

Improved Baselines with Visual Instruction Tuning

Haotian Liu¹ Chunyuan Li² Yuheng Li¹ Yong Jae Lee¹
¹University of Wisconsin–Madison ²Microsoft Research, Redmond
<https://llava-vl.github.io>

Abstract

Large multimodal models (LMM) have recently shown encouraging progress with visual instruction tuning. In this paper, we present the first systematic study to investigate the design choices of LMMs in a controlled setting under the LLaVA framework. We show that the fully-connected vision-language connector in LLaVA is surprisingly powerful and data-efficient. With simple modifications to LLaVA, namely, using CLIP-ViT-L-336px with an MLP projection and adding academic-task-oriented VQA data with response formatting prompts, we establish stronger baselines that achieve state-of-the-art across 11 benchmarks. Our final 13B checkpoint uses merely 1.2M publicly available data, and finishes full training in ~ 1 day on a single 8-A100 node. Furthermore, we present some early exploration of open problems in LMMs, including scaling to higher resolution inputs, compositional capabilities, and model hallucination, etc. We hope this makes state-of-the-art LMM research more accessible. Code and model will be publicly available.

1. Introduction

Large multimodal models (LMMs) have become increasingly popular in the research community, as they are the key building blocks towards general-purpose assistants [2, 29, 42]. Recent studies on LMMs are converging on a central concept known as visual instruction tuning [35]. The results are promising, *e.g.* LLaVA [35] and MiniGPT-4 [60] demonstrate impressive results on natural instruction-following and visual reasoning capabilities. To better understand the capability of LMMs, multiple benchmarks [16, 26, 33, 36, 54] have been proposed. Recent works further demonstrate improved performance by scaling up the pretraining data [3, 13, 53], instruction-following data [13, 17, 28, 57], visual encoders [3], or language models [38], respectively. The LLaVA architecture is also leveraged in different downstream tasks and domains, including region-level [8, 55] and pixel-level [25, 49] understanding, biomedical assistants [30], image generation [5], adversarial studies [6, 58].

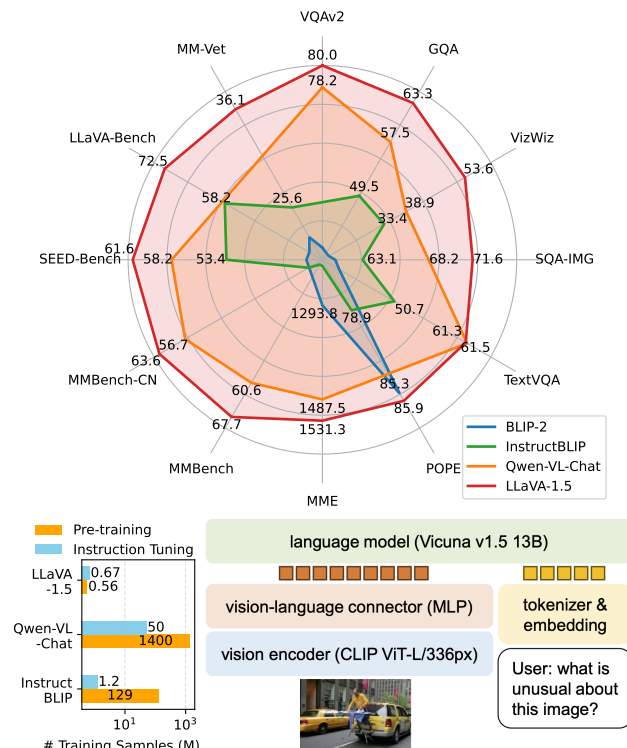


Figure 1. **LLaVA-1.5** achieves SoTA on a broad range of 11 tasks (Top), with high training sample efficiency (Left) and simple modifications to LLaVA (Right): an MLP connector and including academic-task-oriented data with response formatting prompts.

However, despite many benchmarks and developments, it still remains unclear what the best recipe is to train LMMs towards the goal of general-purpose assistants. For example, LLaVA [35] excels in conversational-style visual reasoning and even outperforms later approaches like InstructBLIP [13] on such benchmarks [54], while InstructBLIP excels in traditional VQA benchmarks that demands single-word or short answers. Given the significant differences in the model architecture and training data between them, the root cause of the disparity in their capabilities remains elusive, despite conjectures [36, 54]: the amount of training data, the usage of resamplers like Qformer [31], *etc.* To this

end, we present the first systematic study to investigate the design choices of LMMs in a controlled setting. Our study originates from LLaVA and builds a road map by carefully making effective contributions from the perspectives of the input, model, and data.

First, we unveil that the fully-connected vision-language connector in LLaVA is surprisingly powerful and data-efficient, and we establish stronger and more feasible baselines built upon the LLaVA framework. We report that two simple improvements, namely, an MLP cross-modal connector and incorporating academic task related data such as VQA, are orthogonal to the framework of LLaVA, and when used with LLaVA, lead to better multimodal understanding capabilities. In contrast to InstructBLIP [13] or Qwen-VL [3], which trains specially designed visual resamplers on hundreds of millions or even billions of image-text paired data, LLaVA uses one of the simplest architecture design for LMMs and requires only training a simple fully-connected projection layer on merely 600K image-text pairs. Our final model can finish training in ~ 1 day on a single 8-A100 machine and achieves state-of-the-art results on a wide range of benchmarks. Moreover, unlike Qwen-VL [3] that includes in-house data in training, LLaVA utilizes only publicly available data.

Next, we delve into an early exploration of other open problems of large multimodal models. Our findings include: (1) Scaling to high-resolution image inputs. We show that LLaVA’s architecture is versatile in scaling to higher resolutions by simply dividing images into grids and maintains its data efficiency; with the increased resolution, it improves the model’s detailed perception capabilities and reduces hallucination. (2) Compositional capabilities. We find that large multimodal models are capable of generalizing to compositional capabilities. For example, training on long-form language reasoning together with shorter visual reasoning can improve the model’s writing capability for multimodal questions. (3) Data efficiency. We show that randomly downsampling LLaVA’s training data mixture by up to 75% does not significantly decrease the model’s performance, suggesting that the possibility of a more sophisticated dataset compression strategy can further improve LLaVA’s already efficient training pipeline. (4) Data scaling. We provide empirical evidence for the scaling of data granularity in conjunction with the model’s capability is crucial for an improved capability without introducing artifacts like hallucination.

In sum, we perform a systematic study on the training of large multimodal models, and introduce a simple yet effective approach to balance the multitask learning and effective scaling for large multimodal models. Our improved baselines, LLaVA-1.5, uses only *public* data, achieves the state-of-the-art on a broad range of 11 tasks, and is significantly more data-efficient than previous approaches. By

rethinking the conventional approaches and exploring the open problems in visual instruction tuning, we pave the way for more robust and capable systems for LMMs. We hope these improved and easily-reproducible baselines will provide a reference for future research in open-source LMMs.

2. Related Work

Instruction-following large multimodal models (LMMs).

Common architectures include a pre-trained visual backbone to encode visual features, a pre-trained large language model (LLM) to comprehend the user instructions and produce responses, and a vision-language cross-modal connector to align the vision encoder outputs to the language models. As shown in Fig. 1, LLaVA [35] is perhaps the simplest architecture for LMMs. Optionally, visual resamplers (*e.g.* Qformer [31]) are used to reduce the number of visual patches [3, 13, 60]. Training an instruction-following LMM usually follows a two-stage protocol. First, the vision-language alignment pretraining stage leverages image-text pairs to align the visual features with the language model’s word embedding space. Earlier works utilize relatively few image-text pairs (*e.g.* $\sim 600K$ [35] or $\sim 6M$ [60]), while some recent works pretrain the vision-language connector for a specific language model on a large amount of image-text pairs (*e.g.* 129M [13] and 1.4B [3]), to maximize the LMM’s performance. Second, the visual instruction tuning stage tunes the model on visual instructions [35], to enable the model to follow users’ diverse requests on instructions that involve the visual contents. Dealing with higher resolution with grids in LMM are studied in con-current works [1, 27, 52].

Multimodal instruction-following data. In NLP, studies show that the quality of instruction-following data largely affects the capability of the resulting instruction-following models [59]. For visual instruction tuning, LLaVA [35] is the pioneer to leverage text-only GPT-4 to expand the existing COCO [34] bounding box and caption dataset to a multimodal instruction-following dataset that contains three types of instruction-following data: conversational-style QA, detailed description, and complex reasoning. LLaVA’s pipeline has been employed to expand to textual understanding [56], million-scales [57], and region-level conversations [8]. InstructBLIP [13] incorporates academic-task-oriented VQA datasets to further enhance the model’s visual capabilities. Conversely, [7] identifies that such naive data merging can result in models that tend to overfit to VQA datasets and thus are unable to participate in natural conversations. The authors further propose to leverage the LLaVA pipeline to convert VQA datasets to a conversational style. While this proves effective for training, it introduces added complexities in data scaling. However, in NLP, the FLAN family [12, 50] shows that adding a large number of academic language tasks for instruction tuning can effectively improve the gen-

eralization ability. In light of this, we consider investigating the root cause of the inability to balance between natural conversations and academic tasks in multimodal models.

3. Approach

3.1. Preliminaries

As the seminal work of visual instruction tuning, LLaVA [35] showcases commendable proficiency in visual reasoning capabilities, surpassing even more recent models on diverse benchmarks [4, 54] for real-life visual instruction-following tasks. LLaVA uses a single linear layer to project the visual features to language space, and optimizes the whole LLM for visual instruction tuning. However, LLaVA falls short on academic benchmarks that typically require short-form answers (*e.g.* single-word), and tends to answer *yes* for yes/no questions due to the lack of such data in the training distribution.

On the other hand, InstructBLIP [13] is the pioneer to incorporate academic-task-oriented datasets like VQA-v2 [18] along with LLaVA-Instruct [35], and demonstrates improved performance on VQA benchmarks. It pretrains Qformer [31] on 129M image-text pairs and only finetunes the instruction-aware Qformer for visual instruction tuning. However, recent studies [7, 54] show that it does not perform as well as LLaVA on engaging in real-life visual conversation tasks. More specifically, as shown in Table 1a, it can overfit to VQA training sets with short-answers, even on requests that require detailed responses.

3.2. Response Format Prompting

We find that the inability [7] to balance between short- and long-form VQA for approaches like InstructBLIP [13], which leverages instruction following data that includes both natural responses and short-answers, is mainly due to the following reasons. First, *ambiguous prompts on the response format*. For example, *Q: {Question} A: {Answer}*. Such prompts do not clearly indicate the desired output format, and can overfit an LLM behaviorally to short-form answers even for natural visual conversations. Second, *not finetuning the LLM*. The first issue is worsened by InstructBLIP only finetuning the Qformer for instruction-tuning. It requires the Qformer’s visual output tokens to control the length of the LLM’s output to be either long-form or short-form, as in prefix tuning [32], but Qformer may lack the capability of properly doing so, due to its limited capacity compared with LLMs like LLaMA.

Thus, to enable LLaVA to better handle short-form answers while addressing the issues of InstructBLIP, we propose to use a single response formatting prompt that clearly indicates the output format. It is appended at the end of VQA questions when promoting short answers: *Answer the question using a single word or phrase*. We find that when the

Visual input example, Multitask Balancing Problem:



User	Is this unusual? Please explain in detail.
InstructBLIP	yes

(a) Example of InstructBLIP [13] (Vicuna-13B) having difficulty balancing between short- and long-form answers.

Visual input example, Different Format Prompts:

Normal prompt	What is the color of the shirt that the man is wearing?
Response	The man is wearing a yellow shirt.
Ambiguous prompt	Q: What is the color of the shirt that the man is wearing? A:
Response	The man is wearing a yellow shirt.
Formatting prompt	What is the color of the shirt that the man is wearing? Answer the question using a single word or phrase.
Response	Yellow.

(b) Comparison of how different prompts regularize the output format. The results are obtained zero-shot directly after LLaVA undergoes the first-stage vision-language alignment pretraining, without the second-stage visual instruction tuning.

Table 1. Visual input example to illustrate the challenge of (a) multitask balancing and (b) different format prompts. The same image input is used.

LLM is *finetuned* with such prompts, LLaVA is able to properly adjust the output format according to the user’s instructions (see Table 1b), and does not require additional processing of the VQA answers using ChatGPT [7], which further enables scaling to various data sources. As shown in Table 2, by merely including VQAv2 [18] in training, LLaVA’s performance on MME significantly improves (1323.8 vs 809.6) and outperforms InstructBLIP by 111 points.

3.3. Scaling the Data and Model

MLP vision-language connector. Inspired by the improved performance in self-supervised learning by changing from a linear projection to an MLP [9, 10], we find that improving the vision-language connector’s representation power with a two-layer MLP can improve LLaVA’s multimodal capabilities, compared with the original linear projection.

Academic task oriented data. We further include additional academic-task-oriented VQA datasets for VQA, OCR, and region-level perception, to enhance the model’s capabilities in various ways, as shown in Table 2. We first include

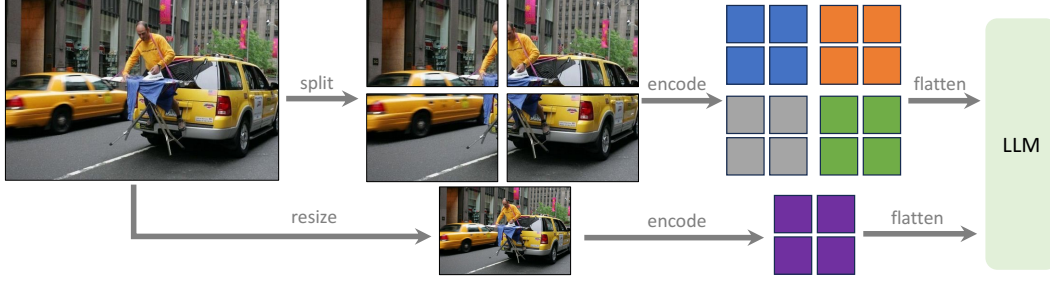


Figure 2. **LLaVA-1.5-HD**. Scaling LLaVA-1.5 to higher resolutions by splitting the image into grids and encoding them independently. This allows the model to scale to any resolution, without performing positional embedding interpolation for ViTs. We additionally concatenate the feature of a downsampled image to provide the LLM with a global context.

Method	LLM	Res.	GQA	MME	MM-Vet
InstructBLIP	14B	224	49.5	1212.8	25.6
<i>Only using a subset of InstructBLIP training data</i>					
0 LLaVA	7B	224	—	809.6	25.5
1 +VQA-v2	7B	224	47.0	1197.0	27.7
2 +Format prompt	7B	224	46.8	1323.8	26.3
3 +MLP VL connector	7B	224	47.3	1355.2	27.8
4 +OKVQA/OCR	7B	224	50.0	1377.6	29.6
<i>Additional scaling</i>					
5 +Region-level VQA	7B	224	50.3	1426.5	30.8
6 +Scale up resolution	7B	336	51.4	1450	30.3
7 +GQA	7B	336	62.0*	1469.2	30.7
8 +ShareGPT	7B	336	62.0*	1510.7	31.1
9 +Scale up LLM	13B	336	63.3*	1531.3	36.1

Table 2. **Scaling results** on data, model, and resolution. We choose to conduct experiments on GQA [20], MME [16], and MM-Vet [54] to examine the representative capabilities of VQA with short answers, VQA with output formatting, and natural visual conversations, respectively. *Training images of GQA were observed during training.

four additional datasets that are used in InstructBLIP: open-knowledge VQA (OKVQA [40], A-OKVQA [44]) and OCR (OCR-VQA [41], TextCaps [46]). A-OKVQA is converted to multiple choice questions and a specific response formatting prompt is used: *Answer with the option’s letter from the given choices directly*. With only a subset of the datasets InstructBLIP uses, LLaVA already surpasses it on all three tasks in Table 2, suggesting LLaVA’s effective design. Furthermore, we find further adding region-level VQA datasets (Visual Genome [24], RefCOCO [23, 39]) improves the model’s capability of localizing fine-grained visual details.

Additional scaling. We further scale up the input image resolution to 336^2 to allow the LLM to clearly “see” the details of images, by swapping the vision encoder to CLIP-ViT-L-336px (the highest resolution available for CLIP). In addition, we add the GQA dataset as an additional visual knowledge source. We also incorporate ShareGPT [45] data and scale up the LLM to 13B as in [3, 8, 38]. Results on MM-

Vet shows the most significant improvement when scaling the LLM to 13B, suggesting the importance of the base LLM’s capability for visual conversations.

LLaVA-1.5. We denote this final model with all the modifications as LLaVA-1.5 (the last two rows in Table 2), which achieves an impressive performance that significantly outperforms the original LLaVA [35].

Computational cost. For LLaVA-1.5, we use the same pretraining dataset, and keep the training iterations and batch size roughly the same for instruction tuning as LLaVA [35]. Due to the increased image input resolution to 336^2 , the training of LLaVA-1.5 is $\sim 2\times$ as long as LLaVA: ~ 6 hours of pretraining and ~ 20 hours of visual instruction tuning, using $8\times$ A100s.

3.4. Scaling to Higher Resolutions

In Sec. 3.3, we observe the advantage that scaling up the input image resolution improves the model’s capabilities. However, the image resolution of the existing open source CLIP vision encoders is limited to 336^2 , preventing the support of higher resolution images by simply replacing the vision encoder as we did in Sec. 3.3. In this section, we present an early exploration of scaling the LLM to higher resolutions, while maintaining the data efficiency of LLaVA-1.5.

When using ViT [14] as the vision encoder, to scale up the resolution, previous approaches mostly choose to perform positional embedding interpolation [3, 31] and adapt the ViT backbone to the new resolution during finetuning. However, this usually requires the model to be finetuned on a large-scale image-text paired dataset [3, 31], and limits the resolution of the image to a fixed size that the LLM can accept during inference.

Instead, as shown in Fig. 2, we overcome this by dividing the image into smaller image patches of the resolution that the vision encoder is originally trained for, and encode them independently. After obtaining the feature maps of individual patches, we then combine them into a single large feature

Method	LLM	Image Size	Sample Size		VQAv2 [18]	GQA [20]	VisWiz [19]	SciQA-IMG [37]	TextVQA [47]
			Pretrain	Finetune					
BLIP-2 [31]	Vicuna-13B	224 ²	129M	-	65.0	41	19.6	61	42.5
InstructBLIP [13]	Vicuna-7B	224 ²	129M	1.2M	-	49.2	34.5	60.5	50.1
InstructBLIP [13]	Vicuna-13B	224 ²	129M	1.2M	-	49.5	33.4	63.1	50.7
Shikra [8]	Vicuna-13B	224 ²	600K	5.5M	77.4*	-	-	-	-
IDEFICS-9B [21]	LLaMA-7B	224 ²	353M	1M	50.9	38.4	35.5	-	25.9
IDEFICS-80B [21]	LLaMA-65B	224 ²	353M	1M	60.0	45.2	36.0	-	30.9
Qwen-VL [3]	Qwen-7B	448 ²	1.4B [†]	50M [†]	78.8*	59.3*	35.2	67.1	63.8*
Qwen-VL-Chat [3]	Qwen-7B	448 ²	1.4B*	50M [†]	78.2*	57.5*	38.9	68.2	<u>61.5*</u>
LLaVA-1.5	Vicuna-7B	336 ²	558K	665K	78.5*	62.0*	50.0	66.8	58.2
LLaVA-1.5	Vicuna-13B	336 ²	558K	665K	80.0*	63.3*	53.6	71.6	<u>61.3</u>
LLaVA-1.5-HD	Vicuna-13B	448 ²	558K	665K	81.8*	64.7*	57.5	<u>71.0</u>	<u>62.5</u>
Specialist SOTA: PaLI-X-55B [11]					86.1*	72.1*	70.9*	-	71.4*

Table 3. **Comparison with SoTA methods on academic-task-oriented datasets.** LLaVA-1.5 achieves the best performance on 4/5 benchmarks, and ranks the second on the other. *The training images/annotations of the datasets are observed during training. [†]Includes in-house data that is not publicly accessible.

Method	POPE [33]			MME [16]	MMBench [36]		SEED-Bench [26]			LLaVA-Wild [35]	MM-Vet [54]
	rand	pop	adv		en	cn	all	img	vid		
BLIP2-14B [31]	89.6	85.5	80.9	1293.8	-	-	46.4	49.7	36.7	38.1	22.4
InstructBLIP-8B [13]	-	-	-	-	36	23.7	53.4	58.8	38.1	60.9	26.2
InstructBLIP-14B [13]	<u>87.7</u>	77	72	1212.8	-	-	-	-	-	58.2	25.6
Shikra-13B [8]	-	-	-	-	58.8	-	-	-	-	-	-
IDEFICS-9B [21]	-	-	-	-	48.2	25.2	-	44.5	-	-	-
IDEFICS-80B [21]	-	-	-	-	54.5	38.1	-	53.2	-	-	-
Qwen-VL [3]	-	-	-	-	38.2	7.4	56.3	62.3	39.1	-	-
Qwen-VL-Chat [3]	-	-	-	1487.5	60.6	56.7	<u>58.2</u>	65.4	37.8	-	-
LLaVA-7B [35]	76.3	72.2	70.1	809.6	38.7	36.4	33.5	37.0	23.8	62.8	25.5
LLaVA-1.5-7B	87.3	<u>86.1</u>	<u>84.2</u>	<u>1510.7</u>	<u>64.3</u>	58.3	<u>58.6</u>	<u>66.1</u>	37.3	65.4	<u>31.1</u>
LLaVA-1.5-13B	87.1	86.2	84.5	1531.3	67.7	63.6	61.6	68.2	42.7	72.5	36.1
LLaVA-1.5-13B-HD	87.5	86.4	85.0	1500.1	68.8	<u>61.9</u>	62.6	70.1	<u>41.3</u>	<u>72.0</u>	39.4

Table 4. **Comparison with SoTA methods on benchmarks for instruction-following LMMs.** LLaVA-1.5 achieves the best overall performance.

map of the target resolution, and feed that into the LLM. To provide the LLM with the global context and to reduce the artifact of the split-encode-merge operation, we additionally concatenate the feature of a downsampled image to the merged feature map. This allows us to scale the input to any arbitrary resolution and maintain the data efficiency of LLaVA-1.5. We call this resulting model LLaVA-1.5-HD.

4. Empirical Evaluation

4.1. Benchmarks

We evaluate LLaVA-1.5 on a collection of both academic-task-oriented benchmarks and recent benchmarks specifically proposed for instruction-following LMMs, totalling 12 benchmarks. For academic-task-oriented benchmarks, VQA-v2 [18] and GQA [20] evaluate model’s visual perception capabilities on open-ended short answers. VizWiz [19] contains 8,000 images to evaluate model’s zero-shot generalization on visual questions asked by visually impaired people. Following InstructBLIP [13], the image subset of

ScienceQA [37] with multiple choice are used to evaluate the zero-shot generalization on scientific question answering. TextVQA [47] contains text-rich visual question answering.

For recent benchmarks proposed for instruction-following LMMs, POPE [33] evaluates model’s degree of hallucination on three sampled subsets of COCO [34]: random, common, and adversarial and we report the F1 score on all three splits. Other benchmarks evaluate the model’s capabilities on a wide range of domains and applications, with different response formats. MME-Perception [16] evaluates model’s visual perception with yes/no questions. MMBench [36] evaluates model’s answer robustness with all-round shuffling on multiple choice answers. MMBench-CN [36] is the Chinese-translated version of MMBench. SEED-Bench [26] evaluates model’s performance on both images and videos with multiple choice, and we sample the frame in the middle to evaluate the accuracy on videos. LLaVA-Bench-in-the-Wild [35] and MM-Vet [54] evaluate model’s capabilities in engaging in visual conversations on a diverse range of tasks, and evaluates the correctness and the helpfulness of

Visual input example, Tricky Question:



User	If there are factual errors in the questions, point it out; if not, proceed to answering the question. What's happening in the desert?
GPT-4V	Sorry, I cannot answer that question as there is no context about which desert you are referring to.
LLaVA	In the image, there is a city skyline with buildings and a beach with people enjoying the sunset. The scene is set in a desert, which is unusual for a beach setting.
LLaVA-1.5	There are no deserts in the image. The image features a beach with palm trees, a city skyline, and a large body of water.

Table 5. LLaVA-1.5 can detect and answer tricky questions when prompted to verify the question.

the response with GPT-4 evaluation.

4.2. Results

We show that LLaVA-1.5 achieves the best overall performance on 12 benchmarks, despite using magnitudes smaller pretraining and instruction tuning data compared with other methods [3, 13]. LLaVA-1.5 significantly outperforms LLaVA on all benchmarks for instruction-following LMMs. Note that it is challenging to evaluate the original LLaVA on academic datasets like VQA-v2 [18] that demand open-ended short answers.

When we continue to scale up the image resolution to 448² with LLaVA-1.5-HD, it further improves the overall performance on all benchmarks, especially on tasks that require perception of details in the images (*e.g.* OCR in MM-Vet, detailed description in LLaVA-Bench-in-the-Wild [35]). Moreover, we find that adding the global context effectively recovers the model from the split-and-merge artifacts and guides the model to more easily locate the relevant regions from the high-resolution features (see appendix).

It is encouraging that *LLaVA-1.5 achieves the best performance with the simplest architecture, academic compute and public datasets, and yields a fully-reproducible and affordable baseline for future research.* The results also suggest that visual instruction tuning plays an important role in improving an LMM’s capabilities, and raises questions upon the common belief that LMMs require significant amount of vision-language alignment pretraining [3, 13, 31], despite that the vision encoders (*e.g.* CLIP [43], OpenCLIP [22], EVA-CLIP [15], *etc.*) are already pretrained on web-scale image-text paired data. LLaVA-1.5 (even the 7B model)

Visual input example, Constrained JSON Output:



User	Please read the text in this image and return the information in the following JSON format (note xxx is placeholder, if the information is not available in the image, put "N/A" instead). { "class": xxx, "DLN": xxx, "DOB": xxx, "Name": xxx, "Address": xxx, "EXP": xxx, "ISS": xxx, "SEX": xxx, "HGT": xxx, "WGT": xxx, "EYES": xxx, "HAIR": xxx, "DONOR": xxx }
GPT-4V	{ "class": "C", "DLN": "1234568", "DOB": "08/31/1977", "Name": "Ima Cardholder", "Address": "2570 24th Street, Anytown, CA 95818", "EXP": "08/31/2014", "ISS": "08/31/2011", "SEX": "F", "HGT": "5-05", "WGT": "125 lbs", "EYES": "BRN", "HAIR": "BRN", "DONOR": "VETERAN" }
LLaVA	{ "class": "Driver License", "DLN": "CA123 456 789", "DOB": "01/01/1990", "Name": "Amy", "Address": "Anytown, CA", "EXP": "08/31/2023", "ISS": "Veteran", "SEX": "F", "HGT": "5'5", "WGT": "130", "EYES": "Brown", "HAIR": "Brown", "DONOR": "Unknown" }
LLaVA-1.5	{ "class": "C", "DLN": "1132456789", "DOB": "08/23/1971", "Name": "Ima Cardholder", "Address": "Anytown, State, Zip", "EXP": "08/23/2014", "ISS": "California", "SEX": "F", "HGT": "5'5", "WGT": "123", "EYES": "Brown", "HAIR": "Brown", "DONOR": "N/A" }

Table 6. LLaVA-1.5 can extract information from the image and answer following the required format, despite a few errors compared with GPT-4V. GPT-4V results are obtained from [51].

outperforms 80B IDEFICS [21], a Flamingo-like LMM with billions of trainable parameters for cross-modal connection. This also makes us rethink the benefits of the vision samplers and the necessity of the additional large-scale pretraining, in terms of multimodal instruction-following capabilities.

Global context. For higher resolution, we pad and resize the image to a single image of 224², and concatenate it with the high resolution features to provide a global context. Ablation on a 7B model shows that the global context effectively boosts performance on all three validation benchmarks.

	GQA	MME	MM-Vet
high-res patch only	62.9	1425.8	31.9
+global context	63.8 (+0.9)	1497.5 (+71)	35.1 (+3.2)

4.3. Emerging Properties

Format instruction generalization. Although LLaVA-1.5 is only trained with a limited number of format instructions, it generalizes to others. First, VizWiz [19] requires the model

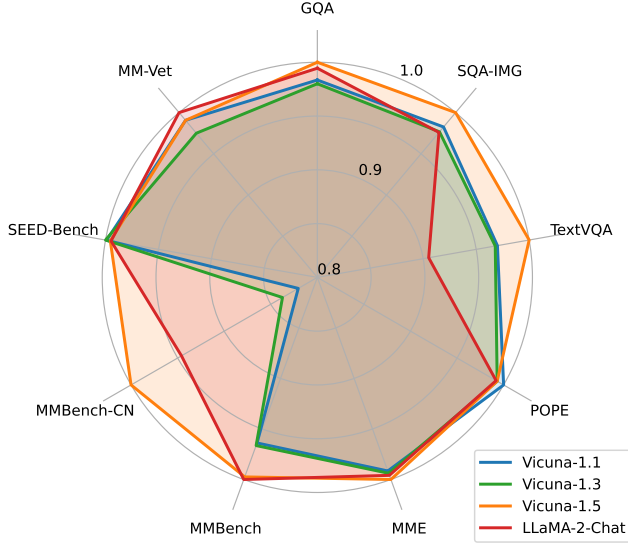


Figure 3. **Ablation on LLM choices.** Data points represent the relative performance of the best performing variant for each dataset.

to output “Unanswerable” when the provided content is insufficient to answer the question, and our response format prompt (see Appendix) effectively instructs the model to do so (11.1% $\rightarrow$ 67.8% on unanswerable questions). We additionally present qualitative examples on instructing LLaVA-1.5 to verify tricky questions (Fig. 5), respond in a constrained JSON format (Fig. 6), and more in appendix.

Multilingual multimodal capability. Though LLaVA-1.5 is *not* finetuned for multilingual multimodal instruction following *at all* (all visual instructions including VQA are in English), we find that it is capable of following multilingual instructions. This is partly due to the multilingual language instructions in ShareGPT [45]. Although ShareGPT does not contain images in its instructions, the model learns from this dataset the behavior of adaptively responding with the language that corresponds to the user’s request. We empirically show that this behavior is transferred to visual conversations. We also quantitatively evaluate the model’s generalization capability to Chinese on MMBench-CN [36], where the questions of MMBench are converted to Chinese. Notably, LLaVA-1.5 outperforms Qwen-VL-Chat by +7.3% (63.6% vs 56.7%), despite Qwen being finetuned on Chinese multimodal instructions while LLaVA-1.5 is not.

4.4. Ablation on LLM Choices

In NLP, findings [48] suggest that the capability of the base LLM can affect its instruction-tuned successors. In this section, we explore two families of LLMs and study their contribution to the final model’s multimodal capability: LLaMA-based (Vicuna-v1.1, Vicuna-v1.3) and LLaMA-2-based (Vicuna-v1.5, LLaMA-2-Chat). Vicuna-v1.3 and Vicuna-v1.5 use the same $\sim 150\text{K}$ ShareGPT [45] data ($2\times$

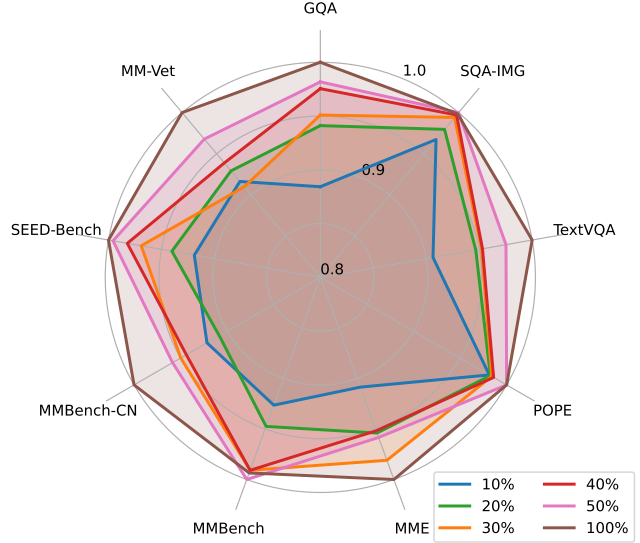


Figure 4. **Ablation on data efficiency.** Data points represent the relative performance of the best performing variant for each dataset.

that used in v1.1). Unlike Vicuna series that is only trained with supervised instruction finetuning (SFT), LLaMA-2-Chat is further optimized with reinforcement-learning from human-feedback (RLHF). We visualize the relative performance of these variants in Fig. 3.

First, we find that Vicuna-v1.5 achieves the best overall performance, and LLaMA-2-based models generally perform better than LLaMA-1-based, suggesting the importance of the base language model. This is further evidenced by the results on MMBench-CN [36]: despite Vicuna-v1.3 and v1.5 using the same ShareGPT data for instruction tuning, the performance in generalization to Chinese of Vicuna-v1.3 is significantly worse than v1.5.

Second, language instruction-tuning matters on specific capabilities that are required by each dataset. For example, although LLaMA-2-Chat and Vicuna-v1.5 achieves almost the same performance on MMBench, the generalization to MMBench-CN [36] of LLaMA-2-Chat is worse than Vicuna-v1.5, which is partly due to that the most SFT/RLHF data of LLaMA-2-Chat is in English and does not contain as many multilingual data as in ShareGPT. Furthermore, TextVQA requires both the model’s capability of identifying the text characters in the images, and also processing the noisy outputs from the OCR engine; such noise *may* be more commonly observed in the ShareGPT data, which is collected in-the-wild from daily usage of ChatGPT.

5. Open Problems in LMMs

Given the successful scaling of LLaVA-1.5, we conduct additional studies on open problems in LMMs using the model design and data mixture of LLaVA-1.5.

5.1. Data Efficiency

Despite the data efficiency of LLaVA-1.5 when compared with approaches like InstructBLIP [13], the training of LLaVA-1.5 still doubles when compared with LLaVA. In this section, we conduct experiments for further improving the data efficiency by randomly sub-sampling the training data mixture of LLaVA-1.5, with a sampling ratio ranging from 0.1 to 0.5. We visualize the relative performance of different sampling variants in Fig. 4.

First, the full data mixture provides the best knowledge coverage, and allows the model to achieve the best overall performance. To our surprise, with only 50% of the samples, the model still maintains more than 98% of the full dataset performance. This suggests that there is room for further improvements in data efficiency.

Second, when downsampling the dataset to 50%, the model’s performance on MMBench, ScienceQA, and POPE does not decrease at all, and it even slightly improves on MMBench. Similarly, the model’s performance remains steady when further downscaling the data from 50% to 30%. These results show promise of having the less-is-more [59] benefit for multimodal models as well.

5.2. Rethinking Hallucination in LMMs

Hallucination is an important issue to tackle for LLMs and LMMs. Often in LMMs, we attribute the model’s hallucination to the errors or hallucinations in the training dataset. For example, the detailed descriptions in LLaVA-Instruct [35] may contain a small amount of hallucinated content, and it is believed that training on such data *may* have caused the model to hallucinate when asked to “describe the image in detail”. However, we find that such hallucination is significantly reduced, when we scale the model’s inputs to higher resolutions like 448².

This finding is interesting as it suggests that the LMMs may be robust to *a few* such errors in the training data. However, when the input resolution is not sufficient for the model to discern all details in the training data, and the amount of data that is at that granularity beyond the model’s capability becomes large enough, the model *learns* to hallucinate. This further suggests that there needs to be a balance between improving the data annotation with more details and the model’s capability to properly process the information at such granularities. We hope this finding provides a reference for future work in terms of dealing with hallucination and the scaling of the models and data.

5.3. Compositional Capabilities

We demonstrate interesting compositional capabilities in LLaVA-1.5: the model trained on a set of tasks independently generalizes to tasks that require a combination of these capabilities without explicit joint training. We note some of the findings below.

First, we observe an improved language capability in visual conversations after including the ShareGPT [45] data, including the multimodal multilingual capability as discussed in Sec. 4.3. Moreover, the model is more capable at providing longer and more detailed responses in visual conversations. Second, the additional visual knowledge from the academic-task-oriented datasets, improves the visual groundness of LLaVA-1.5’s responses in visual conversations, as evidenced quantitatively by the improved results on MM-Vet [54] and LLaVA-Wild [35] in Table 4.

However, there is still difficulty in achieving ideal performance for some tasks that require a certain combination of capabilities. For example, being able to correctly answer the attribute of a certain object in VQA, does not guarantee an accurate depiction of that object attribute in a detailed description of the whole image. Furthermore, the capability of engaging in conversations with certain foreign languages (*e.g.* Korean) still falls behind. See appendix for examples.

These findings suggest that the compositional capabilities of LMMs can be leveraged to improve the model’s performance without significantly increasing the data by exhaustively including all task combinations. Yet, it can be further investigated, and a deeper understanding of the mechanism behind the compositional capabilities of LMMs can further improve the capability and the data efficiency of LLaVA-1.5.

6. Conclusion

In this paper, we take a step towards demystifying the design of large multimodal models, and propose a simple, effective, and data-efficient baseline, LLaVA-1.5, for large multimodal models. In addition, we explore the open problems in visual instruction tuning, scale LMMs to higher resolutions, and present some intriguing findings in terms of model hallucination and compositional capabilities for LMMs. We hope these improved and easily-reproducible baselines as well as the new findings will provide a reference for future research in open-source LMM.

Limitations. Despite the promising results demonstrated by LLaVA-1.5, it still has limitations including prolonged training for high-resolution images, lack of multiple-image understanding, limited problem solving capabilities in certain fields. It is not exempt from producing hallucinations, and should be used with caution in critical applications (*e.g.* medical). See appendix for a detailed discussion.

Acknowledgements. This work was supported in part by NSF CAREER IIS2150012, and Institute of Information & communications Technology Planning & Evaluation(IITP) grants funded by the Korea government(MSIT) (No. 2022-0-00871, Development of AI Autonomy and Knowledge Enhancement for AI Agent Collaboration) and (No. RS-2022-00187238, Development of Large Korean Language Model Technology for Efficient Pre-training).

References

- [1] Adept AI. Fuyu-8b: A multimodal architecture for ai agents. <https://www.adept.ai/blog/fuyu-8b>, 2024. **2**
- [2] Jean-Baptiste Alayrac, Jeff Donahue, Pauline Luc, Antoine Miech, Iain Barr, Yana Hasson, Karel Lenc, Arthur Mensch, Katie Millican, Malcolm Reynolds, et al. Flamingo: a visual language model for few-shot learning. *arXiv preprint arXiv:2204.14198*, 2022. **1**
- [3] Jinze Bai, Shuai Bai, Shusheng Yang, Shijie Wang, Sinan Tan, Peng Wang, Junyang Lin, Chang Zhou, and Jingren Zhou. Qwen-vl: A frontier large vision-language model with versatile abilities. *arXiv preprint arXiv:2308.12966*, 2023. **1, 2, 4, 5, 6**
- [4] Yonatan Bitton, Hritik Bansal, Jack Hessel, Rulin Shao, Wanrong Zhu, Anas Awadalla, Josh Gardner, Rohan Taori, and Ludwig Schimdt. Visit-bench: A benchmark for vision-language instruction following inspired by real-world use, 2023. **3**
- [5] Kevin Black, Michael Janner, Yilun Du, Ilya Kostrikov, and Sergey Levine. Training diffusion models with reinforcement learning. *arXiv preprint arXiv:2305.13301*, 2023. **1**
- [6] Nicholas Carlini, Milad Nasr, Christopher A Choquette-Choo, Matthew Jagielski, Irena Gao, Anas Awadalla, Pang Wei Koh, Daphne Ippolito, Katherine Lee, Florian Tramer, et al. Are aligned neural networks adversarially aligned? *arXiv preprint arXiv:2306.15447*, 2023. **1**
- [7] DeLong Chen, Jianfeng Liu, Wenliang Dai, and Baoyuan Wang. Visual instruction tuning with polite flamingo. *arXiv preprint arXiv:2307.01003*, 2023. **2, 3**
- [8] Keqin Chen, Zhao Zhang, Weili Zeng, Richong Zhang, Feng Zhu, and Rui Zhao. Shikra: Unleashing multimodal llm’s referential dialogue magic. *arXiv preprint arXiv:2306.15195*, 2023. **1, 2, 4, 5**
- [9] Ting Chen, Simon Kornblith, Mohammad Norouzi, and Geoffrey Hinton. A simple framework for contrastive learning of visual representations. In *ICML*, 2020. **3**
- [10] Xinlei Chen, Haoqi Fan, Ross Girshick, and Kaiming He. Improved baselines with momentum contrastive learning. *arXiv preprint arXiv:2003.04297*, 2020. **3**
- [11] Xi Chen, Josip Djolonga, Piotr Padlewski, Basil Mustafa, Soravit Changpinyo, Jialin Wu, Carlos Riquelme Ruiz, Sebastian Goodman, Xiao Wang, Yi Tay, et al. Pali-x: On scaling up a multilingual vision and language model. *arXiv preprint arXiv:2305.18565*, 2023. **5**
- [12] Hyung Won Chung, Le Hou, Shayne Longpre, Barret Zoph, Yi Tay, William Fedus, Yunxuan Li, Xuezhi Wang, Mostafa Dehghani, Siddhartha Brahma, et al. Scaling instruction-finetuned language models. *arXiv preprint arXiv:2210.11416*, 2022. **2**
- [13] Wenliang Dai, Junnan Li, Dongxu Li, Anthony Meng Huat Tiong, Junqi Zhao, Weisheng Wang, Boyang Li, Pascale Fung, and Steven Hoi. Instructblip: Towards general-purpose vision-language models with instruction tuning. *arXiv preprint arXiv:2305.06500*, 2023. **1, 2, 3, 5, 6, 8**
- [14] Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, et al. An image is worth 16x16 words: Transformers for image recognition at scale. *arXiv preprint arXiv:2010.11929*, 2020. **4**
- [15] Yuxin Fang, Wen Wang, Binhui Xie, Quan Sun, Ledell Wu, Xinggang Wang, Tiejun Huang, Xinlong Wang, and Yue Cao. Eva: Exploring the limits of masked visual representation learning at scale. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 19358–19369, 2023. **6**
- [16] Chaoyou Fu, Peixian Chen, Yunhang Shen, Yulei Qin, Mengdan Zhang, Xu Lin, Zhenyu Qiu, Wei Lin, Jinrui Yang, Xiawu Zheng, et al. Mme: A comprehensive evaluation benchmark for multimodal large language models. *arXiv preprint arXiv:2306.13394*, 2023. **1, 4, 5**
- [17] Tao Gong, Chengqi Lyu, Shilong Zhang, Yudong Wang, Miao Zheng, Qian Zhao, Kuikun Liu, Wenwei Zhang, Ping Luo, and Kai Chen. Multimodal-gpt: A vision and language model for dialogue with humans. *arXiv preprint arXiv:2305.04790*, 2023. **1**
- [18] Yash Goyal, Tejas Khot, Douglas Summers-Stay, Dhruv Batra, and Devi Parikh. Making the v in vqa matter: Elevating the role of image understanding in visual question answering. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 6904–6913, 2017. **3, 5, 6**
- [19] Danna Gurari, Qing Li, Abigale J Stangl, Anhong Guo, Chi Lin, Kristen Grauman, Jiebo Luo, and Jeffrey P Bigham. Vizwiz grand challenge: Answering visual questions from blind people. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 3608–3617, 2018. **5, 6**
- [20] Drew A Hudson and Christopher D Manning. Gqa: A new dataset for real-world visual reasoning and compositional question answering. In *CVPR*, 2019. **4, 5**
- [21] IDEFICS. Introducing idefics: An open reproduction of state-of-the-art visual language model. <https://huggingface.co/blog/idefics>, 2023. **5, 6**
- [22] Gabriel Ilharco, Mitchell Wortsman, Ross Wightman, Cade Gordon, Nicholas Carlini, Rohan Taori, Achal Dave, Vaishaal Shankar, Hongseok Namkoong, John Miller, Hannaneh Hajishirzi, Ali Farhadi, and Ludwig Schmidt. Openclip. 2021. If you use this software, please cite it as below. **6**
- [23] Sahar Kazemzadeh, Vicente Ordonez, Mark Matten, and Tamara Berg. Referitgame: Referring to objects in photographs of natural scenes. In *Proceedings of the 2014 conference on empirical methods in natural language processing (EMNLP)*, pages 787–798, 2014. **4**
- [24] Ranjay Krishna, Yuke Zhu, Oliver Groth, Justin Johnson, Kenji Hata, Joshua Kravitz, Stephanie Chen, Yannis Kalantidis, Li-Jia Li, David A Shamma, et al. Visual genome: Connecting language and vision using crowdsourced dense image annotations. *International journal of computer vision*, 123:32–73, 2017. **4**
- [25] Xin Lai, Zhuotao Tian, Yukang Chen, Yanwei Li, Yuhui Yuan, Shu Liu, and Jiaya Jia. Lisa: Reasoning segmentation via large language model. *arXiv preprint arXiv:2308.00692*, 2023. **1**

- [26] Bohao Li, Rui Wang, Guangzhi Wang, Yuying Ge, Yixiao Ge, and Ying Shan. Seed-bench: Benchmarking multimodal llms with generative comprehension. *arXiv preprint arXiv:2307.16125*, 2023. 1, 5
- [27] Bo Li, Peiyuan Zhang, Jingkan Yang, Yuanhan Zhang, Fanyi Pu, and Ziwei Liu. Otterhd: A high-resolution multi-modality model, 2023. 2
- [28] Bo Li, Yuanhan Zhang, Liangyu Chen, Jinghao Wang, Fanyi Pu, Jingkan Yang, Chunyuan Li, and Ziwei Liu. Mimic-it: Multi-modal in-context instruction tuning. *arXiv preprint arXiv:2306.05425*, 2023. 1
- [29] Chunyuan Li, Zhe Gan, Zhengyuan Yang, Jianwei Yang, Linjie Li, Lijuan Wang, and Jianfeng Gao. Multimodal foundation models: From specialists to general-purpose assistants. *arXiv preprint arXiv:2309.10020*, 2023. 1
- [30] Chunyuan Li, Cliff Wong, Sheng Zhang, Naoto Usuyama, Haotian Liu, Jianwei Yang, Tristan Naumann, Hoifung Poon, and Jianfeng Gao. Llava-med: Training a large language-and-vision assistant for biomedicine in one day. *arXiv preprint arXiv:2306.00890*, 2023. 1
- [31] Junnan Li, Dongxu Li, Silvio Savarese, and Steven Hoi. Blip-2: Bootstrapping language-image pre-training with frozen image encoders and large language models. *arXiv preprint arXiv:2301.12597*, 2023. 1, 2, 3, 4, 5, 6
- [32] Xiang Lisa Li and Percy Liang. Prefix-tuning: Optimizing continuous prompts for generation. *arXiv preprint arXiv:2101.00190*, 2021. 3
- [33] Yifan Li, Yifan Du, Kun Zhou, Jinpeng Wang, Wayne Xin Zhao, and Ji-Rong Wen. Evaluating object hallucination in large vision-language models. *arXiv preprint arXiv:2305.10355*, 2023. 1, 5
- [34] Tsung-Yi Lin, Michael Maire, Serge Belongie, James Hays, Pietro Perona, Deva Ramanan, Piotr Dollár, and C Lawrence Zitnick. Microsoft COCO: Common objects in context. In *ECCV*, 2014. 2, 5
- [35] Haotian Liu, Chunyuan Li, Qingyang Wu, and Yong Jae Lee. Visual instruction tuning. In *NeurIPS*, 2023. 1, 2, 3, 4, 5, 6, 8
- [36] Yuan Liu, Haodong Duan, Yuanhan Zhang, Bo Li, Songyang Zhang, Wangbo Zhao, Yike Yuan, Jiaqi Wang, Conghui He, Ziwei Liu, et al. Mmbench: Is your multi-modal model an all-around player? *arXiv preprint arXiv:2307.06281*, 2023. 1, 5, 7
- [37] Pan Lu, Swaroop Mishra, Tanglin Xia, Liang Qiu, Kai-Wei Chang, Song-Chun Zhu, Oyvind Tafjord, Peter Clark, and Ashwin Kalyan. Learn to explain: Multimodal reasoning via thought chains for science question answering. *Advances in Neural Information Processing Systems*, 2022. 5
- [38] Yadong Lu, Chunyuan Li, Haotian Liu, Jianwei Yang, Jianfeng Gao, and Yelong Shen. An empirical study of scaling instruct-tuned large multimodal models. *arXiv preprint arXiv:2309.09958*, 2023. 1, 4
- [39] Junhua Mao, Jonathan Huang, Alexander Toshev, Oana Camburu, Alan L Yuille, and Kevin Murphy. Generation and comprehension of unambiguous object descriptions. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 11–20, 2016. 4
- [40] Kenneth Marino, Mohammad Rastegari, Ali Farhadi, and Roozbeh Mottaghi. Ok-vqa: A visual question answering benchmark requiring external knowledge. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2019. 4
- [41] Anand Mishra, Shashank Shekhar, Ajeet Kumar Singh, and Anirban Chakraborty. Ocr-vqa: Visual question answering by reading text in images. In *2019 international conference on document analysis and recognition (ICDAR)*, pages 947–952. IEEE, 2019. 4
- [42] OpenAI. Gpt-4v(ision) system card. https://cdn.openai.com/papers/GPTV_System_Card.pdf, 2023. 1
- [43] Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, et al. Learning transferable visual models from natural language supervision. *arXiv preprint arXiv:2103.00020*, 2021. 6
- [44] Dustin Schwenk, Apoorv Khandelwal, Christopher Clark, Kenneth Marino, and Roozbeh Mottaghi. A-okvqa: A benchmark for visual question answering using world knowledge. In *European Conference on Computer Vision*, pages 146–162. Springer, 2022. 4
- [45] ShareGPT. <https://sharegpt.com/>, 2023. 4, 7, 8
- [46] Oleksii Sidorov, Ronghang Hu, Marcus Rohrbach, and Amanpreet Singh. Textcaps: a dataset for image captioning with reading comprehension. In *Computer Vision—ECCV 2020: 16th European Conference, Glasgow, UK, August 23–28, 2020, Proceedings, Part II 16*, pages 742–758. Springer, 2020. 4
- [47] Amanpreet Singh, Vivek Natarajan, Meet Shah, Yu Jiang, Xinlei Chen, Dhruv Batra, Devi Parikh, and Marcus Rohrbach. Towards vqa models that can read. In *Proceedings of the IEEE/CVF conference on computer vision and pattern recognition*, pages 8317–8326, 2019. 5
- [48] Hugo Touvron, Louis Martin, Kevin Stone, Peter Albert, Amjad Almahairi, Yasmine Babaei, Nikolay Bashlykov, Soumya Batra, Prajjwal Bhargava, Shruti Bhosale, et al. Llama 2: Open foundation and fine-tuned chat models. *arXiv preprint arXiv:2307.09288*, 2023. 7
- [49] Wenhai Wang, Zhe Chen, Xiaokang Chen, Jiannan Wu, Xizhou Zhu, Gang Zeng, Ping Luo, Tong Lu, Jie Zhou, Yu Qiao, et al. Visionllm: Large language model is also an open-ended decoder for vision-centric tasks. *arXiv preprint arXiv:2305.11175*, 2023. 1
- [50] Jason Wei, Maarten Bosma, Vincent Y Zhao, Kelvin Guu, Adams Wei Yu, Brian Lester, Nan Du, Andrew M Dai, and Quoc V Le. Finetuned language models are zero-shot learners. *arXiv preprint arXiv:2109.01652*, 2021. 2
- [51] Zhengyuan Yang, Linjie Li, Kevin Lin, Jianfeng Wang, Chung-Ching Lin, Zicheng Liu, and Lijuan Wang. The dawn of lmms: Preliminary explorations with gpt-4v (ision). *arXiv preprint arXiv:2309.17421*, 2023. 6
- [52] Jiabo Ye, Anwen Hu, Haiyang Xu, Qinghao Ye, Ming Yan, Guohai Xu, Chenliang Li, Junfeng Tian, Qi Qian, Ji Zhang, Qin Jin, Liang He, Xin Alex Lin, and Fei Huang. Ureader: Universal ocr-free visually-situated language understanding with multimodal large language model, 2023. 2
- [53] Qinghao Ye, Haiyang Xu, Guohai Xu, Jiabo Ye, Ming Yan, Yiyang Zhou, Junyang Wang, Anwen Hu, Pengcheng Shi,

- Yaya Shi, et al. mplug-owl: Modularization empowers large language models with multimodality. *arXiv preprint arXiv:2304.14178*, 2023. [1](#)
- [54] Weihao Yu, Zhengyuan Yang, Linjie Li, Jianfeng Wang, Kevin Lin, Zicheng Liu, Xinchao Wang, and Lijuan Wang. Mm-vet: Evaluating large multimodal models for integrated capabilities. *arXiv preprint arXiv:2308.02490*, 2023. [1](#), [3](#), [4](#), [5](#), [8](#)
- [55] Shilong Zhang, Peize Sun, Shoufa Chen, Min Xiao, Wenqi Shao, Wenwei Zhang, Kai Chen, and Ping Luo. Gpt4roi: Instruction tuning large language model on region-of-interest. *arXiv preprint arXiv:2307.03601*, 2023. [1](#)
- [56] Yanzhe Zhang, Ruiyi Zhang, Jiuxiang Gu, Yufan Zhou, Nedim Lipka, Diyi Yang, and Tong Sun. Llavav: Enhanced visual instruction tuning for text-rich image understanding. *arXiv preprint arXiv:2306.17107*, 2023. [2](#)
- [57] Bo Zhao, Boya Wu, and Tiejun Huang. Svit: Scaling up visual instruction tuning. *arXiv preprint arXiv:2307.04087*, 2023. [1](#), [2](#)
- [58] Yunqing Zhao, Tianyu Pang, Chao Du, Xiao Yang, Chongxuan Li, Ngai-Man Cheung, and Min Lin. On evaluating adversarial robustness of large vision-language models. *arXiv preprint arXiv:2305.16934*, 2023. [1](#)
- [59] Chunting Zhou, Pengfei Liu, Puxin Xu, Srini Iyer, Jiao Sun, Yuning Mao, Xuezhe Ma, Avia Efrat, Ping Yu, Lili Yu, et al. Lima: Less is more for alignment. *arXiv preprint arXiv:2305.11206*, 2023. [2](#), [8](#)
- [60] Deyao Zhu, Jun Chen, Xiaoqian Shen, Xiang Li, and Mohamed Elhoseiny. Minigpt-4: Enhancing vision-language understanding with advanced large language models. *arXiv preprint arXiv:2304.10592*, 2023. [1](#), [2](#)