

Article

A Multi-Modal Foundation Model to Assist People with Blindness and Low Vision in Environmental Interaction

Yu Hao ^{1,†}, Fan Yang ^{1,†}, Hao Huang ¹ , Shuaihang Yuan ¹, Sundeep Rangan ¹, John-Ross Rizzo ^{1,2}, Yao Wang ¹ and Yi Fang ^{1,3,*}

¹ Tandon School of Engineering, New York University, Brooklyn, NY 11201, USA; yh3252@nyu.edu (Y.H.); fy2167@nyu.edu (F.Y.); hh1811@nyu.edu (H.H.); sy2366@nyu.edu (S.Y.); srangan@nyu.edu (S.R.); johnross.rizzo@nyulangone.org (J.-R.R.); yw523@nyu.edu (Y.W.)

² NYU Langone Health, New York University, New York, NY 10016, USA

³ Electrical Engineering and Center for Artificial Intelligence and Robotics, New York University Abu Dhabi, Abu Dhabi 129188, United Arab Emirates

* Correspondence: yfang@nyu.edu

† These authors contributed equally to this work.

Abstract: People with blindness and low vision (pBLV) encounter substantial challenges when it comes to comprehensive scene recognition and precise object identification in unfamiliar environments. Additionally, due to the vision loss, pBLV have difficulty in accessing and identifying potential tripping hazards independently. Previous assistive technologies for the visually impaired often struggle in real-world scenarios due to the need for constant training and lack of robustness, which limits their effectiveness, especially in dynamic and unfamiliar environments, where accurate and efficient perception is crucial. Therefore, we frame our research question in this paper as: *How can we assist pBLV in recognizing scenes, identifying objects, and detecting potential tripping hazards in unfamiliar environments, where existing assistive technologies often falter due to their lack of robustness?* We hypothesize that by leveraging large pretrained foundation models and prompt engineering, we can create a system that effectively addresses the challenges faced by pBLV in unfamiliar environments. Motivated by the prevalence of large pretrained foundation models, particularly in assistive robotics applications, due to their accurate perception and robust contextual understanding in real-world scenarios induced by extensive pretraining, we present a pioneering approach that leverages foundation models to enhance visual perception for pBLV, offering detailed and comprehensive descriptions of the surrounding environment and providing warnings about potential risks. Specifically, our method begins by leveraging a large-image tagging model (i.e., Recognize Anything Model (RAM)) to identify all common objects present in the captured images. The recognition results and user query are then integrated into a prompt, tailored specifically for pBLV, using prompt engineering. By combining the prompt and input image, a vision-language foundation model (i.e., InstructBLIP) generates detailed and comprehensive descriptions of the environment and identifies potential risks in the environment by analyzing environmental objects and scenic landmarks, relevant to the prompt. We evaluate our approach through experiments conducted on both indoor and outdoor datasets. Our results demonstrate that our method can recognize objects accurately and provide insightful descriptions and analysis of the environment for pBLV.

Keywords: assistive technology; multi-model foundation model; vision-language model



Citation: Hao, Y.; Yang, F.; Huang, H.; Yuan, S.; Rangan, S.; Rizzo, J.-R.; Wang, Y.; Fang, Y. A Multi-Modal Foundation Model to Assist People with Blindness and Low Vision in Environmental Interaction. *J. Imaging* **2024**, *10*, 103. <https://doi.org/10.3390/jimaging10050103>

Academic Editor: Elena Casiraghi

Received: 23 January 2024

Revised: 14 April 2024

Accepted: 19 April 2024

Published: 26 April 2024



Copyright: © 2024 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

The prevalence of visual impairment has reached alarming levels, affecting millions of individuals worldwide, as highlighted by recent estimates from the World Health Organization (WHO) [1,2]. The number of people experiencing moderate to severe visual impairment or complete blindness continues to rise steadily, with projections indicating a further surge in these numbers by 2050 [3]. Visual impairment, whether partial or complete,

presents significant challenges that profoundly impact various aspects of daily life for pBLV [4]. Among the critical tasks that pose difficulties for pVLB is visual search, which involves actively scanning the environment and locating a specific target among distracting elements [5]. Even for individuals with normal vision, visual search can be demanding, especially in complex environments. However, for individuals with blindness or low vision, these challenges are further compounded [6]. Those with peripheral vision loss, central vision loss, or hemi-field vision loss often struggle to pinpoint a particular location or search for objects due to reduced fields of view. They often require assistance to accurately identify the environment or locate objects of interest. Similarly, individuals experiencing blurred vision or nearsightedness encounter difficulties in identifying objects at varying distances. Color-deficient vision and low-contrast vision further exacerbate the challenges of distinguishing objects from the background when they share similar colors. In addition to understanding their surroundings and locating objects of interest, assessing potential risks and hazards within the visual environment becomes an intricate task, demanding a comprehensive analysis to ensure personal safety [7]. Therefore, addressing the challenges faced by pBLV in environmental interaction holds profound significance due to the escalating prevalence of visual impairment globally, which substantially affects millions and is projected to increase further. These challenges, which include difficulties in visual search, object identification, and risk assessment in diverse environments, critically impact the independence, safety, and quality of daily life of pBLV. Innovatively enhancing visual perception for these individuals not only promises to mitigate these profound challenges, but also aims to empower them with greater autonomy and confidence in navigating their surroundings, thus fostering inclusivity and accessibility in society.

Current assistive technologies for pBLV [8–10] driven by computer vision approaches have led to the development of assistive systems that utilize object recognition [11], GPS navigation [12], and text-to-speech tools [13]. While these technologies have provided valuable assistance to visually impaired individuals [14], they still face certain challenges and limitations. One of the primary challenges with existing assistive technologies is their limited ability to provide comprehensive scene understanding and guidance to address the specific needs of visually impaired individuals. For instance, while many tools focus on specific functionalities, such as obstacle detection or route planning, they often fall short in delivering detailed descriptions and guidance based on user questions. The current solutions also lack the capability to generate contextually relevant information about objects, scenes, and potential risks in the environment, limiting an in-depth understanding of the environment for visually impaired individuals. Conversational search finds applications in various domains such as basic information retrieval, personal information search, product selection and travel planning, which facilitates information retrieval through conversation [15]. Additionally, these solutions, such as object detection [11], frequently encounter difficulties in real-world scenarios due to the need for constant training and adaptation. They exhibit a lack of robustness, which limits their effectiveness, particularly in dynamic and unfamiliar environments, in which accurate and efficient perception is crucial. This limitation hinders their ability to fully perceive and understand their surroundings, resulting in reduced independence and increased reliance on external assistance.

In this paper, we aim to address the research questions of exploring whether large foundation models can address the limitations of current assistive technologies for pBLV by enhancing comprehensive scene understanding, providing contextually relevant information, and improving adaptability and robustness in dynamic environments. We posit that these large foundation models, through extensive pretraining, can significantly improve the functionality of assistive technologies by providing detailed guidance and accurate environmental perceptions, thus increasing the independence and safety of visually impaired individuals to navigate their surroundings. As shown in Figure 1, we present a novel approach named *VisPercep* that leverages the advanced large vision-language model to enhance visual perception for individuals with blindness and low vision, including scene understanding, object localization, and risk assessment. Our work addresses the challenges

faced by pBLV by providing them with detailed and comprehensive scene descriptions and risk guidance based on user questions, allowing an in-depth understanding of their surroundings, locating objects of interest, and identifying potential risks.

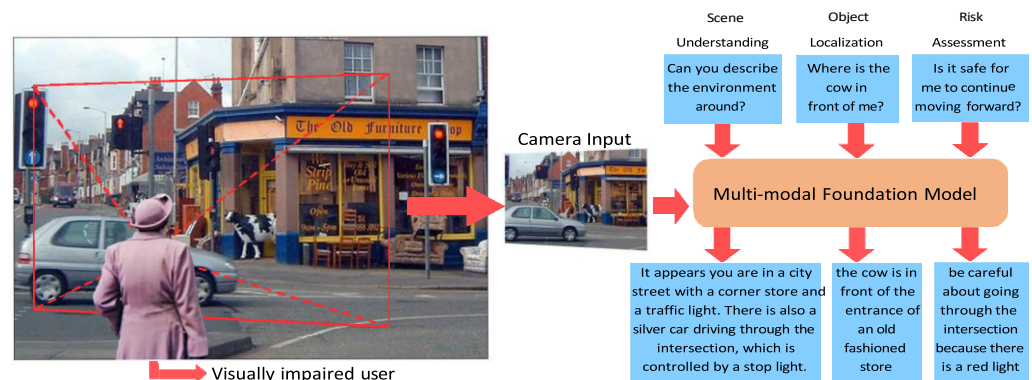


Figure 1. Multi-Modal Foundation Model Sample Illustration.

Our system includes three main modules, as illustrated in Figure 2: image tagging module, prompt engineering module, and vision-language module. The image tagging module, implemented using Recognize Anything Model (RAM) [16], recognizes all objects in the captured image. We then integrate the recognized objects and user questions into a customized prompt designed for visually impaired individuals through prompt engineering. Finally, the vision-language module utilizes InstructBLIP [17] to generate detailed and contextually relevant text, facilitating comprehensive scene understanding, object recognition, and risk assessment for visually impaired individuals. Our experiments demonstrate that our system can recognize objects of interest and provide detailed answers to user questions, significantly enhancing the visual understanding of surroundings.

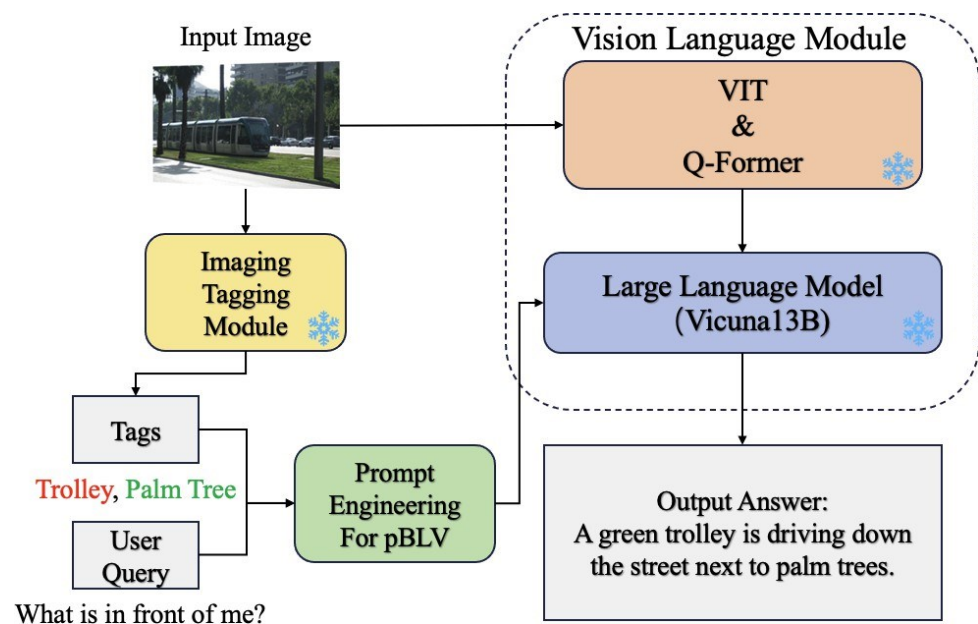


Figure 2. Method Structure Overview.

Our contributions are summarized as follows:

1. In response to the challenges faced by pBLV in achieving comprehensive scene recognition and identifying objects and hazards in unfamiliar environments, we introduce an innovative approach that leverages a multi-modal foundation model. This model is

- designed to significantly enhance environmental understanding by offering detailed and comprehensive descriptions of surroundings and alerting users to potential risks.
2. To directly address the limitations of current assistive technologies, which often lack robustness and the capability to adapt to dynamic scenarios, we have developed a voice-controlled system. This system uniquely combines a large image tagging model with a vision-language foundation model, facilitating intuitive, language-guided question answering that caters specifically to the needs of pBLV.
 3. Our approach's effectiveness is validated through rigorous testing on both indoor and outdoor datasets. These experiments demonstrate the system's superior ability to accurately recognize objects and provide accurate descriptions and analyses of the environment, thereby directly addressing the core research problem of enhancing navigation and interaction for pBLV in diverse settings.

In the following section of the article: (Section 2) Related Work: Summarized existing auxiliary technologies in blind and low vision. (Section 3) Materials and Methods: Proposed to leverage the power of a multi-modal foundation model that integrates image tagging and visual language models to provide detailed environmental descriptions and risk assessments. (Section 4) Experiments and Results: The accuracy and effectiveness of the model were validated through indoor and outdoor datasets, demonstrating its guidance ability for blind and visually impaired individuals. (Section 5) Conclusion: Summarized the contribution of the research, emphasizing the use of new methods to provide better guidance and enhance the independence and safety of visually impaired individuals. (Section 6) Limitations and Future Research: Explained the difficulties encountered by the model at the current stage from different aspects and future research directions.

2. Related Works

Initial research has seen a growth in interest in the development of conversational search systems intended to support users in their information-seeking activities [18]. This work has primarily focused on communication of information exclusively via spoken dialogue. While this is sufficient for simple question-type queries, it is an inefficient means of engagement for more complex or exploratory queries [19]. In the realm of information, conversational search is a relatively new trend [20]. Conversation is the natural mode for information exchange in daily life [18], and conversational approaches to information retrieval are gaining attention [15]. By integrating conversational search approaches with existing assistive technologies, there is potential to enhance the user experience and address the limitations of current systems, providing more natural and efficient interaction for individuals with visual disabilities.

In recent years, several assistive technologies and applications developed to support individuals with visual disabilities in understanding their environment and enhancing their scene understanding [4,21,22]. Traditional tools such as white canes [23] and guide dogs [24] have long been used to aid in mobility and spatial awareness. Additionally, advancements in technology have led to the development of various assistive devices, including wearable cameras [25–27], GPS navigation systems, and object recognition technologies [28].

Wearable camera systems, such as the OrCam MyEye and Seeing AI [29], offer real-time text reading and text-to-speech capabilities to provide auditory feedback to individuals with visual disability. These systems assist in object identification, text reading, and facial recognition, enhancing their ability to interact with their surroundings. GPS navigation systems, such as BlindSquare [30] and Lazarillo [31], utilize location-based services to provide audio instructions and guidance for navigation in both indoor and outdoor environments.

Computer vision-based technologies have also been explored for scene understanding. These include object detection systems using deep learning models like YOLO [32] and Faster R-CNN [33], which provide real-time identification of objects in the environment. Detect and Approach [27] proposes a real-time monocular-based navigation solution based on YOLO for pBLV. Additionally, vision-language models like VizWiz [34] and

SoundScape [35] incorporate natural language processing to describe visual scenes, answer questions, and provide context-aware information.

While these existing assistive technologies have made significant advancements, they still face limitations. Many systems provide partial solutions focused on specific functionalities such as object recognition or detection, but often fall short in delivering comprehensive scene understanding and detailed descriptions. Moreover, these technologies may lack the ability to provide guidance based on user questions, limiting their effectiveness in addressing the specific needs and queries of individuals with visual disability [9]. Furthermore, these technologies often require multiple devices or interfaces, leading to complexity and decreased usability for individuals with visual disability [21]. In contrast to these existing approaches, our proposed method offers a comprehensive and integrated solution. By combining advanced vision-language models, image tagging, and prompt engineering, our approach enhances scene understanding, provides real-time guidance, and offers context-aware prompts tailored specifically for individuals with visual disability.

3. Materials and Methods

In this work, as shown in Figure 1, the proposed model leverages the advanced large vision-language model to assist environmental interaction for individuals with blindness and low vision including scene understanding, object localization, and risk assessment. Our system utilizes a smartphone to capture images and record user questions (left). Based on the input image and user question, our proposed model generates detailed and comprehensive scene descriptions and risk assessments (right). Moreover, the camera input image is from Visual7W dataset [36]. Our work addresses the challenges faced by pBLV by providing them with detailed and comprehensive scene descriptions and risk guidance based on user questions, enabling an in-depth understanding of their surroundings, locating objects of interest, and identifying potential risks.

Our system includes three main modules, as illustrated in Figure 2: image tagging module, prompt engineering module and vision-language module. Firstly, the image tagging module, implemented using Recognize Anything Model (RAM) [16], identifies all common objects present in the captured image. Secondly, using prompt engineering, we integrate the recognized objects and user queries to create customized prompts tailored for individuals with visual disability. Lastly, the vision-language module which utilizes InstructBLIP [17] generates detailed and contextually relevant output text, enabling comprehensive and precise scene understanding, object localization, and risk assessment for individuals with visual disability. [The input image is from Visual7W dataset [36].

Our method aims to overcome the limitations of existing assistive technologies and empower individuals with visual disability with improved guidance. In Section 3.1, we introduce our image tagging module. Section 3.2 illustrates the prompt engineering tailored specifically for individuals with visual disability. We explain the large vision-language module in Section 3.3.

3.1. Image Tagging Module

As shown in the yellow box of Figure 2, the image tagging module is utilized to generate tags for each object present in the captured images, which is crucial as it provides a comprehensive understanding of the visual scene by accurately recognizing various objects. By incorporating the image tagging module, we obtain a catalog of objects present in the environment, facilitating a more precise and comprehensive environment description. We employ the Recognize Anything Model (RAM) [16] as our image tagging module, which has demonstrated the zero-shot ability to recognize any common category with high accuracy.

Specifically, the image tagging module begins with a pre-trained image encoder, which processes an input image I and extracts high-level visual features F , formulated as: $F = \text{Encoder}(I)$. These features capture important characteristics and representations of the objects in the image. After the initial feature extraction stage, an attention mechanism [37]

is employed to focus on the most salient regions within the image. Represented mathematically as $A = \text{Attention}(F)$, this attention mechanism allows the model to pay more attention to relevant objects and suppress irrelevant ones. Thus, the image tagging module can generate accurate and informative tags for the recognized objects. The final stage involves mapping the extracted features to a set of object categories or tags by the image-tag recognition decoder. This mapping, expressed as $T = \text{Decoder}(A)$, is learned through a training process that leverages large-scale annotated datasets, ensuring the model's ability to generalize to various objects and scenes. The trained RAM model can then be applied to new images, accurately recognizing and generating tags for the objects present in the environment.

3.2. Prompt Engineering for pBLV

We incorporate prompt engineering, as shown in the green box of Figure 2, to create customized prompts tailored specifically for individuals with visual disability. This involves integrating the output of the image tagging module with user questions to form contextually relevant and informative prompts. Moreover, the use of prompt engineering eliminates the need for traditional machine learning approaches that require training models on labeled datasets, as prompt engineering focuses on generating effective prompts rather than optimizing model parameters.

The RAM generates a set of tags that represent the recognized objects within the captured images. We utilize these tags to enhance the final prompt. We include the prompt *"The image may contain elements of {tags}"* to seamlessly integrate the object recognition results into a prompt. By incorporating these recognized object tags into the prompt, we ensure that the vision-language module receives specific and accurate information about the objects in their surroundings. This approach significantly enhances the understanding and awareness of the visual environment for the users.

Furthermore, we consider user questions as vital input for prompt engineering. By incorporating user questions into the prompts, we address the individual's specific needs for environmental understanding and ensure that the prompts are highly relevant to their current situation. This personalized approach allows individuals with visual disability to obtain the targeted information about their environment and the objects of interest. For example, in the case of risk assessment, we employ a specific prompt that guides the model to act as an assistant for individuals with visual disability, providing comprehensive analysis. The prompt we use is *"I am visually disabled. You are an assistant for individuals with visual disability. Your role is to provide helpful information and assistance based on my query. Your task is to provide a clear and concise response that addresses my needs effectively. Don't mention that I am visually disabled to offend me. Now, please answer my questions: [user_query]. Your answer should be like a daily conversation with me."* where *[user_query]* is the user question. This prompt enables the model to deliver detailed and accurate explanations regarding potential risks, ensuring that the information is communicated in a respectful and informative manner.

3.3. Vision-Language Module

To generate output text based on the prompts obtained by the prompt-engineering module, we employ InstructBLIP [17], a powerful large vision language model for comprehensive scene understanding and analysis, as shown in the right blue box of Figure 2.

Specifically, InstructBLIP begins by encoding the input image I using the frozen Vision Transformer (ViT) [38], which captures a high-level feature embedding V of the image, represented as $V = \text{ViT}(I)$. The Q-Former [17] in InstructBLIP, distinct from conventional models, employs learnable query embedding Q and the image embedding V from ViT for processing. This is formulated as $C = Q - \text{Former}(Q, V)$, where cross-attention is applied to generate contextualized soft image embedding C [17]. The input prompt is also encoded as the high-dimensional prompt embedding P by the tokenizer. The LLM incorporates both the image embedding and the prompt embedding from the user question to generate

rich and comprehensive textual descriptions. Specifically, given the output of Q-Former as soft image embedding C and prompt embedding P , the goal is to compute the probability of creating the final answer A with a length of N through the transformer model p_θ . The mathematical expression for this process is given by the equation [39]:

$$p(A|V, P) = \prod_{i=1}^N p_\theta(A_i|C, P, A_{<i}) \quad (1)$$

This equation captures the sequential nature of language generation, where each embedding in the answer is dependent on the preceding embeddings, as well as the visual and prompt embeddings. This probabilistic approach ensures that the generated text is not only accurate but also contextually coherent. We demonstrate the algorithm of our proposed model in Algorithm 1.

Algorithm 1: Algorithm of Multi-modal Foundation Model

Input: Image: The captured image

UserQuery: The user question

Output: OutputText: The generated output text

Step 1: Predict Tags

Image $\rightarrow$ Image Tagging Module $\rightarrow$ Tags

Step 2: Prompt Engineering for pBLV

Tags + UserQuery $\rightarrow$ Prompt Engineering for pBLV $\rightarrow$ Prompt

Step 3: Generate OutputText

Image + Prompt $\rightarrow$ Vision-Language Module $\rightarrow$ OutputText

4. Experiments

4.1. Implementation Details

Our system leverages the capabilities of a smartphone, employing a monocular phone camera to capture images and the phone's microphone to receive user voice questions, creating a seamless interaction between the user and the system as shown in Figure 3. The image and voice input are then transferred to our server, where the processing and generation of comprehensive descriptions take place. To convert the user's voice question into text for further processing, we employ Whisper [40], a powerful speech recognition system. This technology accurately transcribes the user's voice question into a textual form, enabling seamless integration with our vision-language model. After the input text is obtained, our system processes the image and text to generate detailed and contextually relevant output descriptions. The system selects the corresponding image frame once the user question is detected, ensuring accurate and timely responses. Here we utilize a LLM, i.e., Vicuna-13B [41] (Model settings include generating sentences with lengths ranging from 1 to 200 using beam search with a width of 5, applying a length penalty of 1, repetition penalty of 3, and temperature of 1), to generate the final output text. The output text is then transformed into audio format to provide a more accessible experience for individuals with visual disability. For text-to-speech conversion, we utilize the robust system Azure [13]. This allows us to transform the output text into clear and natural-sounding audio. The synthesized audio is then sent from the server to the user's phone, enabling real-time delivery of the assist environmental interaction information. By implementing this client-server architecture and incorporating speech recognition and synthesis technologies, our system facilitates seamless interaction between the user and our system.

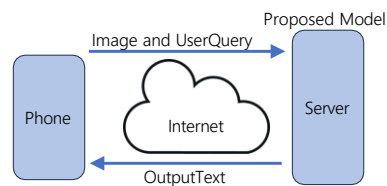


Figure 3. Client-server architecture.

4.2. Tests on Visual7W Dataset

We evaluate our proposed approach to the Visual7W dataset [36]. Compared with previous studies that solely relied on textual answers, Visual7W introduces a novel form of question answering that includes visual answers [36]. This is achieved by establishing a semantic connection through object-level grounding between textual descriptions and image regions [36]. We notice that there are strong connections between objects in images, both in terms of spatial location and meaning of existence. To test our model in assisting people with visual disability, we selected some images from specific perspectives in this dataset. From these perspectives, pBLV often require additional assistance to better understand the current environment. In order to better adapt to the needs, we have set this task into three categories: Scene Understanding, Object Localization and Risk Assessment.

4.2.1. Qualitative Performance Analysis for pBLV

Scene Understanding: We evaluate the effectiveness of our approach on outdoor and indoor scene understanding. Sample results are shown at the top of Figure 4. In our experiment, the user's input is "Can you describe the environment around?". For both indoor and outdoor examples, it is evident that the model's output provides a comprehensive and accurate description of the object composition in the environment depicted in the image. The answer first summarizes names associated with the current place and then gives a specific description of objects and characters in the scene and what is happening at this moment.

Object Localization: We evaluate the effectiveness of our approach in addressing object recognition challenges, as demonstrated in the middle of Figure 4. The user question for this task is "Where is the {giraffe, sheep, bookshelf, rubbish bin} in the image?", where "{" is what the user wants to find out.

In the outdoor scene, the left image is focused on the giraffe. From the answer, we can see that the results are very detailed, not only describing the location of the giraffe on the grass and under the trees, but also providing contextual information "The giraffes appear to be enjoying the shade provided by the tree and the lush green environment around them." for users to better understand the capture environmental images.

Risk Assessment: As shown at the bottom of Figure 4, our model provides safety tips for people with visual disability to help them identify and deal with potential risks according to the current environment. The question is "Is there a risk for me to continue moving forward?".

The first picture depicts a scene where a pedestrian crossing has a red light. The model can provide feedback to the user regarding the risk of crossing the street when the traffic signal is red. In the second scene, a train is approaching, which can be extremely dangerous if proper precautions are not taken. The model can send an alert that it is risky to cross the railway at the current time. It demonstrates that our model can effectively analyze risks and provide necessary alerts for pBLV.

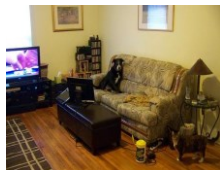
Scene Understanding			
Outdoor Scene		Indoor Scene	
			
Q: Can you describe the environment around?			
A: This place is a park, and the environment around it is surrounded by bushes and hedges. The couple is sitting on a bench in the park while having a conversation.		A: This is a college classroom, where students are sitting at tables and working on their laptops.	
A: This place is a clothing store, and it has a horse cart in front of it. The environment around the store is a busy city street with people walking on the sidewalk.		A: This is an office desk with a computer, monitor, keyboard, mouse, and speakers on it.	
Object Localization			
Outdoor Scene		Indoor Scene	
			
Q: Where is the {giraffe, sheep, bookshelf, rubbish bin} in front of me?			
A: There are several giraffes standing under a large tree in a grassy field. The giraffes are spread out across the scene, with some closer to the tree and others further away. The giraffes appear to be enjoying the shade provided by the tree and the lush green environment around them.		A: A herd of sheep is grazing on a rocky hillside. The sheep are spread out across the hill, with some closer to the bottom and others higher up. There are at least three sheep visible in the scene, with two of them closer to the bottom of the hill and one higher up. The rocky hillside provides a natural habitat for the sheep, and they seem to be comfortable and content in their surroundings.	
A: The bookshelf is located in a living room with a couch, a television, and a cat. The bookshelf is positioned in the corner of the room, near the couch and the television.		A: There is a rubbish bin located in a bathroom. The bin is positioned near the toilet and sink, making it easily accessible for cleaning and disposing of waste.	
Risk Assessment			
			
Q: Is there a risk for me to continue moving forward?			
A: It is risky to cross this street because there is a red light		A: It is risky to walk on a wet sidewalk because it can be slippery	
A: It is risky to cross the train tracks while the train is coming		A: It is risky to walk on the sidewalk because there is a yellow fire hydrant in the middle of the sidewalk	

Figure 4. Examples of scene understanding (top), object localization (middle), and risk assessment (bottom) on Visual7W dataset.

4.2.2. Quantitative Analysis of Inference Time and Helpfulness Scoring for pBLV

In this section, we conduct a statistical quantitative analysis to evaluate our model's performance across three types of tasks. Specifically, we tested the inference time (in seconds) of the image tagging module and the vision-language module. As indicated in Table 1, inference time testing was performed using an NVIDIA RTX A6000 GPU(48 G). In our vision-language model, we utilize a byte-stream transmission approach during answer generation. This method allows users to receive the initial part of the answer even while the model is still generating subsequent segments. Consequently, we measure the inference time when the vision-language module generates the first token. The results demonstrate that our model can respond swiftly to user queries, with the total inference time (Image Tagging time plus Vision-language Inference time) for Object Localization being the fastest, at less than 0.3 s. These times hold potential for further reduction through software optimization.

Furthermore, it's important to assess how well the model's generated answers relate to and assist pBLV. Therefore, we employ a scoring system ranging from 0 to 10, where higher scores indicate greater relevance and helpfulness for pBLV. To determine these

scores, we manually counted the number and names of important objects in each scene and compared them with the answers generated by the model. A point was deducted for each less important object identified in the generated answer. The final scores across all three tasks—Scene Understanding, Object Localization, and Risk Assessment—are reported in the table. Notably, Risk Assessment received the highest average score of 9.4, underscoring the model’s effectiveness in providing relevant and helpful information to pBLV.

Table 1. Quantitative Results of Inference time and Scoring on the Visual7W Dataset

Tasks	Inference Time (Seconds)		Score of 10
	Image Tagging	Vision-Language Inference	
Scene Understanding	0.0419	0.3675	8.85
Object Localization	0.0399	0.2359	8.60
Risk Assessment	0.0356	0.2406	9.40

4.2.3. Ablation Study

We conduct ablation studies to verify the effectiveness of the individual module in our model. The experimental settings are listed in Figure 5 where “✓” denotes the module is enabled. In the first experimental setting, we only utilize the vision-language module, which directly sends user questions and images to InstructBLIP. In the second experimental setting, we employ the image tagging module to generate tags for the input image, which are then integrated into the user question. Then, both the modified question and the input image are fed into the vision-language module. In the third experimental setting, we employ prompt engineering specifically designed for individuals with visual disability to further refine the prompt by incorporating the generated tags and user questions.

Model Setting				
Vision Language Module	Image Tagging Module	Prompt Engineering for pBLV	User Input: Can you describe the environment around?	
✓			In the image, there is a black and white cat sitting inside an open suitcase. The cat appears to be curious about its surroundings as it peers out of the suitcase. The suitcase is placed on a white surface, possibly a table or a shelf. There are other items in the room, such as a pillow, but the cat seems to be more interested in exploring the interior of the suitcase.	In the image, a man in a yellow vest is walking down a sidewalk next to a brick wall. There are several cars parked on the street, and a traffic light can be seen in the background. The scene appears to take place in a residential or commercial area, as there are both houses and businesses visible in the vicinity.
✓	✓		a black and white cat is hiding behind an open suitcase	a man wearing a safety vest walks down the side of a brick building
✓	✓	✓	A black and white cat is hiding behind an open suitcase on a white floor.	The environment is a street corner with a brick building on one side and a road on the other side. There is a guard walking down the pavement, wearing a green safety vest and carrying a baseball hat in his hand.

Figure 5. Ablation study with different model settings on Visual7W dataset.

As shown in Figure 5, for the left scene, when only using the vision-language module, the model provides some answers that do not match the facts, such as “*cat sitting inside an open suitcase*”, “*possibly a table or a shelf*” and “*such as a pillow*”. These are more likely to be inferred from a large language model due to model’s bias (learned from the training data) than what actually exists in the image. After combining with the image tagging module, the model dropped answers that do not match the facts in the image and the generated

answer correctly describes the current scenario. The pipeline should describe the modified user questions after integrating them with the tags. Furthermore, if prompt engineering for pBLV is applied, answers become more precise e.g., it also accurately describes the location of the cat. Here you should describe what is the final prompt generated.

In the case of the right scene, the model that only uses the vision-language module does provide a detailed description of the scene, but there are still errors. The description “There are several cars parked on the street, and a traffic light can be seen in the background.” is inconsistent with the facts shown in the image, as there is no traffic light and only a white street light and an orange fire hydrant. When adding the image tagging module, the model gives a more factual description but lacks details. Again you should provide the new user question based on the tags. In contrast, Prompt Engineering for pBLV makes the answer more precise and detailed. Again you should provide the new user question after prompt engineering.

The example in Figure 5 demonstrates that the integration of the vision-language module, image tagging module, and prompt engineering yield the most accurate and detailed descriptions. In Figure 6, we further present some randomly selected results of the image tagging module. As depicted in the figure, the module successfully recognizes common objects within the images, demonstrating its ability to provide a comprehensive understanding of the visual scene by accurately identifying various objects.



Figure 6. Random selected results of image tagging module on Visual7W dataset.

4.3. Tests on VizWiz Dataset

The VizWiz dataset [34] is a collection of images taken by blind and individuals with visual disability, specifically designed to evaluate computer vision algorithms aimed at assisting individuals with visual disability. In this section, we conduct experiments on the task of Visual Question Answering [42] of VizWiz dataset to verify the effectiveness of our proposed method.

We tested 4319 questions from the validation dataset, which can be categorized into four types, “Unanswerable”, “Other”, “Yes/No” and “Number”. As for evaluation metrics, we use the BLEU [43], ROUGE-L [44], METEOR [45], CIDEr [46] metrics to evaluate our results. Since the answer to Visual Question Answer is usually shorter than 4 words, we evaluate our results on BLEU_1 and BLEU_2. From Table 2, our model averagely achieves quantitative results of 25.43 in BLEU_1 and 53.98 in CIDEr, indicating the effectiveness for pBLV.

Table 2. Quantitative Results on Visual Question Answering. Column Q represents four question types and average results.

Q	BLEU_1	BLEU_2	METEOR	ROUGE_L	CIDEr
Unanswerable	8.53	3.33	7.88	13.18	5.65
Other	34.27	20.75	23.22	44.95	74.72
Yes/No	50.94	9.10	38.24	73.90	112.82
Number	11.41	5.82	12.73	16.95	27.48
Avg.	25.43	14.52	19.09	35.76	53.98

4.4. Real-World Tests

We also conducted experiments to evaluate the proposed system in real-world situations as shown in Figure 7. Specifically, we simulate the walking process of pBLV. The main content is a real video of a person walking on the street entering and then exiting a subway station in New York. Even though this route is quite simple in its geometry and layout, it may engender many risks for ppBLV. We captured several characteristic images from this video and passed them into our model for evaluation. These scenes are on the street, before entering the station, in the subway station, and after exiting the station.

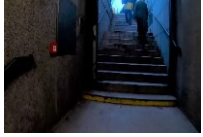
Real-world Tests			
Scene Understanding	Risk Assessment	Object Localization	Risk Assessment
			
Q: Can you describe the environment around?	Q: Is there a risk for me to continue moving forward?	Q: Where is the subway gate in front of me?	Q: Is there a risk for me to continue moving forward?
A: A crowded shopping street is filled with people walking and strolling along the pavement.	A: No, it is not risky for you to go ahead.	A: The subway gate is located inside a subway station. It is positioned at the end of a hallway or corridor within the station. The gate appears to be open, allowing people to enter or exit the subway station through it. There are two turnstiles, one on the left side and one on the right side of the gate, allowing multiple people to pass through the gate at the same time. The turnstile on the left side is closer to the wall, while the turnstile on the right side is closer to the center.	A: be careful about going up the stairs because they are narrow

Figure 7. Examples of scene understanding, object detection and localization, and risk assessment under real-world settings.

The first scene shows a street with a crowd. Moreover, there is a shop on the left of the image. Our model returns the answer *“A crowded shopping street is filled with people walking and strolling along the pavement.”*, which is consistent with the image content. Our model identified the presence of shops on both sides and a large crowd, inferring that this is a bustling street scene. This highlights the importance of caution and slow navigation for individuals with visual disability amidst the bustling crowd.

For the second scene, the user is walking on a straight and empty street, when the user asks *“Is there a risk for me to continue moving forward?”* The model answers that *“No, it is not risky for you to go ahead.”*, which is also in line with the actual situation. From the picture, we can see that the protagonist is walking on the sidewalk. There is mostly an open area on the sidewalk, with only tables and chairs on the right side. Therefore, the model determines that the road on this side can be walked, which means there is no risk.

In the third scene, when the user asks where the subway gate is, the model provides a very detailed explanation of the location of the subway gate with directional adjectives, such as front, back, left, and right. From the answer, it is clear that there are two gates and they are on the left and right sides.

In the last scenario where there is a staircase, the model reminds the user that there currently exists a certain level of risk due to the presence of stairs. Therefore, the model provides this answer *“be careful about going up the stairs because they are narrow”*. This indicates that our model can effectively assess whether the current situation poses some risk to individuals with visual disability in moving forward.

5. Conclusions

In this paper, we present a pioneering approach that addresses the challenges faced by people with blindness and low vision (pBLV) in comprehensive scene understanding, precise object localization, and risk assessment in unfamiliar environments. By leveraging a large vision language model and integrating it with an image tagging module, our method

provides pBLV with detailed and comprehensive descriptions and guidance to address their specific needs. We evaluate our approach through experiments conducted on both indoor and outdoor datasets. Our results demonstrate that our method can recognize objects accurately and provides insightful descriptions and analysis for pBLV.

Theoretical contributions of our work include a thorough exploration of the integration of vision language models with deep learning techniques for assisting individuals with visual impairments. On the other hand, practical contributions encompass the development of a functional system capable of providing detailed scene descriptions and guidance to enhance the mobility and independence of individuals with visual impairments. Enhancing the mobility of individuals with visual impairments is of paramount importance in fostering inclusivity and independence. These advancements will not only improve the practical utility of our system but also contribute to the theoretical understanding of assistive technologies for individuals with visual impairments, paving the way for more inclusive and accessible environments.

6. Limitations and Future Research

Despite these advances, we acknowledge certain limitations inherent to our system. The dynamic and complex nature of real-world environments poses a significant challenge. Factors such as changing lighting conditions, weather variations, and the presence of moving objects can impact the system's ability to accurately interpret scenes and predict potential risks, occasionally leading to false alarms. Moreover, the effectiveness of our system is contingent upon the quality of input data. Images captured by smartphone cameras under suboptimal conditions, such as poor lighting or obstructions, can adversely affect the model's performance, resulting in inaccurate recognitions or missed detections.

Inherent Limitations of AI Models: Despite the advanced capabilities of our Recognize Anything Model (RAM) and the InstructBLIP vision-language model, AI-based systems are not infallible. They operate within the confines of their training data and algorithms, which might not cover every possible real-world scenario or object encountered by people with blindness and low vision (pBLV). This limitation can lead to inaccuracies or false positives in object detection and scene interpretation.

Dynamic and Complex Environments: Real-world environments are highly dynamic and complex, with constant changes that can challenge the model's ability to accurately interpret and predict risks. Factors such as varying lighting conditions, weather changes, and moving objects can affect the system's performance and potentially lead to false alarms.

Quality of Input Data: The effectiveness of our system heavily relies on the quality of the input data, i.e., the images captured by the smartphone camera. Blurred images, poor lighting conditions, or obstructed views can hinder the model's ability to accurately recognize and analyze the scene, leading to potential false alarms or missed detections.

In future work, we aim to refine our approach by expanding training datasets to include a wider variety of environmental conditions, thereby enhancing the model's ability to generalize across different scenarios. We will also focus on developing more sophisticated algorithms that can more effectively distinguish between transient and permanent features in the environment, reducing false alarms. Additionally, integrating multimodal data sources, including auditory and haptic feedback, will be explored to compensate for visual data limitations and improve scene analysis. Incorporating user feedback mechanisms will further allow the system to adapt and learn from real-world applications, continually improving its performance and reliability for assisting individuals with blindness and low vision.

Author Contributions: Conceptualization, Y.H., F.Y. and Y.F.; methodology, Y.H., F.Y., H.H., Y.W. and Y.F.; software, Y.H. and F.Y.; validation, Y.H., F.Y., H.H., S.Y., S.R., J.-R.R., Y.W. and Y.F.; writing—original draft preparation, Y.H. and F.Y.; writing—review and editing, Y.H., F.Y., H.H., S.Y., S.R., J.-R.R., Y.W. and Y.F.; supervision, Y.F.; project administration, Y.F.; funding acquisition, J.-R.R., Y.W., Y.F. and S.R. All authors have read and agreed to the published version of the manuscript.

Funding: This study was supported by the National Science Foundation under Award Number ITE2345139, and number R33EY033689. The content is solely the responsibility of the authors and does not necessarily represent the official views of the National Institutes of Health and the National Science Foundation.

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Data Availability Statement: The raw data supporting the conclusions of this article will be made available by the authors on request.

Conflicts of Interest: The authors declare no conflict of interest.

References

1. Pascolini, D.; Mariotti, S.P. Global estimates of visual impairment: 2010. *Br. J. Ophthalmol.* **2012**, *96*, 614–618. [[CrossRef](#)] [[PubMed](#)]
2. Hakobyan, L.; Lumsden, J.; O'Sullivan, D.; Bartlett, H. Mobile assistive technologies for the visually impaired. *Surv. Ophthalmol.* **2013**, *58*, 513–528. [[CrossRef](#)] [[PubMed](#)]
3. World Health Organization. *Visual Impairment and Blindness Fact Sheet N 282*; World Health Organization: Geneva, Switzerland, 2014.
4. Massiceti, D.; Hicks, S.L.; van Rheede, J.J. Stereoscopic vision: Exploring visual-to-auditory sensory substitution mappings in an immersive virtual reality navigation paradigm. *PLoS ONE* **2018**, *13*, e0199389. [[CrossRef](#)] [[PubMed](#)]
5. Treisman, A.M.; Gelade, G. A feature-integration theory of attention. *Cogn. Psychol.* **1980**, *12*, 97–136. [[CrossRef](#)] [[PubMed](#)]
6. MacKeben, M.; Fletcher, D.C. Target search and identification performance in low vision patients. *Investig. Ophthalmol. Vis. Sci.* **2011**, *52*, 7603–7609. [[CrossRef](#)] [[PubMed](#)]
7. Fernandes, H.; Costa, P.; Filipe, V.; Paredes, H.; Barroso, J. A review of assistive spatial orientation and navigation technologies for the visually impaired. *Univers. Access Inf. Soc.* **2019**, *18*, 155–168. [[CrossRef](#)]
8. Kameswaran, V.; Gupta, J.; Pal, J.; O'Modhrain, S.; Veinot, T.C.; Brewer, R.; Parameshwar, A.; O'Neill, J. 'We can go anywhere' Understanding Independence through a Case Study of Ride-hailing Use by People with Visual Impairments in metropolitan India. *Proc. ACM Hum.-Comput. Interact.* **2018**, *2*, 1–24. [[CrossRef](#)]
9. Roentgen, U.R.; Gelderblom, G.J.; Soede, M.; De Witte, L.P. Inventory of electronic mobility aids for persons with visual impairments: A literature review. *J. Vis. Impair. Blind.* **2008**, *102*, 702–724. [[CrossRef](#)]
10. Loomis, J.M.; Golledge, R.G.; Klatzky, R.L.; Speigle, J.M.; Tietz, J. Personal guidance system for the visually impaired. In Proceedings of the first Annual ACM Conference on Assistive Technologies, Marina Del Rey, CA, USA, 31 October–1 November 1994; pp. 85–91.
11. Redmon, J.; Farhadi, A. Yolov3: An incremental improvement. *arXiv* **2018**, arXiv:1804.02767.
12. Pace, S.; Frost, G.; Lachow, I.; Frelinger, D.; Fossum, D.; Wassem, D.; Pinto, M. *The Global Positioning System*; Prepared for the Executive Office of the President, Office of Science and Technology Policy; RAND: Santa Monica, CA, USA, 1995.
13. Li, N.; Liu, S.; Liu, Y.; Zhao, S.; Liu, M. Neural speech synthesis with transformer network. In Proceedings of the AAAI Conference on Artificial Intelligence, Honolulu, HI, USA, 27 January–1 February 2019; Volume 33, pp. 6706–6713.
14. Montello, D.R. Cognitive research in GIScience: Recent achievements and future prospects. *Geogr. Compass* **2009**, *3*, 1824–1840. [[CrossRef](#)]
15. Radlinski, F.; Craswell, N. A theoretical framework for conversational search. In Proceedings of the 2017 Conference on Conference Human Information Interaction and Retrieval, Oslo, Norway, 7–11 March 2017; pp. 117–126.
16. Zhang, Y.; Huang, X.; Ma, J.; Li, Z.; Luo, Z.; Xie, Y.; Qin, Y.; Luo, T.; Li, Y.; Liu, S.; et al. Recognize Anything: A Strong Image Tagging Model. *arXiv* **2023**, arXiv:2306.03514.
17. Li, J.; Li, D.; Savarese, S.; Hoi, S. Blip-2: Bootstrapping language-image pre-training with frozen image encoders and large language models. *arXiv* **2023**, arXiv:2301.12597.
18. Trippas, J.R.; Spina, D.; Thomas, P.; Sanderson, M.; Joho, H.; Cavedon, L. Towards a model for spoken conversational search. *Inf. Process. Manag.* **2020**, *57*, 102162. [[CrossRef](#)]
19. Kaushik, A.; Bhat Ramachandra, V.; Jones, G.J. An interface for agent supported conversational search. In Proceedings of the 2020 Conference on Human Information Interaction and Retrieval, Virtual, 14–18 March 2020; pp. 452–456.
20. Kaushik, A.; Jacob, B.; Velavan, P. An exploratory study on a reinforcement learning prototype for multimodal image retrieval using a conversational search interface. *Knowledge* **2022**, *2*, 116–138. [[CrossRef](#)]
21. Giudice, N.A.; Legge, G.E. Blind navigation and the role of technology. In *The Engineering Handbook of Smart Technology for Aging, Disability, and Independence*; John Wiley & Sons, Inc.: Hoboken, NJ, USA, 2008; pp. 479–500.
22. Boldini, A.; Rizzo, J.R.; Porfiri, M. A piezoelectric-based advanced wearable: Obstacle avoidance for the visually impaired built into a backpack. In Proceedings of the Nano-, Bio-, Info-Tech Sensors, and 3D Systems IV, Bellingham, WA, USA, 27 April–8 May 2020; Volume 11378, p. 1137806.
23. McDaniel, T.; Krishna, S.; Balasubramanian, V.; Colbry, D.; Panchanathan, S. Using a haptic belt to convey non-verbal communication cues during social interactions to individuals who are blind. In Proceedings of the 2008 IEEE International Workshop on Haptic Audio Visual Environments and Games, Lecco, Italy, 7–8 November 2008; pp. 13–18.

24. Whitmarsh, L. The benefits of guide dog ownership. *Vis. Impair. Res.* **2005**, *7*, 27–42. [\[CrossRef\]](#)
25. Gupta, T.; Li, H. Indoor mapping for smart cities—An affordable approach: Using Kinect Sensor and ZED stereo camera. In Proceedings of the 2017 International Conference on Indoor Positioning and Indoor Navigation (IPIN), Sapporo, Japan, 18–21 September 2017; pp. 1–8.
26. Rizzo, J.R.; Beheshti, M.; Fang, Y.; Flanagan, S.; Giudice, N.A. COVID-19 and visual disability: Can't look and now don't touch. *PM&R* **2021**, *13*, 415–421.
27. Hao, Y.; Feng, J.; Rizzo, J.R.; Wang, Y.; Fang, Y. Detect and Approach: Close-Range Navigation Support for People with Blindness and Low Vision. In *Proceedings of the European Conference on Computer Vision, ECCV 2022*, Tel Aviv, Israel, 23–27 October 2022; Springer: Cham, Switzerland, 2022; pp. 607–622.
28. Boldini, A.; Garcia, A.L.; Sorrentino, M.; Beheshti, M.; Ogedegbe, O.; Fang, Y.; Porfiri, M.; Rizzo, J.R. An inconspicuous, integrated electronic travel aid for visual impairment. *ASME Lett. Dyn. Syst. Control.* **2021**, *1*, 041004. [\[CrossRef\]](#)
29. Granquist, C.; Sun, S.Y.; Montezuma, S.R.; Tran, T.M.; Gage, R.; Legge, G.E. Evaluation and comparison of artificial intelligence vision aids: OrCam myeye 1 and seeing ai. *J. Vis. Impair. Blind.* **2021**, *115*, 277–285. [\[CrossRef\]](#)
30. Kumar, P.A.; Vivek, J.; Senniagiri, N.; Nagarajan, S.; Chandrasekaran, K. A study of added sic powder in kerosene for the blind square hole machining of cfrp using electrical discharge machining. *Silicon* **2022**, *14*, 1831–1849. [\[CrossRef\]](#)
31. Cardoso, Q.; de Melo, A.V.; Orué, A.L. Accessibility analysis for the visually impaired using LazarilloApp. *Int. J. Innov. Educ. Res.* **2019**, *10*, 21–30.
32. Redmon, J.; Divvala, S.; Girshick, R.; Farhadi, A. You only look once: Unified, real-time object detection. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Las Vegas, NV, USA, 26 June–1 July 2016; pp. 779–788.
33. Ren, S.; He, K.; Girshick, R.; Sun, J. Faster r-cnn: Towards real-time object detection with region proposal networks. *Adv. Neural Inf. Process. Syst.* **2015**, *28*, 91–99. [\[CrossRef\]](#) [\[PubMed\]](#)
34. Bigham, J.P.; Jayant, C.; Ji, H.; Little, G.; Miller, A.; Miller, R.C.; Tatarowicz, A.; White, B.; White, S.; et al. Vizwiz: Nearly real-time answers to visual questions. In Proceedings of the 23rd Annual ACM Symposium on User Interface Software and Technology, New York, NY, USA, 3–6 October 2010; pp. 333–342.
35. Aletta, F.; Kang, J.; Axelsson, Ö. Soundscape descriptors and a conceptual framework for developing predictive soundscape models. *Landsc. Urban Plan.* **2016**, *149*, 65–74. [\[CrossRef\]](#)
36. Zhu, Y.; Groth, O.; Bernstein, M.; Fei-Fei, L. Visual7W: Grounded Question Answering in Images. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Las Vegas, NV, USA, 26 June–1 July 2016.
37. Liu, Z.; Lin, Y.; Cao, Y.; Hu, H.; Wei, Y.; Zhang, Z.; Lin, S.; Guo, B. Swin transformer: Hierarchical vision transformer using shifted windows. In Proceedings of the IEEE/CVF International Conference on Computer Vision, Montreal, BC, Canada, 11–17 October 2021; pp. 10012–10022.
38. Dosovitskiy, A.; Beyer, L.; Kolesnikov, A.; Weissenborn, D.; Zhai, X.; Unterthiner, T.; Dehghani, M.; Minderer, M.; Heigold, G.; Gelly, S. An image is worth 16x16 words: Transformers for image recognition at scale. *arXiv* **2020**, arXiv:2010.11929.
39. Zhang, Z.; Zhang, A.; Li, M.; Zhao, H.; Karypis, G.; Smola, A. Multimodal chain-of-thought reasoning in language models. *arXiv* **2023**, arXiv:2302.00923.
40. Radford, A.; Kim, J.W.; Xu, T.; Brockman, G.; McLeavey, C.; Sutskever, I. Robust speech recognition via large-scale weak supervision. In Proceedings of the International Conference on Machine Learning, PMLR, Honolulu, HI, USA, 23–29 July 2023; pp. 28492–28518.
41. Chiang, W.L.; Li, Z.; Lin, Z.; Sheng, Y.; Wu, Z.; Zhang, H.; Zheng, L.; Zhuang, S.; Zhuang, Y.; Gonzalez, J.E.; et al. Vicuna: An Open-Source Chatbot Impressing gpt-4 with 90%* Chatgpt Quality. 2023. Available online: <https://vicuna.lmsys.org> (accessed on 14 April 2023).
42. Gurari, D.; Li, Q.; Stangl, A.J.; Guo, A.; Lin, C.; Grauman, K.; Luo, J.; Bigham, J.P. Vizwiz grand challenge: Answering visual questions from blind people. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Salt Lake City, UT, USA, 18–22 June 2018; pp. 3608–3617.
43. Papineni, K.; Roukos, S.; Ward, T.; Zhu, W.J. Bleu: A method for automatic evaluation of machine translation. In Proceedings of the 40th Annual Meeting of the Association for Computational Linguistics, Philadelphia, PA, USA, 6–12 July 2002; pp. 311–318.
44. Lin, C.Y. Rouge: A package for automatic evaluation of summaries. In Proceedings of the Text Summarization Branches Out, Barcelona, Spain, 25–26 July 2004; pp. 74–81.
45. Banerjee, S.; Lavie, A. METEOR: An automatic metric for MT evaluation with improved correlation with human judgments. In Proceedings of the ACL Workshop on Intrinsic and Extrinsic Evaluation Measures for Machine Translation and/or Summarization, 2005; pp. 65–72. Available online: <https://aclanthology.org/W05-0909/> (accessed on 14 April 2024).
46. Vedantam, R.; Lawrence Zitnick, C.; Parikh, D. Cider: Consensus-based image description evaluation. In Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition, Boston, MA, USA, 7–12 June 2015; pp. 4566–4575.

Disclaimer/Publisher's Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.