek, A., Narayana, V., Pradhan, A., & Fisher, D. L. (2006). Using eye movements to evaluate a PC-based risk awareness and perception training program on a driving simulator. *Human Factors*, 48(3), 447–464. <https://doi.org/10.1518/001872006778606787>
- Poulton, E. C., & Gregory, R. L. (1952). Blinking during visual tracking. *Quarterly Journal of Experimental Psychology*, 4(2), 57–65. <https://doi.org/10.1080/17470215208416604>
- R Core Team. (2017). *R: A language and environment for statistical computing* [Computer software]. R Foundation for Statistical Computing. <https://www.R-project.org/>
- Rayner, K. (2009). Eye movements and attention in reading, scene perception, and visual search. *Quarterly Journal of Experimental Psychology*, 62(8), 1457–1506. <https://doi.org/10.1080/17470210902816461>
- Sampanes, A. C., Tseng, P., & Bridgeman, B. (2008). The role of gist in scene recognition. *Vision Research*, 48(21), 2275–2283. <https://doi.org/10.1016/j.visres.2008.07.011>
- Schmidt, L. J., Belopolsky, A. V., & Theeuwes, J. (2017). The time course of attentional bias to cues of threat and safety. *Cognition & Emotion*, 31(5), 845–857. <https://doi.org/10.1080/02699931.2016.1169998>
- Song, J., Kosovicheva, A., & Wolfe, B. (2024). Road Hazard Stimuli: Annotated naturalistic road videos for studying hazard detection and scene perception. *Behavior Research Methods*, 56(4), 4188–4204. <https://doi.org/10.3758/s13428-023-02299-8>
- Stahl, P., Donmez, B., & Jamieson, G. A. (2019). Eye glances towards conflict-relevant cues: The roles of anticipatory competence and driver experience. *Accident Analysis and Prevention*, 132, Article 105255. <https://doi.org/10.1016/j.aap.2019.07.031>
- Tan, X., & Zhang, Y. (2022). The effects of takeover request lead time on drivers' situation awareness for manually exiting from freeways: A web-based study on level 3 automated vehicles. *Accident Analysis and Prevention*, 168, Article 106593. <https://doi.org/10.1016/j.aap.2022.106593>
- The MathWorks Inc. (2023). *MATLAB version: 9.14.0 (R2023b)* [Computer software]. <https://www.mathworks.com>
- Torralba, A., Oliva, A., Castelano, M. S., & Henderson, J. M. (2006). Contextual guidance of eye movements and attention in real-world scenes: The role of global features in object search. *Psychological Review*, 113(4), 766–786. <https://doi.org/10.1037/0033-295X.113.4.766>
- Underwood, G., Chapman, P., Brocklehurst, N., Underwood, J., & Crundall, D. (2003). Visual attention while driving: Sequences of eye fixations made by experienced and novice drivers. *Ergonomics*, 46(6), 629–646. <https://doi.org/10.1080/0014013031000090116>
- Underwood, G., Crundall, D., & Chapman, P. (2002). Selective searching while driving: The role of experience in hazard detection and general surveillance. *Ergonomics*, 45(1), 1–12. <https://doi.org/10.1080/00140130110110610>
- Vater, C., Wolfe, B., & Rosenholtz, R. (2022). Peripheral vision in real-world tasks: A systematic review. *Psychonomic Bulletin & Review*, 29(5), 1531–1557. <https://doi.org/10.3758/s13423-022-02117-w>
- Ventsislavova, P., Gugliotta, A., Peña-Suarez, E., Garcia-Fernandez, P., Eisman, E., Crundall, D., & Castro, C. (2016). What happens when drivers face hazards on the road? *Accident; Analysis and Prevention*, 91, 43–54. <https://doi.org/10.1016/j.aap.2016.02.013>
- Vlakveld, W., Romoser, M. R. E., Mehranian, H., Diete, F., Pollatsek, A., & Fisher, D. L. (2011). Do crashes and near crashes in simulator-based training enhance novice drivers' visual search for latent hazards? *Transportation Research Record*, 2265, 153–160. <https://doi.org/10.3141/2265-17>
- Võ, M. L.-H., Boettcher, S. E., & Draschkow, D. (2019). Reading scenes: How scene grammar guides attention and aids perception in real-world environments. *Current Opinion in Psychology*, 29, 205–210. <https://doi.org/10.1016/j.copsyc.2019.03.009>
- Wells, P., Tong, S., Sexton, B., Grayson, G., & Jones, E. (2008). *Cohort II: A study of learner and new drivers. Volume 1—Main report*. Road Safety Research Report.
- Wetton, M. A., Horswill, M. S., Hatherly, C., Wood, J. M., Pachana, N. A., & Anstey, K. J. (2010). The development and validation of two complementary measures of drivers' hazard perception ability. *Accident Analysis and Prevention*, 42(4), 1232–1239. <https://doi.org/10.1016/j.aap.2010.01.017>
- Wickham, H., Averick, M., Bryan, J., Chang, W., McGowan, L. D., François, R., Golemund, G., Hayes, A., Henry, L., Hester, J., Kuhn, M., Pedersen, T. L., Miller, E., Bache, S. M., Müller, K., Ooms, J., Robinson, D., Seidel, D. P., Spinu, V., . . . Yutani, H. (2019). Welcome to the tidyverse. *Journal of Open Source Software*, 4(43), Article 1686. <https://doi.org/10.21105/joss.01686>
- Wolfe, B., Dobres, J., Rosenholtz, R., & Reimer, B. (2017). More than the useful field: Considering peripheral vision in driving. *Applied Ergonomics*, 65, 316–325. <https://doi.org/10.1016/j.apergo.2017.07.009>

- Wolfe, B., Sawyer, B. D., & Rosenholtz, R. (2022). Toward a theory of visual information acquisition in driving. *Human Factors*, 64(4), 694–713. <https://doi.org/10.1177/0018720820939693>
- Wolfe, B., Seppelt, B., Mehler, B., Reimer, B., & Rosenholtz, R. (2020). Rapid holistic perception and evasion of road hazards. *Journal of Experimental Psychology: General*, 149(3), 490–500. <https://doi.org/10.1037/xge0000665>
- Wolfe, J. M. (2021). Guided Search 6.0: An updated model of visual search. *Psychonomic Bulletin & Review*, 28(4), 1060–1092. <https://doi.org/10.3758/s13423-020-01859-9>
- Wolfe, J. M., Kosovicheva, A., & Wolfe, B. (2022). Normal blindness: When we look but fail to see. *Trends in Cognitive Sciences*, 26(9), 809–819. <https://doi.org/10.1016/j.tics.2022.06.006>
- Yang, J., Liang, N., Pitts, B. J., Prakah-Asante, K., Curry, R., & Yu, D. (2024). An eye-fixation related electroencephalography technique for predicting situation awareness: Implications for driver state monitoring systems. *Human Factors*, 66(8), 2138–2153. <https://doi.org/10.1177/00187208231204570>
- Yarbus, A. L. (1967). *Eye movements and vision*. Springer. <https://doi.org/10.1007/978-1-4899-5379-7>
- Zhang, T., Yang, J., Liang, N., Pitts, B. J., Prakah-Asante, K., Curry, R., Duerstock, B., Wachs, J. P., & Yu, D. (2023). Physiological measurements of situation awareness: A systematic review. *Human Factors*, 65(5), 737–758. <https://doi.org/10.1177/0018720820969071>

Received September 5, 2025

Revision received March 6, 2026

Accepted April 10, 2026 ■

Reproduced with permission of copyright owner. Further reproduction
prohibited without permission.