

HeedVision: Attention Awareness in Collaborative Immersive Analytics Environments

Arvind Srinivasan , Niklas Elmqvist 

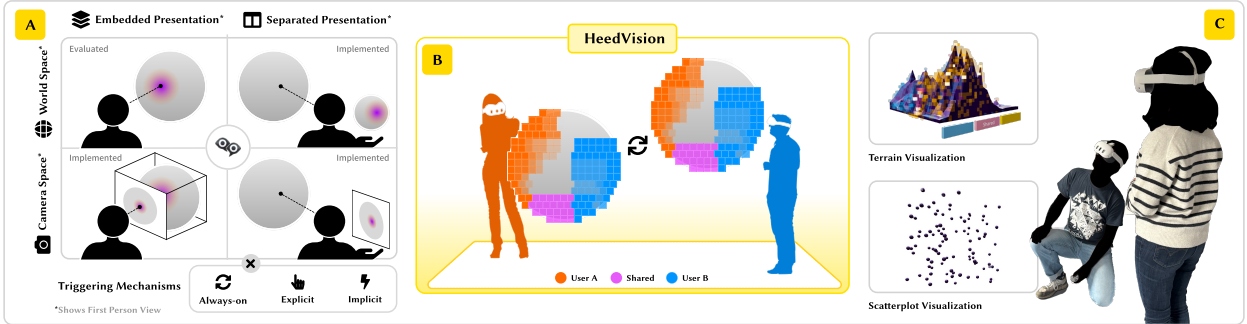


Figure 1: **HeedVision overview.** We propose a (A) Collaborative Attention-Aware Visualization design space covering attention capture, presentation, situatedness, and triggering. (B) HeedVision implements this in a co-located mixed reality setting using voxel-based coverage maps for shared attention. (C) We evaluate it through a treasure hunt task across terrain and scatterplot visualizations with and without attention cues, revealing context-dependent benefits in abstract environments lacking natural reference points.

Abstract—*Group awareness*, the ability to perceive the activities of collaborators in a shared space, is a vital mechanism to support effective coordination and joint data analysis in collaborative visualization. We introduce *collaborative attention-aware visualizations* (CAAVs) that track, record, and revisualize the collective attention of multiple users over time. We implement this concept in HEEDVISION, a standards-compliant WebXR system built with React Three Fiber that runs on modern AR/VR headsets, and complement it with proof-of-concept implementations covering the remaining three quadrants of our design space, varying presentation (embedded vs. separated) and situatedness (world space vs. camera space). Through a mixed-methods exploratory study where pairs of co-located analysts performed visual search tasks in a shared immersive AR environment, we investigate how attention revisualization affects collaborative coordination in immersive analytics. Our results show that CAAVs can improve spatial coordination, search efficiency, and task load distribution in collaborative visual search tasks among collaborators, though benefits vary by context, favoring abstract environments lacking natural landmarks. This work extends attention awareness to multi-user settings and provides empirical evidence for its context-dependent benefits in collaborative visual search tasks within immersive analytics environments.

Index Terms—Attention tracking, eyetracking, immersive analytics, ubiquitous analytics, post-WIMP interaction.

1 INTRODUCTION

Attention is a scarce commodity. Human cognition uses attention as a selection mechanism that determines which sensory inputs are prioritized for processing at any given moment [62]. In visualization, attention helps the analyst focus on areas of interest across charts to address the current analysis goal [58]. This is particularly important in collaborative visualization [31], as it enables multiple analysts to jointly make sense of complex data. However, effective collaboration demands that each participant maintain awareness of what their partners are doing. This *group awareness* (the ability to perceive the activities of collaborators in a shared space [25, 26]) becomes especially challenging in immersive 3D environments, where multiple viewpoints, distributed spatial attention, and real-time coordination create unique obstacles. Prior work has shown how AR/VR environments change the dynamics of collaborative exploration and sensemaking [22, 44], and how gaze, gesture, and embodied cues can support group awareness [5, 52]. More recent efforts have explored representing collaborators' focus of attention on desktop and large-display systems [35, 39]. Yet few systems directly capture and represent collaborators' attention over

time in immersive environments, leaving a gap at the intersection of attention-aware visualization and immersive collaboration.

Attention-aware visualizations (AAVs) [58] address part of this challenge for individual users: they measure a single viewer's attention over time and modify the display accordingly to encourage exploration or support reflection. In this paper, we extend AAVs to multi-analyst settings in what we call *Collaborative Attention-Aware Visualizations* (CAAVs). While this extension builds on the same three-component structure (capture, record, display), it introduces fundamentally new design challenges with no single-user analog. For capture, we must distinguish between users while tracking attention simultaneously. For recording, we maintain identity-aware attention accumulators and introduce aggregation methods (sum, maximum, difference, count) to represent collective patterns. For display, we must encode both attention intensity and user identity while supporting coordination rather than merely self-reflection. These challenges motivate a systematic design space for CAAVs addressing three research questions: how to **capture** collective attention in real time (RQ1), how to **record** it over time (RQ2), and how to **display** it effectively (RQ3).

To probe this design space through a concrete instantiation, we developed HEEDVISION, an immersive analytics environment that tracks and revisualizes the attention of collaborating pairs. HEEDVISION embodies the world-space embedded configuration of our design space: head-based attention capture for accessibility across consumer devices, voxel-based recording with temporal decay, and explicit world-embedded display with color-per-user encoding. Built using React Three Fiber with WebXR support, the system runs on consumer-level AR/VR headsets such as the Meta Quest 3. We com-

• Arvind Srinivasan, and Niklas Elmqvist are with Aarhus University in Aarhus, Denmark. E-mail: {arvind, elm}@cs.au.dk.

Manuscript received xx xxx. 202x; accepted xx xxx. 202x. Date of Publication xx xxx. 202x; date of current version xx xxx. 202x. For information on obtaining reprints of this article, please send e-mail to: reprints@ieee.org. Digital Object Identifier: xx.xxxx/TVCG.202x.xxxxxx

plement HEEDVISION with proof-of-concept implementations of the remaining three quadrants (world-space separated, camera-space embedded, and camera-space separated), illustrating the breadth of the design space. We evaluated HEEDVISION through a mixed-methods exploratory study with eight co-located pairs (16 participants) performing collaborative visual search tasks in a shared augmented reality environment. To understand whether benefits depend on visualization context, we tested CAAVs across both abstract discrete environments (3D scatterplots) and continuous environments with natural landmarks (terrain visualizations). Our findings suggest that CAAV effectiveness depends on visualization context: attention visualization appeared to improve spatial coordination and reduce redundant exploration in abstract environments lacking natural reference points, but introduced potential interference in environments where existing landmarks already supported coordination.

We claim the following contributions:

- (1) A preliminary design space for collaborative attention-aware visualization, extending AAV to multi-user immersive settings through three research questions that organize capture, recording, and display alternatives for collective attention;
- (2) HEEDVISION, a WebXR research prototype built with React Three Fiber that instantiates the world-space embedded configuration, complemented by working proof-of-concept implementations of the remaining three design space quadrants that illustrate the framework’s breadth, though these three use local synchronization rather than networked collaboration and have not been formally evaluated; and
- (3) Empirical evidence from an exploratory study showing when collective attention visualization helps versus hinders collaborative coordination, revealing context-dependent benefits that depend on the structure of the visualization environment.

All design decisions, implementation parameters, and evaluation materials are documented to enable replication and extension. Supplementary materials are available at osf.io/frsp8.

2 BACKGROUND

Our work extends attention-aware visualization to collaborative immersive settings. Below we review the foundations on which we build.

2.1 Group Awareness in Collaborative Visualization

Visual attention determines which elements of a complex visualization are prioritized for cognitive processing [62]. Visualization designers leverage this through pre-attentive features [29] and visual saliency models [34, 46] that predict and guide where viewers direct their focus. Eye tracking has provided empirical insights into these attentional patterns across different visualization types [6, 14, 15, 37, 40].

While individual attention is well studied, collaborative visualization introduces the additional challenge of *group awareness*: understanding others’ activities within a shared workspace [25, 26]. Group awareness is essential for establishing *common ground*: the mutual knowledge, beliefs, and assumptions that underpin effective coordination [16]. The benefits of such collaboration are well documented: Mark and Kobsa [43] and Balakrishnan et al. [2] both found significant improvements when analysts worked together using shared visual representations.

Several techniques support group awareness in collaborative settings: radar overviews [27], embodied cues like collaborators’ arms [61] and fingertips [36], and social navigation techniques [20] that record and visualize user presence and activity [39]. Scented Widgets [64] embed usage traces directly within interface components, while Matejka et al.’s Patina [45] overlays dynamic heatmaps of usage patterns on controls; both demonstrate how aggregated interaction data can convey collective behavior. Collaborative brushing techniques [28, 32] reveal others’ selections across datasets, and Saffo et al.’s “Eyes and Shoes” [52] technique provides asymmetric awareness cues between VR and desktop.

A related dimension of awareness is *collective data coverage*: ensuring the group examines all relevant data without duplication. Sarvghad

and Tory demonstrated that visualizing dimension coverage increases exploration breadth [55] and reduces work duplication in asynchronous collaboration [54]. Badam et al. [1] extended this with a “team-first” approach that balances individual with collective sensemaking. This notion of coverage tracking motivates our recording dimension (Section 3).

RESEARCH GAP *Most collaborative systems focus on discrete interaction events rather than continuous attention, and few leverage attention in immersive 3D where spatial coordination is paramount.*

2.2 Gaze and Attention Tracking

Addressing continuous attention in immersive environments requires reliable head tracking. Gaze serves both as an implicit indicator of visual attention and as an explicit input mechanism. Since Bolt’s seminal gaze-orchestrated windows [7], researchers have explored gaze for target acquisition [59], multimodal interaction [47, 48, 60], and cross-device interaction [63]. A persistent challenge is the “Midas touch” problem [33], where every glance triggers an action; multimodal approaches that require secondary confirmation [60] and explicit triggering mechanisms can mitigate this. Gaze control is now central to consumer XR devices such as the Apple Vision Pro and Android XR.

In immersive settings where eye trackers are unavailable or impractical, head orientation serves as a proxy for visual attention. Studies of viewing behavior in 360° video and VR have shown that head and gaze directions are strongly correlated: Sitzmann et al. [57] found that the mean gaze direction relative to head orientation is $13.85^\circ \pm 11.73^\circ$, indicating strong coupling between head and gaze movements, and David et al. [19] confirmed this coupling in active VR exploration. This correlation is sufficiently strong for tasks where conveying general areas of attention matters more than precise fixation points; as is the case for collaborative spatial coordination, knowing approximately where a partner is looking enables effective division of labor.

RESEARCH GAP *While gaze and head tracking are well established for single users, using these signals to support awareness across multiple collaborators in immersive environments is largely unexplored.*

2.3 Attention-Aware Visualization

Attention-aware visualizations (AAVs) represent a class of systems that track, analyze, and respond to user attention patterns. Srinivasan et al. [58] formalized this concept with a framework organized around three components: how attention is captured, how it is recorded over time, and how the visualization adapts in response. They distinguish between *data-agnostic* approaches that track attention as spatial coordinates without knowledge of the underlying data, and *data-aware* approaches that associate attention with marks or semantic components.

Building on this foundation, Jianu et al. [35] assembled a comprehensive design framework and research agenda for gaze-aware visualizations, distilling research challenges and providing practical guidance for integrating eye-tracking into visualization systems. Earlier systems demonstrated these principles in practice: Exploration Awareness [56] tracked user exploration paths through information spaces, while adaptive annotations [37] adjusted to viewing patterns. In immersive environments, Lu et al. [41] investigated subtle visual cues for directing attention in augmented reality, finding that appropriate cue designs can guide visual search without disrupting the primary task.

Beyond tracking and recording, display techniques can communicate attention history without persistent clutter. Baudisch et al. [4] introduced phosphor transitions, afterglow effects that fade over time to convey interface changes; this ephemeral approach offers relevant design alternatives for attention-aware revisualization.

However, existing AAV systems share a fundamental limitation: they are designed for a single user. They track one viewer’s attention and adapt the display for that viewer alone. In collaborative settings, the challenge shifts from individual reflection to collective coordination; analysts need to know not only where *they* have looked, but where their partners have been, where coverage gaps remain, and how to divide a shared analytical space. These collaborative requirements motivate our extension of AAVs to multi-user settings: Collaborative Attention-Aware Visualizations (CAAVs).

RESEARCH GAP The gap between well-developed single-user AAV approaches and the requirements of collaborative IA motivates extending attention awareness to multi-user settings, where coordination, coverage, and shared understanding become central design goals.

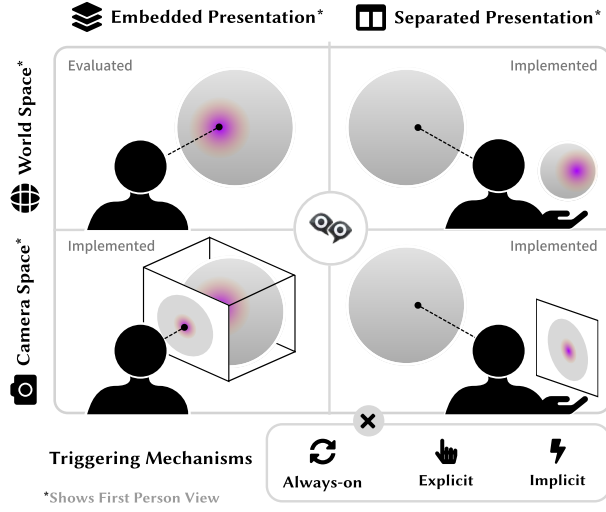


Figure 2: **Revisualizing collective attention (RQ3).** The CAAV design space spans three dimensions (measuring/RQ1, recording/RQ2, revisualizing/RQ3); this figure shows RQ3 only, organized along three sub-dimensions: (1) *presentation* distinguishes between embedding attention directly into the visualization versus separating it into distinct views or overlays; (2) *situatedness* determines whether attention is rendered in world space (anchored to the environment) or camera space (following the user's viewport); and (3) *triggering* specifies when attention visualization appears: continuously (always-on), upon user request (explicit), or based on system-detected context (implicit). The quadrant formed by presentation and situatedness (top) illustrates the four primary configurations, while triggering mechanisms (bottom) can be applied to any. HEEDVISION instantiates the highlighted world-space embedded quadrant.

3 COLLABORATIVE ATTENTION-AWARE VISUALIZATION

Srinivasan et al. [58] defined attention-aware visualizations (AAVs) as displays that track a single viewer's attention and adapt accordingly. We extend this concept to collaborative settings: *Collaborative Attention-Aware Visualizations* (CAAVs) track, record, and revisualize the attention of multiple users engaged in joint analysis tasks (Figure 2). This extension preserves AAV's three-component structure but introduces collaborative dimensions at each level:

- RQ1 How should the attention of multiple viewers on a visualization be **captured in real time**?
- RQ2 How should the attention of multiple viewers on a visualization be **recorded over time**?
- RQ3 How should the attention of multiple viewers on a visualization be **displayed**?

RQ1 maps to Section 3.1 (capture mechanisms); RQ2 maps to Section 3.2 (recording approaches); and RQ3 maps to Section 3.3 (revisualization, encompassing presentation, situatedness, and triggering).

We present this design space not merely as scaffolding for one system, but as a preliminary generative map of unexplored configurations for collaborative attention visualization. We present this as a *preliminary* design space rather than a validated framework; its generative value lies in mapping unexplored configurations for future work, analogous to how task typologies [12] structure design reasoning without themselves being empirically validated in every cell. HEEDVISION (Section 4) instantiates the world-space embedded quadrant and serves

as the focus of our evaluation; we additionally provide implementations of the remaining three quadrants to demonstrate the framework's generative capacity.

3.1 Measuring Collective Attention (RQ1)

If attention represents a person's focus on specific elements, then *collective attention* extends this across multiple analysts. Involving multiple agents yields concepts with no single-user counterpart:

- **Joint Attention**, when multiple users focus on one element;
- **Distributed Attention**, when analysts distribute their focus;
- **Sequential Attention**, when analysts build on each other's focus;
- **Cumulative Attention**, as the collective span of collaborators' attention.

Since attention cannot be directly measured, we must use observable behaviors as proxies. In collaborative settings, these proxies must work for multiple users simultaneously: tapping or clicking; pointer movement; head direction as an indication of gaze; and eye movement for precise focus. Each offers different tradeoffs between precision, implementation complexity, and intrusiveness. Measuring across multiple users also introduces challenges of distinguishing between users, aligning interaction events to a common timeline, managing scalability as team size increases, and accommodating users on different devices.

3.2 Recording Collective Attention (RQ2)

CAAVs can take two fundamentally different approaches to recording attention: *data-agnostic* (no knowledge of the underlying visualization) and *data-aware* (specific knowledge of visualization geometry).

Data-agnostic. This approach tracks attention as coordinates or regions without knowledge of the underlying visualization elements; in other words, the system does not know what the visualization represents, only that the user's attention is moving over it. Data-agnostic approaches are more general but limit attention interpretation.

Data-aware. This method associates attention with specific visualization marks, data points, or semantic components. Because of this knowledge, data-aware approaches enable richer analysis but require deeper integration with the visualization system.

Representing collective attention. Where AAV maintains a single attention accumulator per target, CAAVs must maintain separate accumulators for each user on each target. These accumulators start at zero, increase as a user directs attention to the target, and can be combined to represent collective attention using various aggregation methods: *sum* (total collective attention); *maximum* (highest individual attention); *difference* (balance of attention between members); and *count* (number of analysts who attended to the target). The choice of aggregation affects what collaborative patterns become visible; sum emphasizes total coverage, while difference highlights division of labor.

Collective attention decay. Human memory decays naturally over time, and recorded attention values should gradually decrease to reflect fading relevance. Decay is particularly important in collaborative settings where attention patterns evolve as participants respond to each other's discoveries and insights.

3.3 Revisualizing Collective Attention (RQ3)

As Srinivasan et al. [58] noted, attention does not have to be displayed; it could be used solely for analysis or evaluation. This paper focuses on how the visualization might actively incorporate this data to help coordination, promote exploration, or support reflection; we refer to this process as *revisualizing* the attention.

Presentation. The revisualization can be achieved through two fundamental presentation styles: (a) **embedding** the attention directly into the visualization; or (b) **separating** the attention from the visualization in a view or overlay. In collaborative contexts, the choice between approaches affects how easily participants can develop shared understanding of collective attention. The presentation choice may also be left to each individual collaborator, and the actual implementation depends on whether the approach is data-aware or data-agnostic and whether it is displayed in 2D or 3D.

Situatedness. We also distinguish whether the attention is integrated into the **world space** or the **camera space** of the display. In **world space**, attention visualizations remain fixed in the shared environment, preserving spatial context and enabling deictic reference (“look at that orange region”). In **camera space**, attention visualizations attach to each user’s viewport, ensuring continuous visibility but sacrificing direct spatial correspondence. This distinction matters particularly in immersive 3D where users occupy different positions with independent viewpoints.

Triggering. Following AAV [58], the triggering mechanisms for revisualization may significantly impact user experience: **always-on**, which continuously displays attention (risking Midas touch or reinforcing loops); **explicit**, where users deliberately activate revisualization; and **implicit**, where the system automatically shows attention based on context. In collaborative settings, different triggering mechanisms might apply to personal versus shared displays, and the ability to trigger might depend on user role. Collaboration modality (synchronous vs. asynchronous; co-located vs. remote) is a further cross-cutting influence: asynchronous collaborators may prefer decay-weighted historical overlays for handoff support [54], while always-on triggering can overwhelm remote collaborators on high-latency feeds, making explicit or implicit modes preferable.

Collaborative revisualization patterns. Several patterns emerge as particularly useful for collaborative settings: *Attention Heatmaps* show density of aggregate attention; *Coverage Maps* highlight areas that have received little attention; and *Focus Convergence* emphasizes where multiple users are focusing. These address complementary collaborative needs identified in prior work on group awareness [25] and collaborative visualization [31]: collective interest, thorough examination without redundancy [55], and shared discovery. For this investigation, we prioritize the three patterns as starting points that minimize visual complexity in immersive environments [22].

4 SYSTEM: HEEDVISION

We present HEEDVISION, a collaborative attention-aware visualization system for immersive analytics. Within our design space, HEEDVISION implements attention that is **embedded** in the **world space** and triggered **explicitly**. We additionally provide functional prototype implementations of the remaining three design space quadrants (Section 4.5, Figure 4). We selected the world-space embedded configuration for HEEDVISION because embedded visualization preserves spatial context crucial for collaborative 3D exploration, world-space placement ensures both collaborators share a common reference frame for deictic communication, and explicit triggering gives users agency over when to consult attention information, reducing visual clutter during active exploration while supporting deliberate coordination. These choices prioritize unobtrusive capture with user-controlled display, reflecting our goal of supporting (rather than disrupting) natural collaborative workflows.

4.1 Attention Model

The core architecture features a distributed attention model that captures, records, and visualizes collective attention across multiple visualizations within a shared workspace. Unlike single-user AAV systems, our model maintains separate attention maps for each collaborator on each visualization. The attention model is implemented as a 3D voxel grid overlaid on each visualization, with individual voxels storing attention values for specific regions. This supports both individual and

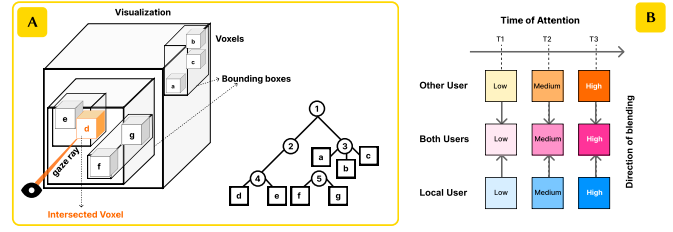


Figure 3: **Voxelization process.** (a) BVH [17] voxelization of a visualization, where the user’s gaze ray intersects with the voxel grid, altering its state on intersection; (b) attention accumulation in individual voxels over time with color blending across two collaborators.

aggregated attention maps, enabling users to review personal focus areas and identify regions of collective interest.

The system updates attention values in real time as collaborators interact. The model follows principles of human short-term memory: attention accumulates when a user focuses on a specific region and decays over time as attention shifts elsewhere. Only voxels that intersect with the actual geometry of a visualization register and accumulate attention data, creating a *data-aware* approach that directly links attention to the visualization’s geometric structure. This same infrastructure can support *data-agnostic* approaches through alternative capture mechanisms (see Section 4.3).

4.2 Voxelization

At initialization, HEEDVISION converts the geometry of a visualization into a discrete grid of invisible voxels at configurable resolution. These voxels serve as the fundamental units for capturing attention, becoming visible only when attention revisualization is triggered. The voxel-based approach enables fine-grained spatial tracking: each voxel independently accumulates attention data, establishing a consistent spatial reference across users, simplifying detection of joint attention when multiple users focus on identical regions, and optimizing computation of attention coverage metrics. We employ Bounding Volume Hierarchies (BVHs) [17] (Figure 3) to accelerate voxelization by efficiently identifying mesh regions that intersect with voxel boundaries.

4.3 Attention Capture

Given gaze information, we capture attention by casting a ray from the user’s view position in their gaze direction, determining the nearest voxel intersection, creating a spherical *region of influence* around it, and incrementing attention values for all voxels within that region with the center receiving the highest increment. Using headsets without eye-tracking capabilities (such as the Meta Quest 3), we use head orientation as a proxy for user attention. For a *data-agnostic* approach, the algorithm would instead accumulate attention in every voxel the gaze ray passes through.

Head orientation introduces imprecision, as users often attend to peripheral regions. Combined with our single-voxel region of influence (Section 4.5), this creates approximate attention maps. However, this suffices for collaborative coordination where conveying *general areas* of partner attention matters more than exact fixation points; this aligns with the strong head-gaze correlation documented in VR viewing studies [19, 57].

4.4 Triggering and Revisualization

HEEDVISION implements an **explicit** triggering mechanism where participants press a controller button while pointing at a visualization to reveal collaborative attention patterns. This prevents Midas touch [33] while giving users control over when to consult the attention display.

For revisualization, HEEDVISION uses an *opacity-driven* approach: each voxel adjusts its opacity based on accumulated attention, becoming more opaque as attention increases and remaining transparent when unattended or below a user-controlled threshold. This creates a real-time coverage map that reveals where attention has been concentrated.

Color distinguishes collaborators: each user is assigned a unique color, and when multiple users attend to the same voxel, their colors blend additively, effectively implementing a per-user aggregation that makes joint attention visually salient. The blue–orange pair was selected for colorblind accessibility (safe for deuteranopia and protanopia), and the resulting additive blend, a warm rose, was chosen deliberately: as a perceptual mix of both user colors, it communicates shared attention without requiring a separate legend entry. In the evaluated AR system, attention voxels were opaque, saturated overlays with a per-user legend, so each user's color stayed legible against the terrain rather than reading as part of it (Figure 4, top-left). This color-based approach works well for small teams (2–4 members) but presents scalability challenges for larger groups; for instance, hierarchical grouping could assign team-level colors to sub-groups and individual colors within each toggle, reducing the palette to N/k distinct colors for k sub-groups. Alternative visual channels (see Section 7.3) such as role-based coloring, texture, outline thickness, or animation could address both scaling and conflicts with data-driven color encodings, since our current implementation assumes color is not the primary data channel. The attention visualization is directly embedded in the 3D world space, allowing users to explore attention patterns from any viewpoint while preserving spatial context.

4.5 Implementation

HEEDVISION is built on open web technologies using React Three Fiber [50] with @react-three/xr for WebXR support, running natively in modern browsers without specialized installation. The collaborative infrastructure uses Spatialstrates [9], DashSpace [8], Webstrates [38], and Varv [10] for real-time cloud-based state synchronization, with an interactive Moveable marker system for positioning visualizations through drag and rotation. The additional implementations of the remaining three quadrants (Section 4.5) are also built with React Three Fiber but operate in VR-only mode and use the BroadcastChannel API for local peer-to-peer synchronization between browser tabs without persistent cloud storage, enabling lightweight simulation of pair collaboration on a single device.

To maintain interactive frame rates, we made several trade-offs: attention capture and decay are computed only for the directly intersected voxel to avoid frame-blocking computations; the region of influence is restricted to a single voxel rather than surrounding neighbors; attention tracking operates at 10 Hz (below the 60 Hz rendering rate); and attention updates are batched and synchronized across devices at configurable intervals (200–500ms). The primary bottleneck lies in network synchronization, as each update must be broadcast and merged across all clients. Thus, a pragmatic tradeoff that sacrifices some spatial precision was necessary to maintain the responsiveness essential for fluid collaborative interaction.

Additional Design Space Implementations To demonstrate the breadth of our design space beyond HEEDVISION's world-space embedded configuration, we built fully functional implementations of the remaining three quadrants (Figure 4). Unlike HEEDVISION, which supports co-located AR with cloud-based synchronization, these implementations operate in VR-only mode and use the BroadcastChannel API to mimic pair collaboration locally; and have not undergone formal user evaluation.

World-space separated. A miniature replica of the visualization floats beside the main view and mirrors attention patterns in real time; users can reposition it via hand tracking or controller input to inspect collective attention from an independent vantage point. This configuration is especially suited to larger environments where the main view is too cluttered; the miniature replica gives an unobstructed overview comparable to a minimap in 2D collaborative interfaces.

Camera-space embedded. Attention is rendered as 2D projections overlaid onto the faces of the visualization's bounding box, computed from the user's current viewpoint, ensuring continuous visibility but sacrificing depth information. Camera-space embedding is most useful when collaborators share a synchronized viewport (e.g., a guided

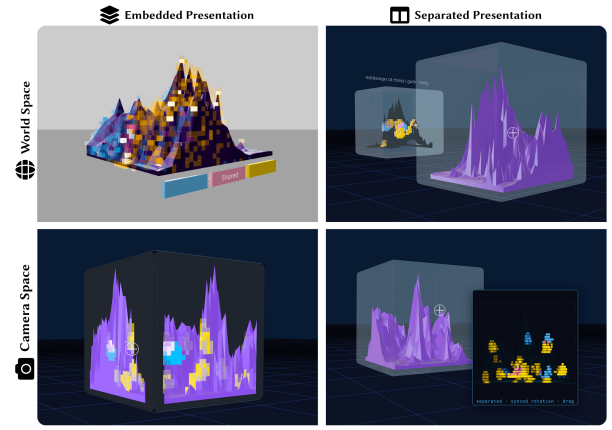


Figure 4: **Implementations across all four design space quadrants:** *Top-left:* World-space embedded (HeedVision), shown as a screenshot of the actual evaluated AR system, with the study's blue/gold/pink attention voxels and the per-user legend. *Top-right:* World-space separated. *Bottom-left:* Camera-space embedded. *Bottom-right:* Camera-space separated. The latter three are later proof-of-concept prototypes that share HEEDVISION's attention model but run VR-only and were not formally evaluated.

tour scenario), as the 2D projection remains interpretable regardless of individual position.

Camera-space separated. A 2D canvas panel in screen space displays a top-down or rotated view of the attention map synchronized with the main visualization's orientation, providing an overview without occluding the 3D scene. A top-down panel suits asynchronous handoffs: a second analyst joining later can review where their partner explored without navigating the full 3D scene.

5 USER STUDY

To evaluate our proposed collaborative attention-aware visualization approach and the HEEDVISION implementation, we conducted a mixed-methods exploratory study. Our goal was to understand how visualizing multiple users' attention patterns impacts coordination and coverage.

This study was conducted in accordance with the research ethics guidelines of Aarhus University. As a behavioral study involving no health data, it falls outside the scope of the Danish Act on Research Ethics Review of Health Research Projects (*Komitéloven*). All participants provided written GDPR informed consent prior to participation.

5.1 Participants

We recruited 16 participants (11 male, 5 female) who formed 8 collaborative pairs (see Table 1). Participants were drawn from graduate-level Computer Science, Data Visualization, and HCI courses, ensuring sufficient background knowledge in data visualization. Ages ranged from 24 to 37 (mean 28.9), and all had normal or corrected-to-normal vision with no self-reported color vision deficiency. The participants varied in their experience with immersive environments, co-located collaboration, and visual analytics (from no experience to experienced). None of the participants reported prior familiarity with the Mt. Bruno terrain dataset used in the study.

5.2 Apparatus

All participants used our implementation running on Meta Quest 3 headsets in a co-located laboratory space (6×3×3m), allowing for verbal communication. The system tracked head position and orientation to determine attention for each participant. Each participant's headset ran the HEEDVISION software on the Quest Browser. The world-embedded attention display could be explicitly triggered using a controller button.

Table 1: **Participant demographics for collaborative pairs** (G1-G8) showing age, gender, vision, and their experience with immersive environments, co-located collaboration, and visual analytics ($N = 16$).

Participant A									Participant B								
			Age Group	Vision	MR Exp.	Coloc. Exp.	VA Exp.					Age Group	Vision	MR Exp.	Coloc. Exp.	VA Exp.	
G1	♂	P1	35-39	6d	★★★	★★	★	♂	P2	25-29	6d	★★	★★	★★	★★★		
G2	♀	P3	25-29	6d	⊙	⊙	★★	♂	P4	25-29	6d	⊙	★★	★★	★★		
G3	♀	P5	35-39	6d	★	★★	★★★	♂	P6	25-29	6d	★	★★★	★★★	★★		
G4	♂	P7	20-24	6d	★★	★★	★★	♂	P8	25-29	6d	★	★★★	★★★	★★		
G5	♀	P9	30-34	6d	★★	★★★	★★	♂	P10	30-34	6d	★★★	★★★	★★★	★★★		
G6	♂	P11	25-29	6d	★★	★★	★★	♂	P12	20-24	6d	★	★★	★★	★★		
G7	♀	P13	30-34	6d	⊙	★★	★★	♂	P14	30-34	6d	⊙	★★	★★	★★		
G8	♂	P15	25-29	6d	★	★★	★★	♀	P16	25-29	6d	★★★	★★★	★★★	★★★		

Summary: $N=16$, 11 ♂ Male, 5 ♀ Female, ages 24–37 ($M=28.9$), 8 6d Corrected-to-normal, 8 6d Normal vision

Legend: Experience: ★★★ Very experienced, ★★ Experienced, ★ Some experience, ★ Novice, ⊙ No experience; Abbreviations: MR = Mixed Reality, Coloc. = Co-located Collaboration, VA = Visual Analytics.

5.3 Experimental Conditions

We employed a 2×2 mixed (split-plot) design with two factors:

- **Visualization Type.** The 3D chart type used.
 - **Terrain:** A continuous 3D terrain visualization of Mt. Bruno’s elevation with targets positioned along the surfaces of the entire volume.
 - **Scatter:** A discrete 3D scatterplot with 100 uniformly distributed data points. We chose this distribution rather than realistic clustered data (through dimensionality reduction techniques like t-SNE or UMAP) to isolate CAAV’s effect from coordination affordances that natural clusters might provide, thus creating a baseline where spatial structure alone offers no guidance for dividing labor.
- **Attention Visualization.** Attention display method.
 - **No-CAAV (control):** No attention display.
 - **CAAV:** World-embedded attention display.

Each pair used CAAV for one visualization type and No-CAAV for the other; thus visualization type varied within pairs and attention visualization varied between pairs (four pairs per cell), with the CAAV-to-geometry assignment and presentation order counterbalanced across the eight pairs to mitigate learning effects.

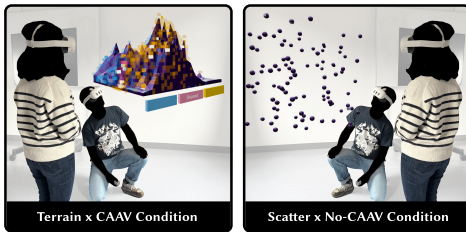


Figure 5: **Experiment with the treasure hunt task.** A pair of participants perform the treasure hunt task. *Left:* Terrain visualization with CAAV enabled (colored voxels visible). *Right:* Scatterplot visualization without CAAV. Each pair experienced both visualization types, one with CAAV and one without.

5.4 Task

All pairs engaged in a collaborative visual search task: a treasure hunt (see Figures 5 and 6). Visual search tasks are well-established in immersive analytics research for evaluating spatial coordination and attention [42, 51]. We designed this task as a simplified probe into key collaborative challenges (spatial navigation, shared attention, and division of labor), without domain- or task-specific complexity that

might confound our analysis of attention-aware mechanisms. During each 10-minute session, pairs searched for white target objects hidden within a 3D visualization. These targets were initially invisible and only became visible (appearing to glow white) when a participant looked at them from the correct angle and within a 10 cm proximity.

In the CAAV condition, attention-based coloring supported collaboration: voxels changed color to indicate where participants had looked (blue for the local user, orange for the partner, and pink for areas explored by both). Participants could toggle the visibility of these attention indicators using a controller button, enabling strategic exploration of unsearched regions and potentially enhancing collective search.

The number of targets was set to five percent of total voxels for each visualization type, with targets distributed pseudo-randomly across the visualization volume using a seeded random generator to ensure reproducibility. For the terrain condition, targets were constrained to lie on or near the terrain surface; for the scatterplot condition, targets were distributed throughout the 3D volume surrounding the data points. This distribution ensured that no single region was disproportionately target-rich, requiring pairs to explore the full visualization space.

Task Rationale. The “treasure hunt” was designed to isolate core coordination challenges while remaining simple enough to produce measurable outcomes within controlled conditions, aligning with established guidance for synthetic tasks [18, 65]. We chose manual spatial scanning because the coordination challenge arises from unknown target *locations* distributed across a shared 3D space. Unlike filter-based approaches that bypass spatial exploration entirely, manual scanning requires pairs to decide how to divide spatial coverage, making attention awareness directly relevant to task success. This design parallels “needle in a haystack” scenarios common in exploratory data analysis [24, 49].

Using both continuous (Terrain) and discrete (Scatterplot) visualizations allowed us to test CAAV across different data representation types. The terrain provides natural landmarks and spatial continuity that may support coordination even without CAAV, while the uniform scatterplot lacks such affordances, allowing us to assess whether CAAV provides greater benefit in abstract environments.

5.5 Procedure

Participants provided informed consent, filled out a demographic form, and received a brief introduction to the system and task. The order of CAAV and No-CAAV conditions was counterbalanced across pairs. Before each condition, participants completed a brief training session to familiarize themselves with the XR environment, controls, and (for CAAV conditions) how to trigger and interpret attention visualizations.

Each pair completed the treasure hunt for both visualization types. After each visualization, participants completed a brief questionnaire about their experience and coordination strategies. Throughout the study, we recorded target counts, spatial coverage, attention overlap

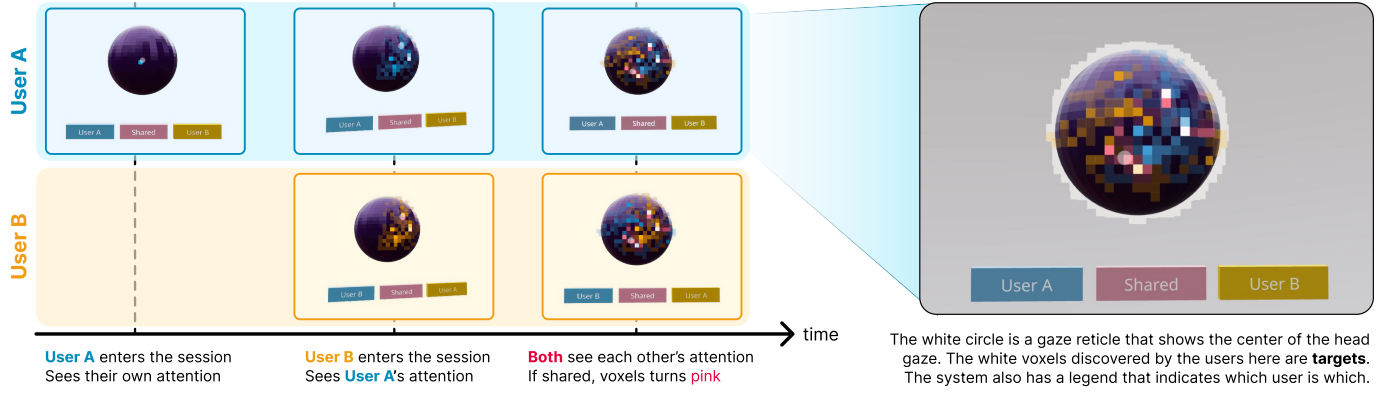


Figure 6: **HeedVision collaborative attention-aware visualization scenario.** (A) How attention visualization evolves: User A initially sees only their own attention; when User B joins, User A's accumulated attention becomes visible; as both discover targets, each sees both trails, with shared regions in pink. Entry order was counterbalanced across pairs; the sequential appearance in (A) is illustrative. (B) Interface detail: gaze reticle (white circle at head-gaze center), discovered targets (white voxels), and the color legend encoding user identity.

between participants, frequency of attention visualization triggering (CAAV condition), and verbal communication patterns.

After completing all conditions, participants were administered a System Usability Scale (SUS) questionnaire [13] and participated in a semi-structured interview about their experience.

6 RESULTS

We evaluated CAAVs to understand their impact on collaborative coordination during visual search tasks. Our findings reveal that attention visualization's effectiveness depends on visualization context, with different patterns emerging across abstract discrete environments (scatterplots) and continuous spatial environments (terrain). In scatterplot environments, pairs using CAAV showed improvements in coordination efficiency and reduced redundant exploration, while terrain environments revealed important boundary conditions for when attention visualization helps versus when existing environmental cues suffice.

6.1 Task Performance

We measured task performance through four complementary metrics. *Target Discovery* (percentage of targets found) captures task success directly. *Coverage* (percentage of space explored) reflects search thoroughness. *Coverage Efficiency* is the ratio of percent targets found to percent space explored; values above 1.0 indicate pairs found targets faster than uniform random exploration would predict, assuming roughly uniform target distribution, met by design for scatterplot but only approximate for terrain, where targets are constrained to the surface. *Gain Metrics* compare pair performance to individual baselines, revealing whether collaboration improved outcomes beyond individuals; target gain is computed as the pair's target count divided by the mean of the two individuals' target counts from the same condition, measuring the multiplicative benefit of working together.

Given our small sample ($N=8$ pairs), we supplement descriptive comparisons with non-parametric bootstrap confidence intervals (10,000 resamples) and Hedges' g effect sizes to contextualize effect magnitudes without assuming normality [18] (see Appendix 9.2). Pairs with CAAV showed improved performance across these measures. For scatterplot, CAAV pairs found 78.5% of targets compared to 62.4% without attention visualization ($\Delta = +16.1$ pp, BCa 95% CI $[-0.012, +0.293]$, $g = 1.12$; BCa just spans zero at $N=4$, treat cautiously), while covering 76.5% of the space versus 67.1%. In terrain, CAAV pairs found 26.0% of targets versus 19.6% without ($\Delta = +6.4$ pp, BCa CI spans zero), and covered 24.6% versus 18.8% of the space.

Coverage efficiency exceeded 1.0 with CAAV in scatterplot conditions (1.027) but dropped to 0.920 without attention visualization ($\Delta = +0.106$, BCa 95% CI $[-0.050, +0.185]$ spans zero; directional, not confirmed), suggesting more effective search behavior with visual feedback. Terrain conditions showed similar efficiency regardless of CAAV presence (0.999 with CAAV, 1.027 without; $\Delta = -0.028$, BCa

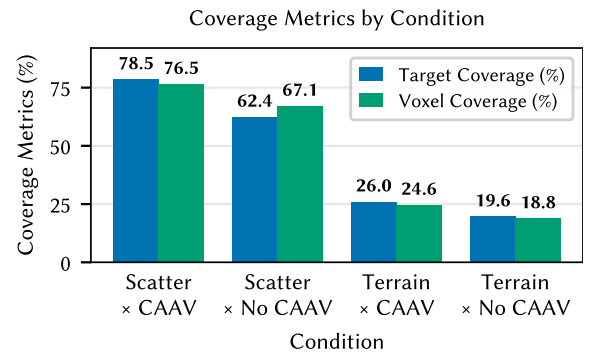


Figure 7: **Target and voxel coverage by condition.** Conditions are grouped by visualization type then by CAAV presence. CAAV conditions show higher target discovery rates, with scatterplot environments achieving substantially better coverage than terrain.

CI $[-0.182, +0.093]$, null), suggesting that terrain geometry itself provides coordination cues. The collaborative advantage was also evident in gain metrics: target gain was higher with CAAV in scatterplot conditions (1.33 vs. 1.19, $\Delta = +0.138$, CI spans zero), while terrain conditions showed collaboration benefits (1.46 with CAAV vs. 1.19 without). The terrain coordination gain is our strongest bootstrap-supported result ($\Delta = +0.277$, BCa 95% CI $[0.112, 0.450]$, $g = 1.67$, permutation $p = 0.029$; all four CAAV pairs outperformed the No-CAAV mean; both BCa and percentile methods agree).

6.2 Spatial Coverage and Coordination

To understand how pairs coordinated their exploration, we examined redundancy patterns alongside coverage. *Redundancy* measures duplicated exploration using normalized Shannon entropy: lower entropy implies more overlap, while higher entropy indicates diverse, non-redundant exploration.

Pairs using CAAV showed reduced redundancy, but patterns diverged by visualization type. For scatterplot tasks, normalized redundancy dropped from 11.3% (control) to 3.3% (CAAV); this represents approximately a 70% reduction ($\Delta = -0.058$, BCa CI $[-0.172, -0.002]$; BCa marginal, exact permutation $p = 0.257$; treat as directional), though we note that some redundancy may be desirable (e.g., deliberate verification of a partner's findings), so this reduction should be interpreted alongside strategy data rather than as an isolated positive outcome. For terrain tasks, however, CAAV *increased* redundancy from 4.0% to 14.9%. This reversal suggests that attention visualization interacts differently with environments that already contain natural coordination cues.

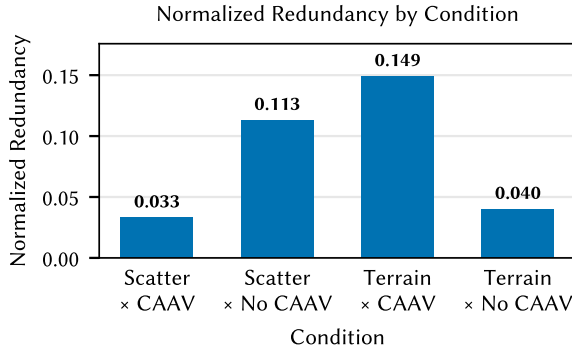


Figure 8: **Redundancy by condition.** Lower values indicate less overlap in exploration. CAAV substantially reduces redundancy in scatterplot environments but increases it in terrain, suggesting that attention visualization benefits depend on whether the environment already provides natural coordination cues.

Behavioral observations help explain this pattern. Participants in terrain + CAAV conditions often made deliberate second passes over explored areas, as if attention visualization made coverage gaps more salient and encouraged thoroughness over efficiency. Terrain landmarks (peaks, valleys) naturally guided both participants toward the same prominent features; attention visualization then highlighted rather than prevented this convergence.

Post-hoc correlation analysis supports this interpretation. We computed terrain gradient magnitude (steepness at prominent features like ridges, peaks, and valleys) across the surface, then correlated it with spatially-aggregated attention heatmaps from all terrain conditions. Terrain gradients correlate strongly with collaborative attention (Fig. 9): Pearson $r = 0.490$ and Spearman $\rho = 0.561$ (both $p < 0.0001$). Attention was also substantially more concentrated in terrain conditions (coefficient of variation: 2.231) than in scatter (1.030).

Together, these findings reveal the mechanism underlying terrain redundancy: prominent features act as natural landmarks that draw collaborative attention regardless of CAAV, reframing our terrain results not as CAAV failure but as evidence that attention visualization’s utility depends on the coordination affordances already present in the environment.

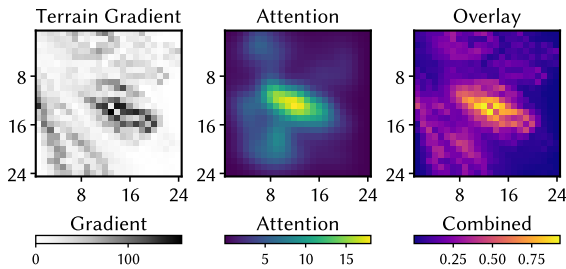


Figure 9: **Terrain gradient-attention correlation.** *Left:* gradient magnitude map (darker regions are steeper features that serve as natural landmarks). *Center:* aggregated attention heatmap from all terrain conditions. *Right:* overlay of the two. The strong correlation (Pearson $r = 0.490$, Spearman $\rho = 0.561$, both $p < 0.0001$) shows participants converged attention on prominent features (peaks, ridges, valleys), explaining the higher redundancy in terrain vs. scatterplot.

Coordination efficiency reflected intentional strategies. For scatterplot, participants described systematic division: “split into front/back faces for each of us” [P7, G4] and “checking the color of CAAV, which matches our initial plan” [P8, G4]. Without CAAV, pairs struggled and “had to double-check everything the other person looked at” [P7, G4].

6.3 Collaborative Strategies

Through systematic thematic analysis [11] of post-study questionnaires, exit interviews, and observed behaviors, we identified six collaborative strategies (Table 2). The analysis involved multiple coding passes: initial familiarization, systematic coding of coordination behaviors, pattern identification across conditions, and theme validation against quantitative behavioral data. CAAV conditions enabled more sophisticated coordination patterns (*spatial division*, i.e., dividing the environment into assigned regions; *systematic search*, i.e., methodical, structured exploration; and *visual feedback utilization*, i.e., leveraging voxel colors for awareness) while reducing reliance on *verbal coordination*. In abstract scatterplot environments, CAAV replaced verbal coordination entirely: P8 noted that “we have no verbal communication during the condition where we have [CAAV]” [P8, G4]. While reduced verbal load may reflect lower coordination cost, it may also indicate less joint discussion and weaker grounding [16]; the shift to visual-channel coordination warrants caution in interpretation. Without CAAV, pairs developed compensatory strategies such as *physical synchronization* (aligning markers in physical space). Terrain showed the richest strategy diversity regardless of CAAV presence, with pairs combining *landmark-based navigation* with other strategies. This explained why CAAV provided less incremental benefit where the environment itself already supported effective coordination through its spatial structure.

Table 2: **Collaborative strategies by experimental condition.** Scatterplot (S) vs. Terrain (T); No CAAV (no) vs. CAAV (yes). Numbers indicate participant counts per strategy (by pair); participants may appear in multiple categories as pairs often employed multiple strategies simultaneously. CAAV conditions supported systematic search and visual feedback, reducing verbal coordination, while non-CAAV conditions relied more on verbal and compensatory strategies in abstract environments. For exact quotes by each pair in each category, see Appendix 9.1.

Strategy	S/No	S/Yes	T/No	T/Yes
<i>Spatial Division</i>	2	4	3	2
<i>Landmark Navigation</i>	2	0	3	5
<i>Systematic Search</i>	0	4	1	2
<i>Verbal Coordination</i>	3	0	1	1
<i>Visual Feedback</i>	0	2	0	5
<i>Physical Synchronization</i>	2	0	0	0

6.4 Attention Visualization Usage

We logged toggle frequency and visibility duration (the toggle was available in CAAV conditions only). Most participants kept the attention visualization active for approximately 75.2% of session duration, and all participants used the toggle feature, indicating active management of the interface rather than passive acceptance (per-participant toggle timelines are provided in the supplementary materials).

Toggle behavior varied: while some rarely disabled the visualization, others toggled it frequently, though “off” periods remained brief, suggesting participants valued the awareness cues despite visual complexity. P6 explained: “it was nice to be able to tell what the other person had seen... but having it turned off made the interface less overwhelming” [P6, G3]. P9 was a notable exception, keeping the visualization inactive for the majority of their session; this individual difference, consistent with P6’s observation above, suggests tolerance for persistent attention overlays varies and warrants consideration in future implicit-triggering designs.

6.5 Subjective Assessment

Post-study ratings showed clear support for attention visualization: 93.8% preferred CAAV conditions, with 81.3% strongly agreeing that CAAV improved collaboration. Preferences split between *Scatter* + CAAV (50.0%) and *Terrain* + CAAV (43.8%), with awareness of partner actions as the most cited reason.

The mean SUS score was 74.38 (SD = 19.79), above the 68 benchmark for good usability [3, 13]; 81.3% of participants scored above

this threshold, with only one outlier (P1, who preferred the non-CAAV condition). NASA TLX ratings showed 24% lower overall workload with CAAV (2.79 vs. 3.69 on a 1–7 scale), with reductions in mental demand (2.50 vs. 3.81), frustration (33% lower), and temporal demand (24% lower). However, bootstrap analysis reveals a directional pattern of higher workload with CAAV in terrain conditions ($\Delta = +0.69$ on the 1–7 TLX scale, CI spans zero), consistent with qualitative reports of visual complexity; though this does not reach reliable statistical support, the direction aligns with the terrain redundancy increase and strategy observations. Overall, CAAV provided clear subjective coordination benefits in scatterplot conditions, even as terrain conditions introduced mixed workload effects.

6.6 Qualitative Experiences

Participants described CAAV as transforming collaboration from effortful coordination to intuitive teamwork (“*seeing my partner’s footsteps or breadcrumbs in real time*”), and frequently cited the color encoding as helpful: “*it was orange where they looked and got pink if I looked at it as well*” [P5, G3]. Communication shifted markedly: without CAAV, detailed spatial descriptions dominated; with CAAV it became strategic (“*you complete my area and I complete your area*” [P5, G3]), and several pairs reported no verbal communication at all. Many treated the task as collaborative gameplay, aiming for “*a nice blue side and a nice orange side with very little overlap*” [P6, G3]. Challenges centered on visual complexity over time, system latency (“*The delay is too much*” [P2, G1]), and occlusion of unexplored marks behind others’ attention color. Despite these concerns, overall perception was positive, with participants envisioning gentler decay: “*It would be nice if color of your attention would fade away [without noticing]*” [P8, G4].

7 DISCUSSION

Our study explored whether and how collaborative attention visualization affects coordination in immersive environments, rather than identifying an optimal CAAV configuration. CAAVs can improve spatial coordination, search efficiency, and task load distribution, though these benefits are strongly context-dependent.

7.1 Explaining the Results

CAAVs address group awareness [25] by making attention patterns visible, enabling what resembles “stigmergy” [23]: coordination through environmental cues rather than direct communication. The attention heatmaps function as “read wear” [30] for collaborative visualization: a form of social navigation [20] where users’ traces become explicit and actionable. This cognitive offloading, no longer needing to track a partner’s activities or verbally coordinate, let participants focus on the analytical task itself, and was especially effective in scatterplots, where pairs naturally distributed efforts.

This mechanism differed by visualization type: our gradient-attention correlation (Section 6.2) suggests prominent terrain features drew both collaborators’ attention regardless of CAAV presence, providing built-in stigmergic cues, so explicit visualization added redundant channels, consistent with natural partitioning in continuous spatial environments [21].

Importantly, our metrics (target discovery, redundancy, verbal communication, strategy shifts) are indirect proxies for group awareness rather than direct measures; we did not assess recall of partner activity, perceived awareness quality, or common ground, so future work should add partner-awareness ratings or recall tasks.

7.2 Generalizing the Results

Within the scope of our visual search evaluation, CAAVs are likely beneficial when the visualization is **abstract and discrete** (lacking landmarks), collaborators must **divide a large search space**, the task requires **exhaustive coverage**, and **verbal descriptions of locations are difficult**. Conversely, within this same scope, CAAVs may be less beneficial when the visualization has **inherent spatial structure**, or if the **visual complexity is already high**. These conditions suggest CAAVs would be particularly valuable for high-dimensional data exploration (t-SNE, UMAP projections), molecular structure analysis,

and 3D network visualization, though further evaluation with domain-specific tasks is needed to confirm this across the broader immersive analytics design space [53]. Our random scatterplot represents an edge case, as real-world dimensionally-reduced data typically exhibits clustering; we used random sampling to create a “pure” abstract environment without emergent structure.

All results derive from synchronous, co-located pairs, so generalizations to other settings should be made with caution. Still, distributed collaboration could plausibly benefit even more, as attention visualization supplies awareness otherwise absent, and asynchronous teams could leverage attention decay for historical handoffs [54].

7.3 Limitations and Future Work

Several technical constraints bound our implementation: the voxel-based approach becomes expensive at high resolutions or with many users, though adaptive resolution or GPU-accelerated tracking could help. Head orientation is a coarser attention proxy than eye tracking; we expect our findings would transfer to eye-tracking with potentially stronger effects. Our evaluation scope also constrains generalizability: we built functional prototypes for all the remaining three quadrants, but formally evaluated only HEEDVISION’s world-space embedded configuration. These proof-of-concepts also use BroadcastChannel local synchronization rather than cloud infrastructure and run VR-only; maturing them for networked, co-located AR remains future work. Our single task type, two-condition, pairs-only design cannot capture the full complexity of collaborative analysis; future work should examine diverse tasks, larger teams, and selective sharing mechanisms that let collaborators choose when their attention is revealed.

A further limitation is our single baseline: comparing CAAV against no attention visualization does not isolate the contribution of accumulated attention *history*. Alternative baselines, including shared gaze rays, pointer trails, manual markers, or minimap overlays, might yield similar benefits, and future work should include at least one such comparison to disentangle attention persistence from other forms of shared awareness. Finally, future CAAVs could incorporate *semantic* attention, extending RQ2 from spatial to data-element-level recording: this would require a data-aware architecture and could enable displays such as “most-attended data clusters” overlays for collaborative sensemaking.

8 CONCLUSION

We introduced Collaborative Attention-Aware Visualizations (CAAVs), extending attention awareness from individual to multi-user immersive analytics. We answer our three questions at two levels: the design space enumerates the *answer space* for capturing (RQ1), recording (RQ2), and revisualizing (RQ3) collective attention, while HEEDVISION supplies one validated point-answer per question. It instantiates the world-space embedded configuration, evaluated with eight co-located pairs in augmented reality, with benefits our study bounds to the contexts where they hold.

Our exploratory study provides initial evidence that CAAV effectiveness depends critically on visualization context, specifically in spatial visual search tasks; whether this extends to richer analytical activities such as hypothesis generation or sensemaking remains open. In abstract, discrete environments such as scatterplots, CAAVs reduced redundant exploration and enabled emergent coordination strategies, while terrain environments yielded a large coordination gain ($\Delta = +0.277$, BCa CI [0.112, 0.450], $g = 1.67$, $p = 0.029$, our strongest result) even as attention visualization sometimes interfered with natural coordination cues, suggesting that CAAV design must carefully consider the affordances of the visualization space itself.

We offer the CAAV design space as a preliminary generative framework rather than a prescription; its remaining configurations await systematic evaluation. As immersive analytics moves toward multi-user settings, attention-aware approaches offer a promising path to collaborative sensemaking in spatial workspaces.

SUPPLEMENTAL MATERIALS

All supplemental materials are available on OSF at osf.io/frsp8, released under a CC BY 4.0 license. They comprise analysis files for

the task performance, NASA TLX, SUS, and qualitative analyses; the collected experimental data, including spatial attention logs, toggle event streams, and survey responses; study materials; and source code for HeedVision and the three proof-of-concept implementations.

ACKNOWLEDGMENTS

This work was supported by Villum Investigator grant VL-54492 by Villum Fonden. Any opinions, findings, and conclusions expressed in this material are those of the authors and do not necessarily reflect the views of the funding agency. The three proof-of-concept implementations were developed with assistance from Claude Sonnet 4.5, which was used for debugging and documentation support.

REFERENCES

- [1] S. K. Badam, Z. Zeng, E. Wall, A. Endert, and N. Elmqvist. Supporting team-first visual analytics through group activity representations. In *Proceedings of the Graphics Interface Conference*, pp. 208–213. ACM, New York, NY, USA, 2017. doi: [10.20380/GI2017.26](https://doi.org/10.20380/GI2017.26) 2
- [2] A. D. Balakrishnan, S. R. Fussell, and S. Kiesler. Do visualizations improve synchronous remote collaboration? In *Proceedings of the ACM Conference on Human Factors in Computing Systems*, pp. 1227–1236. ACM, New York, NY, USA, 2008. doi: [10.1145/1357054.1357246](https://doi.org/10.1145/1357054.1357246) 2
- [3] A. Bangor, P. Kortum, and J. Miller. Determining what individual sus scores mean: adding an adjective rating scale. *J. Usability Studies*, 4(3):114–123, May 2009. 8
- [4] P. Baudisch, D. Tan, M. Collomb, D. Robbins, K. Hinckley, M. Agrawala et al. Phosphor: explaining transitions in the user interface using afterglow effects. In *Proceedings of the ACM Symposium on User Interface Software and Technology*, pp. 169–178. ACM, New York, NY, USA, 2006. doi: [10.1145/1166253.1166280](https://doi.org/10.1145/1166253.1166280) 2
- [5] M. Billingham, M. Cordeil, A. Bezerianos, and T. Margolis. Collaborative immersive analytics. In *Immersive Analytics*, vol. 11190 of *Lecture Notes in Computer Science*, pp. 221–257. Springer, New York, NY, USA, 2018. doi: [10.1007/978-3-030-01388-2_8](https://doi.org/10.1007/978-3-030-01388-2_8) 1
- [6] T. Blaschek, M. John, S. Koch, L. Bruder, and T. Ertl. Triangulating user behavior using eye movement, interaction, and think aloud data. In *Proceedings of the ACM Symposium on Eye Tracking Research & Applications*, pp. 175–182. ACM, New York, NY, USA, 2016. doi: [10.1145/2857491.2857523](https://doi.org/10.1145/2857491.2857523) 2
- [7] R. A. Bolt. Gaze-orchestrated dynamic windows. In *Proceedings of the ACM Conference on Computer Graphics and Interactive Techniques*, pp. 109–119. ACM, New York, NY, USA, 1981. doi: [10.1145/800224.806796](https://doi.org/10.1145/800224.806796) 2
- [8] M. Borowski, P. W. S. Butcher, J. B. Kristensen, J. O. Petersen, P. D. Ritsos, C. N. Klokmoose et al. DashSpace: A live collaborative platform for immersive and ubiquitous analytics. *IEEE Transactions on Visualization and Computer Graphics*, 99(PP):1–13, 2025. doi: [10.1109/TVCG.2025.3537679](https://doi.org/10.1109/TVCG.2025.3537679) 5
- [9] M. Borowski, J. E. S. Grønbaek, P. W. S. Butcher, P. D. Ritsos, C. N. Klokmoose, and N. Elmqvist. Spatialstrates: Cross-reality collaboration through spatial hypermedia. In *Proceedings of the ACM Symposium on User Interface Software and Technology*, pp. 187:1–187:14. ACM, New York, NY, USA, 2025. doi: [10.1145/3746059.3747708](https://doi.org/10.1145/3746059.3747708) 5
- [10] M. Borowski, L. Murray, R. Bagge, J. B. Kristensen, A. Satyanarayan, and C. N. Klokmoose. Varv: Reprogrammable interactive software as a declarative data structure. In *Proceedings of the ACM Conference on Human Factors in Computing Systems*, pp. 492:1–492:20. ACM, New York, NY, USA, 2022. doi: [10.1145/3491102.3502064](https://doi.org/10.1145/3491102.3502064) 5
- [11] V. Braun and V. Clarke. Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3(2):77–101, 2006. doi: [10.1191/1478088706qp0630a](https://doi.org/10.1191/1478088706qp0630a) 8
- [12] M. Brehmer and T. Munzner. A multi-level typology of abstract visualization tasks. *IEEE Transactions on Visualization and Computer Graphics*, 19(12):2376–2385, 2013. doi: [10.1109/TVCG.2013.124](https://doi.org/10.1109/TVCG.2013.124) 3
- [13] J. Brooke et al. SUS – a quick and dirty usability scale. *Usability Evaluation in Industry*, 189(194):4–7, 1996. doi: [10.1201/9781498710411-35](https://doi.org/10.1201/9781498710411-35) 7, 8
- [14] M. Burch. *Eye Tracking and Visual Analytics*. River Publishers, Gistrup, Denmark, 2021. doi: [10.1201/9781003338161](https://doi.org/10.1201/9781003338161) 2
- [15] M. Burch, N. Konevtsova, J. Heinrich, M. Hoeferlin, and D. Weiskopf. Evaluation of traditional, orthogonal, and radial tree diagrams by an eye tracking study. *IEEE Transactions on Visualization and Computer Graphics*, 17(12):2440–2448, 2011. doi: [10.1109/TVCG.2011.193](https://doi.org/10.1109/TVCG.2011.193) 2
- [16] H. H. Clark and S. E. Brennan. Grounding in communication. In L. B. Resnick, J. M. Levine, and S. D. Teasley, eds., *Perspectives on Socially Shared Cognition*, pp. 127–149. American Psychological Association, Washington, D.C., USA, 1991. doi: [10.1037/10096-006](https://doi.org/10.1037/10096-006) 2, 8
- [17] J. H. Clark. Hierarchical geometric models for visible surface algorithms. *Communications of the ACM*, 19(10):547–554, Oct. 1976. doi: [10.1145/360349.360354](https://doi.org/10.1145/360349.360354) 4
- [18] A. Cockburn, C. Gutwin, and J. Alexander. Faster document navigation with space-filling thumbnails. In *Proceedings of the ACM Conference on Human Factors in Computing Systems*, pp. 1–10. ACM, New York, NY, USA, 2006. doi: [10.1145/1124772.1124774](https://doi.org/10.1145/1124772.1124774) 6, 7
- [19] E. J. David, J. Gutiérrez, A. Coutrot, M. P. Da Silva, and P. Le Callet. A dataset of head and eye movements for 360° videos. In *Proceedings of the ACM Conference on Multimedia Systems*, pp. 432–437, 2018. doi: [10.1145/3204949.3208139](https://doi.org/10.1145/3204949.3208139) 2, 4
- [20] P. Dourish and M. Chalmers. Running out of space: Models of information navigation. In *Proceedings of the British Conference on Human Computer Interaction*, vol. 94, pp. 23–26. ACM, New York, NY, USA, 1994. 2, 9
- [21] N. Elmqvist, M. E. Tudoreanu, and P. Tsigas. Evaluating motion constraints for 3D wayfinding in immersive and desktop virtual environments. In *Proceedings of the ACM Conference on Human Factors in Computing Systems*, pp. 1769–1778. ACM, New York, NY, USA, 2008. doi: [10.1145/1357054.1357330](https://doi.org/10.1145/1357054.1357330) 9
- [22] B. Ens, B. Bach, M. Cordeil, U. Engelke, M. Serrano, W. Willett et al. Grand challenges in immersive analytics. In *Proceedings of the ACM Conference on Human Factors in Computing Systems*, art. no. 459, 17 pp. ACM, New York, NY, USA, 2021. doi: [10.1145/3411764.3446866](https://doi.org/10.1145/3411764.3446866) 1, 4
- [23] P.-P. Grassé. La reconstruction du nid et les coordinations interindividuelles chez *Bellicositermes natalensis* et *Cubitermes* sp. la théorie de la stigmergie: essai d'interprétation du comportement des termites constructeurs. *Insectes Sociaux*, 6(1):41–80, 1959. doi: [10.1007/BF02223791](https://doi.org/10.1007/BF02223791) 9
- [24] G. Grinstein, D. Keim, and M. Ward. Information visualization, visual data mining, and its application to drug design. *IEEE Visualization Course #1 Notes*, Oct. 2002. 6
- [25] C. Gutwin and S. Greenberg. Design for individuals, design for groups: Tradeoffs between power and workspace awareness. In *Proceedings of the ACM Conference on Computer-Supported Cooperative Work and Social Computing*, pp. 207–216. ACM, New York, NY, USA, 1998. doi: [10.1145/289444.289495](https://doi.org/10.1145/289444.289495) 1, 2, 4, 9
- [26] C. Gutwin and S. Greenberg. Effects of awareness support on groupware usability. In *Proceedings of the ACM Conference on Human Factors in Computing Systems*, pp. 511–518. ACM, New York, NY, USA, 1998. doi: [10.1145/274644.274713](https://doi.org/10.1145/274644.274713) 1, 2
- [27] C. Gutwin, M. Roseman, and S. Greenberg. A usability study of awareness widgets in a shared workspace groupware system. In *Proceedings of the ACM Conference on Computer-Supported Cooperative Work and Social Computing*, pp. 258–267. ACM, New York, NY, USA, 1996. doi: [10.1145/240080.240298](https://doi.org/10.1145/240080.240298) 2
- [28] A. H. Hajizadeh, M. Tory, and R. Leung. Supporting awareness through collaborative brushing and linking of tabular data. *IEEE Transactions on Visualization and Computer Graphics*, 19(12):2189–2197, 2013. doi: [10.1109/TVCG.2013.197](https://doi.org/10.1109/TVCG.2013.197) 2
- [29] C. G. Healey, K. S. Booth, and J. T. Enns. High-speed visual estimation using preattentive processing. *ACM Transactions on Computer-Human Interaction*, 6(2):107–135, 1999. doi: [10.1145/230562.230563](https://doi.org/10.1145/230562.230563) 2
- [30] W. C. Hill, J. D. Hollan, D. Wroblewski, and T. McCandless. Edit wear and read wear. In *Proceedings of the ACM Conference on Human Factors in Computing Systems*, pp. 3–9. ACM, New York, NY, USA, 1992. doi: [10.1145/142750.142751](https://doi.org/10.1145/142750.142751) 9
- [31] P. Isenberg, N. Elmqvist, J. Scholtz, D. Cernea, K.-L. Ma, and H. Hagen. Collaborative visualization: definition, challenges, and research agenda. *Information Visualization*, 10(4):310–326, 2011. doi: [10.1177/1473871611412817](https://doi.org/10.1177/1473871611412817) 1, 4
- [32] P. Isenberg and D. Fisher. Collaborative brushing and linking for co-located visual analytics of document collections. *Computer Graphics Forum*, 28(3):1031–1038, 2009. doi: [10.1111/j.1467-8659.2009.01444.x](https://doi.org/10.1111/j.1467-8659.2009.01444.x) 2
- [33] R. J. K. Jacob. What you look at is what you get: eye movement-based interaction techniques. In *Proceedings of the ACM Conference on Human Factors in Computing Systems*, pp. 11–18. ACM, New York, NY, USA,

1990. doi: [10.1145/97243.97246](https://doi.org/10.1145/97243.97246) 2, 4
- [34] H. Jänicke and M. Chen. A salience-based quality metric for visualization. *Computer Graphics Forum*, 29(3):1183–1192, 2010. doi: [10.1111/J.1467-8659.2009.01667.X](https://doi.org/10.1111/J.1467-8659.2009.01667.X) 2
- [35] R. Jianu, N. Silva, N. Rodrigues, T. Blascheck, T. Schreck, and D. Weiskopf. Gaze-aware visualisation: Design considerations and research agenda. *Computer Graphics Forum*, 44(6):e70097, 2025. doi: [10.1111/cgf.70097](https://doi.org/10.1111/cgf.70097) 1, 2
- [36] K. Kim, W. Javed, C. Williams, N. Elmqvist, and P. Irani. Hugin: a framework for awareness and coordination in mixed-presence collaborative information visualization. In *Proceedings of the ACM Conference on Interactive Tabletops and Surfaces*, pp. 231–240. ACM, New York, NY, USA, 2010. doi: [10.1145/1936652.1936694](https://doi.org/10.1145/1936652.1936694) 2
- [37] S.-H. Kim, Z. Dong, H. Xian, B. Upatasing, and J. S. Yi. Does an eye tracker tell the truth about visualizations?: Findings while investigating visualizations for decision making. *IEEE Transactions on Visualization and Computer Graphics*, 18(12):2421–2430, 2012. doi: [10.1109/TVCG.2012.215](https://doi.org/10.1109/TVCG.2012.215) 2
- [38] C. N. Klokmoose, J. R. Eagan, S. Baader, W. Mackay, and M. Beaudouin-Lafon. Webstrates: Shareable dynamic media. In *Proceedings of the ACM Symposium on User Interface Software and Technology*, pp. 280–290. ACM, New York, NY, USA, 2015. doi: [10.1145/2807442.2807446](https://doi.org/10.1145/2807442.2807446) 5
- [39] M. Koch, T. Rau, V. Mikheev, S. Öney, M. Becher, X. Wang et al. Group gaze-sharing with projection displays. In *Proceedings of the ACM Symposium on Eye Tracking Research and Applications*, art. no. 92, 7 pp. ACM, New York, NY, USA, 2025. doi: [10.1145/3715669.3725871](https://doi.org/10.1145/3715669.3725871) 1, 2
- [40] K. Kurzhals, B. Fisher, M. Burch, and D. Weiskopf. Eye tracking evaluation of visual analytics. *Information Visualization*, 15(4):340–358, 2016. doi: [10.1177/1473871615609787](https://doi.org/10.1177/1473871615609787) 2
- [41] W. Lu, H. B.-L. Duh, S. Feiner, and Q. Zhao. Attributes of subtle cues for facilitating visual search in augmented reality. *IEEE Transactions on Visualization and Computer Graphics*, 20(3):404–412, 2014. doi: [10.1109/TVCG.2013.241](https://doi.org/10.1109/TVCG.2013.241) 2
- [42] E. Marguet, M. Borowski, J. B. Kristensen, C. N. Klokmoose, and N. Elmqvist. WindowSpace: A web-based XR window manager for interacting with 2D windows in immersive 3D space. *Proceedings of the ACM on Human-Computer Interaction*, 9(8), art. no. ISS006, 26 pp., Nov. 2025. doi: [10.1145/3773063](https://doi.org/10.1145/3773063) 6
- [43] G. Mark and A. Kobsa. The effects of collaboration and system transparency on CIVE usage: an empirical study and model. *Presence: Teleoperators & Virtual Environments*, 14(1):60–80, 2005. doi: [10.1162/1054746053890279](https://doi.org/10.1162/1054746053890279) 2
- [44] K. Marriott, F. Schreiber, T. Dwyer, K. Klein, N. H. Riche, T. Itoh et al., eds. *Immersive Analytics*, vol. 11190 of *Lecture Notes in Computer Science*. Springer, New York, NY, USA, 2018. doi: [10.1007/978-3-030-01388-2](https://doi.org/10.1007/978-3-030-01388-2) 1
- [45] J. Matejka, T. Grossman, and G. Fitzmaurice. Patina: dynamic heatmaps for visualizing application usage. In *Proceedings of the ACM Conference on Human Factors in Computing Systems*, pp. 3227–3236. ACM, New York, NY, USA, 2013. doi: [10.1145/2470654.2466442](https://doi.org/10.1145/2470654.2466442) 2
- [46] L. E. Matzen, M. J. Haass, K. M. Divis, Z. Wang, and A. T. Wilson. Data visualization saliency model: A tool for evaluating abstract data visualizations. *IEEE Transactions on Visualization and Computer Graphics*, 24(1):563–573, 2018. doi: [10.1109/TVCG.2017.2743939](https://doi.org/10.1109/TVCG.2017.2743939) 2
- [47] K. Pfeuffer, J. Alexander, M. K. Chong, and H. Gellersen. Gaze-touch: combining gaze with multi-touch for interaction on the same surface. In *Proceedings of the ACM Symposium on User Interface Software and Technology*, pp. 509–518. ACM, New York, NY, USA, 2014. doi: [10.1145/2642918.2647397](https://doi.org/10.1145/2642918.2647397) 2
- [48] K. Pfeuffer, B. Mayer, D. Mardanbegi, and H. Gellersen. Gaze + pinch interaction in virtual reality. In *Proceedings of the ACM Symposium on Spatial User Interaction*, pp. 99–108. ACM, New York, NY, USA, 2017. doi: [10.1145/3131277.3132180](https://doi.org/10.1145/3131277.3132180) 2
- [49] P. Pirolli and S. Card. The sensemaking process and leverage points for analyst technology as identified through cognitive task analysis. In *Proceedings of the International Conference on Intelligence Analysis*, vol. 5, pp. 2–4. PARC, 3333 Coyote Hill Road, Palo Alto, California, 94304, USA, 2005. 6
- [50] Poimandres. *React Three Fiber*. [Online]. Available <https://docs.pmnd.rs/react-three-fiber/getting-started/introduction>, 2017. Accessed: 07/11/22. 5
- [51] A. Prouzeau, A. Bezerianos, and O. Chapuis. Trade-offs between a vertical shared display and two desktops in a collaborative path-finding task. In *Proceedings of the Graphics Interface Conference*, pp. 214–219. Canadian Human-Computer Communications Society, Waterloo, CAN, 2017. 6
- [52] D. Saffo, A. Batch, C. Dunne, and N. Elmqvist. Through their eyes and in their shoes: Providing group awareness during collaboration across virtual reality and desktop platforms. In *Proceedings of the ACM Conference on Human Factors in Computing Systems*, pp. 383:1–383:15. ACM, New York, NY, USA, 2023. doi: [10.1145/3544548.3581093](https://doi.org/10.1145/3544548.3581093) 1, 2
- [53] D. Saffo, S. Di Bartolomeo, C. Yeh, and C. Dunne. Unraveling the design space of immersive analytics: A systematic review. *IEEE Transactions on Visualization and Computer Graphics*, 30(1):495–506, 2024. doi: [10.1109/TVCG.2023.3327368](https://doi.org/10.1109/TVCG.2023.3327368) 9
- [54] A. Sarvghad and M. Tory. Exploiting analysis history to support collaborative data analysis. In *Proceedings of the Graphics Interface Conference*, pp. 123–130. Canadian Human-Computer Communications Society, Waterloo, ON, Canada, 2015. doi: [10.20380/GI2015.16](https://doi.org/10.20380/GI2015.16) 2, 4, 9
- [55] A. Sarvghad, M. Tory, and N. Mahyar. Visualizing dimension coverage to support exploratory analysis. *IEEE Transactions on Visualization and Computer Graphics*, 23(1):21–30, 2017. doi: [10.1109/TVCG.2016.2598466](https://doi.org/10.1109/TVCG.2016.2598466) 2, 4
- [56] Y. B. Shrinivasan and J. van Wijk. Supporting exploration awareness in information visualization. *IEEE Computer Graphics and Applications*, 29(5):34–43, 2009. doi: [10.1109/MCG.2009.87](https://doi.org/10.1109/MCG.2009.87) 2
- [57] V. Sitzmann, A. Serrano, A. Pavel, M. Agrawala, D. Gutierrez, B. Masia et al. Saliency in VR: How do people explore virtual environments? *IEEE Transactions on Visualization and Computer Graphics*, 24(4):1633–1642, 2018. doi: [10.1109/TVCG.2018.2793599](https://doi.org/10.1109/TVCG.2018.2793599) 2, 4
- [58] A. Srinivasan, J. Ellemose, P. W. S. Butcher, P. D. Ritsos, and N. Elmqvist. Attention-aware visualization: Tracking and responding to user perception over time. *IEEE Transactions on Visualization and Computer Graphics*, 31(1):1017–1027, 2025. doi: [10.1109/TVCG.2024.3456300](https://doi.org/10.1109/TVCG.2024.3456300) 1, 2, 3, 4
- [59] S. Stellmach and R. Dachsel. Look & touch: gaze-supported target acquisition. In *Proceedings of the ACM Conference on Human Factors in Computing Systems*, pp. 2981–2990. ACM, New York, NY, USA, 2012. doi: [10.1145/2207676.2208709](https://doi.org/10.1145/2207676.2208709) 2
- [60] S. Stellmach, S. Stober, A. Nürnberger, and R. Dachsel. Designing gaze-supported multimodal interactions for the exploration of large image collections. In *Proceedings of the Conference on Novel Gaze-Controlled Applications*, pp. 1:1–1:8. ACM, New York, NY, USA, 2011. doi: [10.1145/1983302.1983303](https://doi.org/10.1145/1983302.1983303) 2
- [61] A. Tang, C. Neustaedter, and S. Greenberg. VideoArms: embodiments for mixed presence groupware. In *Proceedings of the British Human-Computer Interaction Conference*, vol. 20, pp. 85–102. ACM, New York, NY, USA, 2006. doi: [10.1007/978-1-84628-664-3_8](https://doi.org/10.1007/978-1-84628-664-3_8) 2
- [62] A. M. Treisman and G. Gelade. A feature-integration theory of attention. *Cognitive Psychology*, 12(1):97–136, 1980. doi: [10.1016/0010-0285\(80\)90005-5](https://doi.org/10.1016/0010-0285(80)90005-5) 1, 2
- [63] J. Turner, J. Alexander, A. Bulling, D. Schmidt, and H. Gellersen. Eye pull, eye push: Moving objects between large screens and personal devices with gaze and touch. In *Proceedings of the IFIP Conference on Human-Computer Interaction*, vol. 8118 of *Lecture Notes in Computer Science*, pp. 170–186. Springer, New York, NY, USA, 2013. doi: [10.1007/978-3-642-40480-1_11](https://doi.org/10.1007/978-3-642-40480-1_11) 2
- [64] W. Willett, J. Heer, and M. Agrawala. Scented widgets: Improving navigation cues with embedded visualizations. *IEEE Transactions on Visualization and Computer Graphics*, 13(6):1129–1136, 2007. doi: [10.1109/TVCG.2007.70589](https://doi.org/10.1109/TVCG.2007.70589) 2
- [65] Y. L. Wong, J. Zhao, and N. Elmqvist. Evaluating social navigation visualization in online geographic maps. *International Journal of Human-Computer Interaction*, 31(2):118–127, 2015. doi: [10.1080/10447318.2014.959106](https://doi.org/10.1080/10447318.2014.959106) 6

9 APPENDIX

9.1 Collaboration Strategy Quotes

Below we present representative participant quotes organized by collaboration strategy and condition. Conditions are abbreviated as **S+C** (Scatter+CAAV), **S** (Scatter+No CAAV), **T+C** (Terrain+CAAV), and **T** (Terrain+No CAAV).

9.1.1 Spatial Division

S+C:

[G2-P3] “We cut the space into top and bottom, and one search for top another search for the bottom.”

[G4-P7] “Split into front/back faces for each of us. Then go from left/right faces.”

[G4-P8] “Separate the task space to have everyone working on different regions.”

[G6-P11] “We split the environment into two halves and each took responsibility for one half.”

S:

[G1-P1] “I was looking at the points within the area agreed upon with my partner (in a ‘I go left, you go right’ manner).”

[G8-P16] “We attempted to divide the Scatter space geometrically.”

T+C:

[G3-P5] “We assigned each other two different regions, back and front, and then top and bottom.”

[G3-P6] “We split the map down the middle, using the geometry as a guide.”

T:

[G2-P3] “Try to separate the area first and then, once I feel like seeing everything, I start working on my partner’s area.”

[G5-P9] “We decided to split the map into different region and each one of us would be responsible for a specific region.”

[G5-P10] “We split the environment equally between us. I searched the back half of the environment.”

9.1.2 Landmark-based Navigation

S:

[G1-P1] “I shared its location with them using relative location of the point to landmark I could identify.”

[G8-P16] “Describe found targets using relative positions to clusters.”

T+C:

[G1-P1] “I was trying to explain positions of the target I had found by relating it to the closest landmarks.”

[G7-P13] “We used the terrain features to divide the environment.”

[G7-P14] “We divided the terrain using natural landmarks.”

[G8-P15] “We used terrain landmarks to divide the space.”

[G8-P16] “Divided terrain using natural features.”

T:

[G4-P7] “We split the graph based on terrain feature.”

[G5-P10] “We used landmarks and described roughly in which area we found the targets.”

[G6-P11] “We tried to divide the terrain using landmarks and verbally described locations.”

9.1.3 Systematic Search

S+C:

[G4-P7] “At last we just randomly check the other CAAV left unexplored.”

[G4-P8] “After first round search, double check the same designated region.”

[G6-P11] “I focused on systematically exploring my assigned region while keeping track of found targets.”

[G6-P12] “I took one half and systematically searched through it.”

T+C:

[G3-P6] “When we found we were looking at the same place one of us moved away.”

[G7-P14] “I worked systematically through my assigned area.”

T:

[G4-P8] “Swap position (relative to the visualization) to double check.”

9.1.4 Verbal Coordination

S:

[G1-P2] “My partner and I gave each other verbal hints as we were in the same physical room.”

[G7-P13] “We tried to coordinate verbally by describing positions relative to the Scatter structure.”

[G8-P15] “We tried to coordinate by describing positions in the Scatter space.”

T+C:

[G1-P2] “In addition to verbal communication we were also using our gaze.”

T:

[G6-P11] “We tried to divide the terrain using landmarks and verbally described locations.”

9.1.5 Visual Feedback Utilization

S+C:

[G2-P3] “The color is easy to distinguish where need to be double check.”

[G6-P12] “The visualization helped us avoid overlapping searches.”

T+C:

[G1-P2] “We were also using our gaze and the highlights that they created to describe positions.”

[G7-P13] “The voxel visualization to track our progress.”

[G7-P14] “Used the voxel system to avoid redundant searches.”

[G8-P15] “Relied on the voxel visualization to track our collective progress.”

[G8-P16] “Used voxel feedback to optimize our search patterns.”

9.1.6 Physical Synchronization

S:

[G3-P5] “We aligned our markers in the physical space, aiming to have the same virtual 3D space.”

[G3-P6] “We placed the Scatter in the same location in the room using the yellow triangle.”

9.2 Bootstrap Analysis

This appendix reproduces the analysis introduced in Section 6: the key findings and methods, the per-geometry bootstrap results (Tables 3 and 4), and the effect-size plots (Figures 10–13). All estimates use $B = 10,000$ BCa bootstrap resamples (seed = 42) and exact permutation tests (all $\binom{8}{4} = 70$ reassignments enumerated), with $N = 4$ pairs per cell.

Key findings.

- **Terrain coordination gain.** CAAV pairs covered 27.7% more ground relative to their stronger individual member ($\Delta = +0.277$, BCa 95% CI [0.112, 0.450], $g = 1.67$ large, exact- $p = 0.029$). This is the only effect whose interval clears zero.
- **Scatter target discovery.** CAAV pairs found 16.1% more targets ($\Delta = +0.161$, BCa 95% CI [-0.012, 0.293], $g = 1.12$ large, exact- $p = 0.114$), suggestive, not confirmed.
- **SUS.** $M = 65$, BCa 95% CI [54, 72], within the “OK” band (51–68).

Statistical methods. *BCa bootstrap confidence intervals* (Efron, 1987): $B = 10,000$ replicates, seed = 42, correcting for both bias (z_0) and skewness/acceleration (a , estimated via jackknife) in the bootstrap distribution. *Exact permutation test:* with $N = 4$ per cell, all $\binom{8}{4} = 70$ reassignments are enumerated exhaustively, giving exact p -values with no Monte Carlo variance (the minimum achievable two-tailed p is $2/70 \approx 0.029$). *Hedges’ g :* small-sample corrected Cohen’s d ($J = 1 - 3/(4df - 1)$, $df = 6$). Both bootstrap and exact permutation p are reported; conclusions follow the more conservative value. Each of the 8 pairs completed exactly one CAAV and one No-CAAV condition, so within-geometry comparisons are between-pairs ($n = 4$ per cell).

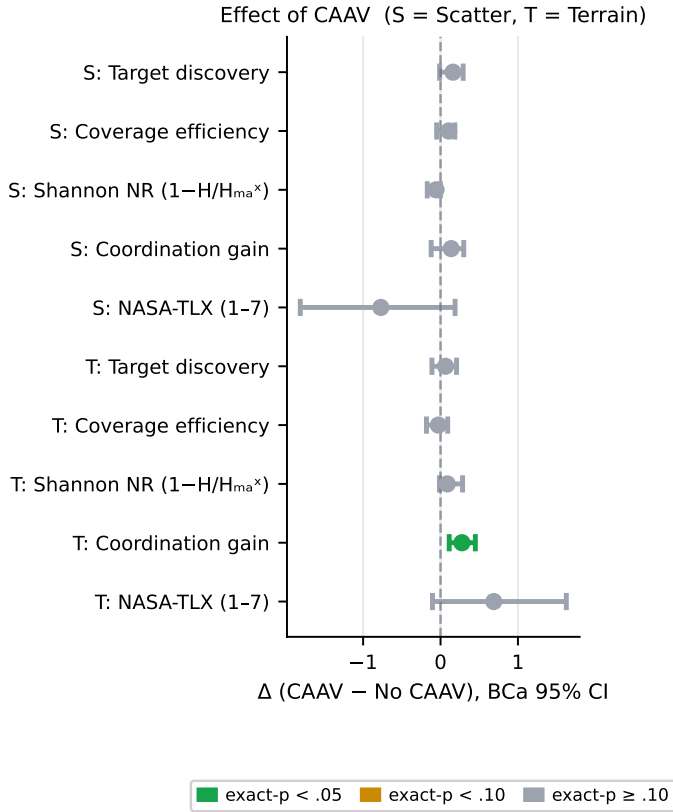


Figure 10: **Effect of CAAV.** Forest plot of Δ (CAAV – No CAAV) with BCa 95% CIs for the primary metrics. Colour encodes the exact permutation p -value: green $p < 0.05$, amber $p < 0.10$, grey $p \geq 0.10$. The dashed line marks $\Delta = 0$; only the terrain coordination gain clears it.

Group means with BCa 95% CIs

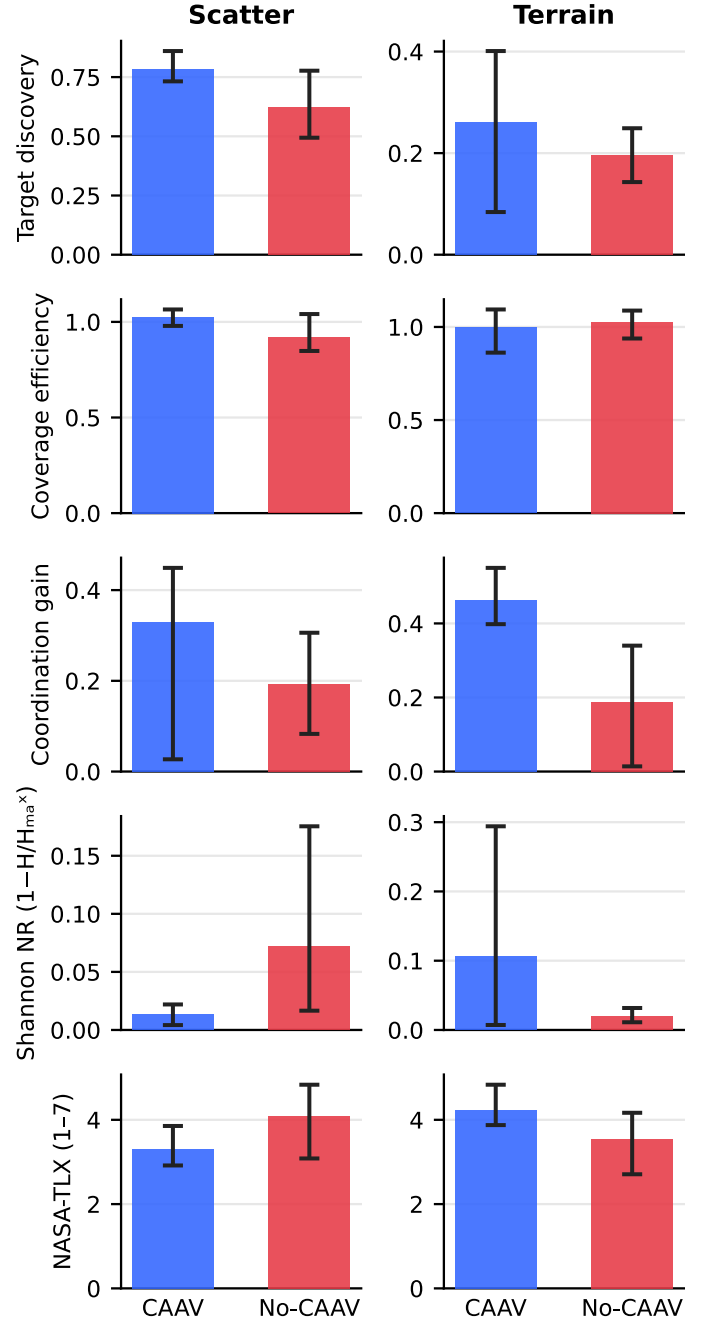


Figure 11: **Group means.** Condition means with BCa 95% CIs for the primary metrics (blue: CAAV, red: No CAAV).

Table 3: **Scatter conditions: bootstrap results** (first results table of the analysis report; $N = 4$ per cell). M = mean; BCa CI = bias-corrected-and-accelerated 95% confidence interval; Δ = CAAV – No CAAV; g = Hedges' g . Significance badges follow exact- p : ** $p < 0.01$, * $p < 0.05$, † $p < 0.10$, ns $p \geq 0.10$. $CI \cap 0$ marks whether the Δ interval contains zero.

Sig	Metric	CAAV M [BCa CI]	No CAAV M [BCa CI]	Δ [BCa CI]	g	boot- p	exact- p	$CI \cap 0$
ns	Target discovery (% found)	0.785 [0.732, 0.860]	0.624 [0.494, 0.777]	+0.161 [−0.012, 0.293]	1.12 (large)	0.022*	0.114	Yes
ns	Coverage efficiency	1.025 [0.979, 1.065]	0.919 [0.848, 1.041]	+0.106 [−0.050, 0.185]	1.01 (large)	0.083	0.143	Yes
ns	Redundancy: Shannon NR	0.014 [0.003, 0.021]	0.071 [0.018, 0.175]	−0.058 [−0.172, −0.002]	−0.69 (med.)	0.038*	0.257	No
ns	Coordination gain (over best indiv.)	0.330 [0.027, 0.449]	0.193 [0.083, 0.306]	+0.138 [−0.121, 0.300]	0.70 (med.)	0.102	0.314	Yes
ns	NASA-TLX composite (1–7)	3.312 [2.917, 3.854]	4.083 [3.083, 4.833]	−0.771 [−1.812, 0.188]	−0.79 (med.)	0.065†	0.314	Yes
ns	Session duration (s)	1022.5 [651.5, 1712.0]	893.1 [679.4, 1214.0]	+129.4 [−352.1, 946.0]	0.22 (small)	0.374	0.857	Yes

Table 4: **Terrain conditions: bootstrap results** ($N = 4$ per cell; columns as in Table 3). The terrain coordination gain is the single effect that reaches significance ($\Delta = +0.277$, exact- $p = 0.029$) and whose Δ interval excludes zero.

Sig	Metric	CAAV M [BCa CI]	No CAAV M [BCa CI]	Δ [BCa CI]	g	boot- p	exact- p	$CI \cap 0$
ns	Target discovery (% found)	0.260 [0.084, 0.401]	0.196 [0.143, 0.249]	+0.064 [−0.111, 0.207]	0.42 (small)	0.224	0.514	Yes
ns	Coverage efficiency	0.999 [0.862, 1.094]	1.027 [0.938, 1.088]	−0.028 [−0.182, 0.093]	−0.22 (small)	0.725	0.743	Yes
ns	Redundancy: Shannon NR	0.106 [0.007, 0.297]	0.019 [0.011, 0.032]	+0.087 [−0.012, 0.285]	0.57 (med.)	0.282	0.771	Yes
*	Coordination gain (over best indiv.)	0.464 [0.398, 0.550]	0.187 [0.014, 0.340]	+ 0.277 [0.112, 0.450]	1.67 (large)	0.000**	0.029*	No
ns	NASA-TLX composite (1–7)	4.229 [3.875, 4.833]	3.542 [2.708, 4.167]	+0.688 [−0.104, 1.625]	0.83 (large)	0.061†	0.257	Yes
ns	Session duration (s)	872.2 [703.1, 1161.6]	760.6 [674.7, 847.0]	+111.5 [−84.9, 424.6]	0.46 (small)	0.198	0.629	Yes

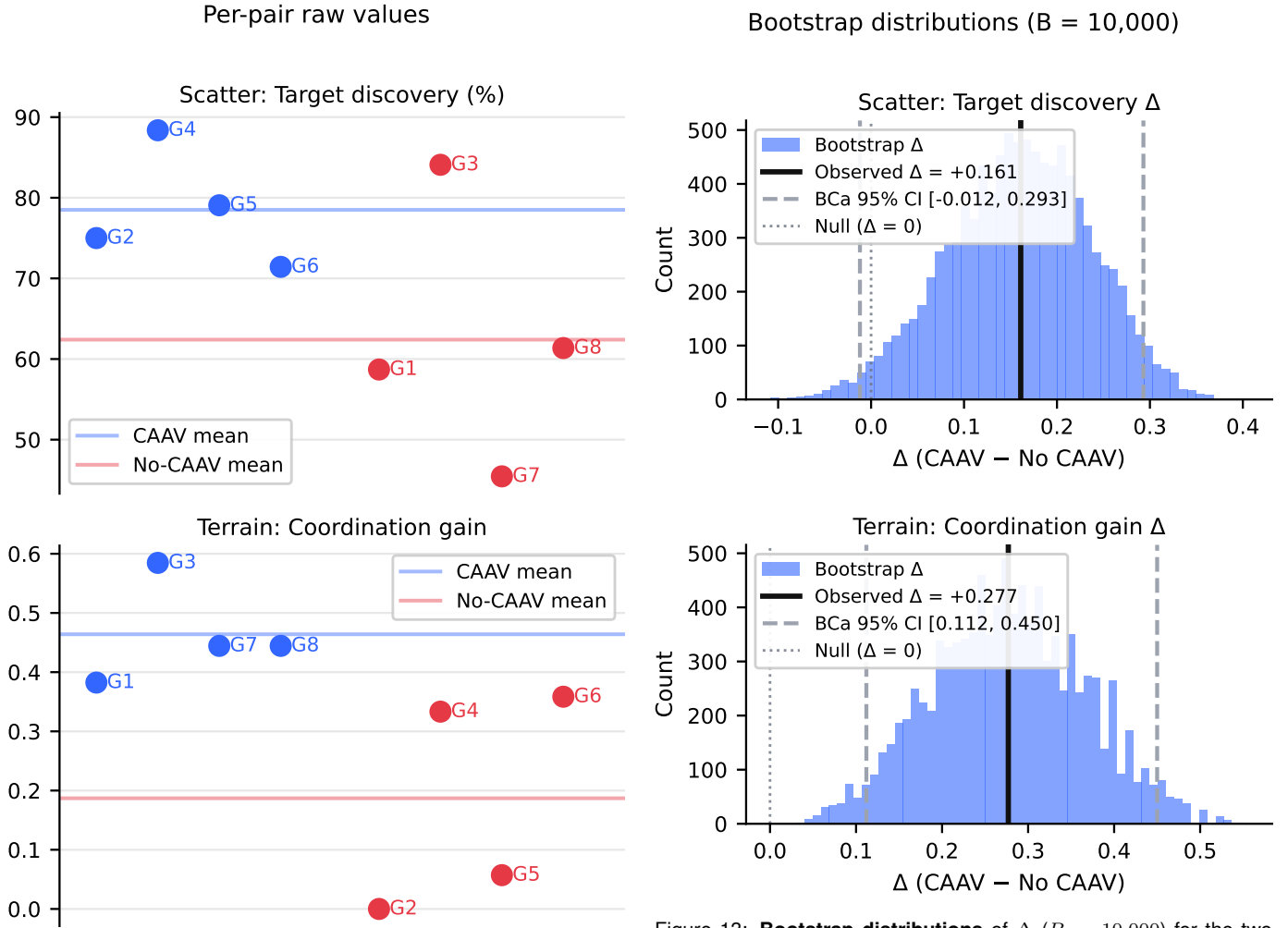


Figure 12: **Per-pair raw values** for the two main findings; horizontal lines mark the condition means.

Figure 13: **Bootstrap distributions** of Δ ($B = 10,000$) for the two primary metrics. Solid line: observed Δ ; dashed lines: BCa 95% CI bounds; dotted line: null ($\Delta = 0$).