

# CCHall: A Novel Benchmark for Joint Cross-Lingual and Cross-Modal Hallucinations Detection in Large Language Models

Yongheng Zhang<sup>1,2\*</sup> Xu Liu<sup>1\*</sup> Ruoxi Zhou<sup>1</sup> Qiguang Chen<sup>1</sup>

Hao Fei<sup>3</sup> Wenpeng Lu<sup>4</sup> Libo Qin<sup>1,2†</sup>

<sup>1</sup>School of Computer Science and Engineering, Central South University, China

<sup>2</sup>Key Laboratory of Data Intelligence and Advanced Computing in Provincial Universities,

<sup>2</sup>Soochow University, China <sup>3</sup>National University of Singapore, Singapore

<sup>4</sup>Key Laboratory of Computing Power Network and Information Security, Ministry of Education

<sup>4</sup>Qilu University of Technology (Shandong Academy of Sciences), China

## Abstract

Investigating hallucination issues in large language models (LLMs) within cross-lingual and cross-modal scenarios can greatly advance the large-scale deployment in real-world applications. Nevertheless, the current studies are limited to a single scenario, either cross-lingual or cross-modal, leaving a gap in the exploration of hallucinations in the joint cross-lingual and cross-modal scenarios. Motivated by this, we introduce a novel joint Cross-lingual and Cross-modal **H**allucinations benchmark (CCHall) to fill this gap. Specifically, CCHall simultaneously incorporates both cross-lingual and cross-modal hallucination scenarios, which can be used to assess the cross-lingual and cross-modal capabilities of LLMs. Furthermore, we conduct a comprehensive evaluation on CCHall, exploring both mainstream open-source and closed-source LLMs. The experimental results highlight that current LLMs still struggle with CCHall. We hope CCHall can serve as a valuable resource to assess LLMs in joint cross-lingual and cross-modal scenarios.

## 1 Introduction

Large language models (LLMs) have made significant progress in recent years (Meta, 2024; Qin et al., 2024; Wang et al., 2025), driving remarkable advancements across a wide range of diverse fields and applications (Wei et al., 2022; Zhang et al., 2025). However, it is disheartening that the issue of hallucinations in LLMs remains unresolved, significantly hindering their large-scale deployment in real-world applications (Park et al., 2024; Yu et al., 2025). This phenomenon is particularly severe in both cross-lingual and cross-modal contexts, potentially causing critical errors in applications such as medical diagnosis, image captioning, and speech-to-text conversion (Guerreiro et al., 2023; Regan et al., 2024; Sriramanan et al., 2025).

\* Equal Contribution. † Corresponding Author.

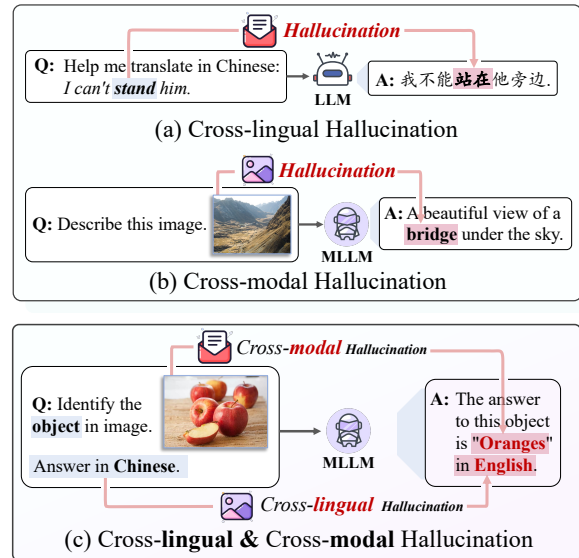
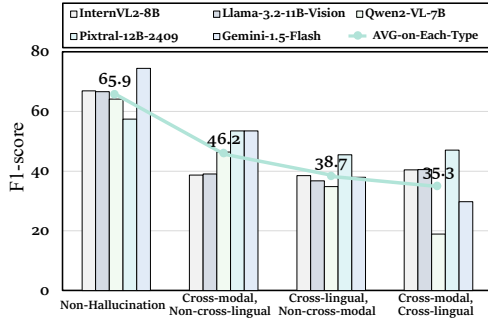


Figure 1: (a) Cross-lingual hallucination: A cross-lingual hallucination emerged: the erroneous translation of “stand” as “站在”. Here it should be “忍受”; (b) Cross-modal hallucination: A cross-modal hallucination occurred, fabricating a “bridge”; (c) Cross-lingual and Cross-modal hallucination: A cross-modal hallucination fabricated “Oranges” and a cross-lingual hallucination did not use Chinese in its Answer.

Recent research has focused on cross-lingual and cross-modal hallucinations in LLMs (Son et al., 2024; Guan et al., 2024; Bai et al., 2024), aiming to understand their causes, including:

(1) **Cross-lingual Hallucinations:** As shown in Figure 1 (a), Cross-lingual hallucination refers to the failure to follow instructions or the generation of inaccurate outputs when processing different languages. Specifically, Qiu et al. (2023) introduce mFACT to evaluate the faithfulness of summaries and observe that LLMs hallucinate more in non-English languages. Dale et al. (2023a) release an annotated dataset covering 18 translation directions, tackling hallucinations. Herrlein et al. (2024) further extend English hallucination detection to German and apply it in long-context scenarios.



(a) F1-score on different hallucination types.

| Benchmark                          | Methods       |             |                         |
|------------------------------------|---------------|-------------|-------------------------|
|                                    | Cross-lingual | Cross-modal | Cross-(Lingual & Modal) |
| XL-Sum (Hasan et al., 2021)        | ✓             | ✗           | ✗                       |
| X-fact (Gupta and Srikumar, 2021)  | ✓             | ✗           | ✗                       |
| HalOmi (Dale et al., 2023b)        | ✓             | ✗           | ✗                       |
| MM-Eval (Son et al., 2024)         | ✓             | ✗           | ✗                       |
| XTRUST (Li et al., 2024)           | ✓             | ✗           | ✗                       |
| CHAIR (Rohrbach et al., 2018)      | ✗             | ✓           | ✗                       |
| POPE (Li et al., 2023)             | ✗             | ✓           | ✗                       |
| M-HalDetect (Gunjal et al., 2024)  | ✗             | ✓           | ✗                       |
| HallusionBench (Guan et al., 2024) | ✗             | ✓           | ✗                       |
| MHalBench (Chen et al., 2024b)     | ✗             | ✓           | ✗                       |
| <b>CCHall (ours)</b>               | ✓             | ✓           | ✓                       |

(b) Comparison of CCHall with related datasets.

Figure 2: (a) Fine-grained performance analysis of MLLMs F1-score for different hallucination types in CCHall. The F1-score of MLLMs on joint cross-lingual and cross-modal hallucinations is 3.4 lower than when addressing cross-lingual hallucinations independently and 10.9 lower than when addressing cross-modal hallucinations independently; (b) A comparative analysis of CCHall with related cross-modal and cross-lingual datasets.

**(2) Cross-modal Hallucinations:** As shown in Figure 1 (b), Cross-modal hallucination refers to the fabrication and inconsistencies that arise in multimodal large language models (MLLMs) when reasoning across different modalities. Specifically, Liu et al. (2023) develop an innovative benchmark for image-context reasoning using image-question pairs created by human experts. Leng et al. (2024) conduct a detailed and systematic study of multimodal hallucinations and introduce the “Curse of Multi-Modalities” benchmark for MLLM evaluation. Yan et al. (2024) introduce a Quality Measurement framework, which aims to evaluate the validity of existing hallucination benchmarks, ensuring that they can accurately assess model outputs.

Despite significant advancements, current research on hallucinations in MLLMs remains overly optimistic, as it *primarily focuses* on either cross-lingual or cross-modal scenarios *in isolation*, leaving a gap in the joint cross-lingual and cross-modal context. In fact, compared to addressing cross-lingual or cross-modal scenarios individually, the joint cross-lingual and cross-modal context is more widely applied in real-world scenarios and presents even *greater challenges* (Qin et al., 2023; Chen et al., 2024a; Zhang et al., 2024a; Wang et al., 2024b; Castillo et al., 2025), as language differences and multimodal inputs increase hallucinations in MLLMs. As illustrated in Figure 1 (c), MLLMs must account not only for the alignment between images and text but also for the alignment across multiple language queries. As shown in Figure 2 (a), numerous models struggle with cross-lingual and cross-modal hallucinations, exhibiting poor performance. More frustratingly, there is currently no research to investigate this critical issue.

Motivated by this, we introduce a novel benchmark, **Cross-lingual and Cross-modal Hallucinations (CCHall)**, to fill the gap. As shown in Figure 2 (b), CCHall provides an evaluation framework that encompasses *not only* individual cross-lingual and cross-modal hallucinations *but also* joint Cross-lingual and Cross-modal Hallucinations. Furthermore, CCHall covers a wide range of topics and nature scenes, aiming to provide a comprehensive evaluation of MLLMs in cross-lingual and cross-modal hallucination scenarios. Through evaluation experiments on CCHall, we derive the following *key takeaways*: (1) Current MLLMs are still struggling in CCHall benchmark; (2) In hallucination mitigation methods, basic strategies are more suitable for MLLMs with fewer parameters (<12B), while powerful MLLMs exhibit advantages when using more advanced strategies; (3) The use of multilingual contexts and tool-assisted invocation can effectively mitigate hallucinations in MLLMs.

Our contributions to our work are as follows:

- (1) We identify and point out that current research on hallucinations tends to be overly optimistic, as it mainly addresses either cross-lingual or cross-modal scenarios separately.
- (2) To address the existing gap in the detection of joint cross-lingual and cross-modal hallucinations, we introduce **Cross-lingual and Cross-modal Hallucinations benchmark (CCHall)**.
- (3) We evaluate CCHall across a range of MLLMs and diverse scenarios, hoping to provide insights that could help mitigate cross-lingual and cross-modal hallucinations.

All data and source code are open-sourced at <https://github.com/BRZ911/CCHall>.

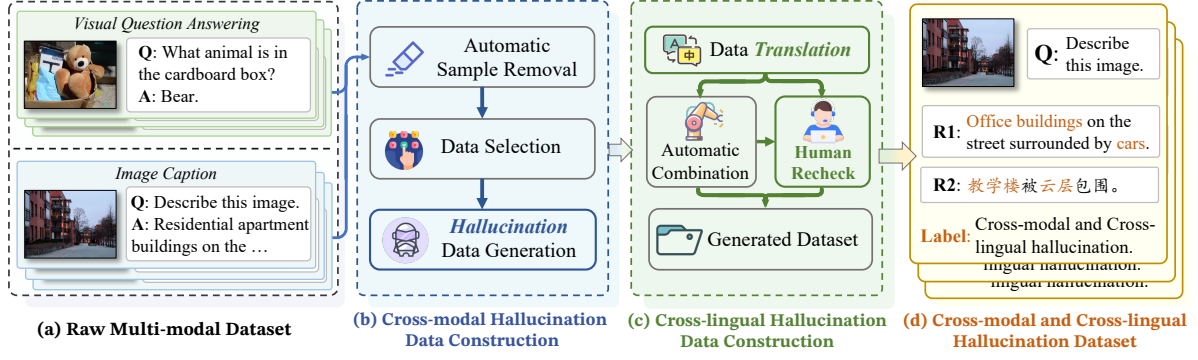


Figure 3: The construction process of CCHa11 includes: (a) Raw Multi-modal Dataset Selection (§3.1), (b) Cross-modal Hallucination Data Construction (§3.2), (c) Cross-lingual Hallucination Data Construction (§3.3), and (d) Cross-modal and Cross-lingual Hallucination Dataset (§3.4).

## 2 Background

### 2.1 Cross-lingual Hallucinations

Cross-linguistic hallucination (Dale et al., 2023b) occurs when the LLMs generate an answer  $\mathcal{A}_t$  in the target language  $L_t$  that deviates from the correct answer  $\mathcal{A}_{gold}$ . As shown in Figure 1 (a), given a question  $Q$  in the source language  $L_s$  and a prompt  $P_{s \rightarrow t}$  to respond in  $L_t$ , which is denoted as:

$$\mathcal{A}_t = \operatorname{argmax}_{\mathcal{A}} P(\mathcal{A} | Q, P_{s \rightarrow t}), \quad (1)$$

where  $P(\mathcal{A} | Q, P_{s \rightarrow t})$  is the probability of generating answer  $\mathcal{A}$  given  $Q$  and  $P_{s \rightarrow t}$ . Hallucination occurs when multilingual instructions are not followed or when incorrect reasoning is provided.

### 2.2 Cross-modal Hallucinations

Cross-modal hallucination (Rohrbach et al., 2018) occurs when the LLMs generate a textual answer  $\mathcal{A}_t$  that does not accurately reflect the content of an input image  $\mathcal{A}_{gold}$ . As shown in Figure 1 (b), given an image  $I$ , a question  $Q$ , and a prompt  $P$ , the model aims to output an answer  $\mathcal{A}$  that aligns with both  $I$  and  $Q$ , which is denoted as:

$$\mathcal{A}_t = \operatorname{argmax}_{\mathcal{A}} P(\mathcal{A} | Q, I, P), \quad (2)$$

where  $P(\mathcal{A} | Q, I, P)$  represents the probability of generating answer  $\mathcal{A}$  given  $I$ ,  $Q$ , and  $P$ . Cross-modal hallucination occurs when there is reasoning that does not correspond to the image.

### 2.3 Joint Cross-Lingual and Cross-Modal Hallucinations

Compared to cross-lingual and cross-modal hallucinations, joint Cross-Lingual and Cross-Modal

Hallucinations involve generating responses in multiple languages that are inconsistent with both the visual content and each other. As shown in Figure 1 (c), given an input image  $I$ , a question  $Q$ , and a prompt  $P_t$  to produce answers in target languages, the MLLMs aim to generate responses  $\mathcal{A}_t$  that accurately represent  $I$  in response to  $Q$  and are semantically consistent with their meanings:

$$\mathcal{A}_t = \operatorname{argmax}_{\mathcal{A}'_t} P(\mathcal{A}'_t | Q, I, P_t), \quad (3)$$

where  $P(\mathcal{A}'_t | Q, I, P_t)$  represents the joint probability of generating the answers given the inputs  $Q$ ,  $I$ , and the prompt  $P_t$ . Joint Cross-Lingual and Cross-Modal Hallucinations occur when the model does not respond in the target language as instructed, and the response includes reasoning that fails to reflect the content of the image.

## 3 The CCHa11 Benchmark

This section provides a detailed overview of the construction process of the CCHa11 benchmark, as illustrated in Figure 3. The process primarily consists of the following four key components:

### 3.1 Raw Multi-modal Dataset Selection

To evaluate the hallucination behaviors of MLLMs from multiple perspectives, we have integrated two widely used multimodal tasks, Visual Question Answering (VQA) and Image Captioning (IC) (Rohrbach et al., 2018; Li et al., 2023; Gunjal et al., 2024), as illustrated in Figure 3 (a).

**Visual Question Answering Task:** Our VQA task extends GQA (Hudson and Manning, 2019) and AMBER (Wang et al., 2023) dataset to assess reasoning and spatial understanding. GQA enriches

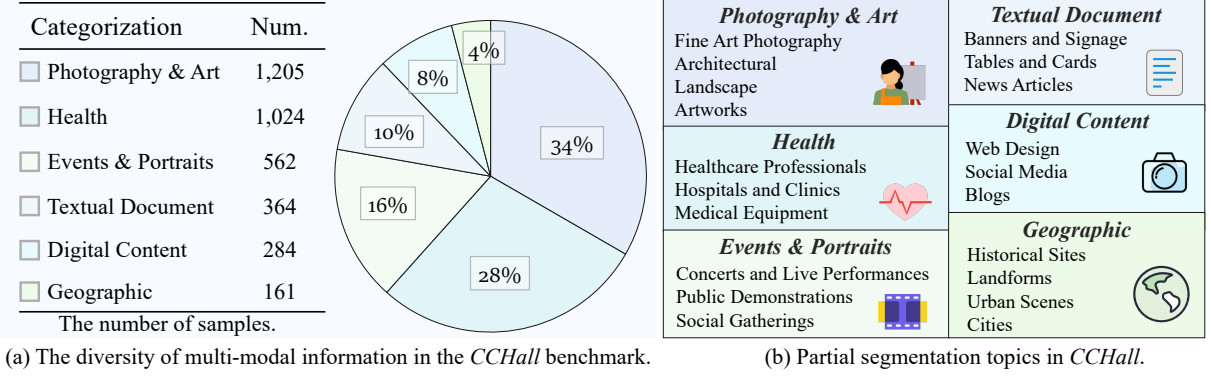


Figure 4: Presentation of data in CCHall: (a) The diversity of multi-modal data as represented by CLIP-based (Radford et al., 2021) classification. (b) Display of part of the detailed topics in CCHall.

scene descriptions, while AMBER covers 337 object categories. These datasets pose challenging VQA questions requiring advanced reasoning.

**Image Captioning Task:** The IC task evaluates hallucination detection in image descriptions, using the XM3600 (Thapliyal et al., 2022) and xFlickr&Co (Bugliarello et al., 2022) datasets. XM3600 features diverse global images, while xFlickr&Co focuses on everyday scenes, ensuring cultural diversity and minimal bias.

### 3.2 Cross-modal Hallucination Data Construction

As shown in Figure 3 (b), the process involves three steps to construct cross-modal hallucination data:

**Automatic Sample Removal:** To prepare the data for hallucination generation, we first filter out abnormal data, including instances of mismatched image-text pairs, images without corresponding text, or text without corresponding images, ensuring the dataset is clean and consistent.

**Data Selection:** We define cross-modal hallucination as significant discrepancies between the outputs of MLLMs and the actual presence of objects in input images. Therefore, we select questions from the GQA (Hudson and Manning, 2019) and AMBER (Wang et al., 2023) datasets related to object existence to assess the model’s accuracy in object reasoning. To minimize redundancy, we ensure that each object appears no more than twice. For the Image Captioning task, we utilize the XM3600 (Thapliyal et al., 2022) and xFlickr&Co (Bugliarello et al., 2022) datasets, with English as the reference language. Finally, we randomly sample 900 entries from each of the filtered datasets, yielding a total of 3600 entries, to construct the dataset for generating hallucinated data.

**Hallucination Data Generation:** We use the Gemini-1.5-Pro (Team et al., 2024) to generate hallucinated data by inserting semantically similar but incorrect entities into image captions, increasing the challenge of CCHall. Real questions and answers were prepared for each image, along with prompts for the model. Gemini-1.5-Pro compared the image with the real answers, embedding misleading nouns not present in the image to ensure natural-sounding hallucinations. For a detailed description, please refer to Appendix A.1.

### 3.3 Cross-lingual Hallucination Data Construction

As shown in Figure 3 (c), the process for constructing cross-lingual hallucination data is as follows:

**Data Translation:** Following Conneau et al. (2018); Heredia et al. (2024); Hasan et al. (2024), we apply machine translation to the text data and then conduct manual verification of the data. Specifically, following Guerreiro et al. (2023), we categorize the languages into low, medium, and high-resource groups based on the availability of resources and select the three languages with the lowest error rates from each group to maximize translation accuracy. The final selections are as follows: Croatian (HR), Welsh (CY), and Swahili (SW) as low-resource languages; Czech (CS), Dutch (NL), and Swedish (SV) as medium-resource languages; and French (FR), Spanish (ES), and Portuguese (PT) as high-resource languages. Google Translate is then used to translate the content into these target languages, resulting in an intermediate dataset.

**Automatic Combination:** After validating the intermediate dataset, we shuffle the data within each subset to eliminate ordering bias and ensure an even distribution of sample types. Next, we re-

fine the selection process by using English as the anchor language, paired with a randomly selected target language for each image. One response is in English, and the other is in a different language.

**Human Recheck:** After generating the hallucinated data and translations, we conduct a human review to ensure data quality and overall accuracy. We focus on two key aspects: (1) *Verification of Hallucinated Data:* We check if the generated data meets the hallucination criteria. This involves verifying that the sentences follow the expected structure, length, and logic and contain the intended elements aligned with the task. (2) *Accuracy of Translated Data:* We verify if the translated data accurately reflects the original meaning and context. This step ensures that translations remain faithful to the intended scenario, avoiding errors from language differences. The details of the check can be found in Appendix A.2.

### 3.4 Cross-modal and Cross-lingual Hallucination Dataset

The Cross-modal and Cross-lingual Hallucination Dataset (CCHall) is shown in Figure 3 (d). Please refer to the Appendix B.4 for specific examples. Specifically, we have retained the original dataset names as the categories. In total, the categories are divided into four types: *AMBER*, *GQA*, *xFlickr&Co*, and *XM3600*. Additionally, we define four types of combinations based on hallucinations:

- **Non-hallucination:** Both the English responses and the answers in the other language are correct, with all answers matching the contents in the images.
- **Cross-lingual, non-cross-modal hallucination:** The English answer is correct, but the non-English answer contains hallucinated objects not present in the images.
- **Cross-modal, non-cross-lingual hallucination:** All answers are hallucinatory, containing objects that do not appear in the images. Meanings in each language are identical.
- **Cross-modal, cross-lingual hallucination:** All the answers are hallucinated, containing various objects that are either absent from the images or incorrectly described. Moreover, the meanings conveyed in the non-English answers differ significantly from those in the English answers.

A more detailed presentation of the CCHall attributes can be found in Appendix A.3.

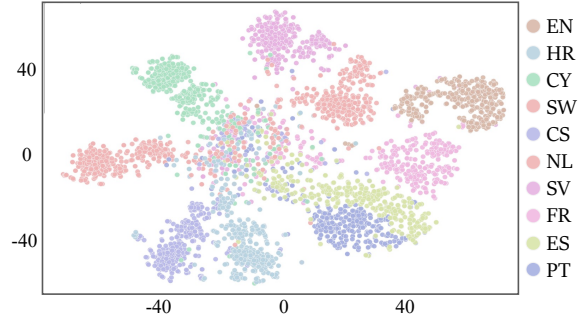


Figure 5: Visualization of the semantic feature coverage of all languages in CCHall, demonstrating the distribution and range of linguistic representations.

## 4 CCHall Analysis

This section presents an analysis of CCHall. **Multi-modal Diversity:** Following Chen et al. (2024a), we employ CLIP (Radford et al., 2021) to classify the images in the CCHall. The results, as shown in Figure 4, reveal that the images in CCHall exhibit a wide variety of types. This diversity enables a more comprehensive coverage of the hallucination scenarios demonstrated by LLMs. **Multi-lingual Diversity:** To demonstrate the diversity of multilingualism, we use the CLIP (Radford et al., 2021) text encoder to encode data in all languages and applied t-SNE for dimensionality reduction. As shown in Figure 5, the language features cover nearly the entire range of semantic feature representations. This indicates that CCHall exhibits broad linguistic richness, enabling comprehensive detection of multilingual hallucinations.

## 5 Experiments

### 5.1 Experiments Setting

We evaluate several MLLMs on the CCHall, including *GPT-4o* (OpenAI, 2023), *Gemini-1.5-Flash* (Team et al., 2024), *Llama-3.2-11B-Vision-Instruct* (Meta, 2024), *Qwen2-VL-7B-Instruct* (Wang et al., 2024a), *Pixtral-12B-2409* (Agrawal et al., 2024), and *InternVL2-8B* (Chen et al., 2024c). In addition to using `Direct` queries with MLLMs, we explore several strategies to mitigate hallucinations. Specifically, we implement `CoT` (Kojima et al., 2022) to promote step-by-step reasoning, `SRO` (Lin et al., 2024) for self-reflection, and `VDGD` (Ghosh et al., 2024) to embed detailed image descriptions and better align reasoning. We apply `HallucimAD` (Lin et al., 2024) to reduce hallucinations via multi-agent debate. All top-p and temperature parameters retain the default of MLLMs values within the specified range of [0, 1].

| Model                                                    | AMBER       |             | GQA         |             | xFlickr&CO  |             | XM3600      |             | AVG         |             |
|----------------------------------------------------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|
|                                                          | Acc         | Macro-F1    | Acc         | Macro-F1    | Acc         | Macro-F1    | Acc         | Macro-F1    | Acc         | Macro-F1    |
| Random                                                   | 25.1        | 30.0        | 25.0        | 29.4        | 24.9        | 30.3        | 25.1        | 29.9        | 25.0        | 29.9        |
| <i>InternVL2-8B (Chen et al., 2024c)</i>                 |             |             |             |             |             |             |             |             |             |             |
| Direct (Chen et al., 2024c)                              | 29.1        | 38.1        | 29.9        | 38.6        | 38.3        | 47.6        | 38.8        | 47.4        | 34.0        | 42.9        |
| CoT (Kojima et al., 2022)                                | 31.3        | <u>40.0</u> | 33.6        | <u>42.1</u> | <u>41.6</u> | 48.0        | <u>40.1</u> | 47.6        | <u>36.7</u> | <u>44.4</u> |
| SRO (Lin et al., 2024)                                   | 30.3        | <u>40.0</u> | 31.1        | 40.7        | 41.2        | <u>48.6</u> | 37.7        | 47.1        | 35.1        | 44.1        |
| VDGD (Ghosh et al., 2024)                                | <u>33.2</u> | <u>40.0</u> | 30.2        | 37.8        | 36.7        | 44.5        | 37.4        | 44.7        | 34.4        | 41.7        |
| HalluciMAD (Lin et al., 2024)                            | 29.9        | 38.6        | 30.0        | 39.0        | 37.9        | 45.9        | 39.6        | <u>47.7</u> | 34.3        | 42.8        |
| <i>Llama-3.2-11B-Vision-Instruct (Meta et al., 2024)</i> |             |             |             |             |             |             |             |             |             |             |
| Direct (Meta et al., 2024)                               | 31.6        | 38.8        | 32.1        | 38.9        | 35.4        | 43.2        | 43.3        | 49.4        | 35.6        | 42.6        |
| CoT (Kojima et al., 2022)                                | 32.0        | 40.6        | 34.3        | 40.9        | 43.6        | 51.6        | <u>46.4</u> | <u>54.0</u> | 39.1        | <u>46.8</u> |
| SRO (Lin et al., 2024)                                   | 32.1        | 40.6        | 34.4        | 40.9        | <u>43.7</u> | <u>51.7</u> | <u>46.4</u> | <u>54.0</u> | <u>39.2</u> | <u>46.8</u> |
| VDGD (Ghosh et al., 2024)                                | <u>34.0</u> | <u>41.2</u> | <u>35.6</u> | <u>42.1</u> | 36.6        | 45.4        | 42.4        | 50.7        | 37.1        | 44.9        |
| HalluciMAD (Lin et al., 2024)                            | 29.7        | 38.7        | 32.6        | 40.3        | 36.4        | 45.1        | 34.8        | 42.8        | 33.4        | 41.7        |
| <i>Qwen2-VL-7B-Instruct (Wang et al., 2024a)</i>         |             |             |             |             |             |             |             |             |             |             |
| Direct (Wang et al., 2024a)                              | 36.2        | 39.5        | 33.3        | 38.3        | 42.9        | 46.7        | 39.9        | 44.4        | 38.1        | 42.2        |
| CoT (Kojima et al., 2022)                                | 38.6        | 43.4        | 33.9        | <u>39.0</u> | 48.3        | 52.1        | <u>48.4</u> | <u>52.5</u> | 42.3        | <u>46.7</u> |
| SRO (Lin et al., 2024)                                   | <u>38.7</u> | <u>43.5</u> | 34.1        | 34.8        | <u>48.4</u> | <u>52.3</u> | <u>48.4</u> | 52.4        | <u>42.4</u> | 45.7        |
| VDGD (Ghosh et al., 2024)                                | 34.8        | 36.3        | <u>36.3</u> | 38.8        | 41.1        | 44.2        | 40.0        | 44.9        | 38.1        | 41.1        |
| HalluciMAD (Lin et al., 2024)                            | 37.0        | 42.5        | 31.4        | 38.1        | 38.0        | 45.4        | 38.6        | 44.7        | 36.3        | 42.7        |
| <i>Pixtral-12B-2409 (Agrawal et al., 2024)</i>           |             |             |             |             |             |             |             |             |             |             |
| Direct (Agrawal et al., 2024)                            | 23.9        | 38.0        | 34.8        | 45.5        | 43.3        | 51.3        | 38.0        | 48.9        | 35.0        | 45.9        |
| CoT (Kojima et al., 2022)                                | 40.7        | 46.0        | 42.7        | 47.9        | 48.7        | 54.2        | 53.3        | 59.0        | 46.3        | 51.8        |
| SRO (Lin et al., 2024)                                   | 43.0        | 47.8        | 41.3        | 48.1        | 50.6        | 55.5        | 51.6        | 57.8        | 46.6        | 52.3        |
| VDGD (Ghosh et al., 2024)                                | 43.4        | 43.1        | 45.1        | 47.3        | 47.3        | 48.4        | 56.4        | 58.2        | 48.1        | 49.3        |
| HalluciMAD (Lin et al., 2024)                            | <u>46.3</u> | <u>49.8</u> | <u>45.2</u> | <u>51.4</u> | <u>57.1</u> | <u>61.4</u> | <u>58.7</u> | <u>63.2</u> | <u>51.8</u> | <u>56.4</u> |
| <i>Gemini-1.5-Flash (Team et al., 2024)</i>              |             |             |             |             |             |             |             |             |             |             |
| Direct (Team et al., 2024)                               | 41.7        | 44.7        | 37.0        | 39.2        | 49.2        | 50.5        | 50.1        | 51.9        | 44.5        | 46.6        |
| CoT (Kojima et al., 2022)                                | 49.2        | 52.3        | 53.2        | 54.7        | 56.4        | 58.0        | 60.6        | 62.2        | 54.9        | 56.8        |
| SRO (Lin et al., 2024)                                   | 49.0        | 52.9        | 51.6        | 53.4        | 55.7        | 58.3        | 58.2        | 60.8        | 53.6        | 56.4        |
| VDGD (Ghosh et al., 2024)                                | <u>58.9</u> | <u>60.3</u> | 50.6        | 52.8        | 52.7        | 54.0        | 60.6        | 62.4        | 55.7        | 57.4        |
| HalluciMAD (Lin et al., 2024)                            | 52.2        | 54.5        | <u>59.0</u> | <u>60.6</u> | <u>61.6</u> | 63.9        | <u>63.7</u> | <u>64.9</u> | <u>59.1</u> | 61.0        |
| <i>GPT-4o (OpenAI, 2023)</i>                             |             |             |             |             |             |             |             |             |             |             |
| Direct (OpenAI, 2023)                                    | 57.1        | 63.2        | 56.1        | 62.1        | 72.2        | 76.4        | 82.9        | 84.3        | 67.1        | 71.5        |
| CoT (Kojima et al., 2022)                                | 68.1        | 70.0        | 66.9        | 69.2        | 81.1        | 83.4        | 83.1        | 84.8        | 74.8        | 76.8        |
| SRO (Lin et al., 2024)                                   | 63.2        | 65.0        | 57.2        | 62.7        | 76.7        | 79.3        | 84.7        | 86.4        | 70.4        | 73.4        |
| VDGD (Ghosh et al., 2024)                                | 64.7        | 66.6        | 56.3        | 62.5        | 72.3        | 77.0        | 83.2        | 85.6        | 69.1        | 73.0        |
| HalluciMAD (Lin et al., 2024)                            | <b>70.9</b> | <b>71.9</b> | <b>68.6</b> | <b>70.3</b> | <b>84.1</b> | <b>85.6</b> | <b>86.4</b> | <b>87.3</b> | <b>77.5</b> | <b>78.8</b> |

Table 1: The experimental results of Acc. (%) and Macro-F1 score on MLLMs. The “Random” refers to the average performance obtained from three separate random selections.   represents the performance of open-source MLLMs, and   represents the performance of closed-source MLLMs. The underline indicates better performance in the MLLM. **1** represents the **Best** performance. The complete results are shown in Table 3.

## 5.2 Results for CCHa11

Results are summarized in the Table 1. We have made the following observations:

**CCHa11 is a challenging hallucination detection benchmark:** Our evaluation results show that the weakest-performing model is *InternVL2-8B*, achieving an accuracy of 34.0% using the *Direct* method. In contrast, the best-performing method, *HalluciMAD*, achieves an accuracy of 77.5% using *GPT-4o*. This suggests that CCHa11 is a highly challenging benchmark and that MLLMs still have significant room for further improvement in cross-lingual and cross-modal tasks.

**The performance of MLLMs depends on their training strategy and parameter size:** The closed-source *GPT-4o* and *Gemini-1.5-Flash* exhibit supe-

rior performance, surpassing several open-source models. Additionally, performance variations exist among open-source models. Notably, the *Qwen-2-VL-7B* outperforms both *Llama-3.2-11B* and *InternVL2-8B*. This highlights the vital role of training strategies in enhancing performance.

**Strategies for mitigating hallucinations work in specific contexts:** Hallucination mitigation methods generally outperform the *Direct* approach, although their effectiveness can vary depending on the specific context. Basic methods, such as CoT and SRO, perform better in less powerful models, including *InternVL2-8B*, *Llama-3.2-11B-Vision*, and *Qwen-2-VL-7B*. In contrast, more advanced methods, such as VDGD and *HalluciMAD*, are more effective in powerful models, including *Pixtral-12B-2409*, *Gemini-1.5-Flash*, and *GPT-4o*.

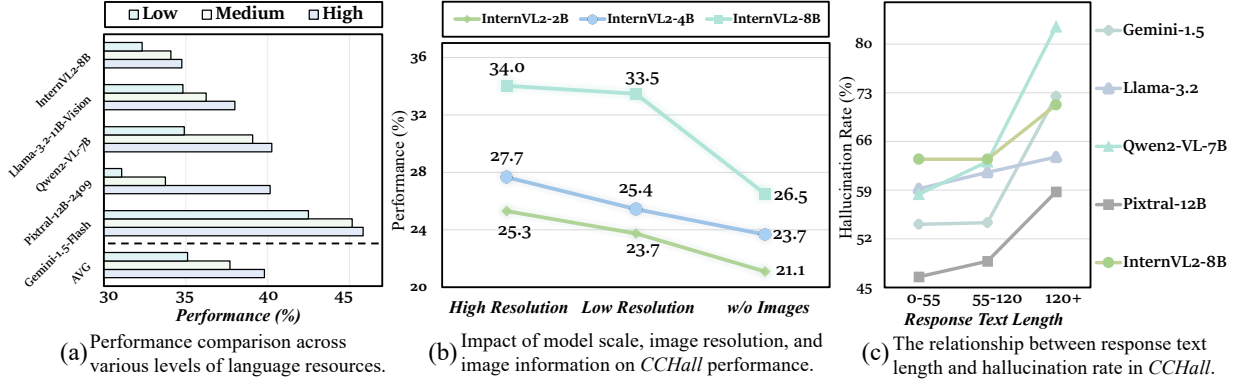


Figure 6: Analysis of the underlying causes of cross-lingual and cross-modal hallucinations in MLLMs.

### 5.3 Analysis

We conduct experiments to analyze how image quality, language resources, response length, and object-specific biases influence hallucination behavior and detection in CCHall.

#### 5.3.1 Models Perform Better in High-resource Languages than in Low-resource Ones

To examine performance variations across languages, we follow [Guerreiro et al. \(2023\)](#) and categorize them into three groups: high-resource languages (FR, ES, PT), medium-resource languages (CS, NL, SV), and low-resource languages (HR, CY, SW). As illustrated in Figure 6 (a), we observe that hallucination detection accuracy is highest for high-resource languages and lowest for low-resource languages. This disparity in performance is likely attributed to the limited availability of training data for low-resource languages, which significantly hampers the model’s ability to understand contextual nuances, leading to increased difficulty in accurate hallucination detection.

#### 5.3.2 High-resolution Images and Larger Models Enhance Task Performance

To explore the impact of image quality and model parameters on hallucination performance, we analyze the effects of high-resolution, low-resolution, and no-image (w/o Images) inputs on model performance. At the same time, we use InternVL of different sizes to examine the effect of model scale on hallucination. As shown in Figure 6 (b), performance declines as image resolution decreases, with a sharp drop when no image is provided, highlighting the importance of visual information in CCHall for reducing hallucinations. Furthermore, model performance declines as the number of parameters decreases (Performance: 8B > 4B > 2B).

#### 5.3.3 Longer Responses Generally Always Lead to Higher Hallucination Rates

We investigate the effect of response length on model performance. As shown in Figure 6 (c), there is a strong positive relationship between how long a response is and how often hallucinations occur. Notably, a substantial surge in hallucination rate is observed when the length of the response surpasses 120 words, indicating a critical threshold beyond which the model’s output becomes significantly less reliable. The objects that are susceptible to hallucination are detailed in Appendix B.2. This is likely because the model generates additional reasoning steps when the answer is uncertain. In the future, incorporating reflective steps into the model may help mitigate hallucinations.

### 5.4 Exploration

This section explores two factors for improving multimodal performance: multi-language prompts and external tools for detecting hallucinations.

#### 5.4.1 Multilingual Prompt Exploration

We investigate the performance enhancement of the MLLMs in cross-lingual and cross-modal scenarios under bilingual context prompts. The results are shown in Figure 7, incorporating the Source Language and English ( $En+SL$ ) consistently improves hallucination detection accuracy across datasets, compared to using only English ( $En$ ). Specifically, in the two VQA tasks, accuracy increases by 2.2% in the AMBER dataset and by 4.8% in GQA. This improvement suggests that bilingual prompts provide better linguistic context and help effectively mitigate biases associated with relying solely on English. Future research should focus on further optimizing bilingual prompts and evaluating their effectiveness across additional language pairs.

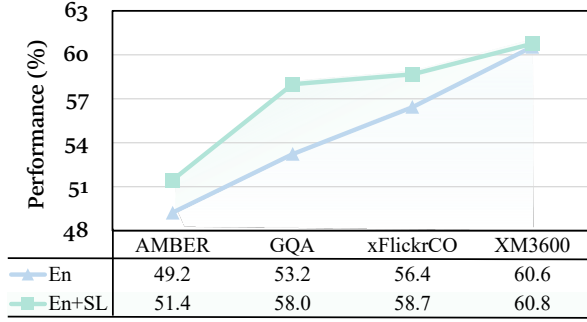


Figure 7: A comparison between using only English (*En*) and using English with an additional Source Language (*En+SL*) as context in *Gemini-1.5-Flash*.

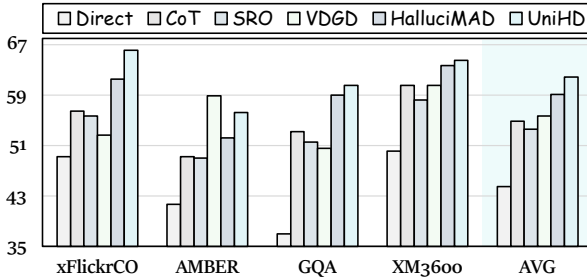


Figure 8: The *UniHD*, which utilizes external tools, is compared with other methods in *Gemini-1.5-Flash*.

#### 5.4.2 Framework Adaptation Exploration

To investigate how external tools and web search improve the mitigation of hallucinations, we integrate the *UniHD* (Chen et al., 2024b) framework. As shown in Figure 8, adapting the *UniHD* framework for *CCHall* allows the model to achieve the highest accuracy across all datasets, surpassing previous methods such as *VDGD* and *HalluciMAD*. Across all datasets, it shows an average improvement of 2.7% over *HalluciMAD*. By validating visual claims using external tools and web resources, *UniHD* bridges the gap between model predictions and real-world data, resulting in significant improvements over methods that rely solely on the internal knowledge reasoning of MLLMs. The experimental details can be found in Appendix B.3.

## 6 Related Work

LLMs have recently made rapid progress (Team et al., 2024; Qin et al., 2025; Chen et al., 2025), achieving outstanding performance (Zhang et al., 2024b). However, LLMs often exhibit hallucinations, producing outputs that fail to align with the given inputs, particularly in cross-lingual or cross-modal settings (Benkirane et al., 2024; Favero et al., 2024). This has prompted several studies exploring hallucinations (Sriramanan et al., 2025).

**Cross-lingual Hallucinations Benchmark:** Qiu et al. (2023) introduce *mFACT* to assess the faithfulness of non-English summaries, revealing that LLMs are more prone to hallucination in languages other than English. Dale et al. (2023a) release an annotated dataset covering hallucinations and omissions across 18 translation directions, with varying levels of hallucination severity. Herrlein et al. (2024) extend English hallucination detection to German and apply it to long-context scenarios. Son et al. (2024) introduce *MM-Eval*, a multilingual benchmark covering 18 languages, to evaluate LLMs as evaluators. Li et al. (2024) introduce *XTRUST*, the multilingual trustworthiness benchmark, covering 10 languages and various topics like hallucination, misinformation, and fairness.

**Cross-modal Hallucinations Benchmark:** Liu et al. (2023) introduce a benchmark for image-context reasoning using image-question pairs created by experts. Leng et al. (2024) propose the “Curse of Multi-Modalities” benchmark to evaluate multimodal hallucinations in MLLMs. Yan et al. (2024) propose a framework for evaluating the reliability and validity of hallucination benchmarks. Gunjal et al. (2024) introduce *M-HalDetect*, a multimodal dataset with 16k VQA annotations for hallucination detection in LVLMs. Guan et al. (2024) introduce *HallusionBench*, a benchmark for image-context reasoning in LVLMs with 346 images and 1129 questions. Chen et al. (2024b) introduce a multi-modal meta-evaluation benchmark for assessing hallucination detection in MLLMs.

Compared to previous work, we introduce the **Cross-lingual and Cross-modal Hallucinations benchmark** (*CCHall*). To the best of our knowledge, this is the first effort to explore joint cross-lingual and cross-modal hallucinations.

## 7 Conclusion

We introduce a benchmark, *CCHall*, for the detection of joint cross-lingual and cross-modal hallucinations. We analyze a wide range of MLLMs and various hallucination mitigation strategies. Our findings demonstrate that current strategies are still insufficient in effectively overcoming cross-lingual and cross-modal hallucinations. Specifically, they often struggle to effectively handle the complexities that arise when both cross-lingual and cross-modal factors interact. We hope *CCHall* can serve as valuable data for evaluating LLMs in joint cross-lingual and cross-modal scenarios.

## Limitations

We introduce CCHall, a cross-lingual and cross-modal hallucination detection benchmark, and perform comprehensive experiments across various MLLMs. Current MLLMs are mainly confined to two modalities: text and image, with only a few extending to the Audio/Speech modality (Borsos et al., 2023; Kuan et al., 2024). Consequently, our work concentrates on the text and image modalities. We hope that future MLLMs will incorporate a broader range of modalities, bringing us one step closer to achieving Artificial General Intelligence (AGI). Additionally, since the hallucinated data in our dataset is generated by MLLMs and the multi-lingual data is obtained through translation, some errors may persist despite manual verification.

## Acknowledgments

This work was supported by the National Natural Science Foundation of China (NSFC) via grant 62306342 and 62376130. This work was sponsored by the Excellent Young Scientists Fund in Hunan Province (2024JJ4070), the Science and Technology Innovation Program of Hunan Province under Grant 2024RC3024 and CCF-Zhipu Large Model Innovation Fund (NO.CCFZhipu202406). This work was supported by Key Laboratory of Data Intelligence and Advanced Computing in Provincial Universities, Soochow University. This work was supported by the Key Laboratory of Computing Power Network and Information Security, Affiliated with Ministry of Education, Qilu University of Technology (Shandong Academy of Sciences) (No.2023ZD032). We are grateful for resources from the High Performance Computing Center of Central South University.

## References

- Pravesh Agrawal, Szymon Antoniak, Emma Bou Hanna, Baptiste Bout, Devendra Chaplot, Jessica Chudnovsky, Diogo Costa, Baudouin De Monicault, Saurabh Garg, Theophile Gervet, et al. 2024. Pixtral 12b. *arXiv preprint arXiv:2410.07073*.
- Zechen Bai, Pichao Wang, Tianjun Xiao, Tong He, Zongbo Han, Zheng Zhang, and Mike Zheng Shou. 2024. Hallucination of multimodal large language models: A survey. *arXiv preprint arXiv:2404.18930*.
- Kenza Benkirane, Laura Gongas, Shahar Pelles, Naomi Fuchs, Joshua Darmon, Pontus Stenertorp, David Ifeoluwa Adelani, and Eduardo Sánchez. 2024. Machine translation hallucination detection for low and high resource languages using large language models. *arXiv preprint arXiv:2407.16470*.
- Zalán Borsos et al. 2023. Audiollm: a language modeling approach to audio generation. *IEEE/ACM transactions on audio, speech, and language processing*, 31:2523–2533.
- Emanuele Bugliarello, Fangyu Liu, Jonas Pfeiffer, Siva Reddy, Desmond Elliott, Edoardo Maria Ponti, and Ivan Vulić. 2022. Iglue: A benchmark for transfer learning across modalities, tasks, and languages. In *International Conference on Machine Learning*, pages 2370–2392. PMLR.
- Ximena Castillo, Lobry Hsu, and Archer Flynn. 2025. Hierarchical cross-lingual representation learning for diverse video contexts.
- Qiguang Chen, Libo Qin, Jinhao Liu, Dengyun Peng, Jiannan Guan, Peng Wang, Mengkang Hu, Yuhang Zhou, Te Gao, and Wanxiang Che. 2025. Towards reasoning era: A survey of long chain-of-thought for reasoning large language models. *Preprint*, arXiv:2503.09567.
- Qiguang Chen, Libo Qin, Jin Zhang, Zhi Chen, Xiao Xu, and Wanxiang Che. 2024a. M<sup>3</sup>CoT: A novel benchmark for multi-domain multi-step multi-modal chain-of-thought. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, Bangkok, Thailand. Association for Computational Linguistics.
- Xiang Chen, Chenxi Wang, Yida Xue, Ningyu Zhang, Xiaoyan Yang, Qiang Li, Yue Shen, Lei Liang, Jinjie Gu, and Huajun Chen. 2024b. Unified hallucination detection for multimodal large language models. *arXiv preprint arXiv:2402.03190*.
- Zhe Chen, Weiyun Wang, Hao Tian, Shenglong Ye, Zhangwei Gao, Erfei Cui, Wenwen Tong, Kongzhi Hu, Jiapeng Luo, Zheng Ma, et al. 2024c. How far are we to gpt-4v? closing the gap to commercial multimodal models with open-source suites. *arXiv preprint arXiv:2404.16821*.
- Alexis Conneau, Rutu Rinott, Guillaume Lample, Adina Williams, Samuel Bowman, Holger Schwenk, and Veselin Stoyanov. 2018. Xnli: Evaluating cross-lingual sentence representations. In *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*, pages 2475–2485.
- David Dale, Elena Voita, Janice Lam, Prangthip Hansanti, Christophe Ropers, Elahe Kalbassi, Cynthia Gao, Loic Barrault, and Marta Costa-jussa. 2023a. Halomi: A manually annotated benchmark for multilingual hallucination and omission detection in machine translation. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 638–653.
- David Dale et al. 2023b. Halomi: A manually annotated benchmark for multilingual hallucination and omission detection in machine translation.

- Alessandro Favero et al. 2024. Multi-modal hallucination control by visual information grounding. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 14303–14312.
- Sreyan Ghosh, Chandra Kiran Reddy Evuru, Sonal Kumar, Utkarsh Tyagi, Oriol Nieto, Zeyu Jin, and Dinesh Manocha. 2024. Vdgd: Mitigating lvlm hallucinations in cognitive prompts by bridging the visual perception gap. *arXiv preprint arXiv:2405.15683*.
- Tianrui Guan, Fuxiao Liu, Xiyang Wu, et al. 2024. Hallusionbench: an advanced diagnostic suite for entangled language hallucination and visual illusion in large vision-language models. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*.
- Nuno M Guerreiro, Duarte M Alves, Jonas Waldendorf, Barry Haddow, Alexandra Birch, Pierre Colombo, and André FT Martins. 2023. Hallucinations in large multilingual translation models. *Transactions of the Association for Computational Linguistics*.
- Anisha Gunjal, Jihan Yin, and Erhan Bas. 2024. Detecting and preventing hallucinations in large vision language models. In *Proceedings of the AAAI Conference on Artificial Intelligence*.
- Ashim Gupta and Vivek Srikumar. 2021. X-factor: A new benchmark dataset for multilingual fact checking. *arXiv preprint arXiv:2106.09248*.
- Md Arif Hasan, Maram Hasanain, Fatema Ahmad, Sahinur Rahman Laskar, Sunaya Upadhyay, Vrunda N Sukhadia, Mucahid Kutlu, Shammur Absar Chowdhury, and Firoj Alam. 2024. Nativqa: Multilingual culturally-aligned natural query for llms. *arXiv preprint arXiv:2407.09823*.
- Tahmid Hasan, Abhik Bhattacharjee, Md Saiful Islam, Kazi Samin, Yuan-Fang Li, Yong-Bin Kang, M. Sohel Rahman, and Rifat Shahriyar. 2021. [Xl-sum: Large-scale multilingual abstractive summarization for 44 languages](#). *Preprint*, arXiv:2106.13822.
- Maite Heredia, Julen Etxaniz, Maitze Zulaika, Xabier Saralegi, Jeremy Barnes, and Aitor Soroa. 2024. Xnlieu: a dataset for cross-lingual nli in basque. In *Proceedings of the 2024 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies (Volume 1: Long Papers)*.
- Janek Herrlein, Chia-Chien Hung, and Goran Glava. 2024. Anhalten: Cross-lingual transfer for german token-level reference-free hallucination detection. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 4: Student Research Workshop)*, pages 186–194.
- Drew A Hudson and Christopher D Manning. 2019. Gqa: A new dataset for real-world visual reasoning and compositional question answering. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 6700–6709.
- Takeshi Kojima, Shixiang Shane Gu, Machel Reid, Yutaka Matsuo, and Yusuke Iwasawa. 2022. Large language models are zero-shot reasoners. *Advances in neural information processing systems*, 35.
- Chun-Yi Kuan, Wei-Ping Huang, and Hung-yi Lee. 2024. Understanding sounds, missing the questions: The challenge of object hallucination in large audio-language models. *arXiv preprint arXiv:2406.08402*.
- Michelle A Lee. 2021. *Examining Machine Translation Systems and Translation Quality using the Back-Translation Method*. Ph.D. thesis, University of Wisconsin–Stout.
- Sicong Leng, Yun Xing, Zesen Cheng, Yang Zhou, Hang Zhang, Xin Li, Deli Zhao, Shijian Lu, Chunyan Miao, and Lidong Bing. 2024. The curse of multi-modalities: Evaluating hallucinations of large multimodal models across language, visual, and audio. *arXiv preprint arXiv:2410.12787*.
- Yahan Li, Yi Wang, Yi Chang, and Yuan Wu. 2024. [Xtrust: On the multilingual trustworthiness of large language models](#). *Preprint*, arXiv:2409.15762.
- Yifan Li, Yifan Du, Kun Zhou, Jinpeng Wang, Wayne Xin Zhao, and Ji-Rong Wen. 2023. Evaluating object hallucination in large vision-language models. *arXiv preprint arXiv:2305.10355*.
- Zheng Lin, Zhenxing Niu, Zhibin Wang, and Yinghui Xu. 2024. Interpreting and mitigating hallucination in mlms through multi-agent debate. *arXiv preprint arXiv:2407.20505*.
- Fuxiao Liu, Tianrui Guan, Zongxia Li, Lichang Chen, Yaser Yacoob, Dinesh Manocha, and Tianyi Zhou. 2023. Hallusionbench: You see what you think? or you think what you see? an image-context reasoning benchmark challenging for gpt-4v (ision), llava-1.5, and other multi-modality models. *arXiv preprint arXiv:2310.14566*.
- Shilong Liu, Zhaoyang Zeng, Tianhe Ren, Feng Li, Hao Zhang, Jie Yang, Qing Jiang, Chunyuan Li, Jianwei Yang, Hang Su, Jun Zhu, and Lei Zhang. 2024. [Grounding dino: Marrying dino with grounded pre-training for open-set object detection](#). *Preprint*, arXiv:2303.05499.
- Abhimanyu Meta, Dubey, Abhinav Jauhri, Abhinav Pandey, Abhishek Kadian, Ahmad Al-Dahle, Aiesha Letman, Akhil Mathur, Alan Schelten, Amy Yang, Angela Fan, et al. 2024. The llama 3 herd of models. *arXiv preprint arXiv:2407.21783*.
- AI Meta. 2024. Llama 3.2: Revolutionizing edge ai and vision with open, customizable models. *Meta AI*.
- Mai Miyabe and Takashi Yoshino. 2015. Evaluation of the validity of back-translation as a method of assessing the accuracy of machine translation. In *2015 International Conference on Culture and Computing (Culture Computing)*. IEEE.

- OpenAI. 2023. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*.
- Dongmin Park, Zhaofang Qian, Guangxing Han, and Ser-Nam Lim. 2024. Mitigating dialogue hallucination for large multi-modal models via adversarial instruction tuning. *arXiv preprint arXiv:2403.10492*.
- Libo Qin, Qiguang Chen, Xiachong Feng, Yang Wu, Yongheng Zhang, Yinghui Li, Min Li, Wanxiang Che, and Philip S Yu. 2024. Large language models meet nlp: A survey. *ArXiv preprint*.
- Libo Qin, Qiguang Chen, Fuxuan Wei, Shijue Huang, and Wanxiang Che. 2023. [Cross-lingual prompting: Improving zero-shot chain-of-thought reasoning across languages](#). *Preprint*, arXiv:2310.14799.
- Libo Qin, Qiguang Chen, Yuhang Zhou, Zhi Chen, Yinghui Li, Lizi Liao, Min Li, Wanxiang Che, and Philip S Yu. 2025. A survey of multilingual large language models. *Patterns*, 6(1).
- Yifu Qiu, Yftah Ziser, Anna Korhonen, Edoardo Ponti, and Shay B Cohen. 2023. Detecting and mitigating hallucinations in multilingual summarisation. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*.
- Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, et al. 2021. Learning transferable visual models from natural language supervision. In *International conference on machine learning*, pages 8748–8763. PMLR.
- Michael Regan, Shira Wein, George Baker, and Emilio Monti. 2024. Massive multilingual abstract meaning representation: A dataset and baselines for hallucination detection. *arXiv preprint arXiv:2405.19285*.
- Anna Rohrbach, Lisa Anne Hendricks, Kaylee Burns, Trevor Darrell, and Kate Saenko. 2018. Object hallucination in image captioning. *arXiv preprint arXiv:1809.02156*.
- Guijin Son, Dongkeun Yoon, Juyoung Suk, Javier Aula-Blasco, Mano Aslan, Vu Trong Kim, Shayekh Bin Islam, Jaume Prats-Cristia, Lucia Tormo-Banuelos, and Seungone Kim. 2024. [Mm-eval: A multilingual meta-evaluation benchmark for llm-as-a-judge and reward models](#). *Preprint*, arXiv:2410.17578.
- Gaurang Sriramanan, Siddhant Bharti, Vinu Sankar Sadasivan, Shoumik Saha, Priyatham Kattakinda, and Soheil Feizi. 2025. Llm-check: Investigating detection of hallucinations in large language models. *Advances in Neural Information Processing Systems*, 37:34188–34216.
- Gemini Team, Petko Georgiev, Ving Ian Lei, Ryan Burnell, Libin Bai, Anmol Gulati, Garrett Tanzer, Damien Vincent, Zhufeng Pan, Shibo Wang, et al. 2024. Gemini 1.5: Unlocking multimodal understanding across millions of tokens of context. *arXiv preprint arXiv:2403.05530*.
- Ashish V Thapliyal, Jordi Pont-Tuset, Xi Chen, and Radu Soricut. 2022. Crossmodal-3600: A massively multilingual multimodal evaluation dataset. *arXiv preprint arXiv:2205.12522*.
- Junyang Wang, Yuhang Wang, Guohai Xu, Jing Zhang, Yukai Gu, Haitao Jia, Ming Yan, Ji Zhang, and Jitao Sang. 2023. An llm-free multi-dimensional benchmark for mllms hallucination evaluation. *arXiv preprint arXiv:2311.07397*.
- Kun Wang, Guibin Zhang, Zhenhong Zhou, Jiahao Wu, Miao Yu, Shiqian Zhao, Chenlong Yin, Jinhu Fu, Yibo Yan, Hanjun Luo, et al. 2025. A comprehensive survey in llm (-agent) full stack safety: Data, training and deployment. *arXiv preprint arXiv:2504.15585*.
- Peng Wang, Shuai Bai, Sinan Tan, Shijie Wang, Zhihao Fan, Jinze Bai, Keqin Chen, Xuejing Liu, Jialin Wang, Wenbin Ge, et al. 2024a. Qwen2-vl: Enhancing vision-language model’s perception of the world at any resolution. *arXiv preprint arXiv:2409.12191*.
- Yabing Wang, Le Wang, Qiang Zhou, Zhibin Wang, Hao Li, Gang Hua, and Wei Tang. 2024b. Multimodal llm enhanced cross-lingual cross-modal retrieval. In *Proceedings of the 32nd ACM International Conference on Multimedia*.
- Jason Wei, Xuezhi Wang, and Dale et al. Schuurmans. 2022. Chain-of-thought prompting elicits reasoning in large language models. *Advances in neural information processing systems*, 35.
- Bei Yan, Jie Zhang, Zheng Yuan, Shiguang Shan, and Xilin Chen. 2024. Evaluating the quality of hallucination benchmarks for large vision-language models. *arXiv preprint arXiv:2406.17115*.
- Miao Yu, Fanci Meng, Xinyun Zhou, Shilong Wang, Junyuan Mao, Linsey Pang, Tianlong Chen, Kun Wang, Xinfeng Li, Yongfeng Zhang, et al. 2025. A survey on trustworthy llm agents: Threats and countermeasures. *arXiv preprint arXiv:2503.09648*.
- Guibin Zhang, Luyang Niu, Junfeng Fang, Kun Wang, Lei Bai, and Xiang Wang. 2025. Multi-agent architecture search via agentic supernet. *arXiv preprint arXiv:2502.04180*.
- Yongheng Zhang, Qiguang Chen, Min Li, Wanxiang Che, and Libo Qin. 2024a. Autocap: Towards automatic cross-lingual alignment planning for zero-shot chain-of-thought. In *Findings of the Association for Computational Linguistics ACL 2024*, pages 9191–9200.
- Yongheng Zhang, Qiguang Chen, Jingxuan Zhou, Peng Wang, Jiasheng Si, Jin Wang, Wenpeng Lu, and Libo Qin. 2024b. Wrong-of-thought: An integrated reasoning framework with multi-perspective verification and wrong information. In *Findings of the Association for Computational Linguistics: EMNLP 2024*, pages 6644–6653.

## Appendix

### A CCHa11 Construction Details

#### A.1 Generating Hallucinated Data

Before generating data, we first prepare corresponding real questions and answers for each image and meticulously design specific prompts to guide Gemini-1.5-Pro in generating hallucinated sentences. As shown in Figure 10, the prompt includes the following elements:

1. **Image:** The inclusion of the image ensures that the model effectively recognizes and processes visual content, thereby facilitating more accurate multi-modal alignment.
2. **Question and Real Answer:** These provide the model with an accurate semantic context, enabling it to understand the requirements for generating hallucinated data.
3. **Examples:** By providing examples, the prompt clarifies the generation target for the model, ensuring that the hallucinated sentences maintain logical consistency and follow the intended structure and content. As shown in Figure 10, these are the provided examples used to construct A and B.

Incorporating these elements helps the model better comprehend the input context, ensuring that the generated hallucinated sentences appropriately combine information from the image, question, and real answer to produce the expected output. This design prevents the hallucinations from deviating from the actual context, ensuring high quality in the generated content.

During the generation process, we require the hallucinated sentences to closely match the real answers in both length and structure. Specifically, the Gemini-1.5-Pro, when processing images, is instructed to compare the real entities in the answers with the content in the images. It then integrates entities that do not exist in the image into the generated hallucinated sentence.

This approach ensures that the hallucinated data maintains a fundamental level of naturalness and consistency with the real answer. If the structure of the hallucinated sentence deviates significantly from the real answer, it becomes more easily identifiable as an anomaly by both models and human users. Thus, maintaining similarity in structure and length makes the hallucinated sentence appear

more natural and harder to detect. Moreover, this strategy ensures that the hallucinated data we generate is highly deceptive. Structurally similar hallucinated sentences can mask the inserted false entities, ensuring that the hallucinated content maintains fluency similar to the real answer, thus increasing the difficulty for models during evaluation.

Overall, this process ensures that the hallucinated sentences semantically approximate the real answers, while preserving structural consistency, thus enhancing their deceptiveness and improving evaluation effectiveness.

#### A.2 Cross-lingual Hallucination Data Construction

**Human Recheck:** After generating the hallucinated data and performing translations, we carefully conduct a human recheck to ensure the quality and accuracy of the data. Specifically, we consistently employ the back-translation method (Miyabe and Yoshino, 2015; Lee, 2021) to thoroughly review all the data, focusing on two key aspects:

1. **Verification of Hallucination Data:** We check whether the generated hallucinated data meets our requirements. This involves ensuring that the hallucinated sentences adhere to the expected structure, length, and logic, while also verifying that they contain the intended hallucinated elements in alignment with the task.
2. **Accuracy of Translated Data:** We also check whether the translated data accurately reflects the original meaning and context. This step ensures that the translations remain faithful to the intended hallucination scenario, avoiding any misinterpretations or errors that arise from language differences.

In terms of details, we provide the reviewers with a specific scoring rubric as a clear guide for evaluation. The scoring criteria are as follows:

- 0-60 points: Only a small portion of the original meaning is conveyed. The word choice is inaccurate, and there are a significant number of severe language errors.
- 60-70 points: The meaning of the original text is conveyed. Word choice is somewhat inaccurate, and there are quite a few language errors, some of which are serious.

- 70-80 points: The meaning of the original text is generally conveyed. The text is smooth and coherent, with no major language errors.
- 80-90 points: The text conveys the original meaning clearly, with minor errors that don't hinder comprehension. The writing is smooth and coherent, though some word choices may be slightly imprecise or awkward.
- 90-100 points: The original meaning is accurately conveyed. The word choice is appropriate, the writing flows smoothly, and there are virtually no language errors.

Each data is reviewed by three evaluators, and the final score for the data is obtained by averaging the scores from all three reviewers. If a data point receives a score below 80, we regenerate and re-translate it. The data will then be re-evaluated and re-scored until it passes the review. A total of only 11 data points received a score below 80, and we have already made corrections to them.

In our review, the average score across all data points is 87.1. The manual check guarantees the accuracy, consistency, and quality of the generated data and translations, significantly minimizing the risk of errors, inconsistencies, or biases that could affect the model's overall performance.

### A.3 Four Hallucination Types

Based on the definitions of four hallucination types, we carefully select and combine the intermediate dataset after randomization and balancing. Specifically, we use English as the anchor language and select two answers for each image in different languages. For each image, one English answer is randomly chosen, along with another in a different language. The selected English answer can be either the correct answer or a hallucinated one generated by a previous model. According to the hallucination definitions, we have the following four combinations:

1. **Non-hallucination:** If both the English answer and the answer in another language are correct, meaning that the objects mentioned in the sentence match those in the image and are consistent with it, the answer is classified as non-hallucination.

2. **Cross-lingual, non-cross-modal hallucination:** If the English answer is correct but the answer in another language is hallucinated — meaning the English answer is consistent with the image while the objects in the non-English answer do not appear in the image — this is classified as cross-lingual non-cross-modal hallucination.
3. **Cross-modal, non-cross-lingual hallucination:** If the English answer is hallucinated and the answer in another language is identical, meaning both answers contain the same hallucinated objects not present in the image, this is classified as cross-modal non-cross-lingual hallucination.
4. **Cross-modal, cross-lingual hallucination:** If the English answer is hallucinatory and the answer in another language is inconsistent with it, meaning that the two hallucinatory answers differ, with each containing different objects that are either not actually present in the image or are incorrectly described, this is classified as a cross-modal, cross-lingual hallucination.

This combination process ensures that we generate a diverse set of hallucination types, while maintaining consistency in image-object relationships, thus providing a robust and high-quality dataset for training and evaluation.

### A.4 Prompt Construction for Experimentation

During the experimental phase, we test different models on our CCHall. As shown in Figure 11, our prompt consists of four key components: Intention Description, Hallucination Type Explanation, Task Description, and Output Format.

1. **Intention Description:** This section establishes the model's role, clearly instructing it that its task is to detect hallucinations based on the given image, question, and answers in two different languages, and classify the hallucinations. The model selects the correct option from the provided choices. This helps set expectations for the model's behavior, ensuring that it understands the specific goal of identifying and classifying hallucinations.
2. **Hallucination Type Explanation:** This part defines four types of hallucinations: "Non-hallucination", "Cross-lingual, non-cross-

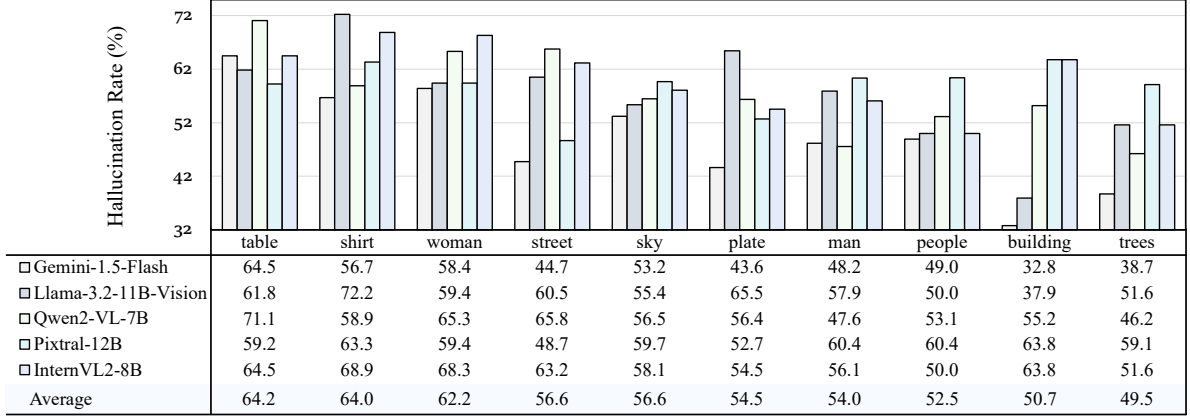


Figure 9: Analysis of the relationship between object presence and hallucination rate in the image captioning task.

modal hallucination”, “Cross-modal, non-cross-lingual hallucination”, and “Cross-modal, cross-lingual hallucination”. Explaining the four types of hallucinations is critical because it provides the model with a clear framework to understand the specific nature of each hallucination type. By making these definitions explicit, the model can more effectively distinguish between different hallucination scenarios and improve its accuracy in detecting and classifying hallucinations. Without this explanation, the model struggles to identify subtle distinctions between hallucination types, which impacts its overall performance and classification accuracy.

3. **Task Description:** This section provides concrete task details, including the image, question, two answers in different languages, and the available options. This supplies the model with all necessary input data, ensuring that it has everything needed to perform the task effectively and make informed decisions when selecting the correct answer.
4. **Output Format:** The output format section specifies the required structure of the model’s response, designed for easy extraction using regular expressions. This standardization allows efficient evaluation of the model’s performance, facilitates accurate calculations, and ensures consistent output formatting.

This structured approach to prompt construction ensures that the model receives clear, unambiguous instructions at every stage of the task, improving consistency and reliability in hallucination detection and classification across different models.

## B Experiment Details

### B.1 Main Result Details

To conduct a more thorough and comprehensive evaluation of CCHall, we assess its performance on the following extensive models: *GPT-4o* (OpenAI, 2023), *Gemini-1.5-Flash* (Team et al., 2024), *Llama-3.2-11B-Vision-Instruct* (Meta, 2024), *Qwen2-VL-7B-Instruct* (Wang et al., 2024a), *Pixtral-12B-2409* (Agrawal et al., 2024), and *InternVL2-(2B, 4B, 8B)* (Chen et al., 2024c). The complete results are shown in Table 3.

### B.2 Certain objects are more prone to hallucination.

To gain deeper insights into the types of entities that are more prone to triggering hallucinations in MLLMs, we conducted a thorough statistical analysis of hallucination-associated entities produced by five representative MLLMs: *Gemini-1.5-Flash*, *Llama-3.2-11B-Vision-Instruct*, *Qwen2-VL-(2B, 7B)-Instruct*, *Pixtral-12B-2409*, and *InternVL2-8B*. For each model, we systematically extracted and aggregated entities from the hallucinated outputs, allowing us to identify recurring patterns and quantify their frequency. Based on these findings, we established a comparative ranking of entities according to their tendency to induce hallucinations across different models. The statistical results are shown in Figure 9, objects such as “table”, “shirt”, and “woman” frequently appear in erroneous captions, indicating their potential role in inducing hallucinations. This is likely due to the inherent visual ambiguity or semantic complexity of objects, which makes them harder for models to accurately identify and categorize.

| Model                         | AMBER       |             | GQA         |             | xFlickr&CO  |             | XM3600      |             | AVG         |             |
|-------------------------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|
|                               | Acc         | Macro-F1    | Acc         | Macro-F1    | Acc         | Macro-F1    | Acc         | Macro-F1    | Acc         | Macro-F1    |
| Direct (Team et al., 2024)    | 41.7        | 44.7        | 37.0        | 39.2        | 49.2        | 50.5        | 50.1        | 51.9        | 44.5        | 46.6        |
| CoT (Kojima et al., 2022)     | 49.2        | 52.3        | 53.2        | 54.7        | 56.4        | 58.0        | 60.6        | 62.2        | 54.9        | 56.8        |
| SRO (Lin et al., 2024)        | 49.0        | 52.9        | 51.6        | 53.4        | 55.7        | 58.3        | 58.2        | 60.8        | 53.6        | 56.4        |
| VDGD (Ghosh et al., 2024)     | <b>58.9</b> | <b>60.3</b> | 50.6        | 52.8        | 52.7        | 54.0        | 60.6        | 62.4        | 55.7        | 57.4        |
| HalluciMAD (Lin et al., 2024) | 52.2        | 54.5        | 59.0        | <b>60.6</b> | 61.6        | <b>63.9</b> | 63.7        | 64.9        | 59.1        | 61.0        |
| UniHD (Chen et al., 2024b)    | 56.2        | 52.4        | <b>60.6</b> | <b>60.6</b> | <b>66.1</b> | 59.5        | <b>64.5</b> | <b>65.7</b> | <b>61.8</b> | <b>62.1</b> |

Table 2: Comparison of the UniHD framework and other methods on Gemini. **Bold** represents the Best performance.

### B.3 Framework Adaptation Exploration Experiment Details

We adopt the Unified Multimodal Hallucination Detection (UniHD) framework, which integrates *Claim Extraction*, *Autonomous Tool Selection*, *Tool Execution*, and *Hallucination Verification* for image-to-text and text-to-image hallucination detection. Using tools like object detection, attribute detection, scene text recognition, and common-sense knowledge search, the framework aims to improve performance on CCHall by bridging internal reasoning with real-world context.

Given that the CCHall focuses on image-to-text hallucinations, particularly object hallucinations, we adapt the UniHD framework to better suit the specific requirements of our benchmark. These adjustments are as follows:

1. **Input Adaptation:** For VQA tasks, we combine the question and answer into a declarative statement for clarity, e.g., “What animal is in the box?” and “A bear” become “The animal in the box is a bear”. For Image Captioning, the input is the caption itself. Since CCHall is multilingual and some tools support only specific languages, all non-English inputs are translated to English for compatibility.
2. **Claim Extraction and Query Generation:** The adjusted framework begins by using *Gemini-1.5-Flash* to perform claim extraction, breaking complete sentences into sub-claims for detailed hallucination analysis. For each extracted claim, *Gemini-1.5-Flash* generates specific queries focused on object existence and commonsense verification, aligning with the object-focused nature of our benchmark.
3. **Tool Execution:** The framework automatically invokes *Grounding DINO* (Liu et al., 2024) model to detect objects and return

bounding box coordinates. Additionally, it utilizes the *Serper Google Search API*<sup>1</sup> to perform internet-based fact verification by analyzing search results, comparing sources, and validating the truthfulness of each claim.

4. **Integrated Evaluation:** Evidence collected from the tools is integrated with the original image and its corresponding claims into a comprehensive prompt. *Gemini-1.5-Flash* then evaluates each claim’s hallucination likelihood and provides reasoning, resulting in a final hallucination detection outcome.

During experiments, the adapted UniHD framework executes 3,532 queries via the Serper API, averaging one query per call. As shown in Table 2, UniHD outperforms all other methods, achieving the best results in hallucination detection. These findings demonstrate the effectiveness of external tools and web-based verification for object existence validation, offering a clear advantage over traditional prompting and reasoning strategies. The results highlight the importance of integrating external resources to enhance hallucination detection, particularly for complex object hallucinations.

### B.4 Samples & Error Analysis

In this section, as shown in Figures 12-15, we conduct a detailed analysis of samples and error cases from the *GPT-4o* across the four sub-datasets of the CCHall Benchmark: AMBER, xFlickrCO, GQA, and XM3600. For each of the four sub-datasets, we systematically collect and analyze error cases observed in *GPT-4o*’s test-time performance. Through this detailed evaluation, we identify recurring failure patterns, domain-specific challenges, and systematic shortcomings. We aim to surface not only the nature of current limitations but also to quantify the actionable headroom.

<sup>1</sup><https://serper.dev>

| Model                                             | AMBER       |             | GQA         |             | xFlickr&CO  |             | XM3600      |             | AVG         |             |
|---------------------------------------------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|
|                                                   | Acc         | Macro-F1    | Acc         | Macro-F1    | Acc         | Macro-F1    | Acc         | Macro-F1    | Acc         | Macro-F1    |
| Random                                            | 25.1        | 30.0        | 25.0        | 29.4        | 24.9        | 30.3        | 25.1        | 29.9        | 25.0        | 29.9        |
| <i>Qwen2-VL-2B-Instruct (Wang et al., 2024a)</i>  |             |             |             |             |             |             |             |             |             |             |
| Direct (Wang et al., 2024a)                       | 25.1        | 33.8        | 27.1        | 35.8        | 23.9        | 33.1        | 23.7        | 33.2        | 24.9        | 34.0        |
| CoT (Kojima et al., 2022)                         | <u>25.2</u> | 43.4        | <u>27.2</u> | 35.0        | 24.1        | 52.1        | <u>24.3</u> | <u>52.5</u> | <u>25.2</u> | <u>45.7</u> |
| SRO (Lin et al., 2024)                            | 24.7        | <u>43.5</u> | 24.5        | 34.8        | <u>24.6</u> | <u>52.3</u> | 24.2        | 52.4        | 24.5        | <u>45.7</u> |
| VDGD (Ghosh et al., 2024)                         | 23.7        | 36.3        | 19.8        | <u>38.8</u> | 20.2        | 44.2        | 23.7        | 41.9        | 21.8        | 40.3        |
| HalluciMAD (Lin et al., 2024)                     | 20.9        | 42.5        | 20.1        | 38.1        | 19.9        | 45.4        | 22.8        | 44.7        | 20.9        | 42.7        |
| <i>InternVL2-2B (Chen et al., 2024c)</i>          |             |             |             |             |             |             |             |             |             |             |
| Direct (Chen et al., 2024c)                       | <u>24.4</u> | 28.1        | <u>24.1</u> | 26.0        | <u>24.6</u> | 26.9        | <u>28.1</u> | 27.3        | <u>25.3</u> | 27.1        |
| CoT (Kojima et al., 2022)                         | 20.6        | 26.1        | 20.1        | 24.9        | 23.8        | 27.6        | 22.4        | 27.7        | 21.7        | 26.6        |
| SRO (Lin et al., 2024)                            | 9.7         | <u>28.7</u> | 10.3        | <u>28.9</u> | 4.3         | <u>28.4</u> | <u>28.1</u> | <u>30.9</u> | 8.9         | <u>29.2</u> |
| VDGD (Ghosh et al., 2024)                         | 14.9        | 25.0        | 18.4        | 28.1        | 16.3        | 27.1        | 18.4        | 26.9        | 17.0        | 26.8        |
| HalluciMAD (Lin et al., 2024)                     | 13.0        | 25.6        | 11.7        | 25.6        | 12.4        | 25.4        | 13.6        | 27.2        | 12.7        | 26.0        |
| <i>InternVL2-4B (Chen et al., 2024c)</i>          |             |             |             |             |             |             |             |             |             |             |
| Direct (Chen et al., 2024c)                       | <u>25.0</u> | <u>34.9</u> | <u>25.2</u> | <u>35.5</u> | <u>30.2</u> | 38.2        | <u>30.2</u> | 38.3        | <u>27.7</u> | <u>36.7</u> |
| CoT (Kojima et al., 2022)                         | 14.2        | 31.1        | 13.4        | 29.9        | 17.9        | 33.5        | 16.4        | 33.3        | 15.5        | 31.9        |
| SRO (Lin et al., 2024)                            | 12.3        | <u>34.9</u> | 11.8        | 34.7        | 6.6         | 37.8        | 13.2        | <u>38.9</u> | 11.0        | 36.6        |
| VDGD (Ghosh et al., 2024)                         | 20.3        | <u>32.6</u> | 21.8        | 32.4        | 29.7        | <u>39.1</u> | 24.8        | <u>35.9</u> | 24.2        | 35.0        |
| HalluciMAD (Lin et al., 2024)                     | 16.7        | 31.9        | 16.2        | 33.0        | 21.1        | 37.0        | 29.4        | 34.9        | 20.8        | 34.2        |
| <i>InternVL2-8B (Chen et al., 2024c)</i>          |             |             |             |             |             |             |             |             |             |             |
| Direct (Chen et al., 2024c)                       | 29.1        | 38.1        | 29.9        | 38.6        | 38.3        | 47.6        | 38.8        | 47.4        | 34.0        | 42.9        |
| CoT (Kojima et al., 2022)                         | 31.3        | <u>40.0</u> | <u>33.6</u> | <u>42.1</u> | <u>41.6</u> | 48.0        | <u>40.1</u> | 47.6        | <u>36.7</u> | <u>44.4</u> |
| SRO (Lin et al., 2024)                            | 30.3        | <u>40.0</u> | 31.1        | 40.7        | 41.2        | <u>48.6</u> | 37.7        | 47.1        | 35.1        | 44.1        |
| VDGD (Ghosh et al., 2024)                         | <u>33.2</u> | <u>40.0</u> | 30.2        | 37.8        | 36.7        | 44.5        | 37.4        | 44.7        | 34.4        | 41.7        |
| HalluciMAD (Lin et al., 2024)                     | 29.9        | 38.6        | 30.0        | 39.0        | 37.9        | 45.9        | 39.6        | <u>47.7</u> | 34.3        | 42.8        |
| <i>Llama-3.2-11B-Vision-Instruct (Meta, 2024)</i> |             |             |             |             |             |             |             |             |             |             |
| Direct (Meta, 2024)                               | 31.6        | 38.8        | 32.1        | 38.9        | 35.4        | 43.2        | 43.3        | 49.4        | 35.6        | 42.6        |
| CoT (Kojima et al., 2022)                         | 32.0        | 40.6        | 34.3        | 40.9        | 43.6        | 51.6        | <u>46.4</u> | <u>54.0</u> | 39.1        | <u>46.8</u> |
| SRO (Lin et al., 2024)                            | 32.1        | 40.6        | 34.4        | 40.9        | <u>43.7</u> | <u>51.7</u> | <u>46.4</u> | <u>54.0</u> | <u>39.2</u> | <u>46.8</u> |
| VDGD (Ghosh et al., 2024)                         | <u>34.0</u> | <u>41.2</u> | <u>35.6</u> | <u>42.1</u> | 36.6        | 45.4        | 42.4        | 50.7        | 37.1        | 44.9        |
| HalluciMAD (Lin et al., 2024)                     | 29.7        | 38.7        | 32.6        | 40.3        | 36.4        | 45.1        | 34.8        | 42.8        | 33.4        | 41.7        |
| <i>Qwen2-VL-7B-Instruct (Wang et al., 2024a)</i>  |             |             |             |             |             |             |             |             |             |             |
| Direct (Wang et al., 2024a)                       | 36.2        | 39.5        | 33.3        | 38.3        | 42.9        | 46.7        | 39.9        | 44.4        | 38.1        | 42.2        |
| CoT (Kojima et al., 2022)                         | 38.6        | 43.4        | 33.9        | <u>39.0</u> | 48.3        | 52.1        | <u>48.4</u> | <u>52.5</u> | 42.3        | <u>46.7</u> |
| SRO (Lin et al., 2024)                            | 38.7        | <u>43.5</u> | 34.1        | 34.8        | 48.4        | <u>52.3</u> | <u>48.4</u> | 52.4        | <u>42.4</u> | 45.7        |
| VDGD (Ghosh et al., 2024)                         | 34.8        | 36.3        | <u>36.3</u> | 38.8        | 41.1        | 44.2        | 40.0        | 44.9        | 38.1        | 41.1        |
| HalluciMAD (Lin et al., 2024)                     | 37.0        | 42.5        | 31.4        | 38.1        | 38.0        | 45.4        | 38.6        | 44.7        | 36.3        | 42.7        |
| <i>Pixtral-12B-2409 (Agrawal et al., 2024)</i>    |             |             |             |             |             |             |             |             |             |             |
| Direct (Agrawal et al., 2024)                     | 23.9        | 38.0        | 34.8        | 45.5        | 43.3        | 51.3        | 38.0        | 48.9        | 35.0        | 45.9        |
| CoT (Kojima et al., 2022)                         | 40.7        | 46.0        | 42.7        | 47.9        | 48.7        | 54.2        | 53.3        | 59.0        | 46.3        | 51.8        |
| SRO (Lin et al., 2024)                            | 43.0        | 47.8        | 41.3        | 48.1        | 50.6        | 55.5        | 51.6        | 57.8        | 46.6        | 52.3        |
| VDGD (Ghosh et al., 2024)                         | 43.4        | 43.1        | 45.1        | 47.3        | 47.3        | 48.4        | 56.4        | 58.2        | 48.1        | 49.3        |
| HalluciMAD (Lin et al., 2024)                     | <u>46.3</u> | <u>49.8</u> | <u>45.2</u> | <u>51.4</u> | <u>57.1</u> | <u>61.4</u> | <u>58.7</u> | <u>63.2</u> | <u>51.8</u> | <u>56.4</u> |
| <i>Gemini-1.5-Flash (Team et al., 2024)</i>       |             |             |             |             |             |             |             |             |             |             |
| Direct (Team et al., 2024)                        | 41.7        | 44.7        | 37.0        | 39.2        | 49.2        | 50.5        | 50.1        | 51.9        | 44.5        | 46.6        |
| CoT (Kojima et al., 2022)                         | 49.2        | 52.3        | 53.2        | 54.7        | 56.4        | 58.0        | 60.6        | 62.2        | 54.9        | 56.8        |
| SRO (Lin et al., 2024)                            | 49.0        | 52.9        | 51.6        | 53.4        | 55.7        | 58.3        | 58.2        | 60.8        | 53.6        | 56.4        |
| VDGD (Ghosh et al., 2024)                         | <u>58.9</u> | <u>60.3</u> | 50.6        | 52.8        | 52.7        | 54.0        | 60.6        | 62.4        | 55.7        | 57.4        |
| HalluciMAD (Lin et al., 2024)                     | 52.2        | 54.5        | <u>59.0</u> | <u>60.6</u> | <u>61.6</u> | <u>63.9</u> | <u>63.7</u> | <u>64.9</u> | <u>59.1</u> | <u>61.0</u> |
| <i>GPT-4o (OpenAI, 2023)</i>                      |             |             |             |             |             |             |             |             |             |             |
| Direct (OpenAI, 2023)                             | 57.1        | 63.2        | 56.1        | 62.1        | 72.2        | 76.4        | 82.9        | 84.3        | 67.1        | 71.5        |
| CoT (Kojima et al., 2022)                         | 68.1        | 70.0        | 66.9        | 69.2        | 81.1        | 83.4        | 83.1        | 84.8        | 74.8        | 76.8        |
| SRO (Lin et al., 2024)                            | 63.2        | 65.0        | 57.2        | 62.7        | 76.7        | 79.3        | 84.7        | 86.4        | 70.4        | 73.4        |
| VDGD (Ghosh et al., 2024)                         | 64.7        | 66.6        | 56.3        | 62.5        | 72.3        | 77.0        | 83.2        | 85.6        | 69.1        | 73.0        |
| HalluciMAD (Lin et al., 2024)                     | 🔴 70.9      | 🔴 71.9      | 🔴 68.6      | 🔴 70.3      | 🔴 84.1      | 🔴 85.6      | 🔴 86.4      | 🔴 87.3      | 🔴 77.5      | 🔴 78.8      |

Table 3: The experimental results of Acc. (%) and Macro-F1 score on MLLMs. The “Random” refers to the average performance obtained from three separate random selections.   represents the performance of open-source MLLMs, and   represents the performance of closed-source MLLMs. The underline indicates better performance in the MLLM. 🔴 represents the **Best** performance.

### Intention Description:

I want you to act as a **hallucination answer generator**. Given a picture, a question, and the correct answer. The key objective is to **create responses that mistakenly identify objects that do not exist in the image**, focusing specifically on errors in the existence of tangible objects (e.g., people, animals, vehicles). Do not replace non-object nouns like events, actions, or abstract concepts.

You must provide two answers based on the given image, the given question and the correct answer. For each answer, **randomly replace 1-2 object nouns** (e.g., person, tree, car) in the correct answer with similar but different object names that are not present in the image. Ensure that the replaced objects in the two answers are distinct and do not accurately describe the content of the image.

### Example (XM3600 and xFlickr&Co) :

#Question#: Describe this image.

#Right Answer #: **Plumbing part** lying on a wooden table.

#Hallucinated Answer1#: **Camera** lying on a wooden table.

#Hallucinated Answer2#: **Robot** lying on a wooden table.

You should focus on making the answers convincingly hallucinated, ensuring that #Hallucinated Answer1# and #Hallucinated Answer2# clearly misidentify tangible objects not present in the image. Only provide the two hallucinated answers in the required format, nothing else.



### Example (GQA) :

#Question#: What is the woman to the right of the boat holding?

#Right Answer #: **umbrella**

#Hallucinated Answer1#: **stick**

#Hallucinated Answer2#: **handbag**

You should focus on making the answers convincingly hallucinated, ensuring that #Hallucinated Answer1# and #Hallucinated Answer2# clearly misidentify tangible objects not present in the image. Only provide the two hallucinated answers in the required format, nothing else.



### Example (AMBER) :

#Question#: What objects are in this picture?

#Right Answer #: sky, forest, grass, **person**, **lake**, mountain, road

#Hallucinated Answer1#: sky, forest, grass, **monkey**, lake, mountain, road

#Hallucinated Answer2#: sky, forest, grass, person, **river**, mountain, road

You should focus on making the answers convincingly hallucinated, ensuring that #Hallucinated Answer1# and #Hallucinated Answer2# clearly misidentify tangible objects not present in the image. Only provide the two hallucinated answers in the required format, nothing else.



### Task :

#Question#: Describe this image.

#Right Answer #:

#Hallucinated Answer1#:

#Hallucinated Answer2#:

### Output Format :

Your response must only include the following two lines:

#Hallucinated Answer1#: (Hallucinated Answer1 You Generated)

#Hallucinated Answer2#: (Hallucinated Answer2 You Generated)

Each answer should be a comma-separated list of objects. Do not include any additional text, explanations, or formatting outside of these lines.



Figure 10: Structure of the Prompt for Generating Hallucinated Sentences

### ***Intention Description:***

You are a **hallucination detection system** tasked with categorizing hallucinations based on images , questions and answers in two different languages. Your job is to analyze the provided image, question, and both language answers to determine the type of hallucination, if any.

### ***Hallucination Type Explanation***

1. **Non-hallucination**: Both the English and the other language answers are consistent with the content of the image, and they correctly describe what is present in the image.
2. **Cross-modal, non-cross-lingual hallucination**: Both the English and the other language answers are inconsistent with the image content, specifically regarding the existence of objects. The key point is that the both languages make the same error, mentioning the existence of the same object(s) that do not appear in the image.
3. **Cross-lingual, non-cross-modal hallucination**: The English answer is consistent with the image, but the answer in the other language contains an error regarding the existence of objects.
4. **Cross-modal, cross-lingual hallucination**: Both the English and the other language answers are inconsistent with the image, and they differ in the objects they describe as being present. In other words, both languages make errors, but they reference different nonexistent objects.

### ***Task :***

Based on the following information, identify the type of hallucination:

Image: [The provided Image]

Question: [The provided Question]

Answer1: [The provided Answer1]

Answer2: [The provided Answer2]

Options: [The provided Options]

Let's think step by step!

### ***Output Format :***

You MUST provide the reasoning steps and conclude with the final answer in the format: "Option. Hallucination Type". For example: A. Hallucination Type

Figure 11: Example of the Prompt Used for Hallucination Detection and Classification

### AMBER Task

#### Question:

What objects are in this picture?

#### Answer1 (in English):

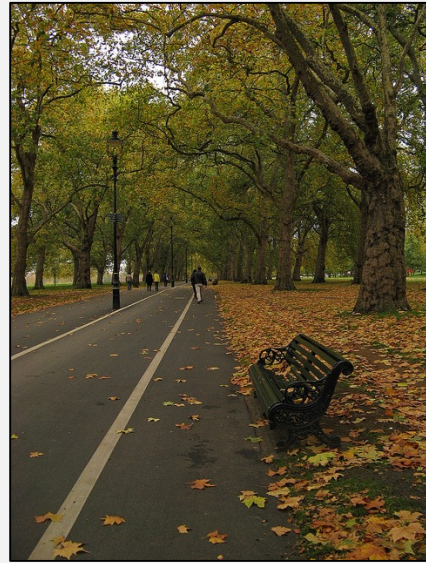
tree, bench, **person**, ground, leave, sky, streetlamp, grass, road

#### Answer2 (in Welsh):

coeden, mainc, **delw**, daear, gadael, awyr, lamp stryd,  
glaswellt, ffordd (tree, bench, **statue**, ground, leave, sky, street  
lamp, grass, road)

#### Options:

- A. Cross-modal, non-cross-lingual hallucination
- B. Non-hallucination
- C. Cross-modal, cross-lingual hallucination
- D. Cross-lingual, non-cross-modal hallucination



### OpenAI GPT-4o :

#### STEP 1:

Examine the provided image to identify the objects visible.

In the image, I can see trees, a bench, a person, the ground, leaves, the sky, streetlamps, grass, and a road.

#### STEP 2:

Analyze Answer1 (English): "tree, bench, person, ground, leave, sky, streetlamp, grass, road".

This answer lists the same objects that are visible in the image.

#### STEP 3:

Analyze Answer2 (Another Language - Welsh): "coeden, mainc, delw, daear, gadael, awyr, lamp stryd, glaswellt, ffordd". Translate these words to English:

- coeden = tree

- mainc = bench

**[Hallucination]**

**- delw = person**

- daear = ground

- gadael = leaves

- awyr = sky

- lamp stryd = streetlamp

- glaswellt = grass

- ffordd = road

**The words are consistent with the objects seen in the image.**

#### STEP 4:

Both Answer1 and Answer2 are consistent with the content of the image, as they both correctly describe objects present in the image.

**Final Answer: B. Non-hallucination.**

#### Label:

D. Cross-lingual, non-cross-modal hallucination

Figure 12: Error Analysis of Gemini on the AMBER

### *xFlickr&CO Task*

#### **Question:**

Describe this image.

#### **Answer1 (in English):**

A **horse** in a red shirt and jeans a toddler on a rock near the water.

#### **Answer2 (in Portuguese):**

Um **cavalo** de camisa vermelha e jeans e uma criança em uma pedra perto da água. (A **horse** in a red shirt and jeans a toddler on a rock near the water.)

#### **Options:**

- A. Cross-modal, cross-lingual hallucination
- B. Cross-lingual, non-cross-modal hallucination
- C. Non-hallucination
- D. Cross-modal, non-cross-lingual hallucination



### *OpenAI GPT-4o :*

#### **STEP 1:**

Image Analysis: The image shows a person (possibly an adult) lying on a rock near water, and there is a child in front of them. The child is wearing a life vest, which suggests that they are near water, possibly preparing for an activity like swimming or boating.

#### **STEP 2:**

Answer1 (English): "A horse in a red shirt and jeans a toddler on a rock near the water." Answer2 (Portuguese): "Um cavalo de camisa vermelha e jeans e uma criança em uma pedra perto da água." Both answers describe a "horse," which is incorrect. There is no horse in the image; instead, it depicts a person and a child near water. The correct description should not mention a horse, as the image clearly shows a human (adult) and a child.

#### **STEP 3:**

**[Hallucination]** Hallucination Type Determination: Both answers contain a major error: they both incorrectly mention the presence of a horse, which does not exist in the image. The error is consistent in both languages (English and Portuguese), but the object (horse) is mentioned in both. This type of hallucination is **\*\*cross-modal\*\*** (both answers describe a non-existent object), **and it is \*\*cross-lingual\*\* (both languages make the same mistake).**

**Final Answer: A. Cross-modal, cross-lingual hallucination.**

#### *Label*

D. Cross-modal, non-cross-lingual hallucination

Figure 13: Error Analysis of Gemini on the xFlickr&CO

### GQA Task

#### Question:

What kind of baked good is the powdered sugar on?

#### Answer1 (in English):

*donuts*

#### Answer2 (in French):

*beignets*

#### Options:

- A. Cross-lingual, non-cross-modal hallucination
- B. Non-hallucination
- C. Cross-modal, cross-lingual hallucination
- D. Cross-modal, non-cross-lingual hallucination



### OpenAI GPT-4o :

#### STEP 1:

**[Hallucination]** Analyze the Image: The image shows a display case with trays containing various types of donuts. The donuts on the lower tray have powdered sugar on them. **No beignets are visible in the image.**

#### STEP 2:

Evaluate Answer1 (English): Answer1 states "donuts," which is consistent with the image. There are indeed donuts with powdered sugar in the display case.

#### STEP 3:

**[Hallucination]** Evaluate Answer2 (French): Answer2 states "beignets," which is inconsistent with the image, as **there are no beignets visible.**

#### STEP 4:

Determine Hallucination Type: Since Answer1 (English) is correct and consistent with the image, while Answer2 (French) is incorrect and inconsistent, this fits the definition of a Cross-lingual, non-cross-modal hallucination. The content error occurs in the French answer but not in the English one.

**Final Answer: A. Cross-lingual, non-cross-modal hallucination.**

#### Label :

- B. Non-hallucination

Figure 14: Error Analysis of Gemini on the GQA

### XM3600 Task

#### Question:

Describe this image.

#### Answer1 (in English):

A macro shot of a *gazelle* in a field.

#### Answer2 (in Swahili):

Risasi kubwa ya *kulungu jangwani* (A big *deer* shot in the *desert*)

#### Options:

- A. Non-hallucination
- B. Cross-modal, non-cross-lingual hallucination
- C. Cross-lingual, non-cross-modal hallucination
- D. Cross-modal, cross-lingual hallucination



### OpenAI GPT-4o :

#### STEP 1:

Analyze the Image: The image depicts a close-up of an antelope with distinct horns. Based on the features, it appears to be a type of gazelle or similar species, with a background that suggests it is in a natural setting like a field or savannah.

#### STEP 2:

**[Hallucination]** English Answer Analysis: "A macro shot of a gazelle in a field.". *This description aligns with the content of the image*, which shows a close-up (macro shot) of an animal that *resembles a gazelle* in a natural setting (field).

#### STEP 3:

**[Hallucination]** Other Language Answer Analysis: "Risasi kubwa ya kulungu jangwani." *The translation of this is "A macro shot of a gazelle in the savanna."* This answer also appears consistent with the image's content, as the background and the animal suggest a setting that could indeed be described as "savanna".

#### STEP 4:

Comparison of Answers: Both the English and the other language answers are consistent with the content of the image. There is no discrepancy or incorrect identification of objects in relation to the image.

**Final Answer: A. Non-hallucination.**

#### Label :

- D. Cross-modal, cross-lingual hallucination

Figure 15: Error Analysis of Gemini on the XM3600