

UrbanGazeVis: A Visualization System for Analyzing Eye-Tracking Data on Urban Safety Perception

Andres De La Puente^a, Luis Sante^{a,1}, Felipe Moreno-Vera^{a,1}, Mauro Diaz^{a,1}, Jorge Poco^{a,1}

^aFundação Getúlio Vargas, Praia de Botafogo, 190 - Botafogo, Rio de Janeiro - RJ, 22250-145, Brazil

ARTICLE INFO

Article history:

Received August 26, 2026

Keywords: Urban Safety Perception, Visual attention, Eye-tracking, Urban Physical Disorder, Street-level imagery

ABSTRACT

Perceived safety in streetscapes depends on where people look, yet how gaze relates to visual cues of urban disorder remains poorly understood. Prior work treats safety as an image-level label, offering little insight into how attention to specific elements (e.g. buildings, greenery, people, signs of decay) shapes these judgments. We present a head-mounted eye-tracking study in which 30 participants viewed and rated the safety of 150 street-view images from Rio de Janeiro using a HoloLens 2 headset. Gaze traces were mapped onto semantic segments and disorder cues (e.g., damaged walls, graffiti, overhead cables), yielding a multimodal dataset linking gaze dynamics, scene semantics, and safety scores. To analyze it, we introduce URBANGAZEVis, an interactive visual analytics system with image- and participant-centric views that connects the spatial, temporal, and semantic dimensions of gaze to perceived safety, supporting comparisons between safe and unsafe scenes, inspection of divergent ratings for similar images, and region-of-interest analysis via glyph-based summaries. Statistical models show that sustained attention to physical disorder is associated with lower perceived safety, while the visual analysis reveals context-specific effects often masked by global aggregation. Together, these analyses offer actionable insights for urban design and planning.

© 2026 Elsevier B.V. All rights reserved.

1. Introduction

Perceived safety in streetscapes is strongly shaped by visual attention, yet the relationship between gaze behavior and urban disorder remains poorly understood. Broken Windows Theory [1] and subsequent empirical work [2, 3] suggest that visible signs of neglect, such as deteriorated façades and damaged infrastructure, reduce perceived safety and reinforce cycles of disorder, with downstream effects on behavior, crime, and mental health. Yet most approaches model safety at the image level, even though recent studies [4] suggest that individual disorder cues, such as overhead cables or structural damage, may affect perception differently. Real-world streetscapes are also structurally and socially heterogeneous: as Fig. 1 shows, some areas

feature well-maintained, orderly infrastructure (Fig. 1.a), while others show pronounced disorder, such as degraded walls, litter, and irregular architecture (Fig. 1.b). Understanding how attention navigates these conflicting cues requires moving beyond static, image-level safety scores toward a spatiotemporal visual analytics approach. Street View Imagery (SVI) has enabled large-scale safety studies pairing street scenes with crowd-sourced judgments, as in StreetScore [5], Scenic-or-Not [6], and Place Pulse [7]. However, these approaches mainly capture overall safety or aesthetic appeal, treating scenes holistically rather than revealing where people look or how attention unfolds over time during evaluation [8, 9, 10]. Recent immersive and eye-tracking studies begin to address this: virtual and mixed-reality setups allow controlled study of factors like lighting, greenery, and density in relation to safety, beauty, and liveliness [11], and eye-tracking in urban contexts [12] identifies

e-mail: andres.puente@fgv.br (Andres De La Puente)

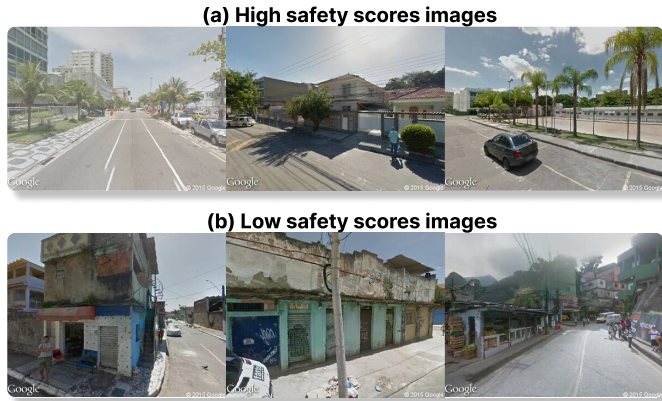


Figure 1: **Visual and structural heterogeneity within the URBANGAZE dataset.** These representative Street View panoramas from Rio de Janeiro illustrate the contrasting urban stimuli evaluated in our study. (a) Well-maintained environments characterized by orderly infrastructure, clean sidewalks, and preserved vegetation, providing examples associated with higher perceived safety. (b) Streetscapes exhibiting pronounced physical disorder, including degraded façades, uncollected litter, and disorganized overhead cables.

regions that draw attention. Still, results are typically summarized with static, aggregated fixation maps, obscuring the dynamic, temporal, and semantic dimensions of attention.

In this work, we adopt a visualization-first perspective. We conduct a head-mounted eye-tracking experiment using 150 street-view scenes from Rio de Janeiro, which participants rate in terms of perceived safety. We use a head-mounted optical see-through (OST) headset as a self-contained apparatus for controlled 2D stimulus presentation and gaze capture. Gaze traces are projected onto semantically segmented scenes (e.g., buildings, sidewalks, and trees) and annotated with disorder cues, including deteriorated walls, overhead cables, garbage, and graffiti. This produces a multimodal dataset integrating scene content, visual attention patterns, and perceived safety ratings. To analyze these data, we introduce URBANGAZEVis, a visualization system for exploring eye-tracking data in urban safety studies. URBANGAZEVis supports post hoc inspection of spatiotemporal gaze behavior in relation to semantic regions and safety judgments, linking gaze trajectories, fixation patterns, and scene elements to support both hypothesis generation and evidence assessment for researchers and practitioners. Our contributions are threefold:

Head-mounted eye-gaze collection methodology. A head-mounted eye-tracking protocol for selecting, filtering, and assigning SVI, including calibration, trial design, and quality control, to collect curated safety judgments.

URBANGAZE dataset. A multimodal dataset of Rio de Janeiro street-view images with safety ratings, eye-tracking data, semantic segmentations, and participant metadata for urban perception research.

URBANGAZEVis visualization system. An interactive system that overlays dynamic gaze information on semantically segmented scenes, enabling exploration of gaze paths and fixation patterns in relation to safety scores through domain-inspired case studies.

All datasets, supplemental materials, and source code will be made available upon acceptance, in accordance with the venue’s

anonymization guidelines.

2. Related Work

This research lies at the intersection of five pillars: (i) Urban Perception and Computer Vision, (ii) Environmental Criminology, (iii) Eye-Tracking for Urban Perception, (iv) Immersive and Mixed Reality in Urban Visualization, and (v) Visual Analytics for Eye-Tracking Data.

Urban Perception and Computer Vision. Urban perception research examines how the built environment is visually characterized and how people judge properties such as safety from street-level imagery [13, 14, 15]. Large-scale image collections paired with crowdsourced judgments, such as Place Pulse and StreetScore [16, 5], support questions such as “*What makes Paris look like Paris?*” [17] and “*What makes a place feel safe?*” [18]. Early work relied on hand-crafted features; later approaches adopted CNNs to predict perceived attributes [19, 20, 7], and more recent methods relate semantic scene content to safety and other judgments [21, 22], bridging perception studies and visual analysis. Most of these approaches treat perception as an *image-level* label, learning from global ratings without observing how viewers explore a scene. Our work extends this line by linking safety judgments to *where* observers look, using semantic classes and disorder cues in an interactive visual analytics workflow.

Environmental Criminology. While computer vision frameworks traditionally model urban perception through passive image-level feature extraction, environmental criminology establishes safety perception as a dynamic cognitive tension between physical infrastructure and social activity. This perspective draws on foundational frameworks in environmental criminology—routine activity theory [23], crime pattern theory [24], and situational crime prevention [25]—which we invoke not as models of crime occurrence but as an interpretive lens for how everyday spatial routines and the configuration of the built environment shape where people perceive safety or threat. Jane Jacobs’s “eyes on the street” [26] and Oscar Newman’s “defensible space” [27] further assert that human presence and street vitality fundamentally alter how an environment is monitored and interpreted, principles underlying Crime Prevention Through Environmental Design (CPTED), which identifies “natural surveillance” as a primary psychological mitigator of perceived threat. Rather than treating pedestrians, crowds, and street activity as mere geometric or baseline semantic classes, our approach conceptualizes human presence as an active social cue that can dynamically modulate how observers interpret co-occurring signals of physical disorder.

Eye-Tracking for Urban Perception. Eye-tracking methodologies offer a window into spatial cognition, decomposing visual attention into fixations, saccades, and scanpaths [12, 28].

Desktop and head-mounted eye-tracking systems involve different trade-offs in sampling rate, head-movement tolerance, calibration, and stimulus presentation, and the appropriate choice depends on the study design [11, 12]. As synthesized by Moreno-Arjonilla et al. [29], integrating head-mounted eye

trackers within spatial computing frameworks demands rigorous calibration and geometric uncertainty modeling to yield reproducible attentional datasets. Contemporary work increasingly intersects spatiotemporal gaze with deep feature spaces and explainable AI to isolate visual correlates of safety and aesthetic judgment [30, 31]. Although these studies show gaze provides rich information about safety judgments, results are usually presented as static heatmaps, aggregate statistics, or a few example scanpaths—summaries that flatten a dynamic process and obscure how attention shifts over time and across semantic elements. Our work treats the HoloLens and calibration as infrastructure, focusing on structuring and visualizing spatiotemporal gaze data through a safety-judgment protocol and URBANGAZEVis for multimodal exploration.

Immersive and Mixed Reality in Spatial Visualization. The convergence of spatial computing and visual analytics has shifted data exploration from desktop interfaces toward immersive environments. Recent advances emphasize situated visualizations, where geometric data and semantic layouts are embedded within the user’s egocentric frame to enhance cognitive processing and environmental presence [32, 33]. Immersive analytics has evolved to leverage advanced view computation and cross-scale interaction, letting analysts explore large-scale urban representations and complex environmental phenomena fluidly [34, 35]. Furthermore, contemporary frameworks integrate realistic 3D modeling with interactive simulations to evaluate spatial contexts and urban risks in real time [36, 37].

URBANGAZEVis sits within this frontier but with a deliberately narrow scope. We do not use Mixed Reality (MR) for situated 3D visualization; the headset serves only as a self-contained eye tracker that presents 2D SVI and records gaze. Our contribution is not the display medium but the protocol and visual analytics that relate the resulting spatiotemporal gaze data to semantic scene content and fine-grained urban decay.

Visual Analytics for Eye-Tracking Data. Eye-tracking visualization is a mature research vector within visual analytics. Established studies [38, 39] categorize techniques across spatial, temporal, and hybrid spatiotemporal dimensions, using in-context gaze overlays, timelines, and coordinated views to reduce cognitive overhead. Early paradigms such as space-time cubes [40, 41] and AOI-centric encodings like gaze stripes and rivers [42, 43, 44] optimized localized sequence extraction. More recently, the VETA system [45] provides diagnostic interfaces for spatial gaze distributions, while SVATA [46] introduces infrastructure for multi-user fixation analysis, implementing an Area-Fixation Weight (AFW) metric that balances spatial distribution bias—conceptually mirroring the area-normalized attention intensity metric (AttIn) we propose in Sec. 6.2. Both frameworks confront spatial dominance artifacts by discount-weighting raw fixation durations relative to bounding visual area, emphasizing genuinely salient features over structurally dominant background regions.

Building on these ideas, URBANGAZEVis targets a different setting: rather than abstract AOIs or screen regions, we project gaze onto semantic segments (e.g., buildings, greenery, sky) and urban disorder cues (e.g., damaged façades, exposed cables) in SVI, linking these encodings to time-resolved

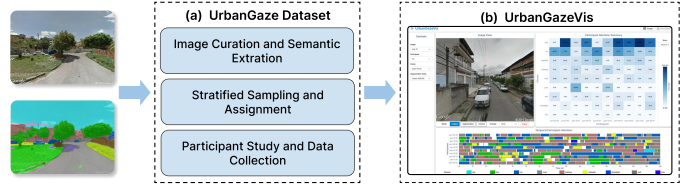


Figure 2: URBANGAZEVis workflow. (a) Dataset construction: street-view scenes are semantically processed, sampled, and used in the participant study to collect gaze data and safety ratings. (b) System analysis: the resulting multimodal data are explored interactively to analyze spatial and temporal patterns of visual attention and perceived safety.

views and safety scores. This lets analysts explore *which* urban elements attract attention, *when* they are inspected, and *how* patterns differ between scenes perceived as safe or unsafe—adapting general eye-tracking visualization concepts to a domain-specific, semantics-aware analysis of urban streetscapes.

3. System Overview

Eye-tracking research shows that gaze data are high-volume and spatiotemporal, yet perceptual judgments are often reduced to image-level scores, leaving interactive exploration of gaze–semantics–rating relationships limited [38, 39]. Drawing on “overview first, zoom and filter, details-on-demand” [47] and established gaze-analysis task abstractions [38], we formulated URBANGAZEVis’s design goals and analytical tasks.

3.1. Design Goals

Informed by eye-tracking taxonomies [38], spatiotemporal visual analytics [39], and urban perception studies interpretability [21, 22], our design goals are: **DG1: Support image- and participant-centric analysis** of how multiple individuals explore one streetscape and how one individual behaves across scenes, using consistent encodings [38]; **DG2: Characterize spatiotemporal gaze behavior**, revealing temporal dynamics and scanpath differences hidden by static heatmaps and aggregates [39]; and **DG3: Relate gaze, semantics, and safety judgments** by linking gaze traces to semantic segments, disorder cues, and ratings [21, 22, 18]. These goals motivate dual views (*By-Image* and *By-Participant*) with coordinated summaries of spatial, temporal, and semantic gaze patterns.

3.2. Analytical Tasks

Grounded in these goals and eye-tracking task taxonomies [38], we define analytical tasks following an overview-to-detail pattern [47]: from a global summary of gaze and ratings, to inspecting images or participants, to detailed traces:

T1: Spatial attention analysis (DG1, DG2) summarizes which regions and semantic elements attract attention across participants for a selected image (**T1.1**), and where a selected participant tends to look across all assessed images (**T1.2**).

T2: Temporal gaze pattern exploration (DG1, DG2) visualizes how attention moves across an image over time to compare scanpaths between participants (**T2.1**), and how a participant’s

gaze evolves across a session to detect learning or fatigue effects (T2.2).

T3: Linking gaze, semantics, and safety assessments (DG3) relates fixations on semantic categories and disorder cues to the safety rating for a specific image (T3.1), and compares a participant’s semantic focus and gaze coverage with their score distribution to inspect rating consistency across visually similar images (T3.2).

These tasks inform URBANGAZEVis’s two main components: a *By-Image* view for image-centric tasks (T1.1, T2.1, T3.1) and a *By-Participant* view for participant-centric tasks (T1.2, T2.2, T3.2), detailed in Sec. 6.

4. URBANGAZE Dataset

Fig. 2 summarizes our end-to-end workflow, from multimodal dataset curation using SVI and head-mounted eye-tracking to the interactive visual analytics environment detailed in Sec. 6. This section covers the first phase: constructing URBANGAZE.

4.1. Image Curation and Semantic Extraction

We curated a compact but diverse set of streetscapes, distributed evenly across participants to support the analyses enabled by URBANGAZEVis. We used the Urban Physical Disorder-4k (UrbanPD4k) dataset [4], which contains Rio de Janeiro SVIs with pixel-level annotations of physical disorder cues. We applied the OneFormer [48] model pre-trained on ADE20K dataset to obtain urban elements. To support fine-grained analysis in URBANGAZEVis, we define four cases based on data:

ADE20K (baseline): 150 segmentation classes (e.g., car, building, tree, road).

Grouped-ADE (7 macro-classes): 7 segmentation group classes: *City Elements, Construction, Floor, Human, Sky, Terrain Vehicle, and Vegetation*.

ADE20K+UrbanPD4k: Adds physical disorder classes: *Overhead cables, Graffiti, Broken/Damaged Bricks Wall, Damaged Pavement/Road, Car Taxi, Damaged Traffic Sign, Garbage Bag, Garbage Box, Kiosk, Exposed Overhead Cables, and Trashcan*.

Grouped-ADE + UrbanPD4k: Groups and disorder classes.

4.2. Multivariate Stratified Sampling

We used *multivariate stratified sampling* to select a subset suitable for a controlled experiment while preserving the visual diversity of the full dataset. Rather than manually normalizing for varying traffic or pedestrian counts, we formulated the selection as an Integer Linear Programming (ILP) problem over the baseline ADE20K segmentation layouts. For each structural and social class (e.g., buildings, roads, vegetation, cars, pedestrians), pixel-ratio distributions were partitioned into quartiles, and the ILP solved for a balanced sample of 150 images with equal representation across all baseline strata (e.g., vehicles, pedestrians, vegetation), though it does not model all acquisition conditions such as time of day. This balancing mitigates the risk of any single layout category acting as an uncontrolled

confound, giving our post-hoc regression models the statistical stability to isolate the perceptual weights of general infrastructure and co-occurring disorder features. Mathematical details of the ILP objective and constraints are in Appendix A.

4.3. Image Assignment Protocol

To balance the experimental design without participant fatigue, the 150 images were distributed via randomized block assignment. The 30 participants were split into three blocks of 10 raters each, with each block evaluating a unique subset of 50 streetscapes drawn from the ILP strata. Within each evaluation session, these 50 images were presented to the participant in a completely randomized sequence. This guarantees every image was evaluated by exactly 10 distinct participants, mitigating systematic presentation-order bias while maintaining structural symmetry across the sparse measurement matrix (full matrix in Appendix B).

4.4. Participant Study and Data Collection

To build URBANGAZE, we conducted a controlled experiment in which participants rated the safety of SVI while their eye movements were recorded.

Participant demographics and inclusion criteria. We recruited 30 volunteers via institutional mailing lists, with informed consent obtained prior to participation. The cohort included 16 undergraduate, 9 master, and 5 doctoral students in fields such as urban planning, computer science, and social sciences. Twenty participants were Brazilian (from seven states), seven Peruvian, and three Paraguayan. Ages ranged from 19 to 43 (median = 23.8, SD = 5.2), with 20 men and 10 women. Participants with conditions incompatible with head-mounted eye tracking were excluded. In total, the study yielded 1,500 trials. We used a Microsoft HoloLens 2 solely as a head-mounted eye tracker, chosen to integrate calibration, stimulus presentation, and gaze capture in one setup. Each SVI was shown as a flat 2D stimulus on a frontoparallel virtual plane at 2.0 m (800×600 px); using the standard Eye Tracking API [49], we recorded the single eye-gaze ray (origin and direction) at approximately 30 Hz and intersected it with this plane to obtain 2D coordinates. Dedicated desktop trackers offer higher sampling rates and remain viable alternatives; we account for the resulting spatial uncertainty through the smoothing model in Sec. 5. Each session began with a nine-point gaze calibration and a brief familiarization phase with reference images. During main trials, participants viewed each image for 15 seconds while gaze data were recorded continuously, then rated perceived safety on a 1–10 Likert scale via a handheld controller. Sessions lasted about 25 minutes; additional protocol details are in Appendix C.

Experimental Controls and Noise Mitigation. All sessions were conducted in a dedicated, climate-controlled laboratory. Participants faced a blank, matte-white wall, and stimuli were rendered at maximum opacity so that the physical background could not show through or trigger exogenous saccades and peripheral distractions. Each entry in the resulting URBANGAZE dataset includes:

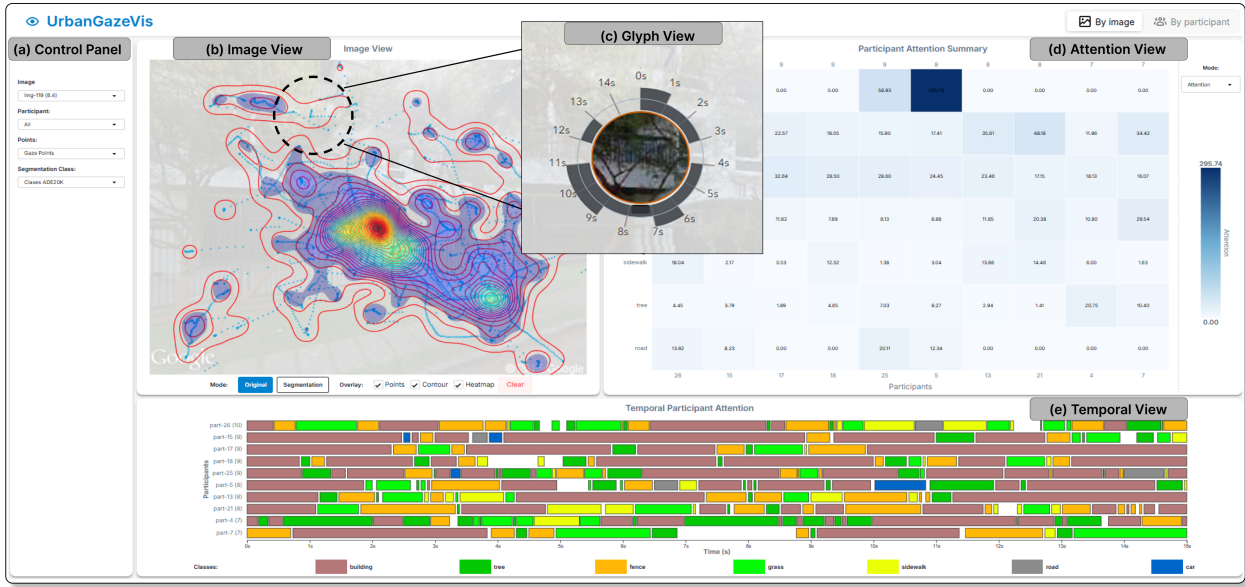


Figure 3: By-Image component of URBANGAZEVis. **(a)** Control panel for selecting the image, participant filters, point type (gaze or fixation), and semantic representation. **(b)** Image View showing the selected streetscape with optional overlays for gaze points, density contours, and heatmap, and access to the Glyph View for selected regions. **(c)** Glyph View, triggered from a circular region of interest in the Image View, summarizing visual exploration within the selected bounds over the 15 s trial: stacked radial bars encode participant counts per 1 s segment. **(d)** Attention View (*Participant Attention Summary*) summarizing *Total Attention* or *normalized Attention Intensity (AttIn)* by participant (columns) and semantic class (rows). **(e)** Temporal View (*Temporal Participant Attention*) using a scarf-timeline chart to show, for each participant, which semantic class is fixated at each moment during the 15 s viewing period. Participant traceability is maintained interactively via tooltips and single-user filtering.

- **A street-view image** and its geographic coordinates (latitude and longitude);
- **A temporal sequence of gaze points** (x, y, t) recorded during the 15-seconds observation interval;
- **A Likert-scale safety rating (1–10)**, where 1 denotes “very unsafe” and 10 denotes “very safe”.

5. Dataset Validation and Robustness

Before integrating the collected multimodal data into URBANGAZEVis, we conducted a quantitative validation of the dataset, assessing signal stability, inter-rater reliability, and structural validity.

5.1. Hardware Robustness and Spatial Uncertainty

Robustness was first addressed at the hardware level to minimize coordinate mapping uncertainty and segmentation-boundary noise. Due to the inherent physiological and hardware jitter of head-mounted eye tracking, raw fixation points were modeled as attentional catchment areas rather than absolute pixel coordinates. We applied a Gaussian smoothing kernel with a standard deviation of $\sigma = 10.68$ pixels. This smoothing parameter was not empirically tuned; it was analytically derived from the Microsoft HoloLens 2 empirical angular precision (0.24°) and the effective pixel-per-degree resolution of our display setup to map hardware uncertainty into the stimulus space (see Appendix F for the geometric derivation). This provides a reliable foundation for evaluating attention on small, complex urban cues.

5.2. Inter-Rater Reliability Analysis

To evaluate the consistency of urban safety perceptions across participants, we performed a rater calibration analysis tailored for sparse rating matrices. Given our III-Structured Measurement Design (ISMD)—where 30 participants evaluated overlapping but incomplete subsets of the 150 images—traditional Intraclass Correlation Coefficients (ICC) based on balanced ANOVA are mathematically biased.

Following Putka et al. [50], we used a Linear Mixed-Effects Model (LMM) with crossed random effects to isolate the variance components of the 1,500 observations, partitioning total variance into the urban scene (signal), the participant (subjective bias), and residual noise. The model-based reliability coefficient $G(q, k)$ was computed as follows:

$$G(q, k) = \frac{\sigma_{\text{img}}^2}{\sigma_{\text{img}}^2 + \frac{q\sigma_{\text{rat}}^2 + \sigma_{\text{res}}^2}{k}} \quad (1)$$

Where $k = 10$ is the number of raters per image, and $q \approx 0.67$ is the coefficient of non-overlap derived from our sparse matrix layout ($q = 1 - k/N_{\text{total}}$), reflecting that each image was evaluated by one-third of the pool. The analysis yielded a high inter-rater reliability of $G(q, k) = 0.812$. In behavioral research, values above 0.80 indicate robust consensus, indicating that 10 raters per image provide sufficient reliability to capture collective perception in this study.

5.3. Baseline Variance Decomposition

To validate the dataset structure, a baseline LMM was fitted with participants as random effects to account for variance in baseline severity or leniency. This model achieved an ICC of

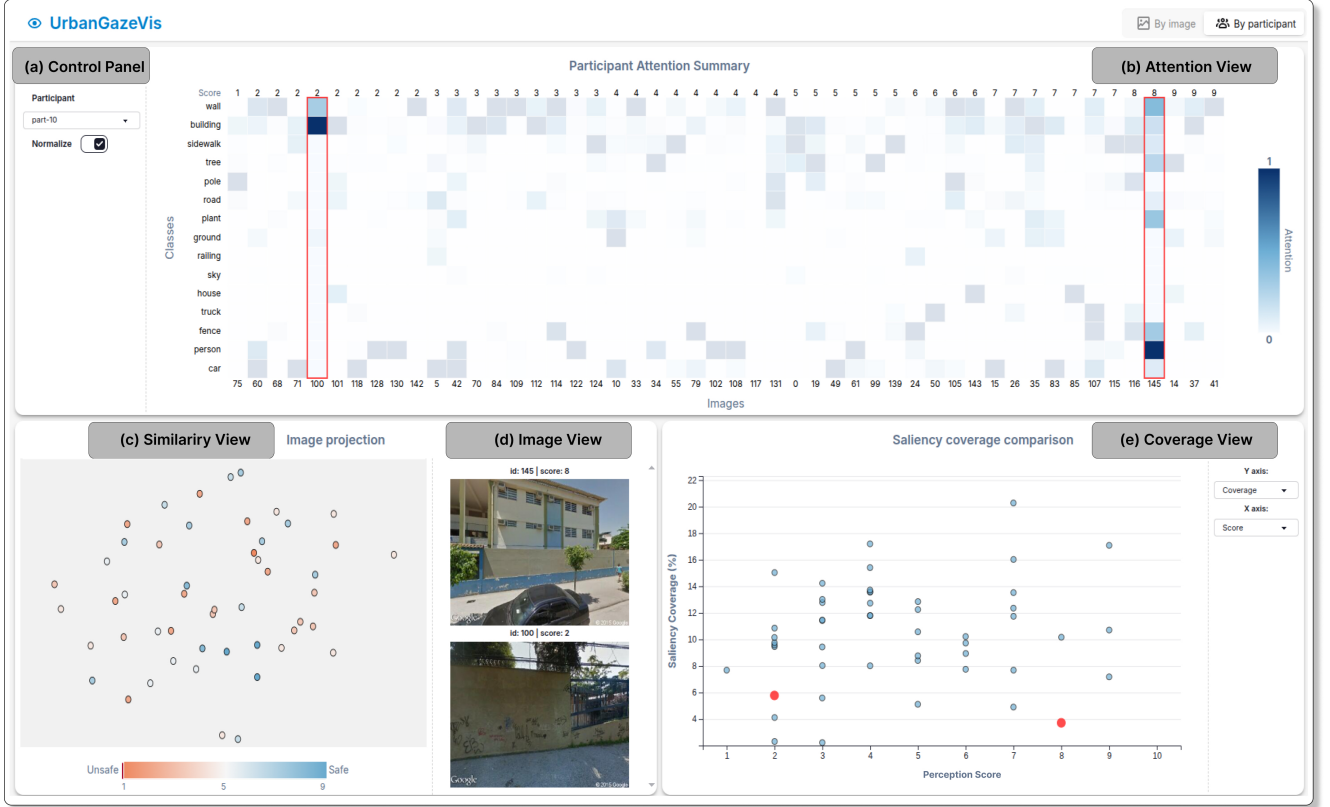


Figure 4: By-Participant component of URBANGAZEVis. (a) Control panel for selecting a participant, choosing the semantic representation, and enabling normalization in the attention heatmap. (b) Attention View showing, for the selected participant, *normalized Attention Intensity* (*AttIn*) by image (columns, ordered by the participant’s safety score) and semantic class (rows). (c) Similarity View (*Image Projection*), where each point represents a Places365 scene embedding projected via t-SNE, colored by the participant’s safety score (red = low, blue = high). (d) Image View showing the images selected from the projection. (e) Coverage View (*Saliency Coverage Comparison*) relating *Saliency Coverage* or entropy to safety scores or image order.

0.203, indicating that approximately one-fifth of the variance in safety scores is attributable to participant-specific baselines, with the remainder attributable to image-level differences and residual factors. These findings confirm that while safety perception has a subjective component, a substantial portion of the variation is linked to image-level differences, justifying the development of an interactive visual analytics platform to isolate these signals.

6. UrbanGazeVis

URBANGAZEVis is an interactive visual analytics system addressing the tasks outlined in Sec. 3, offering two coordinated analysis modes: *By-Image* and *By-Participant*. These modes let analysts alternate between image- and participant-centric perspectives (DG1), characterize spatiotemporal gaze behavior (DG2), and relate gaze patterns to semantic content and safety ratings (DG3); Table 1 summarizes how each view supports tasks T1–T3.

6.1. Data Abstraction and Preprocessing

Raw eye-tracking data form a continuous stream of gaze coordinates comprising fixations and saccades. Since visual information is primarily acquired during fixations, saccadic samples are excluded prior to analysis using the Identification by Velocity Threshold (I-VT) algorithm [51]: velocity is computed

between consecutive samples, and sequences below 1.15 px/ms are classified as fixations, consistent with established protocols for head-mounted displays [52], while samples exceeding this threshold are labeled saccades and excluded. To establish a baseline for *Saliency Coverage*, we construct Empirical Saliency Maps (ESMs) from aggregated human gaze rather than computational saliency models. Gaze points are smoothed with a Gaussian kernel ($\sigma = 10.68$ pixels), analytically derived from the HoloLens 2 empirical angular precision (0.24°) and screen resolution specifications to map hardware uncertainty into the stimulus pixel space (see Appendix F for the geometric derivation).

Spatial dispersion of attention is quantified using *spatial entropy*, defined as Shannon entropy over normalized fixation density [53], a standard measure of gaze dispersion [54, 55]. Fixations are then mapped onto semantic segmentations to estimate *Attention Intensity* for each urban class.

6.2. Attention Intensity Metric

To examine the relationship between the built environment and safety perception, we quantify the attention participants allocate to each visual element [4]. Total attention is computed as accumulated fixation time; however, large objects (e.g., façades or sky) naturally receive more fixations due to their greater pixel area.

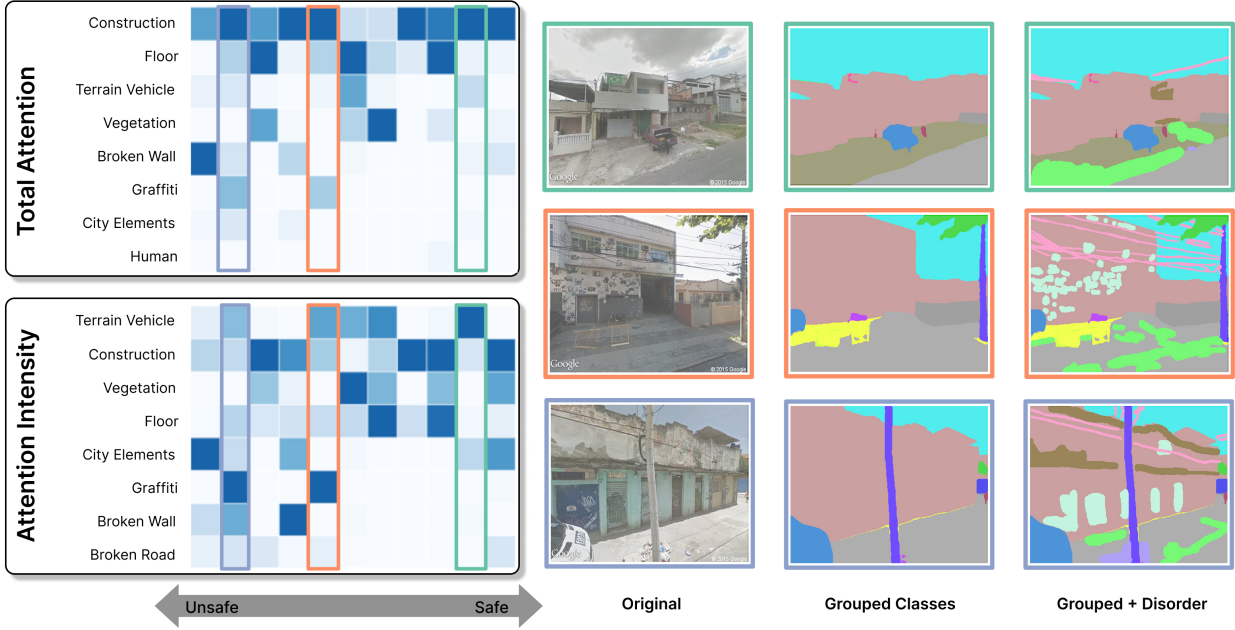


Figure 5: Usage scenario 1: semantic comparison in the By-Participant view for Participant 29. **(Left)** The Attention View shows two stacked matrices—**Total Attention** (top) and **normalized Attention Intensity (AttIn)** (bottom)—with semantic classes as rows (independently ranked by magnitude in each matrix) and the participant’s assessed images as columns, ordered left to right from unsafe to safe. Three representative scenes are highlighted with colored outlines (blue, orange, green). **(Right)** The three highlighted scenes (Image IDs 60, 38, and 109) are shown one per row—each as the original street-view stimulus (**left**), the Grouped classes (**center**), and the Grouped + Disorder classes (**right**)—with row outline colors matching the highlighted columns on the left. Contrasting Total Attention with AttIn enables comparison between attention to general urban structure and to specific disorder cues.

Table 1: URBANGAZEVis views and the analytical tasks they support.

Analysis	View	<i>T1</i>	<i>T2</i>	<i>T3</i>
By-Image	Image View	✓	✓	
	Attention View	✓		✓
	Temporal View		✓	✓
By-Participant	Attention View	✓		✓
	Similarity View			✓
	Coverage View		✓	✓

To reduce this bias, we introduce an area-normalized attention intensity (*AttIn*) that discounts each class by its image area, where $t_{p,i,c}$ denotes the total time (ms) participant p fixated on class c in image i , and $r_{i,c} \in (0, 1]$ is the fraction of pixels belonging to class c .

To avoid instability when normalizing by very small regions ($r_{i,c} \rightarrow 0$), a lower-bound threshold τ is applied; given the eye tracker’s spatial uncertainty, τ is set to 0.005 (0.5%), and regions below this threshold are excluded as indistinguishable from hardware noise (detailed formulation in Appendix F). The resulting metric is defined as follows:

$$AttIn_{i,c} = \frac{1}{P_i} \sum_{p=1}^{P_i} \frac{t_{p,i,c}}{\max(r_{i,c}, \tau)} \quad (2)$$

This formulation highlights genuinely salient small objects, such as graffiti or cables, while maintaining robustness to segmentation noise.

6.3. By-Image Component

The By-Image component supports image-centric analysis, visualizing how participants explore a selected streetscape via a coordinated control panel and four views (Fig. 3). The **Control Panel** (Fig. 3a) enables image selection, sorting by safety scores, participant filtering, and data toggling (e.g., ADE20K or disorder cues) to coordinate linked views (**T1–T3**). The central **Image View** (Fig. 3b) shows the selected streetscape with optional overlays for raw gaze, density contours, or heatmaps, and toggles between original and semantic representations (**T1**). For detailed ROI-centered analysis, a circular Region of Interest (ROI) is instantiated via a single click; analysts can drag and scale this selector to trigger the **Glyph View** (Fig. 3c), summarizing visual exploration within the designated bounds over time (**T2, T3**). The Glyph View’s ring functions as a 15-second clock: stacked radial bars represent participant counts per 1-second segment, revealing early versus sustained attention. Hovering over a segment reports the participant IDs contributing to that interval; filtering the Control Panel to a single participant turns the glyph into an individual 15-second timeline, making sustained attention and revisits explicit. The **Attention View** (Participant Attention Summary, Fig. 3d) relates participants (columns) to semantic classes (rows), encoding total time or our AttIn metric (Eq. 2) via heatmap intensity. Finally, the bottom **Temporal View** (Temporal Participant Attention, Fig. 3e) uses a scarf-timeline chart, where rows represent participants and colored segments encode semantic class fixations (**T2, T3**), revealing scanpath diversity and transitions be-



Figure 6: Usage scenario 2: divergent safety ratings for visually similar scenes in the By-Participant view. Using the Similarity View, we identify image pairs (red boxes) that are close in the projection space but received markedly different scores. The side panels summarize the participant's **Attention Intensity (AttIn)** across visual categories such as vegetation, graffiti, and vehicles. Heatmap intensity is normalized according to the AttIn metric defined in Eq. 2.

tween urban elements.

6.4. By-Participant Component

This component offers a complementary perspective, aggregating behavior across images viewed by a selected participant (Fig. 4). Its **Control Panel** (Fig. 4a) specifies attention types and data, including normalization toggles for the attention heatmap. The participant-level **Attention View** (Fig. 4b) transposes the previous heatmap: columns represent images (sorted by safety score) and rows represent semantic classes (**T1**, **T3**). The **Similarity View** (Image Projection, Fig. 4c) shows a 2D t-SNE projection of the 50 images viewed by the participant, using 4096-dimensional Places365 embeddings [56], colored by safety score to facilitate comparisons across visually similar scenes (**T3**); the **Image View** (Fig. 4d) displays the streetscenes selected from that projection alongside their safety scores. Finally, the **Coverage View** (Saliency coverage comparison, Fig. 4e) relates *Saliency Coverage* or *Stationary Entropy* to safety scores or image order (**T2**, **T3**), letting analysts assess whether perceived safety correlates with broader versus focused exploratory visual strategies.

6.5. Implementation Details

URBANGAZEVis employs a decoupled client-server architecture: client-side coordinated views are implemented in JavaScript with D3.js [57], while the server uses Python and

Flask [58] for data indexing. To ensure responsive interactions, computationally expensive feature extraction, dimensionality reduction, and segmentation are performed offline.

Our preprocessing pipeline uses OneFormer [48] (pretrained on ADE20K via HuggingFace [59]), selected for its superior label correctness and boundary alignment over alternatives (e.g., DeepLabv3+, SegFormer; see Appendix I). For similarity-driven exploration in the By-Participant component, we extract 4,096-dimensional scene embeddings with a pretrained Places365Net [56]. Unlike raw semantic maps or general-purpose models such as CLIP, Places365Net better captures urban layouts and structural decay, preventing visually distinct but semantically similar scenes from appearing artificially identical. The embeddings are projected into 2D using t-SNE [60] with PCA initialization, preserving local neighborhoods and enabling stable comparisons within each participant's image subset ($N = 50$).

7. Usage Scenarios

Three usage scenarios illustrate how URBANGAZEVis supports the analysis of eye-gaze data in urban safety studies.

7.1. Associations between Safety Perceptions and the Built Environment

Fig. 5 presents the analytical workflow within the *By-Participant* component for Participant 29. Three streetscenes

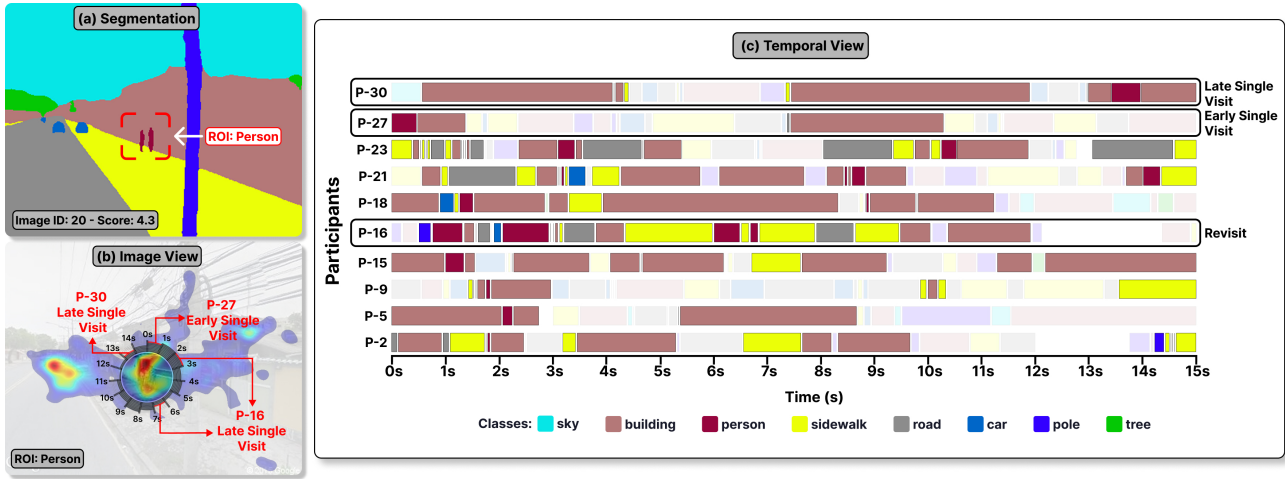


Figure 7: Usage scenario 3: Temporal divergence in exploring the “person” class (Image ID 20, safety score: 4.3). (a) **Semantic segmentation** of Image ID 20. (b) **Image View** heatmap reveals a high fixation density (visual hotspot) over the pedestrian region, prompting the placement of the ROI Glyph. (c) **Temporal View** details how individual participants interact with this specific area over the 15-second trial: Participant 27 makes an early single visit, Participant 30 makes a late single visit, and Participant 16 exhibits a recurrent attention (last single visit or revisit) pattern. These distinct spatiotemporal strategies are effectively disentangled, revealing behavioral nuances that the aggregated spatial heatmap alone would conceal.

are selected for detailed inspection, denoted by their unique dataset identifiers: Image ID 60 (mint green), Image ID 38 (coral), and Image ID 109 (soft blue), sorted left to right by mean perceived safety score, from unsafe to safe. The default view displays *total attention* (top heatmap), which emphasizes large background elements such as construction and floor. Switching to *attention intensity* (AttIn, defined in Sec. 6.2) instead highlights elements with low spatial coverage but strong attentional focus, particularly disorder cues such as graffiti and damaged walls. In Image 60 (top row), total attention emphasizes construction, whereas AttIn reveals stronger emphasis on smaller elements such as vehicles and city features like street lamps. In Image 38 (middle row), total attention highlights construction, floor, and graffiti, while AttIn concentrates on graffiti, linking infrastructural decay to reduced perceived safety. In Image 109 (bottom row), AttIn again emphasizes graffiti and damaged walls over the broader streetscape. By isolating these subtle yet influential elements, URBANGAZEVis demonstrates that combining AttIn with disorder cues helps identify key correlates of safety perception.

7.2. Analyzing Divergent Scoring Patterns

Assessing how consistently participants score visually similar images reveals the stability of their judgments; large rating differences for similar scenes may indicate sensitivity to specific visual cues. Using AttIn, an analyst selects a participant in the *By-Participant* component and uses the **Similarity View** to find nearby images in the projection space that received different scores. The **Attention View** and **Coverage View** summarize visual focus, while gaze points are overlaid in the **Image View**.

Fig. 6 presents three examples for Participants 10, 22, and 2. In Fig. 6.a, Participant 10 rates Image 145 highly (8) but Image 100 low (2): attention in Image 145 is balanced between *Construction* and *Vegetation*, whereas in Image 100 dense fixations concentrate on *Grffiti*, contributing to the lower score. In Fig. 6.b, Participant 22 rates Image 72 as 6 and Image 38 as

3. Although both scenes contain graffiti, gaze patterns differ: Image 72 attracts attention to *Construction*, *Grffiti*, and *Human/pedestrians*, yielding a moderate score, while gaze in Image 38 stays largely confined to *Construction* and *Grffiti*. This suggests that while physical disorder is associated with lower perceived safety, the presence of people can partially mitigate this negative association. Finally, Fig. 6.c shows Participant 2’s evaluations of Images 72 and 10. Despite prominent graffiti, Image 72 (score 7) draws substantial attention to pedestrians, while Image 10 (score 3), lacking human presence, directs attention toward vehicles and *Broken Road* cues. Across these cases, URBANGAZEVis reveals a consistent pattern: the visual presence of people often shows a stronger positive association with perceived safety than graffiti alone.

7.3. Spatiotemporal Dynamics and Participant Divergence

In this scenario, we investigate the spatiotemporal dynamics of attention to understand how different observers process the same streetscape during the 15-second trial (Fig. 7). Aggregate heatmaps reveal where attention is concentrated but not whether a region was fixated continuously, intermittently, early, or late. We analyze Image ID 20 with a safety score of 4.3 (Fig. 7.a), whose segmentation identifies all semantic objects. Although the *person* class (pedestrian area) occupies only a small portion of the image, the **Image View** heatmap (Fig. 7.b) highlights it as a visual hotspot.

To investigate this hotspot, the ROI Glyph links the selected region to the coordinated **Temporal View** (Fig. 7.c), which decomposes the 15-second scanpaths of individual participants. This view reveals distinct exploration strategies. Participant 27 (P-27) fixates the *person* class only during the initial seconds (early single visit), whereas Participant 30 (P-30) attends to it only near the end of the trial (late single visit). In contrast, Participant 16 (P-16) repeatedly returns to the *person* class after exploring other semantic elements (revisit).

Together, the ROI Glyph and **Temporal View** expose temporal behaviors—early visits, late visits, and revisits—that are

not visible in conventional spatial heatmaps. These differences suggest that human presence is attended not only as a semantic category but also at different stages of the visual assessment process, providing insights into how social cues may influence perceived safety.

8. System Evaluation and Value-Driven Assessment

A user study with 22 participants from diverse academic and professional backgrounds evaluated the analytical effectiveness and value-driven utility of URBANGAZEVIS. The study intentionally focused on individuals without formal expertise in eye-tracking analysis, to test whether the system makes complex visual behavior data accessible to a broader audience.

8.1. Study Setup and Participants

We recruited 22 participants with experience in general data analysis but no prior exposure to eye-tracking data. The evaluation design and questionnaire were explicitly inspired by the ICE-T framework [61], which assesses visualization systems along four dimensions: Insight, Confidence, Essence, and Time. Prior research indicates that relatively small samples suffice to capture exploratory feedback and assess visualization effectiveness [61], supporting this sample size. Although participants were familiar with system evaluations, none had previously used URBANGAZEVIS.

8.2. Tasks and Procedure

The study was conducted in person. Participants received a 10-minute introduction to the URBANGAZEVIS interface, covering the By-Image and By-Participant views and their primary interactions, then completed four visual analytics tasks (U1–U4) engineered to operationalize and cover the high-level analytical design tasks (T1–T3) defined in Sec. 3.2:

U1: Identify dominant attention elements required toggling attention metrics to identify initial and sustained focus, mapping to T1.1 and serving as a baseline for evaluating category-specific distributions (T3.1).

U2: Reconstruct visual trajectories used the scarf plot and gaze overlays to describe sequential paths, operationalizing T2.1 to validate the deconstruction of temporal scanpaths and transition episodes.

U3: Analyze specific urban features involved inspecting regions of urban decay via glyphs and switching data (e.g., to Disorder Cues), addressing T3.1 by linking localized fixations on physical disorder with safety judgments.

U4: Discover global viewing patterns had users explore multiple images in the By-Participant view; this scenario covered the participant-centric tasks by evaluating cross-image attention (T1.2), temporal strategy shifts like fatigue (T2.2), and judgment consistency across visually similar embeddings (T3.2).

For each task, participants were encouraged to think aloud [62], verbalizing their reasoning and observations before answering quantitative (QT) and qualitative (QL) questions. Because the tasks were designed to support open-ended,

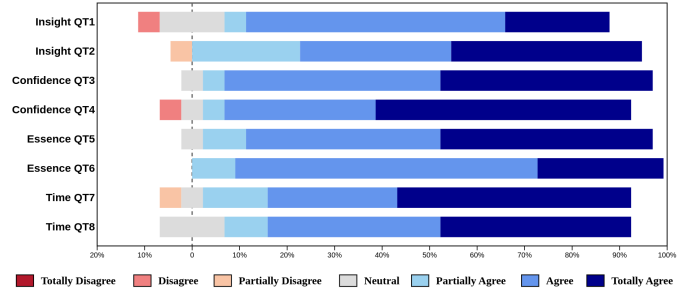


Figure 8: Distribution of participant responses for questions QT1–QT8, grouped by ICE-T-inspired dimensions (Insight, Confidence, Essence, and Time). Bar length indicates the percentage of participants selecting each Likert level on a 7-point scale from “Totally Disagree” to “Totally Agree”.

autonomous discovery and hypothesis generation, we intentionally omitted conventional task-performance metrics, such as task completion rates and error frequencies. Such metrics do not adequately capture the qualitative synthesis and insight validation that are central to value-driven evaluation frameworks. The quantitative questions operationalized the ICE-T-inspired dimensions through eight statements (two per dimension), each rated on a 7-point Likert scale. The qualitative questions elicited open-ended feedback on the usefulness of the visual components and suggestions for improvement. The full questionnaire is provided in Appendix D.

8.3. Results

Both quantitative results from our ICE-T-inspired questionnaire and qualitative feedback are reported below.

8.3.1. Quantitative Results

Averaged over the 22 participants, all four ICE-T dimensions scored above the neutral midpoint (4) on the 7-point Likert scale, with **Confidence** highest (6.3), followed by **Essence** (6.2), **Time** (6.1), and **Insight** (5.9); individual results are reported in Appendix E. We qualify these ratings by noting that our cohort consisted of general data analysts with no prior domain exposure to eye-tracking data; the high Confidence and Essence scores thus reflect the system’s ability to make complex, aggregated spatiotemporal gaze data immediately accessible and trustworthy for non-expert users, rather than an expert-level verification of visualization heuristics.

Insight (QT1, QT2) showed more varied responses, suggesting deeper interpretation may rely on extended user familiarity. **Confidence (QT3, QT4)** showed strong agreement, as coordinated views and consistent color encodings enhanced trust in spatial and temporal interpretations. **Essence (QT5, QT6)** results highlighted the system’s ability to summarize complex visual behavior concisely. The **Time (QT7, QT8)** dimension was rated positively, though temporal visualizations required slightly more cognitive effort. Fig. 8 confirms most responses fall on the positive end of the Likert scale, with minimal neutral or negative feedback.

8.3.2. Qualitative feedback

Participant feedback (QL1, QL2) consistently highlighted the *Participant Attention Summary* and the *Temporal Participant Attention* scarf plot as the most intuitive components for

comparing focus, interpreting numerical metrics, and tracking chronological scanpaths. Interactive features like cross-view filtering were deemed essential for reducing visual clutter; however, users noted that while the *By-Image* mode was visually accessible, the *By-Participant* mode required higher cognitive effort due to its numerical density. Improvement requests centered on expanding global search (e.g., quickly locating specific images or participants) and sorting results by attention-based metrics. Users also suggested minor refinements such as resizable panels, and proposed automated insights to highlight dominant attention correlations across scenes for higher-level analytical reasoning.

9. Discussion

The quantitative validation established in Sec. 5 confirms that the URBANGAZE dataset possesses a stable signal-to-noise ratio. In this section, we interpret the domain-specific implications of these patterns, discussing how visual attention aligns with foundational urban sociology and how fine-grained granularity reveals context-specific nuances.

9.1. Alignment with Broken Windows Theory

Consistent with the Broken Windows Theory [1], our hybrid regression metrics, detailed in Table 2, demonstrate that physical disorder cues from UrbanPD4k are the strongest negative predictors of perceived urban safety. Specifically, sustained visual attention to broken or damaged brick walls ($\beta = -0.442$, $p < .001$) demonstrated the most pronounced negative association with safety scores, closely followed by damaged pavement ($\beta = -0.229$) and graffiti ($\beta = -0.225$). These results confirm that the interactive visual patterns highlighted by analysts within the URBANGAZEVis **Attention View** correspond directly to statistically meaningful and predictable human perceptual mechanisms of decay.

9.2. General Urban Infrastructure vs. Semantic Granularity

Similar results for the other three data cases (detailed in the Supplementary Material Appendix H) underscore the importance of semantic granularity. Baseline models (such as raw ADE20K) merge deterioration into broad, generic categories like *wall*, whereas our ADE20K+UrbanPD4k separates general infrastructure from decay signals: the baseline absorbs structural decay into background elements, while our fine-grained representation distinguishes blind walls ($\beta = -0.326$) from physically damaged walls ($\beta = -0.442$), the latter showing a substantially stronger negative correlation.

Conversely, macro-level Grouped-ADE proved too coarse and statistically unstable for safety analysis. Variance Inflation Factor (VIF) analysis revealed severe multicollinearity when fine-grained elements were merged into broad clusters—for example, aggregating infrastructure into *Construction* or ground elements into *Floor* produced VIF values of 43.29 and 29.78, respectively, inflating standard errors enough to render several predictors insignificant. These findings support the necessity of fine-grained, disorder-aware semantic cues in urban visual analytics.

Table 2: Fixed effects of visual attention on perceived safety (ADE20K + Disorder, $N = 1500$).

Predictor	Coef. (β)	S.E.	z-val	VIF	p-val	95% CI
(Intercept)	5.091	0.159	31.940	—	<.001	[4.780, 5.400]
<i>Positive (Safe)</i>						
Car / Taxi	0.197	0.047	4.230	1.160	<.001	[0.110, 0.290]
Door	0.155	0.051	3.060	1.360	0.002	[0.060, 0.250]
<i>Negative (Unsafe)</i>						
Broken/Dam Wall	-0.442	0.090	-4.930	4.270	<.001	[-0.620, -0.270]
Wall	-0.326	0.121	-2.690	7.780	0.007	[-0.560, -0.090]
Road	-0.277	0.117	-2.360	7.340	0.018	[-0.510, -0.050]
Damaged Pavement	-0.229	0.056	-4.100	1.650	<.001	[-0.340, -0.120]
Graffiti	-0.225	0.075	-2.990	2.990	0.003	[-0.370, -0.080]
Sky	-0.200	0.071	-2.830	2.170	0.005	[-0.340, -0.060]
Bridge	-0.188	0.051	-3.690	1.370	<.001	[-0.290, -0.090]
Mountain	-0.184	0.046	-4.020	1.120	<.001	[-0.270, -0.090]
Signboard	-0.152	0.055	-2.770	1.590	0.006	[-0.260, -0.040]
Garbage Bag	-0.142	0.049	-2.900	1.280	0.004	[-0.240, -0.050]
Pole	-0.142	0.048	-2.990	1.210	0.003	[-0.240, -0.050]
Field	-0.138	0.045	-3.070	1.070	0.002	[-0.230, -0.050]
<i>Random Effects</i>						
Participant	0.707	0.841		0.203		
Residual	2.774	1.666				

9.3. The Mitigating Role of Social Cues

Beyond physical decay, URBANGAZEVis captures localized contextual patterns that global statistical models often obscure. As shown in our second usage scenario in Sec. 7.2, the visual presence of pedestrians acts as a powerful social mediator: even amid prominent disorder cues like graffiti, scenes with active human presence (Image ID 72) maintained substantially higher safety ratings than visually identical, isolated spaces (Image ID 10). This suggests that visual attention to human presence is consistent with natural surveillance interpretations, modifying how physical decay is interpreted. Across these qualitative cases, the presence of people often exerts a stronger positive effect on perceived safety than visual clutter or graffiti alone.

10. Limitations and Future Work

Although this study provides interpretable insights into urban safety perception, several limitations affect the scope and generalizability of the findings. First, regarding **participant diversity and experimental context**, our sample of 30 participants from a single institution is sufficient for methodological evaluation but does not represent the socioeconomic diversity of Rio de Janeiro residents. Moreover, the analysis relies on static SVI viewed for 15 seconds in a controlled head-mounted setting, whereas real-world perception is shaped by movement, sound, and unconstrained exploration. **Future work** should therefore include more diverse populations and multimodal, in-the-wild data collection.

Second, URBANGAZEVis has limitations in **visual scalability, representation, and the correlational nature of attention**. Designed for a moderate dataset (150 images, 30 participants), visualizations such as heatmaps and scarf plots may require filtering or aggregation for larger cohorts. Similarly, Places365 embeddings capture coarse scene semantics but may

miss fine-grained disorder cues, while t-SNE projections remain sensitive to initialization with only 50 images per participant. **Future work** should investigate scalable visualization and dimensionality-reduction strategies, as well as domain-specific visual representations that better capture fine-grained urban disorder cues. Finally, although attention to physical disorder is associated with lower safety ratings, these results are correlational rather than causal. URBANGAZEVis is therefore intended as an exploratory, hypothesis-generating tool, motivating future neurocognitive studies with manipulated stimuli.

Third, regarding **spatial mapping**, URBANGAZEVis does not include a macro-geographic map layer. Because the dataset prioritizes structural diversity over geographic coverage, mapping the sampled locations could imply unsupported spatial trends. **Future work** should integrate micro-level gaze analysis with city-scale geospatial representations to better connect perception and urban planning.

11. Conclusions

We presented URBANGAZEVis, a visual analytics system for exploring eye-movement behavior in urban safety perception. Built on a head-mounted eye-tracking study with 30 participants and 150 SVIs from Rio de Janeiro, the system integrates gaze dynamics, semantic segments, disorder annotations, and safety ratings into a unified analytical framework.

Our results show that perceived safety cannot be fully understood through image-level judgments alone. By linking gaze behavior, temporal attention patterns, and inspected urban elements, URBANGAZEVis reveals how disorder cues, infrastructure, vegetation, and human presence contribute differently to safety perception, complementing statistical analyses by exposing local, context-specific patterns often hidden by global aggregation. More broadly, this work demonstrates the value of integrating eye tracking, semantic scene understanding, and visual analytics to support more interpretable studies of urban perception, providing a foundation for future research in urban analytics, human-centered planning, and perception-aware urban design.

ACKNOWLEDGMENTS

This work was supported by the National Council for Scientific and Technological Development (CNPq, grant #313454/2026-4 and #132349/2025-6), Brazilian Federal Agency for Support and Evaluation of Graduate Education (CAPES, grant #88887.684234/2022-00), Carlos Chagas Filho Foundation for Research Support of Rio de Janeiro State (FAPERJ, grant #E-26/210.585/2025), São Paulo Research Foundation (FAPESP, grant #2021/07012-0 and #2024/05760-8), and the Fundação Getulio Vargas (FGV).

References

- [1] Wilson, JQ, Kelling, GL. Broken windows. *Atlantic monthly* 1982;249(3):29–38. doi:10.4135/9781412959193.n281.
- [2] Sampson, RJ, Morenoff, JD, Gannon-Rowley, T. Assessing “neighborhood effects”: Social processes and new directions in research. *Annual review of sociology* 2002;28(1):443–478. doi:10.1146/annurev.soc.28.110601.141114.
- [3] Keizer, K, Lindenberg, S, Steg, L. The spreading of disorder. *Science* (New York, NY) 2008;322:1681–5. doi:10.1126/science.1161405.
- [4] Moreno-Vera, F, De-la Puente, A, Poco, J. UrbanPhysicalDisorder-4K: Understanding urban perception via counterfactuals and street view signs of physical disorder. In: 2025 IEEE International Conference on Big Data (BigData). 2025, p. 5194–5200. doi:10.1109/BigData66926.2025.11401786.
- [5] Naik, N, Philipoom, J, Raskar, R, Hidalgo, C. Streetscore-predicting the perceived safety of one million streetscapes. In: Proceedings of the IEEE conference on computer vision and pattern recognition workshops. 2014, p. 779–785. doi:10.1109/CVPRW.2014.121.
- [6] Seresinhe, CI, Preis, T, Moat, HS. Using deep learning to quantify the beauty of outdoor places. *Royal Society Open Science* 2017;4. doi:10.1098/rsos.170170.
- [7] Dubey, A, Naik, N, Parikh, D, Raskar, R, Hidalgo, CA. Deep learning the city: Quantifying urban perception at a global scale. In: European Conference on Computer Vision (ECCV). Springer; 2016, p. 196–212.
- [8] Zhang, F, Salazar-Miranda, A, Duarte, F, Vale, L, Hack, G, Chen, M, et al. Urban visual intelligence: Studying cities with artificial intelligence and street-level imagery. *Annals of the American Association of Geographers* 2024;114(5):876–897. doi:10.1080/24694452.2024.2313515.
- [9] Moreno-Vera, F, Poco, J. Assessing urban environments with vision-language models: A comparative analysis of ai-generated ratings and human volunteer evaluations. In: 2025 International Joint Conference on Neural Networks (IJCNN). IEEE; 2025, p. 1–8.
- [10] Lavi, B, Tokuda, EK, Moreno-Vera, F, Nonato, LG, Silva, CT, Poco, J. 17k-graffiti: Spatial and crime data assessments in são paulo city. In: Proceedings. 2022.
- [11] Li, Y, Yabuki, N, Fukuda, T. Measuring visual walkability perception using panoramic street view images, virtual reality, and deep learning. *Sustainable Cities and Society* 2022;86:104140. doi:10.1016/j.scs.2022.104140.
- [12] Yang, N, Deng, Z, Hu, F, Chao, Y, Wan, L, Guan, Q, et al. Urban perception by using eye movement data on street view images. *Transactions in GIS* 2024;doi:10.1111/tgis.13172.
- [13] Santani, D, Ruiz-Correa, S, Gatica-Perez, D. Looking south: Learning urban perception in developing cities. *ACM Transactions on Social Computing* 2018;1(3):1–23. doi:10.1145/3224182.
- [14] Moreno-Vera, F, Lavi, B, Poco, J. Quantifying urban safety perception on street view images. In: International Conference on Web Intelligence and Intelligent Agent Technology (WI-IAT). 2021;doi:10.1145/3486622.3493975.
- [15] Moreno-Vera, F, Lavi, B, Poco, J. Urban perception: Can we understand why a street is safe? In: Mexican International Conference on Artificial Intelligence. Springer; 2021, p. 277–288. doi:10.1007/978-3-030-89817-5_21.
- [16] Salesses, P, Schechtner, K, Hidalgo, CA. The collaborative image of the city: mapping the inequality of urban perception. *PLoS one* 2013;8(7):e68400. doi:10.1371/journal.pone.0068400.
- [17] Doersch, C, Singh, S, Gupta, AK, Sivic, J, Efros, AA. What makes paris look like paris? *ACM Transactions on Graphics (TOG)* 2012;31:1–9. doi:10.1145/2830541.
- [18] Moreno-Vera, F, Brandoli, B, Poco, J. What makes a place feel safe? analyzing street view images to identify relevant visual elements. In: International Conference on Web Intelligence and Intelligent Agent Technology (WI-IAT). 2024;doi:10.1109/WI-IAT62293.2024.00062.
- [19] Ordonez, V, Berg, TL. Learning high-level judgments of urban perception. *European Conference on Computer Vision (ECCV)* 2014;doi:10.1007/978-3-319-10599-4_32.
- [20] Porzi, L, Rota Bulò, S, Lepri, B, Ricci, E. Predicting and understanding urban perception with convolutional neural networks. In: ACM international conference on Multimedia. 2015;doi:10.1145/2733373.2806273.
- [21] Zhang, F, Zhou, B, Liu, L, Liu, Y, Fung, HH, Lin, H, et al. Measuring human perceptions of a large-scale urban region using machine learning. *Landscape and Urban Planning* 2018;180:148–160. doi:10.1016/j.landurbplan.2018.08.020.
- [22] Min, W, Mei, S, Liu, L, Wang, Y, Jiang, S. Multi-task deep relative attribute learning for visual urban perception. *IEEE Transactions on Image*

- Processing 2020;29:657–669. doi:10.1109/TIP.2019.2932502.
- [23] Cohen, LE, Felson, M. Social change and crime rate trends: A routine activity approach. *American sociological review* 1979;588–608.
- [24] Brantingham, PL, Brantingham, PJ. Nodes, paths and edges: Considerations on the complexity of crime and the physical environment. vol. 13. 1993, p. 3–28. URL: <https://www.sciencedirect.com/science/article/pii/S0272494405802129>. doi:https://doi.org/10.1016/S0272-4944(05)80212-9.
- [25] Clarke, RV. Situational crime prevention: Its theoretical basis and practical scope. *Crime and justice* 1983;4:225–256.
- [26] Jacobs, J. *The Death and Life of Great American Cities*. Vintage Books ed; Vintage Books; 1961. ISBN 9780679741954. URL: <https://books.google.com.br/books?id=deqqU0jBTnUC>.
- [27] Newman, O. *Defensible Space: Crime Prevention Through Urban Design*. New York: Macmillan; 1972.
- [28] Yang, N, Deng, Z, Hu, F, Guan, Q, Chao, Y, Wan, L. Urban perception assessment from street view images based on a multifeature integration encompassing human visual attention. *Annals of the American Association of Geographers* 2024;doi:10.1080/24694452.2024.2363783.
- [29] Moreno-Arjonilla, J, et al., . Eye-tracking integration in virtual reality environments: A comprehensive methodological survey. *Computers & Graphics* 2024;118:45–58. doi:10.1016/j.cag.2024.01.002.
- [30] Kang, Y, Chen, J, Liu, L, Sharma, K, Mazzarello, M, Mora, S, et al. Decoding human safety perception with eye-tracking systems, street view images, and explainable ai. *Computers, Environment and Urban Systems* 2026;123:102356. doi:10.1016/j.compenvurbsys.2025.102356.
- [31] Oki, T, Kizawa, S. Evaluating visual impressions based on gaze analysis and deep learning: a case study of attractiveness evaluation of streets in densely built-up wooden residential area. *The International Archives of the Photogrammetry, Remote Sensing and Spatial Information Sciences* 2021;doi:10.5194/isprs-archives-XLIII-B3-2021-887-2021.
- [32] Boorboor, S, Kim, Y, Hu, P, Moses, JM, Colle, BA, Kaufman, AE. Submerse: Visualizing storm surge flooding simulations in immersive display ecologies. *IEEE Transactions on Visualization and Computer Graphics* 2023;30(9):6365–6377.
- [33] Perez-Martin, E, Medina, SLC, Herrero-Tejedor, TR, Ezquerra-Canalejo, A, Lopez-Fernandez, D. Using virtual reality in the learning of geomatic engineering education. *IEEE Computer Graphics and Applications* 2024;44(6):77–88.
- [34] Cobeli, S, Omar, KS, Valena, R, Ferreira, N, Miranda, F. A neural field-based approach for view computation & data exploration in 3d urban environments. *IEEE Transactions on Visualization and Computer Graphics* 2025;.
- [35] Ouyang, Y, Wu, Y, Wang, X, Xie, L, Cheng, W, Gan, J, et al. Oceanvive: An immersive visualization system for communicating complex oceanic phenomena. In: *2025 IEEE Visualization and Visual Analytics (VIS)*. IEEE; 2025, p. 326–330.
- [36] Banno, T, Takenawa, M, Wöhler, L, Ikehata, S, Aizawa, K. 360citygml: Realistic and interactive urban visualization system integrating citygml model and 360 videos. *IEEE Transactions on Visualization and Computer Graphics* 2025;.
- [37] Gonzalez, E, Sajja, R, Sermet, Y, Demir, I. Floodsim sandbox: An immersive interactive simulation framework for urban flood risk management. *EarthArXiv* 2025;.
- [38] Blascheck, T, Kurzahls, K, Raschke, M, Burch, M, Weiskopf, D, Ertl, T. Visualization of eye tracking data: A taxonomy and survey. In: *Computer graphics forum*; vol. 36. Wiley Online Library; 2017, p. 260–284. doi:10.1111/cgf.13079.
- [39] Andrienko, G, Andrienko, N, Burch, M, Weiskopf, D. Visual analytics methodology for eye movement studies. *IEEE Transactions on Visualization and Computer Graphics* 2012;18(12):2889–2898. doi:10.1109/TVCG.2012.276.
- [40] Kurzahls, K, Weiskopf, D. Space-time visual analytics of eye-tracking data for dynamic stimuli. *IEEE Transactions on Visualization and Computer Graphics* 2013;19(12):2129–2138. doi:10.1109/TVCG.2013.194.
- [41] Kurzahls, K, Heimerl, F, Weiskopf, D. Iseecube: Visual analysis of gaze data for video. In: *Proceedings of the symposium on eye tracking research and applications*. 2014, p. 43–50. doi:10.1145/2578153.2578158.
- [42] Burch, M, Kull, A, Weiskopf, D. Aoi rivers for visualizing dynamic eye gaze frequencies. In: *Computer Graphics Forum*; vol. 32. Wiley Online Library; 2013, p. 281–290. doi:10.1111/cgf.12115.
- [43] Kurzahls, K, Hlawatsch, M, Heimerl, F, Burch, M, Ertl, T, Weiskopf, D. Gaze stripes: Image-based visualization of eye tracking data. *IEEE Transactions on Visualization and Computer Graphics* 2015;22(1):1005–1014. doi:10.1109/TVCG.2015.2468091.
- [44] Tula, AD, Kurauchi, A, Coutinho, F, Morimoto, C. Heatmap explorer: an interactive gaze data visualization tool for the evaluation of computer interfaces. In: *Proceedings of the 15th Brazilian Symposium on Human Factors in Computing Systems*. Association for Computing Machinery; 2016;doi:10.1145/3033701.3033725.
- [45] Goodwin, S, et al., . Veta: Visual analytics for spatial eye-gaze diagnostic data. *IEEE Transactions on Visualization and Computer Graphics* 2022;28(1):230–240. doi:10.1109/TVCG.2022.3110001.
- [46] Ren, X, et al., . Svata: Visual analytics for multi-user fixation and spatial attention optimization. *IEEE Transactions on Visualization and Computer Graphics* 2026;32(1):112–122. doi:10.1109/TVCG.2026.3332511.
- [47] Shneiderman, B. The eyes have it: A task by data type taxonomy for information visualizations. In: *The craft of information visualization*. Elsevier; 2003;doi:10.1016/B978-155860915-0/50046-9.
- [48] Jain, J, Li, J, Chiu, MT, Hassani, A, Orlov, N, Shi, H. Oneformer: One transformer to rule universal image segmentation. In: *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*. 2023, p. 2989–2998. doi:10.1109/CVPR52729.2023.00292.
- [49] Microsoft, . Eye tracking on HoloLens 2. 2022. URL: <https://learn.microsoft.com/en-us/windows/mixed-reality/design/eye-tracking>; mixed Reality documentation.
- [50] Putka, DJ, Le, H, McCloy, RA, Diaz, T. Ill-structured measurement designs in organizational research: Implications for estimating interrater reliability. *Journal of Applied psychology* 2008;93(5):959.
- [51] Salvucci, DD, Goldberg, JH. Identifying fixations and saccades in eye-tracking protocols. In: *Proc. of the 2000 Symposium on Eye Tracking Research & Applications*. ETRA '00; Association for Computing Machinery. ISBN 158113246X; 2000, p. 71–78. doi:10.1145/355017.355028.
- [52] Karthik, G, Amudha, J, Jyotsna, C. A custom implementation of the velocity threshold algorithm for fixation identification. In: *2019 International Conference on Smart Systems and Inventive Technology (ICSSIT)*. IEEE; 2019, p. 488–492. doi:10.1109/ICSSIT46314.2019.8987791.
- [53] Shannon, CE. A mathematical theory of communication. *SIG-MOBILE Mob Comput Commun Rev* 2001;5(1):3–55. doi:10.1145/584091.584093.
- [54] Le Meur, O, Baccino, T. Methods for comparing scanpaths and saliency maps: strengths and weaknesses. *Behavior research methods* 2013;45(1):251–266. doi:10.3758/s13428-012-0226-9.
- [55] Borji, A, Itti, L. State-of-the-art in visual attention modeling. *IEEE Transactions on Pattern Analysis and Machine Intelligence* 2012;35(1):185–207. doi:10.1109/TPAMI.2012.89.
- [56] Zhou, B, Lapedriza, A, Khosla, A, Oliva, A, Torralba, A. Places: A 10 million image database for scene recognition. *IEEE Transactions on Pattern Analysis and Machine Intelligence* 2017;40(6):1452–1464. doi:10.1109/TPAMI.2017.2723009.
- [57] Bostock, M, Ogievetsky, V, Heer, J. D³ data-driven documents. *IEEE Transactions on Visualization and Computer Graphics* 2011;17(12):2301–2309. doi:10.1109/TVCG.2011.185.
- [58] Grinberg, M. Flask web development. "O'Reilly Media, Inc."; 2018. URL: <https://books.google.com/books?id=cV1PDwAAQBAJ>.
- [59] Wolf, T, Debut, L, Sanh, V, Chaumond, J, Delangue, C, Moi, A, et al. Transformers: State-of-the-art natural language processing. In: *Proc. of the 2020 Conference on Empirical Methods in Natural Language Processing: System Demonstrations*. Online: Association for Computational Linguistics; 2020;doi:10.18653/v1/2020.emnlp-demos.6.
- [60] Pedregosa, F, Varoquaux, G, Gramfort, A, Michel, V, Thirion, B, Grisel, O, et al. Scikit-learn: Machine learning in python. *the Journal of machine Learning research* 2011;12:2825–2830. URL: <https://arxiv.org/abs/1201.0490>.
- [61] Wall, E, Agnihotri, M, Matzen, L, Divis, K, Haass, M, Endert, A, et al. A heuristic approach to value-driven evaluation of visualizations. *IEEE Transactions on Visualization and Computer Graphics* 2018;25(1):491–500. doi:10.1109/TVCG.2018.2865146.
- [62] Lewis, C. Using the "thinking aloud" method in cognitive interface design. *Research Report RC9265*, IBM TJ Watson Research Center 1982;URL: <https://books.google.com.br/books?id=F5AKHQAAQAAJ>.
- [63] Khowaja, S, Ghufuran, S, Ahsan, M. Estimation of population means in multivariate stratified random sampling. *Communications in*

- Statistics—Simulation and Computation® 2011;40(5). doi:10.1080/03610918.2010.551014.
- [64] Tatler, BW, Hayhoe, MM, Land, MF, Ballard, DH. Eye guidance in natural vision: Reinterpreting salience. *Journal of vision* 2011;11(5):5–5. doi:10.1167/11.5.5.
- [65] Clay, V, König, P, König, S. Eye tracking in virtual reality. *Journal of eye movement research* 2019;12(1):10–16910. doi:10.16910/jemr.12.1.3.
- [66] Kapp, S, Barz, M, Mukhametov, S, Sonntag, D, Kuhn, J, Arett: Augmented reality eye tracking toolkit for head mounted displays. *Sensors* 2021;21(6):2234. doi:10.3390/s21062234.

Supplementary Material

Supplementary material that may be helpful in the review process.

Appendix A. Multivariate Stratified Sampling Optimization

To preserve the diversity of visual element proportions, we used a multivariate stratified sampling scheme [63]. For each visual element $v \in V$, we compute its pixel-ratio distribution over all images and divide it into quartile groups $g \in G_v$. Let $S_{v,g}$ be the set of images whose proportion of element v falls into quartile g , and let s be the desired sample size. We select a binary vector $\mathbf{x} = (x_1, \dots, x_n)$, where $x_i = 1$ if image i is kept, by solving:

$$\text{objective: } \max \sum_{i=1}^n x_i \quad (\text{A.1})$$

$$\text{subject to: } \left\lfloor \frac{s}{\sum_{v \in V} |G_v|} \right\rfloor \leq \sum_{i \in S_{v,g}} x_i \quad \forall v \in V, \forall g \in G_v, \quad (\text{A.2})$$

which enforces a minimum number of selected images in each quartile across all visual elements, yielding the final 150-image subset.

Appendix B. Randomized Image Assignment Matrix

Given n_p participants $\mathcal{P}$ and n_i selected images $\mathcal{I}$, we require that every image be rated by exactly $K_{\mathcal{P}}$ distinct participants and every participant view exactly $K_{\mathcal{I}}$ images, with $n_i K_{\mathcal{P}} = n_p K_{\mathcal{I}}$. In our study, we targeted $K_{\mathcal{P}} = 10$ ratings per image. We implement a randomized assignment procedure that iteratively fills an $n_p \times n_i$ binary matrix V such that:

$$\sum_{p=1}^{n_p} V_{p,i} = K_{\mathcal{P}} \quad \forall i \in \mathcal{I}, \quad (\text{B.1})$$

$$\sum_{i=1}^{n_i} V_{p,i} = K_{\mathcal{I}} \quad \forall p \in \mathcal{P}, \quad (\text{B.2})$$

where $V_{p,i} = 1$ denotes that participant p views image i .

Appendix C. Extended Experimental Protocol and Hardware Setup

This study was conducted with approval from the Institutional Ethics Committee of our university. The experiment followed a standardized protocol with three main phases, lasting at most 25 minutes to avoid fatigue [64, 65]. Fig. C.1 shows the complete experimental protocol for eye-gaze data collection.

Calibration: The HoloLens 2 device was cleaned between sessions. We performed a nine-point gaze calibration routine,

monitoring accuracy in real time. Following established benchmarks, our setup maintained a quantitative spatial accuracy of $0.77^\circ \pm 0.35^\circ$ at the 2.0m focal distance used for stimuli projection.

Task familiarization: Participants viewed two reference images—one clearly perceived as *unsafe*, one as *safe*—for 15 s each to confirm stable tracking.

Safety rating trials: To prevent visual carry-over effects and ensure stable semantic scenes, we implemented a 5-second interval buffer between stimuli. Any trial failing to meet the 15-second viewing duration due to system interruption was automatically discarded.

Appendix D. ICE-T Questionnaire and Open-Ended Questions

The quantitative questions were directly *inspired* by the ICE-T framework [61], covering four dimensions: Insight, Confidence, Essence, and Time. Each dimension was assessed through two statements using a 7-point Likert scale (from “Totally Disagree” to “Totally Agree”).

Insight.

QT1: “*In the By-Image view, the ability to switch between attention metrics and segmentation classes helped me identify relevant elements and generate new interpretations of visual attention.*”.

QT2: “*The combination of visual representations in the By-Image view with the metrics in the By-Participant view allowed me to discover patterns and better understand participants’ exploration behavior.*”.

Confidence.

QT3: “*In the By-Image view, the consistency of colors associated with semantic classes across visualizations, together with tooltips for inspecting exact values, increased my confidence in interpreting the data accurately.*”.

QT4: “*The synchronized coordination between visualizations and the overlay of gaze/fixation points on the image increased my confidence in the correctness of spatial and temporal alignment.*”.

Essence.

QT5: “*Both image-centered analysis and participant-centered analysis facilitated a quick and overall understanding of visual behavior.*”.

QT6: “*The system effectively summarizes complex eye-tracking data by integrating multiple coordinated views, while maintaining clarity.*”.

Time.

QT7: “*The use of temporal visualization tools allowed me to quickly relate when participants observed specific elements and identify dominant attention patterns over time.*”.

QT8: “*Filtering options and smooth navigation significantly reduced the effort and time required to derive insights from eye-tracking data.*”.

The qualitative questions collected open-ended feedback about the usefulness of the system and potential improvements:

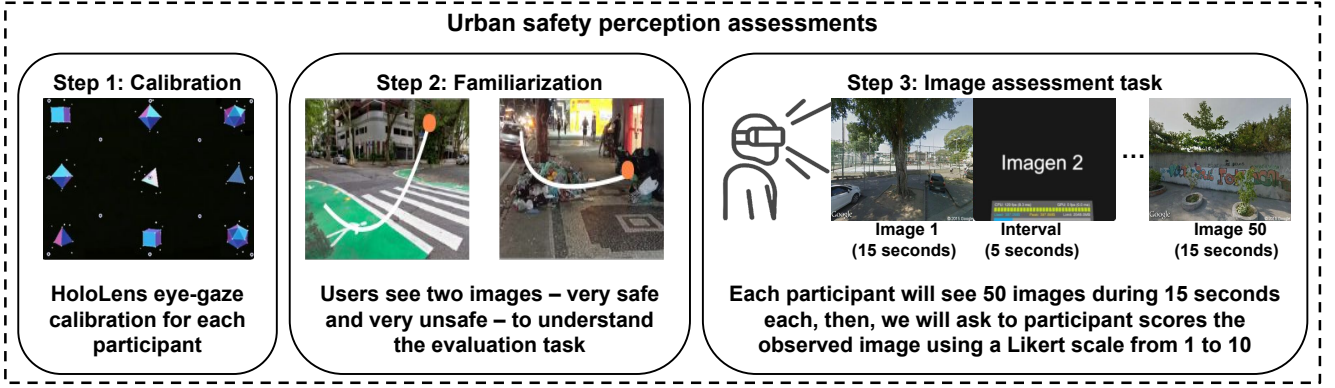


Figure C.1: Urban safety perception protocol. Step 1: gaze calibration with the HoloLens 2 headset. Step 2: task familiarization using one clearly safe and one clearly unsafe example image. Step 3: image assessment, in which each participant views 50 assigned images for 15 s, followed by a 5 s rating phase in which perceived safety is reported on a 1–10 Likert scale.

Table E.1: ICE-T evaluation scores per participant across the four dimensions.

Participant	<i>Insight</i>	<i>Confidence</i>	<i>Essence</i>	<i>Time</i>
P1	6.5	7.0	7.0	7.0
P2	6.0	6.0	6.0	6.0
P3	5.5	7.0	6.0	5.5
P4	3.5	4.0	6.0	5.0
P5	5.0	6.5	5.0	6.5
P6	6.5	7.0	6.0	6.5
P7	5.0	6.5	6.5	7.0
P8	5.5	5.5	5.5	6.5
P9	6.0	6.0	6.0	6.0
P10	5.5	7.0	6.5	6.5
P11	6.0	6.0	5.5	5.0
P12	5.5	6.5	6.0	4.5
P13	6.5	6.5	6.0	7.0
P14	7.0	7.0	6.5	7.0
P15	6.0	5.5	6.0	6.5
P16	5.5	5.5	6.0	5.5
P17	7.0	6.5	6.0	6.0
P18	7.0	7.0	7.0	5.0
P19	7.0	6.0	7.0	7.0
P20	6.5	7.0	7.0	7.0
P21	4.5	7.0	7.0	7.0
P22	6.0	5.0	6.5	4.0
Avg.	5.9	6.3	6.2	6.1

“Considering both the *By-Image* and *By-Participant* views, which visual component, chart, or interaction did you find most useful for your analysis, and why?” (QL1); “What would you change, add, or remove from the interface to make it more efficient? Please feel free to include any additional comments or suggestions about your experience.” (QL2).

Appendix E. ICE-T Results Per Participants

Table E.1 presents the ICE-T scores per participant, indicating consistently positive evaluations across all four dimensions.

Appendix F. Spatial Uncertainty and Kernel Calculation

To ensure the statistical validity of the *Attention Intensity* metric, it is necessary to calibrate the gaze data to account for the inherent hardware noise of the HoloLens 2. This appendix details the step-by-step derivation of the spatial kernel (σ) and the area exclusion threshold (τ) used to prevent false-positive gaze allocations on minute urban elements (e.g., distant overhead cables or small debris).

Step 1: Angular Precision of the Device

In eye-tracking methodology, spatial *precision* is the standard parameter used to define the dispersion of a Gaussian kernel, as it models the physiological and hardware jitter of a fixation. According to benchmark evaluations [66], the HoloLens 2 exhibits a spatial precision of 0.24° of visual angle at a focal distance of 2.0 meters, which is the exact distance used in our experimental setup.

Step 2: Conversion to Pixels Per Degree (PPD)

To translate this angular precision into the 2D pixel space of our stimuli, we calculated the device’s angular resolution. Based on the manufacturer’s specifications¹ and hardware evaluations, the HoloLens 2 displays a resolution of $2,048 \times 1,080$ pixels per eye with a diagonal field of view (FOV) of 52° and a diagonal of 2,315 pixels. Dividing this by the diagonal FOV yields an angular resolution of approximately 44.5 Pixels Per Degree ($2,315 \text{ px} / 52^\circ \approx 44.5 \text{ PPD}$).

Step 3: Deriving the Gaussian Kernel (σ)

Multiplying the device’s angular precision by the calculated PPD directly provides the operational standard deviation (σ) for our heatmap generation. Thus, we applied a Gaussian kernel with $\sigma \approx 10.68$ pixels ($0.24^\circ \times 44.5 \text{ PPD}$). This ensures that the spatial spread of each fixation accurately reflects the empirical uncertainty of the device.

Step 4: Defining the Area Exclusion Threshold (τ)

Finally, to establish a robust noise threshold (τ) for semantic segmentation, we considered an effective noise radius of 2σ ($2 \times 10.68 = 21.36$ pixels). In a 2D Gaussian distribution, a 2σ radius encompasses over 95.4% of the spatial uncertainty. Geometrically, this margin of error forms a circular area of roughly

¹<https://learn.microsoft.com/en-us/hololens/hololens2-hardware>

1,433 squared pixels ($A = \pi \times 21.36^2$).

When evaluated against our 800×600 stimulus resolution (480,000 total pixels), this hardware noise floor occupies exactly 0.3% of the visual field ($1,433/480,000 \approx 0.003$). To ensure an even more robust and conservative analysis, we padded this noise floor with a safety margin, establishing a strict area threshold of $\tau = 0.005$ (0.5%). Any semantic mask or disorder cue occupying less than this threshold was considered statistically indistinguishable from hardware noise and was excluded from the analysis.

Appendix G. Inter-Rater Reliability (IRR) and Rater Calibration

To validate the consistency of urban safety perceptions and ensure that our findings reflect environmental features rather than individual biases, we performed a reliability analysis tailored for sparse rating matrices. Given our *Ill-Structured Measurement Design* (ISMD)—where 30 participants evaluated overlapping but incomplete subsets of 150 images—traditional Intraclass Correlation Coefficients (ICC) based on balanced ANOVA are mathematically biased.

Following the framework proposed by Putka et al. (2008), we utilized a Linear Mixed-Effects Model (LMM) with crossed random effects to isolate the variance components of our dataset ($N = 1,500$ observations). This approach allows for explicit rater calibration by partitioning the total variance into three distinct sources: the urban scene (signal), the participant (subjective bias), and residual noise. The variance components estimated via Restricted Maximum Likelihood (REML) are detailed in Table G.2.

Table G.2: Variance Components for Safety Perception Scores.

Source of Variation	Variance (σ^2)	Interpretation
Urban Scene (Image)	0.876	True variance in streetscape quality (Signal)
Participant (Rater)	0.394	Subjective severity/leniency bias (Calibration)
Residual (Noise)	1.766	Unexplained variance and interaction

To calculate the reliability of the aggregated safety scores, we applied the $G(q, k)$ estimator, which accounts for the partial overlap between raters in ISMD designs:

$$G(q, k) = \frac{\sigma_{img}^2}{\sigma_{img}^2 + \frac{q \cdot \sigma_{part}^2 + \sigma_{res}^2}{k}} \quad (G.1)$$

Where $k = 10$ represents the number of raters per image, and $q \approx 0.67$ is the coefficient of non-overlap for our design. This coefficient is mathematically derived from the ratio of raters per image to the total participant pool ($q = 1 - k/N_{total}$), reflecting that each image was evaluated by exactly one-third of our 30-participant panel ($1 - 10/30$).

The analysis yielded a high inter-rater reliability of $G(q, k) = 0.812$. In the context of behavioral research, a value exceeding 0.80 indicates a robust and high level of consensus. This result confirms that the safety metrics used in URBANGAZEVis are highly stable and that the sample of 10 raters per image provides sufficient statistical power to capture the collective perception of the urban environment.

Appendix H. Detailed Statistical Results for LMM

This section provides the comprehensive statistical outputs for the four semantic data evaluated in this study: *ADE20K Base*, *Grouped-ADE*, *ADE20K + UrbanPD4k*, and *Grouped-ADE + UrbanPD4k*. While the main text focuses on the most predictive features, the following tables provide the full set of fixed effects, including non-significant predictors, to ensure transparency and allow for meta-analytical comparisons.

To focus the analysis on strictly structural and urban elements, broad contextual categories (such as indoor, outdoor, nature, and miscellaneous) have been excluded from these summary tables. Each table reports the estimated coefficients (β), standard errors (S.E.), z -values, Variance Inflation Factors (VIF, where applicable), p -values, and 95% confidence intervals (CI) for the visual elements. Predictors are categorized into elements that positively correlate with perceived safety (Safe) and those that correlate negatively (Unsafe).

Appendix H.1. ADE20K Base Model Results

Table H.1 presents the results for the baseline ADE20K classes. This model serves as the primary benchmark for evaluating the impact of semantic granularity and disorder cues.

Appendix H.2. Grouped-ADE Model Results

Table H.2 details the results for the macro-level Grouped-ADE where elements are clustered into broad infrastructure categories. As discussed in Sec. 9, this model exhibits significant multicollinearity in certain categories.

Appendix H.3. ADE20K + UrbanPD4k Model Results

Table H.3 provides the complete output for the hybrid model integrating specific physical disorder cues. This model incorporates predictors such as “Broken/Damaged Wall” and “Graffiti”, which emerged as critical indicators of perceived insecurity.

Appendix H.4. Grouped-ADE + UrbanPD4k Model Results

Finally, Table H.4 presents the results for the grouped classes augmented with disorder categories. This model evaluates whether the predictive power of disorder cues persists when environmental infrastructure is analyzed at a macro level.

Table H.1: Report: ADE20K (baseline)

Predictor	Coef. (β)	S.E.	z-val	VIF	p-val	95% CI
(Intercept)	5.091***	0.155	32.790	—	< .001	[4.790, 5.400]
<i>Positive (Safe)</i>						
Tree	0.195	0.084	2.310	3.400	0.056	[0.030, 0.360]
Door	0.182**	0.051	3.570	1.300	0.010	[0.080, 0.280]
Palm	0.130	0.056	2.320	1.560	0.056	[0.020, 0.240]
Ashcan	0.113	0.045	2.500	1.020	0.055	[0.020, 0.200]
Fence	0.099	0.093	1.060	4.290	0.398	[-0.080, 0.280]
Plant	0.097	0.055	1.760	1.520	0.162	[-0.010, 0.200]
Bus	0.096	0.050	1.910	1.250	0.135	[-0.000, 0.190]
Minibike	0.078	0.047	1.650	1.130	0.190	[-0.010, 0.170]
Streetlight	0.063	0.046	1.390	1.040	0.262	[-0.030, 0.150]
Car	0.036	0.076	0.470	2.910	0.719	[-0.110, 0.190]
Sidewalk	0.004	0.098	0.040	4.790	0.970	[-0.190, 0.200]
<i>Negative (Unsafe)</i>						
House	-0.003	0.054	-0.060	1.470	0.970	[-0.110, 0.100]
Van	-0.016	0.048	-0.330	1.170	0.803	[-0.110, 0.080]
Person	-0.030	0.055	-0.550	1.500	0.687	[-0.140, 0.080]
Truck	-0.062	0.066	-0.940	2.180	0.448	[-0.190, 0.070]
Stairs	-0.064	0.046	-1.390	1.050	0.262	[-0.150, 0.030]
Grass	-0.072	0.060	-1.190	1.820	0.352	[-0.190, 0.050]
Earth	-0.102	0.097	-1.050	4.740	0.398	[-0.290, 0.090]
Signboard	-0.105	0.056	-1.880	1.560	0.135	[-0.220, 0.000]
Field	-0.130*	0.046	-2.800	1.070	0.027	[-0.220, -0.040]
Mountain	-0.153*	0.047	-3.260	1.110	0.013	[-0.250, -0.060]
Pole	-0.154*	0.049	-3.180	1.170	0.013	[-0.250, -0.060]
Bridge	-0.159*	0.052	-3.060	1.360	0.015	[-0.260, -0.060]
Building	-0.179	0.219	-0.820	24.020	0.509	[-0.610, 0.250]
Sky	-0.188	0.077	-2.430	2.450	0.055	[-0.340, -0.040]
Road	-0.198	0.122	-1.620	7.520	0.190	[-0.440, 0.040]
Wall	-0.370	0.154	-2.400	11.820	0.055	[-0.670, -0.070]
<i>Random Effects</i>						
	Var.	SD	ICC			
Participant	0.665	0.815	0.184			
Residual	2.941	1.715				

* $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$

Table H.2: Report: Grouped-ADE

Predictor	Coef. (β)	S.E.	z-val	VIF	p-val	95% CI
(Intercept)	5.091***	0.159	32.080	—	< .001	[4.780, 5.400]
<i>Positive (Safe)</i>						
Vegetation	0.262	0.149	1.760	10.210	0.288	[-0.030, 0.550]
Terrain Vehicle	0.097	0.136	0.710	8.550	0.781	[-0.170, 0.360]
Human	0.029	0.065	0.450	1.930	0.851	[-0.100, 0.160]
<i>Negative (Unsafe)</i>						
Construction	-0.057	0.306	-0.190	43.290	0.851	[-0.660, 0.540]
Floor	-0.073	0.254	-0.290	29.780	0.851	[-0.570, 0.420]
City Elements	-0.107	0.070	-1.530	2.260	0.346	[-0.240, 0.030]
Sky	-0.178	0.099	-1.790	4.110	0.288	[-0.370, 0.020]
<i>Random Effects</i>						
	Var.	SD	ICC			
Participant	0.692	0.832	0.179			
Residual	3.180	1.783				

* $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$

Table H.3: Report: ADE20K + UrbanPD4k

Predictor	Coef. (β)	S.E.	z-val	VIF	p-val	95% CI
(Intercept)	5.091***	0.159	31.940	—	< .001	[4.780, 5.400]
<i>Positive (Safe)</i>						
Car / Taxi	0.197***	0.047	4.230	1.160	< .001	[0.110, 0.290]
Door	0.155*	0.051	3.060	1.360	0.010	[0.060, 0.250]
Tree	0.100	0.083	1.200	3.510	0.302	[-0.060, 0.260]
Plant	0.078	0.053	1.460	1.520	0.238	[-0.030, 0.180]
Trashcan	0.076	0.044	1.720	1.040	0.164	[-0.010, 0.160]
Palm	0.076	0.055	1.390	1.580	0.245	[-0.030, 0.180]
Streetlight	0.052	0.045	1.170	1.050	0.308	[-0.040, 0.140]
Motorcycle	0.041	0.046	0.890	1.140	0.412	[-0.050, 0.130]
Bus	0.041	0.049	0.840	1.270	0.429	[-0.060, 0.140]
Fence	0.013	0.089	0.150	4.180	0.884	[-0.160, 0.190]
<i>Negative (Unsafe)</i>						
Van	-0.047	0.047	-1.010	1.170	0.370	[-0.140, 0.040]
House	-0.050	0.053	-0.940	1.500	0.395	[-0.150, 0.050]
Sidewalk	-0.061	0.094	-0.650	4.600	0.532	[-0.240, 0.120]
Stairs	-0.068	0.045	-1.520	1.060	0.224	[-0.150, 0.020]
Car	-0.083	0.074	-1.130	2.870	0.315	[-0.230, 0.060]
Overhead Cable	-0.088	0.054	-1.640	1.500	0.184	[-0.190, 0.020]
Truck	-0.089	0.065	-1.370	2.250	0.245	[-0.220, 0.040]
Person	-0.095	0.054	-1.770	1.530	0.159	[-0.200, 0.010]
Grass	-0.116	0.060	-1.940	1.880	0.116	[-0.230, 0.000]
Ground	-0.117	0.090	-1.300	4.330	0.268	[-0.290, 0.060]
Field	-0.138*	0.045	-3.070	1.070	0.010	[-0.230, -0.050]
Garbage Bag	-0.142*	0.049	-2.900	1.280	0.012	[-0.240, -0.050]
Pole	-0.142*	0.048	-2.990	1.210	0.010	[-0.240, -0.050]
Signboard	-0.152*	0.055	-2.770	1.590	0.015	[-0.260, -0.040]
Mountain	-0.184***	0.046	-4.020	1.120	< .001	[-0.270, -0.090]
Bridge	-0.188**	0.051	-3.690	1.370	0.001	[-0.290, -0.090]
Sky	-0.200*	0.071	-2.830	2.170	0.014	[-0.340, -0.060]
Graffiti	-0.225*	0.075	-2.990	2.990	0.010	[-0.370, -0.080]
Damaged Pavement	-0.229***	0.056	-4.100	1.650	< .001	[-0.340, -0.120]
Road	-0.277*	0.117	-2.360	7.340	0.043	[-0.510, -0.050]
Building	-0.293	0.204	-1.430	22.120	0.238	[-0.690, 0.110]
Wall	-0.326*	0.121	-2.690	7.780	0.018	[-0.560, -0.090]
Broken/Dam. Wall	-0.442***	0.090	-4.930	4.270	< .001	[-0.620, -0.270]
<i>Random Effects</i>						
	Var.	SD	ICC			
Participant	0.707	0.841	0.203			
Residual	2.774	1.666				

* $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$

Table H.4: Report: Grouped-ADE + UrbanPD4k

Predictor	Coef. (β)	S.E.	z-val	VIF	p-val	95% CI
(Intercept)	5.091***	0.161	31.700	—	< .001	[4.780, 5.410]
<i>Positive (Safe)</i>						
Car / Taxi	0.231***	0.050	4.590	1.270	< .001	[0.130, 0.330]
Vegetation	0.156	0.138	1.130	9.500	0.425	[-0.110, 0.430]
Trashcan	0.085	0.046	1.850	1.050	0.128	[-0.000, 0.170]
<i>Negative (Unsafe)</i>						
Terrain Vehicle	-0.024	0.124	-0.190	7.710	0.896	[-0.270, 0.220]
Human	-0.038	0.061	-0.620	1.870	0.690	[-0.160, 0.080]
Overhead Cable	-0.082	0.059	-1.380	1.740	0.299	[-0.200, 0.030]
Construction	-0.104	0.275	-0.380	37.730	0.848	[-0.640, 0.440]
City Elements	-0.128	0.066	-1.930	2.190	0.120	[-0.260, 0.000]
Garbage Bag	-0.142*	0.054	-2.630	1.460	0.031	[-0.250, -0.040]
Floor	-0.158	0.226	-0.700	25.460	0.690	[-0.600, 0.290]
Sky	-0.180	0.086	-2.100	3.260	0.107	[-0.350, -0.010]
Graffiti	-0.183	0.095	-1.930	4.450	0.120	[-0.370, 0.000]
Damaged Pavement	-0.204*	0.067	-3.040	2.240	0.011	[-0.330, -0.070]
Broken/Dam. Wall	-0.382*	0.122	-3.130	7.440	0.011	[-0.620, -0.140]
<i>Random Effects</i>						
	Var.	SD	ICC			
Participant	0.715	0.846	0.196			
Residual	2.943	1.716				

* $p < 0.05$, ** $p < 0.01$, *** $p < 0.001$

Appendix I. Segmentation Model Comparison

Fig. I.2 shows the segmentation mask obtained from these three models. We note that SegFormer, PSPNet, and DeepLab do not segment the light pole; they fail to differentiate among walls, sidewalks, and gates, and they fail to segment cars correctly. However, while it requires more computational resources, the superior accuracy and versatility of OneFormer often justify the additional effort.

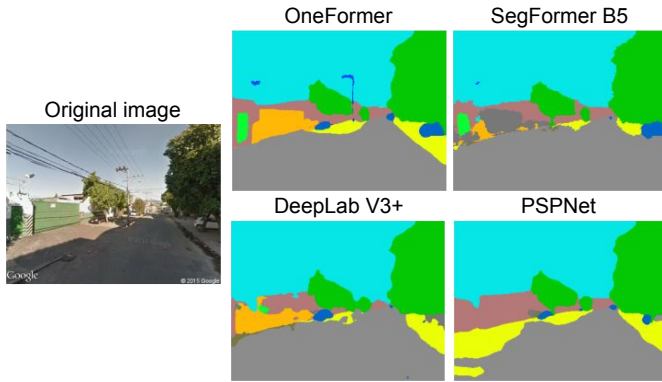


Figure I.2: Comparison of segmentation outputs produced by OneFormer, SegFormer B5, PSPNet, and DeepLabV3+. Some visual elements, such as light poles, cars, gates, and walls, are incorrectly classified, illustrating the differences in spatial accuracy across models.