

Making Street View Accessible Using Context-Aware Multimodal AI: A Demo of StreetReaderAI

Jon E. Froehlich
Google Research
Seattle, Washington, USA
jfroehlich@google.com

Alex Fiannaca
Google DeepMind
Seattle, Washington, USA
afiannaca@google.com

Nimer Jaber
Google
Mountain View, California, USA
nimer@google.com

Victor Tsaran
Google
Mountain View, California, USA
vtsaran@google.com

Shaun Kane
Google Research
Boulder, Colorado, USA
shaunkane@google.com



Figure 1: We introduce StreetReaderAI, an accessible streetscape mapping prototype for blind users using context-aware AI, accessible controls, and audio UI. Above, a screenshot showing an interactive view of San Francisco from atop a grassy hill. The right pane shows (a) a screen-reader accessible AI-generated description and (b) a context-aware AI chat where the user can ask open-ended questions about the scene and local geography—in this case, about how to get to the nearby houses and available pedestrian infrastructure (e.g., the presence of crosswalks and curb ramps).¹

Abstract

Interactive streetscape mapping tools, such as *Google Street View* (GSV) and *Meta Mapillary*, allow users to virtually navigate the world and plan travel in unprecedented ways, yet remain fundamentally inaccessible to blind users. We introduce *StreetReaderAI*, a new accessible street view prototype that uses multimodal AI, accessible controls, and audio UI, enabling users to examine destinations, engage in open-world exploration, and virtually tour the

world. We present the design of StreetReaderAI and preliminary findings from a lab evaluation with 11 blind participants.

Keywords

Accessible maps, Street view, Multimodal LLMs, AI chat

ACM Reference Format:

Jon E. Froehlich, Alex Fiannaca, Nimer Jaber, Victor Tsaran, and Shaun Kane. 2025. Making Street View Accessible Using Context-Aware Multimodal AI: A Demo of StreetReaderAI. In *The 27th International ACM SIGACCESS Conference on Computers and Accessibility (ASSETS '25)*, October 26–29, 2025, Denver, CO, USA. ACM, New York, NY, USA, 6 pages. <https://doi.org/10.1145/3663547.3759762>

1 Introduction

¹Some streetscape figures have been slightly altered to help protect privacy.

Permission to make digital or hard copies of all or part of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for third-party components of this work must be honored. For all other uses, contact the owner/author(s).

ASSETS '25, Denver, CO, USA

© 2025 Copyright held by the owner/author(s).

ACM ISBN 979-8-4007-0676-9/2025/10

<https://doi.org/10.1145/3663547.3759762>

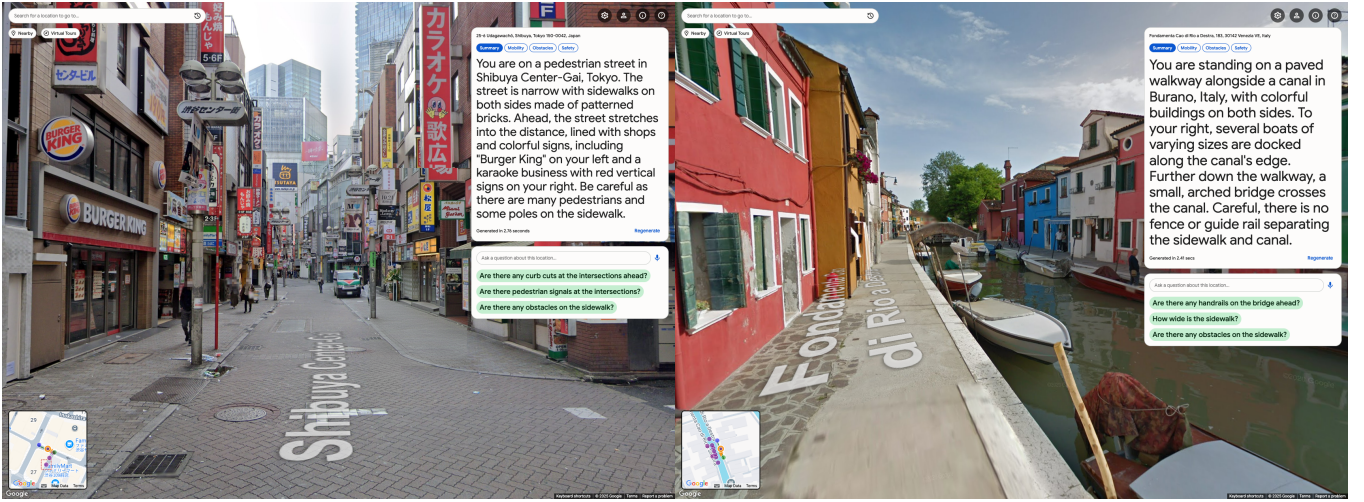


Figure 2: StreetReaderAI’s AI Descriptor provides context-aware descriptions across diverse recreational, urban, and residential scenes. In the screenshots, notice how AI Descriptor concisely overviews the scene, including geographic and visual details but also offers mobility-related information for blind navigation (e.g., road type, potential obstacles, lack of railing at the canal).

Interactive, digital maps have transformed how people plan travel and move about the world yet are historically inaccessible [8]. Recent work has enhanced the accessibility of traditional two-dimensional maps using audio descriptions and spatialized audio [10, 21, 23], tactile representations [6, 14, 15, 24], and dynamic haptic feedback [7, 26, 37]. Despite significant progress, an entire class of digital maps remains inaccessible: streetscape mapping tools such as *Google Street View* (GSV) [11], *Meta Mapillary* [20] and *Apple Look Around* [1]. These interactive tools enable users to virtually navigate and experience real-world environments via immersive 360° imagery but remain fundamentally inaccessible to blind users due to their reliance on visual panoramas, lack of textual descriptions, and inaccessible controls.

In this demo paper, we introduce *StreetReaderAI* (Figure 1), a new, accessible street view prototype using context-aware, real-time AI and accessible navigation controls. With *StreetReaderAI*, users can interactively pan and move between panoramic images, learn about nearby roads, intersections, and places, hear real-time AI descriptions, and dynamically converse with a live, multimodal AI agent about the scene and local geography. *StreetReaderAI* was designed iteratively, drawing on the experiences of two blind team members and literature in accessible first-person gaming [3, 27, 34], mixed-reality [5, 13], and mapping [17, 21, 23, 28]. To evaluate *StreetReaderAI*, we conducted a lab study with 11 blind participants finding that participants could effectively use the accessible controls and AI to virtually navigate streetscapes. Participants favored using the context-aware *AI Chat Agent* to ask situated questions rather than triggering general, AI-generated scene descriptions. Key challenges also emerged, including reconciling users’ mental models of pedestrian navigation vs. car-based streetscape imagery, a tendency to over-trust AI output, and the difficulty of synthesizing rich, panoramic, spatial information into concise audio.

During the ASSETS’25 session, we will showcase *StreetReaderAI* and allow attendees to virtually visit locations of their choice and interact with the multimodal AI (via text or speech). As the first

accessible street view tool, our work advances research in accessible maps, contributes new methods to accessibly converse with a context-aware AI agent about street scenes, and helps identify emergent challenges with accessible geospatial AI.

2 The StreetReaderAI Prototype

To design and build *StreetReaderAI*, we followed an iterative, human-centered design process that included nine co-design sessions with our two blind co-authors and feedback from professional designers and engineers with expertise in interactive maps and AI-based mixed-reality. We built *StreetReaderAI* to be natively self-voicing but also follow web standards for screen reader accessibility. Below, we describe key components. Data used with permission of Google.

2.1 Virtual Navigation with StreetReaderAI

Users interact with *StreetReaderAI* either via keyboard interactions or voice. At any point, the user can trigger an AI description or chat with the AI agent. For example, after taking a virtual step, the user can hit **[Alt] + [D]** to hear an AI-generated summary of their current view. We support two types of avatar-based control within GSV: panning and movement. See video demo.

Panning. The user can pan left and right in 45° increments using the **[←]** **[→]** arrow keys, respectively. We fix the pitch (vertical pan) to 0° so the user view is roughly at eye level. As the user pans, *StreetReaderAI* immediately voices the current heading as a cardinal or intercardinal direction (e.g., “Now facing: North” or “Northeast”), expresses whether the user can move forward along that heading and, if so, to what road address, and also explains whether the user is now facing a nearby place. When the user’s heading shifts, we describe places in front of the user—within a 45° angle of the current heading and a maximum distance of 35 meters.

Taking a step. If a nearby pano is available at the current heading, the user can take a virtual “step” using the **[↑]** arrow or move backward with **[↓]**. These steps are approximately 5-15 meters

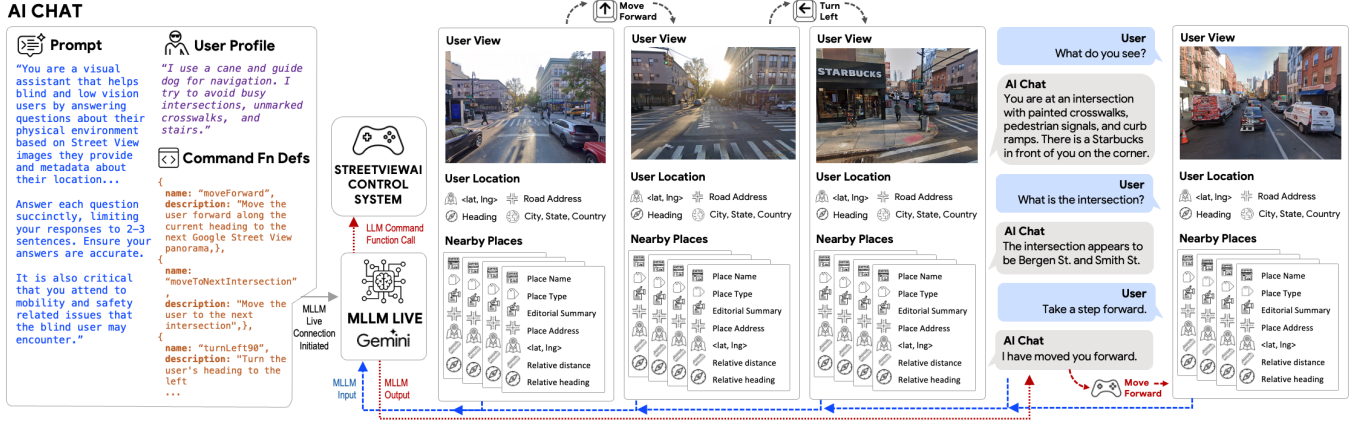


Figure 3: A system diagram of *AI Chat* showing the multimodal input, AI chat interaction, and AI-based command subsystem.

depending on the GSV pano distribution in that region. If no pano exists along the current heading, StreetReaderAI voices an *available movements* status message (also triggerable by **[Alt] + [M]**), which expresses possible movements. After taking a step, StreetReaderAI describes the movement type (e.g., step, jump), how far the user traveled, key geographic information (e.g., nearby places), and previous visits to the specific pano. For brevity, we only voice information that changes from the previous pano. The list of nearby places uses relative positions (e.g., “The Braille Library is now on your left and 25 meters away.”)

Jumping. In addition to “stepping”, the user can also *jump* along their current heading by 70 meters (230 feet) or to the nearest intersection, whichever comes first. After jumping, we generate a similar status message as a “step.”

Teleporting. Finally, the user can teleport by searching for and selecting a specific place or address in the “search box” field. Once a new destination is selected, we virtually “teleport” the user to the closest nearby pano. Crucially, to help orient the user upon landing, we automatically adjust the POV to point to the destination.

2.2 AI Subsystems

In addition to navigation, the user can invoke one of three AI subsystems built with Google Gemini [31]: *AI Descriptor*, *AI Chat*, and *AI Tour Guide*, each which use custom prompts. Unlike traditional AI-based image description tools [18, 22, 35], our AI models are multimodal and input relevant contextual information about the user, surrounding geography, and the image itself. Specifically, our AI models accept three inputs: (1) an optional user profile describing vision level and mobility needs; (2) a formatted list of auto-generated geographic elements, including current heading, address, and nearby places; (3) the user’s current view as a 640x640 image.

AI Descriptor. Invoked by **[Alt] + [D]**, *AI Descriptor* combines the three data components above with a custom multimodal LLM prompt that asks the model to “describe street view scenes for people who are blind or have low vision” and to focus on eight specific areas, including spatial relationships and navigational cues relevant to a blind pedestrian. The prompt also enumerates guidelines such as

using clear and concise language, a consistent form of reference, and to limit descriptions to two or three sentences. See Figure 2.

AI Chat Agent. The *AI Chat Agent* allows for conversational interactions about the user’s current and past views as well as nearby geography (Figure 3). The agent uses Google’s *Multimodal Live API* [12], which supports real-time interaction, function calling, and retains memory of all interactions within a single session. When the user initiates a chat either via typing (**[Alt] + [C]**) or speaking (**[Alt] + [Spacebar]**), we track and send each pan or movement interaction along with the user’s current view and geographic context (e.g., nearby places, current heading). The context window is set to a maximum of 1,048,576 input tokens, which is roughly equivalent to over 4k input images. To help the user discover features in the scene, we also seed the chat interface with three scene-based questions (e.g., see Figure 2). Additionally, we support a small list of avatar controls via *AI Chat commands* such as turning, stepping, and jumping. These commands can be spoken or typed.

AI Tour Guide. Finally, the *AI Tour Guide* functions similarly to *AI Descriptor* but changes the MLLM prompt to act as an “expert tour guide for blind or low-vision virtual tourists.”, including historical facts, cultural significance, architectural styles, interesting anecdotes, and nearby popular attractions. See Figure 4.

3 User Study

To evaluate StreetReaderAI and explore the potential for accessible, AI-driven street view experiences, we conducted an in-person lab study with eleven blind participants (6 men, 5 women, aged 20–66+). All participants used white canes for navigation and screen readers for computing. While most had experience with digital mapping tools and some familiarity with AI applications like *Be My AI* [4], none had previously used streetscape imagery tools. The 1.5-2 hour sessions comprised a formative interview on navigation and technology use, a system tutorial followed by point-of-interest (POI) investigations at three distinct unfamiliar locations (a bus stop, playground, and restaurant), two open-world navigation tasks in pre-selected areas, and a concluding debrief interview. Data collection methods included audio/video recordings, client-side interaction logs, researcher observations, and Likert-scale questionnaires.



Figure 4: The AI Tour Guide employs a specialized prompt directing the MLLM to act as an “expert tour guide for blind virtual tourists.” In this screenshot, the user is virtually touring the *Treasury of Petra* also known as “Al-Khazneh.” The AI Tour Guide describes the look of the treasury, its architecture and history, the presence of two camels adorned with colorful saddles, and recommends two other nearby places (the Royal Tombs and the Monastery).

Overall, participants reacted positively to StreetReaderAI and used it to move to 356 panos, make 568 heading changes, and 1,053 AI requests (136 AI Descriptor and 917 AI Chats). All eleven participants completed the POI investigations and 10/11 completed at least one open-world navigation task. As P10 said, “Most navigation systems can get you to the last 5-10 feet but this helps you get to a door and even describes that door” and P1, “If I’m going to a place, I can familiarize myself first from my home.” During the post-study debrief, participants rated the overall usefulness of StreetReaderAI as 6.4/7 ($Median=7$; $SD=0.9$) and all wanted to use the system as a product, if available ($Avg=6.6$; $Med=7$; $SD=0.8$). While participants found value in StreetReaderAI, relied heavily on its AI features, and adeptly combined virtual world navigation with AI interactions, they struggled with maintaining spatial orientation, distinguishing the veracity of AI responses, and determining StreetReaderAI’s knowledge limits (e.g., access to transit schedules, restaurant menus). See our full paper for details [9].

4 Discussion and Conclusion

In this demo paper, we introduced StreetReaderAI, a new, accessible street view prototype with context-aware, real-time AI and accessible navigation controls. Below, we reflect on limitations and opportunities for future work.

Accuracy and Trust. Participants exhibited a high-level of trust in StreetReaderAI. However, as other scholars have emphasized [16, 25, 36], it is difficult to discern hallucinations from truth—StreetReaderAI seems just as confident about both. Future work should explore how to frame AI-generated descriptions, informed by guidelines in *Explainable AI*.

Data Discrepancies. StreetReaderAI draws primarily on two data sources: geographic databases of road, place, and address information (which is typically up-to-date and trustworthy) and AI-generated descriptions about the scene and local geography (which can depend on outdated street view imagery and imperfect inferences). However, when contradictions arose, users again had no way to discern truth. Such discrepancies can impact spatial mental models and travel planning.

Future Work. Beyond the above, we propose five key areas of future work: (1) support origin-to-destination route previewing with easy-to-use controls that provide BLV users with an accessible overview of the walking experience; (2) enhancing the AI chat agent to support questions beyond the current (and past) views and by providing additional geographic data sources (e.g., transit schedules); (3) using spatialized audio or haptics to provide non-verbal cues about distances, object locations, or other information (drawing on accessible navigation tools [21, 23, 32], VR techniques [19, 29, 30, 33], and audio games [2, 27]); (4) more robustly evaluating the MLLM’s spatial understanding across geographic contexts and scenarios; (5) a deployment study to examine *what* questions BLV users ask during open-ended natural usage.

Conclusion. In conclusion, our work introduces AI-driven accessibility support for an emergent class of previously inaccessible mapping tools: immersive streetscape imagery. Our ASSETS’25 demo will invite attendees to use StreetReaderAI, visit locations of their choice, and interact with our context-aware AI subsystems.

Acknowledgments

Diagram icons from Noun Project, including: “prompt icon” by Firdaus Faiz, “command functions” by Kawalan Icon, “<lat,lng>” by Didik Darmanto, “heading” and “relative heading” by IronCV, “road address” and “place address” by IGraphics, “city, state, country” by Karyative, “café” by loviana, “hospital” by loviana, “park” by Made x Made, “school” by Omah Icon, “storefront” by Sutiayah, “place type” by BEARicons, “editorial summary” by ilham firmansyah, “relative distance” by Sandy Walsh, “blind person” by Nuno Sequeira, “game controller” by PMSO BEMFEBUNUD, “up arrow key” by Se-hui Jo, “left arrow key” by Se-hui Jo, “MLLM icon” by Funtasticon.

References

- [1] Apple Inc. 2025. Apple Look Around. <https://www.apple.com/maps/>
- [2] Matthew Tylee Atkinson. 2018. AGRIP: Accessible Gaming Rendering Independence Possible / AudioQuake. <https://github.com/matatk/agrip>
- [3] Matthew T. Atkinson, Sabahattin Gucukoglu, Colin H. C. Machin, and Adrian E. Lawrence. 2006. Making the mainstream accessible: redefining the game. In *Proceedings of the 2006 ACM SIGGRAPH Symposium on Videogames* (Boston, Massachusetts) (Sandbox '06). Association for Computing Machinery, New York, NY, USA, 21–28. doi:10.1145/1183316.1183321
- [4] Be My Eyes. 2025. Be My AI. <https://www.bemyeyes.com/blog/be-my-ai-launches-globally>
- [5] Ruei-Che Chang, Yuxuan Liu, and Anhong Guo. 2024. WorldScribe: Towards Context-Aware Live Visual Descriptions. In *Proceedings of the 37th Annual ACM Symposium on User Interface Software and Technology* (Pittsburgh, PA, USA) (UIST '24). Association for Computing Machinery, New York, NY, USA, Article 140, 18 pages. doi:10.1145/3654777.3676375
- [6] Kirk Andrew Crawford, Jennifer Posada, Yetunde Esther Okueso, Erin Higgins, Laura Lachin, and Foad Hamidi. 2024. Co-designing a 3D-Printed Tactile Campus Map With Blind and Low-Vision University Students. In *Proceedings of the 26th International ACM SIGACCESS Conference on Computers and Accessibility* (St. John's, NL, Canada) (ASSETS '24). Association for Computing Machinery, New York, NY, USA, Article 77, 6 pages. doi:10.1145/3663548.3688537
- [7] Julie Ducasse, Anke M. Brock, and Christophe Jouffrais. 2018. *Accessible Interactive Maps for Visually Impaired Users*. Springer International Publishing, Cham, 537–584. doi:10.1007/978-3-319-54446-5_17
- [8] Jon E. Froehlich, Anke M. Brock, Anat Caspi, João Guerreiro, Kotaro Hara, Reuben Kirkham, Johannes Schöning, and Benjamin Tannert. 2019. Grand challenges in accessible maps. 26, 2 (Feb. 2019), 78–81. doi:10.1145/3301657
- [9] Jon E. Froehlich, Alex Fiannaca, Nimer Jaber, Victor Tsaran, and Shaun Kane. 2025. StreetReaderAI: Making Street View Accessible Using Context-Aware Multimodal AI. In *The 38th Annual ACM Symposium on User Interface Software and Technology* (UIST '25) (Busan, Republic of Korea) (UIST '25). ACM, New York, NY, USA, 22. doi:10.1145/3746059.3747756
- [10] Google. 2023. Voice Guidance in Maps, Built for People with Impaired Vision. <https://blog.google/products/maps/better-maps-for-people-with-vision-impairments/> Official Google blog post about accessibility features in Google Maps.
- [11] Google. 2025. Google Street View. <https://www.google.com/streetview/>
- [12] Google. 2025. Vertex AI Multimodal Live API. Google Cloud. <https://cloud.google.com/vertex-ai/generative-ai/docs/multimodal-live-api>
- [13] Google DeepMind. 2024. Project Astra. <https://deepmind.google/technologies/project-astra/>. Accessed: 2025-04-08.
- [14] Leona Holloway, Matthew Butler, and Kim Marriott. 2023. TactIcons: Designing 3D Printed Map Icons for People who are Blind or have Low Vision. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems* (Hamburg, Germany) (CHI '23). Association for Computing Machinery, New York, NY, USA, Article 543, 18 pages. doi:10.1145/3544548.3581359
- [15] Leona Holloway, Kim Marriott, Matthew Butler, and Samuel Reinders. 2019. 3D Printed Maps and Icons for Inclusion: Testing in the Wild by People who are Blind or have Low Vision. In *Proceedings of the 21st International ACM SIGACCESS Conference on Computers and Accessibility* (Pittsburgh, PA, USA) (ASSETS '19). Association for Computing Machinery, New York, NY, USA, 183–195. doi:10.1145/3308561.3353790
- [16] Ilkka Kaate, Joni Salminen, Soon-Gyo Jung, Trang Thi Thu Xuan, Essi Häyhänen, Jinan Y. Azem, and Bernard J. Jansen. 2025. “You Always Get an Answer”: Analyzing Users’ Interaction with AI-Generated Personas Given Unanswerable Questions and Risk of Hallucination. In *Proceedings of the 30th International Conference on Intelligent User Interfaces* (IUI '25). Association for Computing Machinery, New York, NY, USA, 1624–1638. doi:10.1145/3708359.3712160
- [17] Minchu Kulkarni, Chu Li, Jaye Jungmin Ahn, Katrina Oi Yau Ma, Zhihan Zhang, Michael Saugstad, Kevin Wu, Yochai Eisenberg, Valerie Novack, Brent Chamberlain, and Jon E. Froehlich. 2023. BusStopCV: A Real-time AI Assistant for Labeling Bus Stop Accessibility Features in Streetscape Imagery. In *Proceedings of the 25th International ACM SIGACCESS Conference on Computers and Accessibility* (New York, NY, USA) (ASSETS '23). Association for Computing Machinery, New York, NY, USA, Article 91, 6 pages. doi:10.1145/3597638.3614481
- [18] Jaewook Lee, Jaylin Herskovitz, Yi-Hao Peng, and Anhong Guo. 2022. Image-Explorer: Multi-Layered Touch Exploration to Encourage Skepticism Towards Imperfect AI-Generated Image Captions. In *Proceedings of the 2022 CHI Conference on Human Factors in Computing Systems* (New Orleans, LA, USA) (CHI '22). Association for Computing Machinery, New York, NY, USA, Article 462, 15 pages. doi:10.1145/3491102.3501966
- [19] Keenan R. May, Brianna J. Tomlinson, Xiaomeng Ma, Phillip Roberts, and Bruce N. Walker. 2020. Spotlights and Soundscapes: On the Design of Mixed Reality Auditory Environments for Persons with Visual Impairment. *ACM Trans. Access. Comput.* 13, 2, Article 8 (April 2020), 47 pages. doi:10.1145/3378576
- [20] Meta Platforms, Inc. 2025. Mapillary by Meta. <https://www.mapillary.com/>
- [21] Microsoft. 2023. *Microsoft Soundscape: Empowering people who are blind or have low vision to explore the world around them*. Microsoft. <https://blogs.microsoft.com/accessibility/soundscape/> Microsoft Accessibility Blog.
- [22] Microsoft Corporation. 2025. Seeing AI. <https://www.microsoft.com/en-us/ai/seeing-ai>
- [23] MIPsoft. 2025. BlindSquare. <https://www.blindsquare.com/>
- [24] Ruth G Nagassa, Matthew Butler, Leona Holloway, Catagay Goncu, and Kim Marriott. 2023. 3D Building Plans: Supporting Navigation by People who are Blind or have Low Vision in Multi-Storey Buildings. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems* (Hamburg, Germany) (CHI '23). Association for Computing Machinery, New York, NY, USA, Article 539, 19 pages. doi:10.1145/3544548.3581389
- [25] Mahjabin Nahar, Haeseung Seo, Eun-Ju Lee, Aiping Xiong, and Dongwon Lee. 2024. Fakes of Varying Shades: How Warning Affects Human Perception and Engagement Regarding LLM Hallucinations. arXiv:2404.03745 [cs.HC] <https://arxiv.org/abs/2404.03745>
- [26] Maria Teresa Paratore and Barbara Leporini. 2024. Exploiting the haptic and audio channels to improve orientation and mobility apps for the visually impaired. *Universal Access in the Information Society* 23, 2 (6 2024), 859–869. doi:10.1007/s10209-023-00973-4
- [27] Jamie Pauls. 2020. Vintage Games Series, Part 4: Immerse Yourself in the World of Shades of Doom. <https://www.afb.org/aw/21/12/17336>
- [28] Manaswi Saha, Michael Saugstad, Hanuma Teja Maddali, Aileen Zeng, Ryan Holland, Steven Bower, Aditya Dash, Sage Chen, Anthony Li, Kotaro Hara, and Jon Froehlich. 2019. Project Sidewalk: A Web-based Crowdsourcing Tool for Collecting Sidewalk Accessibility Data At Scale. In *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems* (Glasgow, Scotland UK) (CHI '19). Association for Computing Machinery, New York, NY, USA, 1–14. doi:10.1145/3290605.3300292
- [29] David W. Schloerb, Orly Lahav, Joseph G. Desloge, and Mandayam A. Srinivasan. 2010. BlindAid: Virtual environment system for self-reliant trip planning and orientation and mobility training. In *2010 IEEE Haptics Symposium*. 363–370. doi:10.1109/HAPTIC.2010.5444631
- [30] Alexa F. Siu, Mike Sinclair, Robert Kovacs, Eyal Ofek, Christian Holz, and Edward Cutrell. 2020. Virtual Reality Without Vision: A Haptic and Auditory White Cane to Navigate Complex Virtual Worlds. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems* (Honolulu, HI, USA) (CHI '20). Association for Computing Machinery, New York, NY, USA, 1–13. doi:10.1145/3313831.3376353
- [31] Gemini Team, Rohan Anil, Sebastian Borgeaud, Jean-Baptiste Alayrac, Jiahui Yu, Radu Soricut, Johan Schalkwyk, and Andrew M. Dai et al. 2024. Gemini: A Family of Highly Capable Multimodal Models. arXiv:2312.11805 [cs.CL] <https://arxiv.org/abs/2312.11805>
- [32] Bruce N. Walker and Jeffrey Lindsay. 2006. Navigation Performance With a Virtual Auditory Display: Effects of Beacon Sound, Capture Radius, and Practice. *Human Factors* 48, 2 (2006), 265–278. doi:10.1518/001872006777724507 arXiv:https://doi.org/10.1518/001872006777724507 PMID: 16884048
- [33] Ryan Wedoff, Lindsay Ball, Amelia Wang, Yi Xuan Khoo, Lauren Lieberman, and Kyle Rector. 2019. Virtual Showdown: An Accessible Virtual Reality Game with Scaffolds for Youth with Visual Impairments. In *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems* (Glasgow, Scotland UK) (CHI '19). Association for Computing Machinery, New York, NY, USA, 1–15. doi:10.1145/3290605.3300371
- [34] Thomas Westin. 2004. Game accessibility case study: Terraformers—a real-time 3D graphic game. In *Proceedings of the 5th International Conference on Disability, Virtual Reality and Associated Technologies, ICDVRAT*, Vol. 1. 95–100. https://www.researchgate.net/publication/260853226_Proceedings_of_the_5th_International_Conference_on_Disability_Virtual_Reality_and_Associated_Technologies_ICDVRAT_2004/links/0deec53282e0b1593600000/Proceedings-of-the-5th-International-Conference-on-Disability-Virtual-

- Reality-and-Associated-Technologies-ICDVRAT-2004.pdf#page=121
- [35] Shaomei Wu, Jeffrey Wieland, Omid Farivar, and Julie Schiller. 2017. Automatic Alt-text: Computer-generated Image Descriptions for Blind Users on a Social Network Service. In *Proceedings of the 2017 ACM Conference on Computer Supported Cooperative Work and Social Computing* (Portland, Oregon, USA) (CSCW '17). Association for Computing Machinery, New York, NY, USA, 1180–1192. doi:10.1145/2998181.2998364
- [36] Zhangyue Yin, Qiushi Sun, Qipeng Guo, Jiawen Wu, Xipeng Qiu, and Xuanjing Huang. 2023. Do Large Language Models Know What They Don't Know? arXiv:2305.18153 [cs.CL] <https://arxiv.org/abs/2305.18153>
- [37] Limin Zeng, Mei Miao, and Gerhard Weber. 2014. Interactive Audio-haptic Map Explorer on a Tactile Display. *Interacting with Computers* 27, 4 (02 2014), 413–429. doi:10.1093/iwc/iwu006 arXiv:<https://academic.oup.com/iwc/article-pdf/27/4/413/9644221/iwu006.pdf>