

# Learning Video Representations from Large Language Models

Yue Zhao<sup>1,2\*</sup> Ishan Misra<sup>1</sup> Philipp Krähenbühl<sup>2</sup> Rohit Girdhar<sup>1</sup>

<sup>1</sup>FAIR, Meta AI <sup>2</sup>University of Texas, Austin

[facebookresearch.github.io/LaViLa](https://facebookresearch.github.io/LaViLa)

## Abstract

We introduce **LAVILA**, a new approach to learning video-language representations by leveraging Large Language Models (LLMs). We repurpose pre-trained LLMs to be conditioned on visual input, and finetune them to create automatic video narrators. Our auto-generated narrations offer a number of advantages, including dense coverage of long videos, better temporal synchronization of the visual information and text, and much higher diversity of text. The video-language embedding learned contrastively with these narrations outperforms the previous state-of-the-art on multiple first-person and third-person video tasks, both in zero-shot and finetuned setups. Most notably, LAVILA obtains an absolute gain of **10.1%** on EGTEA classification and **5.9%** Epic-Kitchens-100 multi-instance retrieval benchmarks. Furthermore, LAVILA trained with only half the narrations from the Ego4D dataset outperforms models trained on the full set, and shows positive scaling behavior on increasing pre-training data and model size.

## 1. Introduction

Learning visual representation using web-scale image-text data is a powerful tool for computer vision. Vision-language approaches [31, 49, 80] have pushed the state-of-the-art across a variety of tasks, including zero-shot classification [49], novel object detection [87], and even image generation [52]. Similar approaches for videos [4, 39, 46], however, have been limited by the small size of paired video-text corpora compared to the billion-scale image-text datasets [31, 49, 84]—even though access to raw video data has exploded in the past decade. In this work, we show it is possible to *automatically* generate text pairing for such videos by leveraging Large Language Models (LLMs), thus taking full advantage of the massive video data. Learning video-language models with these automatically generated annotations leads to stronger representations, and as Figure 1 shows, sets a new state-of-the-art on six popular first

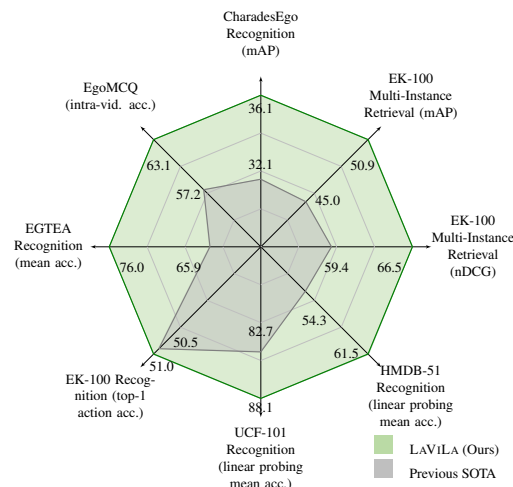


Figure 1. **LAVILA sets a new state-of-the-art** across a number of first and third-person video understanding tasks (*cf.* Table 1 for details), by learning a video-language representation using supervision from large language models as narrators.

and third-person video benchmarks.

Our method, called **LAVILA: Language-model augmented Video-Language** pre-training, leverages pre-trained LLMs, *e.g.* GPT-2 [50], which encode within their weights a treasure trove of factual knowledge and conversational ability. As shown in Figure 2, we repurpose these LLMs to be “visually-conditioned narrators”, and finetune on all accessible paired video-text clips. Once trained, we use the model to densely annotate thousands of hours of videos by generating rich textual descriptions. This pseudo-supervision can thus pervade the entire video, in between and beyond the annotated snippets. Paired with another LLM trained to rephrase existing narrations, LAVILA is able to create a much larger and more diverse set of text targets for video-text contrastive learning. In addition to setting a new state-of-the-art as noted earlier, the stronger representation learned by LAVILA even outperforms prior work using only half the groundtruth annotations (Figure 5).

LAVILA’s strong performance can be attributed to a number of factors. First, LAVILA can provide temporally dense supervision for long-form videos, where the associ-

\*Work done during an internship at Meta.

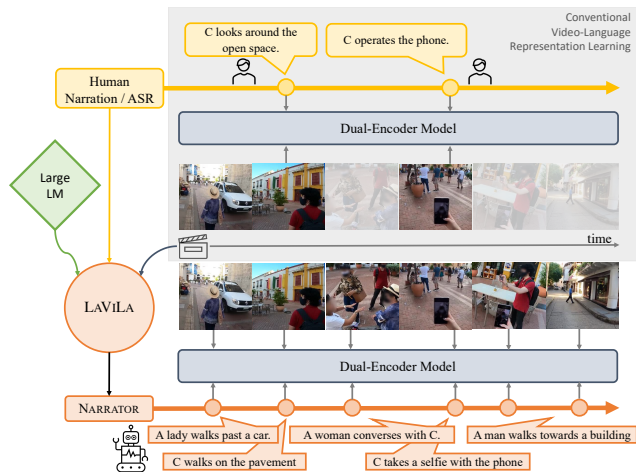


Figure 2. **LAViLA** leverages Large Language Models (LLMs) to densely narrate long videos, and uses those narrations to train strong dual-encoder models. While prior work uses sparsely labeled text by humans, or weakly aligned text transcribed from speech, LAViLA is able to leverage dense, diverse, and well-aligned text generated by a LLM.

ated captions are either too sparse, or the video-level “Alt-Text” (in the case of web videos) does not describe all the nuanced activities happening in it. Second, the generated text is well-aligned with the visual input. Although prior work has leveraged automatic speech transcription on HowTo videos [45] to automatically extract clips paired with text from the speech, such datasets have relatively poor alignment between the visual and textual content ( $\leq 50\%$ , *cf.* [25, 45]), limiting the quality of the learned representations. Third, LAViLA can significantly expand annotations when only a little is available. For instance, videos of mundane day-to-day activities, especially from an egocentric viewpoint, could be very useful for assistive and augmented reality applications. Such videos, however, are rare on the internet, and hence do not readily exist with associated web text. Recent work [24] instead opted to manually capture and narrate such video data. These narrations however required significant manual effort: 250K hours of annotator time spent in narrating 3.6K hours of video. In contrast, LAViLA is able to automatically narrate each video multiple times and far more densely, and hence learns much stronger representations.

We extensively evaluate LAViLA across multiple video-text pre-training datasets and downstream tasks to validate its effectiveness. Specifically, after being pre-trained on Ego4D, the largest egocentric video datasets with narrations, LAViLA can re-narrate the whole dataset  $10\times$  over. The resulting model learned on these expanded narrations sets a new state-of-the-art on a wide range of downstream tasks across challenging datasets, including multi-instance

video retrieval on Epic-Kitchens-100 (**5.9%** absolute gain on mAP), multiple-choice question answering on Ego4D (**5.9%** absolute gain on intra-video accuracy), and action recognition on EGTEA (**10.1%** absolute gain on mean accuracy). It obtains gains both when evaluated for zero-shot transfer to the new dataset, as well as after fine-tuning on that dataset. Similar gains are shown in third-person video data. When training LAViLA after densely re-narrating HowTo100M, we outperform prior work on downstream action classification on UCF-101 and HMDB-51. In a case study of semi-supervised learning, we show that our model, which only ever sees 50% of the human-labeled data, is capable of outperforming the baseline model trained with all the narrations. Moreover, the gains progressively increase as we go to larger data regimes and larger backbones, suggesting the scalability of our method.

## 2. Related Work

**Vision-language representation learning** maps visual and textual embeddings into a common space using metric-learning techniques [21, 73]. Recently, different pretext tasks are proposed to learn a finer-grained association between visual and textual modality, *e.g.* masked language modeling (MLM) [10, 41, 62] and captioning [16, 80]. Another line of research focuses on scaling up both models and pre-training data. For instance, CLIP [49] is pre-trained on 400M image-text pairs with a contrastive loss (InfoNCE [48, 59]) while CoCa [80] unifies contrastive and generative approaches with a single foundation model. Similar trends are also witnessed in the video-text domain [36, 64, 88]. However, collecting high-quality video-text data is more difficult than image-text. Therefore, efforts are made to learn from uncuration videos with machine-generated audio transcripts via contrastive learning [44, 77, 82] or unsupervised alignment [25] while other works focus on either adapting well-performing image-text models to videos [32, 40, 47, 78], or curriculum learning from a single frame to multiple frames [4]. In contrast, our approach leverages language models to generate temporally dense textual supervision on long-form videos.

**Generative Visual Language Models (VLM)** were first used for image/video captioning using recurrent networks [17, 68] and Transformer-based architectures [42, 56]. More recently, generative VLMs have unified multiple vision tasks [11, 89] by training multi-modal Transformers on visual-text pairs [30, 81]. Meanwhile, generative VLMs also excel at multimodal tasks via zero-shot or few-shot prompting [1, 65, 83] by leveraging multi-billion-parameter LLMs pre-trained on massive text corpus [7, 28, 50]. In our work, we demonstrate that generative VLMs can narrate long videos and the resulting video-text data benefits video-language representation learning.

**Large-scale multimodal video datasets** are crucial for

video understanding tasks but are hard to collect. Conventional video-text datasets [8, 53, 86] either have limited scenarios, *e.g.* cooking, or are not large enough to learn generic video representation. Miech *et al.* [45] scrape over 100 million video-text pairs via automatic audio transcription from long-form How-To videos. However, ASR introduces textual noise and visual-text unalignment [25]. WebVid [4] contains 10 million short videos with textual descriptions. But it is still several orders of magnitude smaller than the image counterparts [49, 54] and is harder to scale up since it is sourced from stock footage sites. The recently released Ego4D [24] dataset offers 3,600 hours of egocentric videos in which written sentence narrations are manually annotated every few seconds but requires significant manual effort. In contrast, our method shows a promising alternative by automatically narrating videos using supervision from LLM.

**Data augmentation techniques in NLP**, including word-level replacement based on synonyms [72, 85] or nearest-neighbor retrieval [19, 70], improve text classification accuracy. We refer readers to [20] for a comprehensive survey. In this paper, we show that sentence-level paraphrasing based on text-to-text models [51] is helpful for video-language pre-training.

### 3. Preliminaries

A video  $V$  is a stream of moving images  $I$ . The number of frames  $|V|$  can be arbitrarily long while video models typically operate on shorter clips, which are often in the range of a few seconds. Therefore, we skim through a long-form video and represent it by a set of  $N$  short clips, *i.e.*  $\mathcal{X}$ . Each clip  $x_i$  is defined by a specific start and end frame  $x_i = \{I_{t_i}, \dots, I_{e_i}\}$ , where  $0 < t_i < e_i \leq |V|$ , and is typically associated with some annotation  $y_i$ . This annotation could be a class label or free-form textual description of the clip. We denote a video by the set of annotated clips with their corresponding annotations, *i.e.*  $(\mathcal{X}, \mathcal{Y}) = \{(x_1, y_1), \dots, (x_N, y_N)\}$ . Note that the annotated clips often cannot densely cover the entire video due to the annotation cost and visual redundancy, *i.e.*  $\bigcup_i [t_i, e_i] \subsetneq [0, |V|]$ .

A typical video model  $\mathcal{F}(\mathcal{X}, \mathcal{Y})$  learns from these clip-level annotations using a standard training objective such as a cross-entropy loss when the annotations are class labels with a fixed vocabulary. However, more recently, dual-encoder-based contrastive approaches like CLIP [49, 77] have gained popularity. They work with free-form textual annotations which are tokenized [55] into sequences of discrete symbols, *i.e.*  $y = (s_1, s_2, \dots, s_L) \in \{1, 0\}^{|S| \times L}$ . The model consists of a visual encoder  $f_v : \mathbb{R}^{T \times 3 \times H \times W} \mapsto \mathbb{R}^{D_v}$  plus a projection head  $h_v : \mathbb{R}^{D_v} \mapsto \mathbb{R}^d$  and a textual encoder  $f_t : \{1, 0\}^{|S| \times L} \mapsto \mathbb{R}^{D_t}$  plus a projection head  $h_t : \mathbb{R}^{D_t} \mapsto \mathbb{R}^d$  in parallel to obtain the global visual and textual embedding respectively:

$$\mathbf{v} = h_v(f_v(x)), \quad \mathbf{u} = h_t(f_t(y)).$$

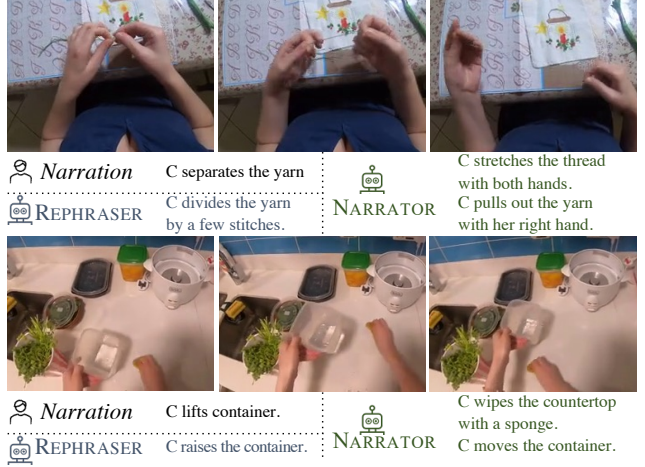


Figure 3. **Generated samples by NARRATOR and REPHRASER.** NARRATOR generates new descriptions of the action taking place, potentially focusing on other objects being interacted with. REPHRASER not only changes the word order of the human narration but also diversifies it by using related verbs or nouns.

A contrastive loss, such as InfoNCE [48], learns global embeddings that associate corresponding video and text embeddings within a batch of samples  $\mathcal{B}$ ,

$$\frac{1}{|\mathcal{B}|} \sum_{(x, y) \in \mathcal{B}} (\text{InfoNCE}(\mathbf{v}, \mathbf{u}) + \text{InfoNCE}(\mathbf{u}, \mathbf{v})). \quad (1)$$

### 4. LAVILA

In LAVILA, we leverage large language models (LLMs) as supervision to train the dual-encoder model, where the LLMs serve as vision-conditioned narrators and automatically generate textual descriptions from video clips. In particular, we exploit supervision from two LLMs: (1) **NARRATOR** (§ 4.1) is a *visually-conditioned* LLM that pseudo-labels existing and new clips with narrations, generating new annotations  $(\mathcal{X}', \mathcal{Y}')$ . (2) **REPHRASER** (§ 4.2) is a standard LLM that paraphrases narrations in existing clips, augmenting those annotations to  $(\mathcal{X}, \mathcal{Y}'')$ . As illustrated in Figure 3, NARRATOR generates new descriptions of the action taking place, potentially focusing on other objects being interacted with. REPHRASER serves to augment the text input, *e.g.*, changes word order of the human narration and additionally replaces common verbs or nouns, making annotations more diverse. Finally, we train the **dual-encoders** (§ 4.3) on all these annotations combined, *i.e.*  $(\mathcal{X}, \mathcal{Y}) \cup (\mathcal{X}', \mathcal{Y}') \cup (\mathcal{X}, \mathcal{Y}'')$ .

#### 4.1. NARRATOR

Traditional LLMs, such as GPT-2 [50], are trained to generate a sequence of text tokens  $(s_1 \dots s_L)$  from scratch

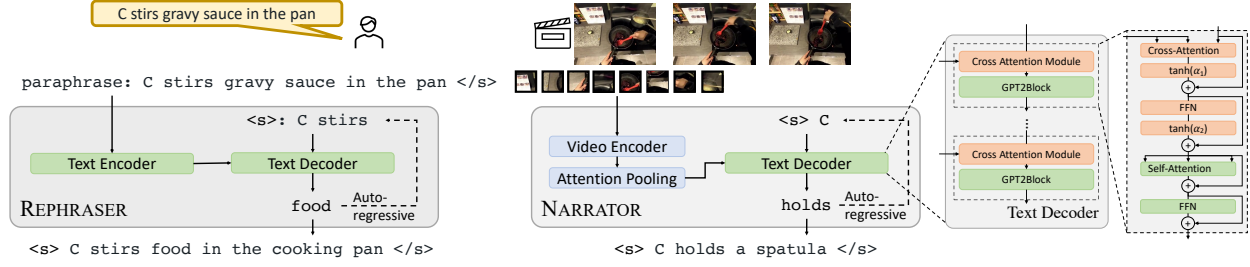


Figure 4. **Language supervision from REPHRASER and NARRATOR.** REPHRASER(*left*) takes the narration as input, passes it through a text encoder and uses a text decoder to autoregressively generate the rephrased output. NARRATOR(*right*) takes video frames as input and obtains the visual embeddings through a video encoder followed by attentional pooling. Equipped with a few additional cross-attention modules, the text decoder autoregressively generates new narrations for those new frames.

by modeling the probability of the next token given all tokens seen so far:  $p(s_t | s_{<t})$ . NARRATOR repurposes existing LLMs to be conditioned on the visual input and is trained on the original annotations  $(\mathcal{X}, \mathcal{Y})$ . The resulting model produces dense new annotations  $(\mathcal{X}', \mathcal{Y}')$  on the full video. Following the formulation of factorized probabilities in language models [5], we model the visually conditioned text likelihood as follows:

$$p_{\text{NARRATOR}}(y' | x') = \prod_{\ell=1}^L p(s'_\ell | s'_{<\ell}, x'). \quad (2)$$

**Architecture.** We design NARRATOR to closely follow the architecture of standard LLMs, with only a few additional cross-attention modules added to provide visual conditioning, as illustrated in Figure 4 (*right*). This enables NARRATOR to be initialized from pre-trained weights, which is crucial for our task as the data we use to train NARRATOR (narrations associated with video clips) are far smaller in scale compared to the large text corpus typically used to train LLMs. Moreover, video narrations are less diverse and noisier because they are either collected by only a few annotators or automatically transcribed from speech. Similar “frozen-LM” approaches have shown effectiveness in multimodal few-shot adaptation in recent work [1, 65].

Specifically, we take a frozen pre-trained LLM and add a cross-attention module before each Transformer decoder layer so that the text input can attend to visual information. The cross-attended output then sums with the input text feature via residual connection [26] and goes to the Transformer decoder layer. Each cross-attention module comprises a cross-attention layer, which takes textual tokens as queries and visual embedding as keys and values, followed by a feed-forward network (FFN). Layer Normalization [3] is applied at the beginning of both cross-attention and FFN. We add  $\tanh$ -gating [27], with an initial value of zero, such that the output of the new model is the same as that from the original language model at the beginning.

While features from any video model are applicable for conditioning, for convenience we adopt the video encoder from  $\mathcal{F}$  in § 3, trained contrastively on the ground-truth data

$(\mathcal{X}, \mathcal{Y})$ . We use features before global pooling to allow the LLM to leverage fine-grained spatial-temporal information. **Training.** We train NARRATOR on all of, or a subset of, the ground-truth annotations  $(\mathcal{X}, \mathcal{Y})$ . For each pair  $(x, y)$ , the captioning loss is the sum of the negative log-likelihood of the correct word at each step,

$$\mathcal{L}_{\text{NARRATOR}}(x, y) = - \sum_{\ell=1}^L \log p(s_\ell | s_{<\ell}, x). \quad (3)$$

**Inference.** At inference time, we query NARRATOR by feeding visual input  $x$  plus a special start-of-sentence token  $\langle s \rangle$ . We sample from the distribution recursively, *i.e.*  $\tilde{s}_\ell \sim p(s | \langle s \rangle, \dots, \tilde{s}_{\ell-1}, x)$  until an end-of-sentence token  $\langle /s \rangle$  is reached. Following [29], at each step we sample from a subset of tokens that contain the vast majority of the probability mass, which is known as nucleus sampling.

The effect of nucleus sampling is two-fold. On the one hand, it generates more diverse, open-ended, and human-like text than maximum-likelihood-based methods such as beam search and its variants [67]. On the other hand, the generated text may contain irrelevant or noisy information due to sampling without post-processing based on sentence-level likelihood. To address this, we repeat the sampling process for  $K$  times on the same visual input. We later demonstrate that the contrastive pre-training objective is robust to the noise caused by sampling, and the final performance benefits from a more diverse set of narrations.

To sample video clips for captioning, we start by simply re-captioning the existing clips labeled in the dataset  $\mathcal{X}$ , resulting in expanded annotations. Furthermore, long-form videos are typically sparsely narrated, meaning that the temporal union of all labeled clips cannot cover the entire video. Hence, we use NARRATOR to annotate the remainder of the video to obtain additional annotations by pseudo-captioning. With a simple assumption that video is a stationary process, we uniformly sample clips from the unlabeled intervals. The clip duration is equal to the average of all ground-truth clips, *i.e.*  $\Delta = \frac{1}{N} \sum_{i=1}^N (e_i - t_i)$  while the sampling stride is computed likewise. Finally, by

combining both re-captioned and pseudo-captioned narrations, we refer to the final set of annotations generated by NARRATOR as  $(\mathcal{X}', \mathcal{Y}')$ .

**Post-processing.** Exhaustive pseudo-captioning may contain some uninformative visual clips and generate text that is not useful. Thus, we add a filtering process to eliminate low-quality clips and their associated descriptions. We use the baseline dual-encoder model  $\mathcal{F}$ , which is trained on the ground-truth paired clips, to compute the visual and textual embedding of pseudo-labeled pairs and filter based on the similarity score, *i.e.*  $\text{Filter}(f_v(x'_j)^\top \cdot f_t(y'_j))$ , where  $\text{Filter}(\cdot)$  can be either top- $k$  of all generated text or a threshold filtering. In the experiments, we use a threshold of 0.5.

## 4.2. REPHRASER

The data generated by NARRATOR is several times larger than the ground-truth pairs. To ensure that we do not overfit the pseudo-labeled data, we increase the number of ground-truth narrations by paraphrasing. In particular, we use a text-to-text LLM which models conditional text likelihood:

$$p_{\text{REPHRASER}}(y''|y) = \prod_{\ell=1}^L p(s''_\ell | s''_{<\ell}, y).$$

The text-to-text model is implemented by an encoder-decoder architecture, *e.g.* T5 [51], to auto-regressively generate a new sentence given the original one. We observe that REPHRASER is able to do basic manipulations such as replacing synonyms or changing word order, which serves as an efficient way of automatic data augmentation. The resulting annotations are referred to as  $(\mathcal{X}, \mathcal{Y}'')$ .

## 4.3. Training the Dual-Encoders

We train the dual-encoders as described in Algorithm 1 in Appendix E. In each iteration, we first sample a batch  $\mathcal{B}$  of video clips. It comprises a subset of clips  $\mathcal{B}_l$  with labeled timestamps as well as narrations, and a subset  $\mathcal{B}_u$  whose clips are randomly sampled from videos without narrations. For clip  $x_i \in \mathcal{B}_u$ , we obtain the pseudo-caption  $y'_i$  by querying the NARRATOR  $y'_i \sim p_{\text{NARRATOR}}(y'|x)$ , resulting in a set of clips with LLM-generated narrations  $\tilde{\mathcal{B}}_u$ . For clip  $(x_i, y_i) \in \mathcal{B}_l$ , the text supervision is obtained from either the REPHRASER or the NARRATOR, with a probability of 0.5. Hence, the effective number of iterations per epoch for LAViLA is the same as that for the baseline Dual-Encoder. We denote the resulting set of pairs to be  $\tilde{\mathcal{B}}_l$  similarly. Following CLIP [49], we use the symmetric cross-entropy loss over the similarity scores of samples in the batch  $\tilde{\mathcal{B}}_l \cup \tilde{\mathcal{B}}_u$ .

In practice, we run REPHRASER and NARRATOR in advance and cache the resulting video-narration pairs so that there is no computational overhead during pre-training. Therefore, training a dual-encoder in LAViLA is as fast as training a standard dual-encoder contrastive model.

| Datasets         | Task | Ego? | Metrics                 | Eval. Prot. |
|------------------|------|------|-------------------------|-------------|
| EK-100 [14]      | MIR  | ✓    | mAP, nDCG               | ZS, FT      |
| EK-100 [14]      | CLS  | ✓    | top-1 action acc.       | FT          |
| Ego4D [24]       | MCQ  | ✓    | inter-/intra-video acc. | ZS          |
| Ego4D [24]       | NLQ  | ✓    | Recall@N                | FT          |
| EGTEA [37]       | CLS  | ✓    | top-1, mean acc.        | ZS, FT      |
| CharadesEgo [58] | CLS  | ✓    | video-level mAP         | ZS, FT      |
| UCF-101 [60]     | CLS  | ✗    | mean acc.               | LP          |
| HMDB-51 [35]     | CLS  | ✗    | mean acc.               | LP          |

Table 1. **Downstream datasets** and metrics used for evaluation. We evaluate LAViLA on a wide range of tasks including Multi-Instance Retrieval (MIR), Multiple-Choice Question (MCQ), Natural Language Query (NLQ), and Action Recognition (CLS). The evaluation protocols include zero-shot (ZS), fine-tuning (FT), and linear-probing (LP). Please refer to Appendix C for more details.

## 5. Experiments

**Dual-Encoder Architecture.** The video-language model follows a dual-encoder architecture as CLIP [49]. The Visual encoder is a TimeSformer (TSF) [6], whose spatial attention modules are initialized from a ViT [18] which is contrastively pre-trained on large-scale paired image-text data as in CLIP [49]. We sample 4 frames per clip during pre-training and 16 when finetuning on downstream tasks. The text encoder is a 12-layer Transformer [50, 66]. We use BPE tokenizer [55] to pre-process the full sentence corresponding to the video clip and keep at most 77 tokens.

NARRATOR’s architecture is a visually conditioned auto-regressive Language Model. The visual encoder is by default TimeSformer-L while the text decoder is a GPT-2 XL. During inference, we use nucleus sampling [29] with  $p = 0.95$  and return  $K = 10$  candidate outputs.

**REPHRASER.** We use an open-source paraphraser [23] based on T5-large [51]. It is pre-trained on C4 [51] and then finetuned on a cleaned subset of ParaNMT [74]. During inference, we use Diverse Beam Search [67] with group number the same as beam number ( $G = B = 20$ ) and set the diversity penalty to be 0.7. We keep 3 candidates per sentence, remove punctuations, and do basic de-duplication.

**Pre-training dataset.** We train on the video-narration pairs from Ego4D [13, 24], the largest egocentric video dataset to date. We exclude videos that appear in the validation and test sets of the Ego4D benchmark and determine each clip’s interval using the same pairing strategy in [39]. This results in around 4M video-text pairs with an average clip length of 1 second. We also experiment with third-person videos by pre-training on HowTo100M [45] in § 5.2.

**Evaluation protocols.** We evaluate the learned video-text encoders using three evaluation protocols. (1) *Zero-Shot* (ZS), meaning that we apply the pre-trained video-text encoders directly on the downstream validation dataset to perform video↔text retrieval tasks, without any tuning on the downstream dataset. Zero-shot classification is performed

| Method       | Backbone | mAP  |      |      | nDCG |      |      |
|--------------|----------|------|------|------|------|------|------|
|              |          | V→T  | T→V  | Avg. | V→T  | T→V  | Avg. |
| (ZERO-SHOT)  |          |      |      |      |      |      |      |
| EgoVLP [39]  | TSF-B    | 19.4 | 13.9 | 16.6 | 24.1 | 22.0 | 23.1 |
| EgoVLP* [39] | TSF-B    | 26.0 | 20.6 | 23.3 | 28.8 | 27.0 | 27.9 |
| LAViLA       | TSF-B    | 35.1 | 26.6 | 30.9 | 33.7 | 30.4 | 32.0 |
| LAViLA       | TSF-L    | 40.0 | 32.2 | 36.1 | 36.1 | 33.2 | 34.6 |
| (FINETUNED)  |          |      |      |      |      |      |      |
| MME [75]     | TBN      | 43.0 | 34.0 | 38.5 | 50.1 | 46.9 | 48.5 |
| JPoSe [75]   | TBN      | 49.9 | 38.1 | 44.0 | 55.5 | 51.6 | 53.5 |
| EgoVLP [39]  | TSF-B    | 49.9 | 40.5 | 45.0 | 60.9 | 57.9 | 59.4 |
| LAViLA       | TSF-B    | 55.2 | 45.7 | 50.5 | 66.5 | 63.4 | 65.0 |
| LAViLA       | TSF-L    | 54.7 | 47.1 | 50.9 | 68.1 | 64.9 | 66.5 |

Table 2. **EK-100 MIR**. LAViLA outperforms prior work across all settings, metrics and directions of retrieval, with larger gains when switching to a larger model. Specifically, our best model achieves over 10% absolute gain in the zero-shot setting and 5.9 ~ 7.1% gain in the finetuned setting. EgoVLP\* refers to our improved version of [39], details of which are given in Appendix F.

similarly, where we compute the similarity score between the video clip and the textual description of all possible categories. (2) *Finetuned (FT)*, where we take the pre-trained video-text model and perform end-to-end fine-tuning on the training split of the target downstream dataset. (3) *Linear-Probe (LP)*, where we compute the video features from a frozen encoder and train a linear SVM on top of the training split of the downstream dataset.

**Downstream benchmarks.** We use multiple benchmarks across four first-person (egocentric) and two third-person datasets, as enumerated in Table 1. We summarize them here and refer the reader to Appendix C for details on datasets and metrics. (1) Two tasks on Epic-Kitchens-100: Multi-Instance Retrieval (**EK-100 MIR**) and Action Recognition (**EK-100 CLS**) [14]. EK-100 is a very popular and challenging egocentric video recognition benchmark. The MIR task requires retrieving the text given videos (V→T) and videos given text (T→V). The CLS task requires classifying each video into one of 97 verbs and 300 nouns each, resulting in a combination of 3,806 action categories. (2) Two downstream tasks of Ego4D: Multiple-Choice Questions (**EgoMCQ**) and Natural Language Query (**EgoNLQ**). EgoMCQ requires selecting the correct textual description from five choices given a query video clip while EgoNLQ asks the model to output the relevant temporal intervals of video given a text query. We select these two benchmarks because they require reasoning about both visual and textual information. (3) Action Recognition on **EGTEA** [37]. It requires classifying into 106 classes of fine-grained cooking activities. (4) Action Recognition on **CharadesEgo** [58]. It requires classification into 157 classes of daily indoor activities. Note that CharadesEgo is very different from EK-100, Ego4D and EGTEA since its videos are captured by head-mounted phone cameras in a crowd-sourcing way.

| Method        | EgoMCQ Accuracy (%) |             | EgoNLQ       |              |              |              |
|---------------|---------------------|-------------|--------------|--------------|--------------|--------------|
|               | Inter-video         | Intra-video | mIOU@0.3 R@1 | mIOU@0.5 R@5 | mIOU@0.3 R@1 | mIOU@0.5 R@5 |
| SlowFast [24] | -                   | -           | 5.45         | 10.74        | 3.12         | 6.63         |
| EgoVLP [39]   | 90.6                | 57.2        | <b>10.84</b> | 18.84        | <b>6.81</b>  | 13.45        |
| LAViLA (B)    | <b>93.8</b>         | <b>59.9</b> | <u>10.53</u> | <b>19.13</b> | 6.69         | <b>13.68</b> |
| LAViLA (L)    | <b>94.5</b>         | <b>63.1</b> | <b>12.05</b> | <b>22.38</b> | <b>7.43</b>  | <b>15.44</b> |

Table 3. **Ego4D EgoMCQ and EgoNLQ**. LAViLA outperforms prior work on both Multiple-Choice Questions and Natural Language Questions on Ego4D, with nearly 6% absolute gain on the challenging intra-video MCQ task that requires reasoning over multiple clips from the same video to answer a question.

In all tables, we bold and underline the best and second-best performing methods with comparable backbones architectures. We highlight the overall best performing method, which typically uses a larger backbone, if applicable.

## 5.1. Main Results

**EK-100.** We compare LAViLA with prior works on EK-100 MIR in Table 2. In the zero-shot setup, LAViLA remarkably surpasses an improved version of EgoVLP [39] under similar model complexity: we use TSF-Base+GPT-2 as the dual-encoder architecture while EgoVLP uses TSF-Base+Distil-BERT. With a stronger video encoder, *i.e.* TSF-Large, the performance improves further. In the finetuned setting, LAViLA significantly outperforms all previous supervised approaches including MME, JPOSE [75] and EgoVLP [39]. We also compare LAViLA on EK-100 CLS in Appendix E, and establish a new state-of-the-art.

**Ego4D.** We evaluate the pre-trained LAViLA model on EgoMCQ and EgoNLQ tasks and compare the results in Table 3. On EgoMCQ, our method achieves 93.8% inter-video accuracy and 59.9% intra-video accuracy, outperforming EgoVLP by a noticeable margin. Note that EgoVLP’s performance reported in Table 3 is obtained by using EgoNCE loss [39], a variant of InfoNCE specialized for Ego4D while ours uses a standard InfoNCE loss. EgoVLP with InfoNCE has lower performance (89.4% inter-video and 51.5% intra-video accuracy). On EgoNLQ, LAViLA achieves comparable results with EgoVLP with similar model complexity.

**EGTEA.** We evaluate the learned video representation by finetuning the video encoder for action classification in Table 4 on another popular egocentric dataset, EGTEA [37]. Our method surpasses the previous state-of-the-art which takes multiple modalities including visual, auditory and textual inputs [33] by a more than 10% absolute margin on the mean accuracy metric. Since previous methods are based on different backbones, we experiment with a TSF-Base (“Visual only”) model pre-trained on Kinetics [9] as a fair baseline for LAViLA. We observe that its accuracy is comparable to previous methods but much lower than LAViLA, implying the effectiveness of learning visual representation

| Method                | Backbone         | Pretrain       | Top-1 Acc.   | Mean Acc.    |
|-----------------------|------------------|----------------|--------------|--------------|
| Li <i>et al.</i> [37] | I3D              | K400           | -            | 53.30        |
| LSTA [63]             | ConvLSTM         | IN-1k          | 61.86        | 53.00        |
| IPL [71]              | I3D              | K400           | -            | 60.15        |
| MTCN [33]             | SlowFast (V+A+T) | K400+VGG-Sound | 73.59        | 65.87        |
| Visual only           | TSF-B            | IN-21k+K400    | 65.58        | 59.32        |
| LAViLA                | TSF-B            | WIT+Ego4D      | <b>77.45</b> | <b>70.12</b> |
| LAViLA                | TSF-L            | WIT+Ego4D      | <b>81.75</b> | <b>76.00</b> |

Table 4. **EGTEA Classification.** LAViLA obtains significant gains on this task, outperforming prior work with over 10% mean accuracy. Since the backbones used are not all comparable, we also report a comparable baseline with TSF-B (“Visual only”).

| Method                | Backbone      | mAP (ZS)    | mAP (FT)    |
|-----------------------|---------------|-------------|-------------|
| ActorObserverNet [57] | ResNet-152    | -           | 20.0        |
| SSDA [12]             | I3D           | -           | 25.8        |
| Ego-Exo [38]          | SlowFast-R101 | -           | 30.1        |
| EgoVLP [39]           | TSF-B         | 25.0        | 32.1        |
| LAViLA                | TSF-B         | <b>26.8</b> | <b>33.7</b> |
| LAViLA                | TSF-L         | <b>28.9</b> | <b>36.1</b> |

Table 5. **CharadesEgo Action Recognition.** LAViLA sets new state-of-the-art in both zero-shot (ZS) and finetuned (FT) settings. Note that CharadesEgo videos are visually different compared to Ego4D videos, on which LAViLA is pretrained.

on large-scale egocentric videos and using LLM as textual supervision during pre-training.

**CharadesEgo.** Next, we compare LAViLA’s representation on the CharadesEgo action classification task. As shown in Table 5, LAViLA’s representation excels on this task as well, which is notable as CharadesEgo videos are significantly different compared to Ego4D, being captured by crowdsourced workers using mobile cameras.

## 5.2. Application to Third-Person Video Pre-training

We apply LAViLA to third-person videos by experimenting with the HowTo100M [45] dataset. Specifically, we use the temporally aligned subset provided by [25], which contains 3.3M sentences from 247k videos. We evaluate the video representation on two third-person video datasets, *i.e.*

| Method             | Vis. Enc. | UCF-101     | HMDB-51     |
|--------------------|-----------|-------------|-------------|
| MIL-NCE [44]       | S3D       | 82.7        | 54.3        |
| TAN [25]           | S3D       | 83.2        | 56.7        |
| Baseline (w/o LLM) | TSF-B     | <u>86.5</u> | <b>59.4</b> |
| LAViLA             | TSF-B     | <b>87.4</b> | <u>57.2</u> |
| LAViLA             | TSF-L     | <b>88.1</b> | <b>61.5</b> |

Table 6. **LAViLA on third-person videos.** We measure the linear-probing action classification performance of the video model after pre-training on HowTo100M [45].

UCF-101 [60] and HMDB-51 [35] for action classification using the linear probing protocol. For more details, please refer to Appendix D. From Table 6, we see that LAViLA outperforms previous methods such as MIL-NCE [44] and TAN [25] by a large margin. Since we use a different backbone, we report a baseline without LLM and show that LAViLA indeed benefits from the language supervision.

## 5.3. Application to Semi-supervised Learning

While LAViLA is very effective at leveraging existing narrations to augment them, we now show that it is also applicable when only a limited number of narrations are available to begin with. We first divide each long video from Ego4D into 15-second chunks and assume only the annotated clips within every  $N$  chunks is available during pre-training, leading to approximately  $\frac{100}{N}\%$  of the full set, where  $N \in \{2, 5, 10\}$ . This can be considered a practical scenario when we want to annotate as many videos as possible for diversity when the annotation budget is limited. In the remainder  $(1 - \frac{100}{N}\%)$  part that is skipped, we uniformly sample the same number of the clips per chunk with the same clip length as that in the seen chunks. Both the dual-encoder model and NARRATOR are trained on the  $\frac{100}{N}\%$  available annotations.

We plot the zero-shot performance curve of pre-training with different proportions in Figure 5. We can see that LAViLA consistently outperforms the ground-truth-only baseline at all points (10, 20, 50, and 100%). The improvement

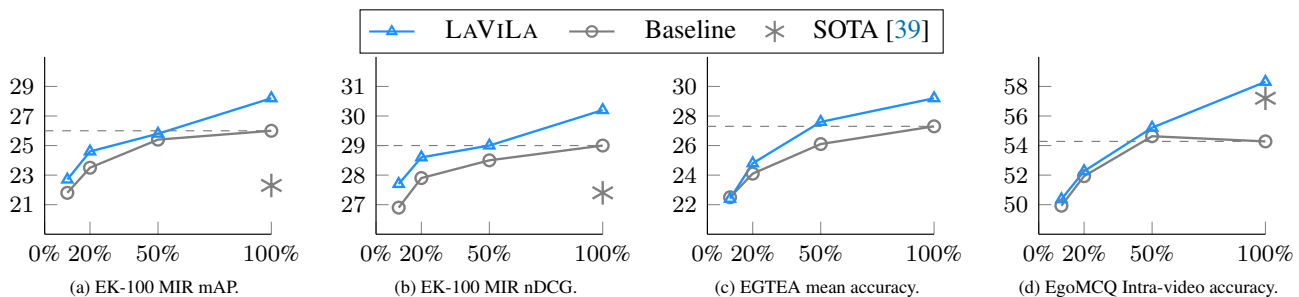
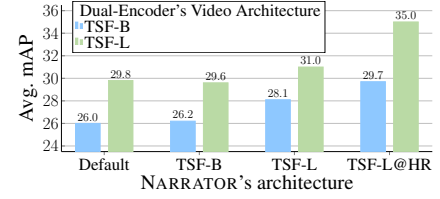


Figure 5. **LAViLA is effective in a semi-supervised setting where only a limited amount of narrations are given.** Comparing zero-shot performance of pre-training, LAViLA consistently outperforms the groundtruth-only baseline when 10, 20, 50, 100% data is used. We also achieve comparable result with state-of-the-art with only 50% of the annotated data.

| Text Dec. arch. | Text Dec. init. | Freeze LM | METEOR       | ROUGE-L      | CIDEr        | Avg. mAP    |
|-----------------|-----------------|-----------|--------------|--------------|--------------|-------------|
| (baseline)      |                 |           | -            | -            | -            | 26.0        |
| GPT-2           | random          | ✗         | 0.284        | 0.524        | 0.882        | 24.3        |
| GPT-2           | WebText         | ✓         | 0.270        | 0.505        | 0.804        | 24.0        |
| GPT-2 XL        | WebText         | ✓         | <b>0.289</b> | <b>0.530</b> | <b>0.940</b> | <b>26.2</b> |

| Sampling method | # of sentences | Avg. mAP    |
|-----------------|----------------|-------------|
| N/A (baseline)  | -              | 26.0        |
| Beam search     | 1              | 27.9        |
| Nucleus         | 1              | 29.6        |
| Nucleus         | 10             | <b>29.7</b> |



(a) **Generation Quality.** Using a sufficiently large language model as the text decoder is crucial for good text generation quality and downstream performance.

(b) **Sampling.** LAViLA benefits more from narrations produced by nucleus sampling than beam search.

(c) **Scaling effect of LAViLA.** Gains increase on scaling the video encoder in NARRATOR. Default refers to only using the original narrations.

Table 7. **Ablations of NARRATOR.** We report zero-shot average mAP on EK-100 MIR for comparing downstream performance. We study NARRATOR from the perspective of generation quality (*left*), sampling techniques (*middle*), and scaling effect (*right*).

| Rephr. Recap. | Pseudo cap. | EK-100 MIR  |             | EgoMCQ      |             | EGTEA       |             |
|---------------|-------------|-------------|-------------|-------------|-------------|-------------|-------------|
|               |             | Avg. mAP    | Avg. nDCG   | inter-video | Intra-video | Mean        | Top-1       |
|               |             | 26.0        | 28.8        | <b>93.6</b> | 54.3        | 27.3        | 30.1        |
| ✓             |             | 28.0        | 30.1        | 93.5        | 56.9        | 29.8        | 30.8        |
|               | ✓           | 27.1        | 29.9        | 93.2        | <b>59.2</b> | 26.8        | 31.2        |
| ✓             | ✓           | <u>29.7</u> | <b>31.5</b> | <b>93.6</b> | 58.3        | 29.4        | <b>36.6</b> |
| ✓             | ✓           | <b>29.9</b> | <u>31.4</u> | <b>93.6</b> | <u>59.1</u> | <b>31.1</b> | <u>36.0</u> |

Table 8. **Contributions of different Language Supervision.** We can see that (1) using REPHRASER (“Rephr.”) and NARRATOR (“Recap.”) improve downstream zero-shot performance complementarily, (2) dense pseudo-captioning further improves performance on 3 out of 6 metrics.

tends to be larger when more data is available, indicating the method’s scalability as more videos are narrated in the future. Furthermore, we observe our method can achieve a similar level of performance with the baseline often using less than 50% data. We also achieve a comparable result with the state-of-the-art using much fewer data.

## 5.4. Ablation Studies

**Contributions of Different Language Supervisions.** We ablate different language supervisions in Table 8 on EK-100 MIR (zero-shot), EgoMCQ and EGTEA. Using the text-only REPHRASER (“rephr.”) or visually conditioned NARRATOR (“recap.”) separately improves the ground-truth baseline noticeably. Combining both REPHRASER and NARRATOR gives an improvement of 3.5% average mAP on EK-100 MIR. We see that dense captioning on the entire video (“pseudo-cap.”) is also helpful. Though the gain on EK-100 MIR is not as significant, it shows nontrivial improvements on EgoMCQ intra-video accuracy and EGTEA mean accuracy. Our conjecture for this marginal gain is that informative clips are mostly covered in Ego4D because all videos are inspected by two annotators.

**Generation Quality of NARRATOR.** We study how the NARRATOR’s configurations affect the quality of generated text and the downstream performance. The generation quality is measured by standard unsupervised automatic metrics including METEOR, ROUGE, and CIDEr [43]. We use a NARRATOR with a smaller GPT-2 as the text decoder and

consider two scenarios in Table 7a: (1) LM is randomly initialized but jointly trained with the gated cross-attention modules, and (2) LM is initialized from the original GPT-2. The generation quality decreases compared to GPT-2 XL in both cases and the zero-shot retrieval result on EK-100 MIR is worse. This indicates that the language model should be sufficiently large and pre-trained on web text data.

**Sampling.** In Table 7b, we investigate different sampling methods for text generation from NARRATOR. We see that nucleus sampling works much better than beam search while repetitive sampling shows marginal improvement.

**Scaling effect.** In Table 7c, we compare the zero-shot retrieval result by progressively increasing the size of NARRATOR’s video encoder from TSF-B to TSF-L and TSF-L@HR, which increases the input resolution to be narrated from 224 to 336 while fixing the dual-encoder architecture. The retrieval performance steadily increases while NARRATOR becomes stronger. We conduct this experiment by varying the dual-encoder architecture, namely TSF-Base and TSF-Large, and show similar trends. Both phenomena suggest that LAViLA can scale to larger models.

## 6. Conclusion and Future Work

In this paper, we proposed LAViLA, a new approach to video-language representation learning by automatically narrating long videos with LLMs. We achieve strong improvements over baselines trained with the same amount of human-narrated videos and set new state-of-the-art on six popular benchmark tasks across first- and third-person video understanding benchmarks. LAViLA also shows positive scaling behavior when adding more training narrations, using larger visual backbones, and using stronger LLMs, all of which are promising areas for future work.

**Acknowledgements:** We thank Naman Goyal, Stephen Roller and Susan Zhang for help with language models, Kevin Qinghong Lin for help with EgoVLP, and the Meta AI team for helpful discussions and feedback. This material is based upon work in-part supported by the National Science Foundation under Grant No. IIS-1845485.

## References

- [1] Jean-Baptiste Alayrac, Jeff Donahue, Pauline Luc, Antoine Miech, Iain Barr, Yana Hasson, Karel Lenc, Arthur Mensch, Katie Millican, Malcolm Reynolds, Roman Ring, Eliza Rutherford, Serkan Cabi, Tengda Han, Zhitao Gong, Sina Samangooei, Marianne Monteiro, Jacob Menick, Sebastian Borgeaud, Andrew Brock, Aida Nematzadeh, Sahand Sharifzadeh, Mikolaj Binkowski, Ricardo Barreira, Oriol Vinyals, Andrew Zisserman, and Karen Simonyan. Flamingo: a visual language model for few-shot learning. In *NeurIPS*, 2022.
- [2] Anurag Arnab, Mostafa Dehghani, Georg Heigold, Chen Sun, Mario Lučić, and Cordelia Schmid. Vivit: A video vision transformer. In *ICCV*, 2021.
- [3] Jimmy Lei Ba, Jamie Ryan Kiros, and Geoffrey E Hinton. Layer normalization. *arXiv preprint arXiv:1607.06450*, 2016.
- [4] Max Bain, Arsha Nagrani, Gül Varol, and Andrew Zisserman. Frozen in time: A joint video and image encoder for end-to-end retrieval. In *ICCV*, 2021.
- [5] Yoshua Bengio, Réjean Ducharme, and Pascal Vincent. A neural probabilistic language model. In *NIPS*, 2000.
- [6] Gedas Bertasius, Heng Wang, and Lorenzo Torresani. Is space-time attention all you need for video understanding? In *ICML*, 2021.
- [7] Tom Brown, Benjamin Mann, Nick Ryder, Melanie Subbiah, Jared D Kaplan, Prafulla Dhariwal, Arvind Neelakantan, Pranav Shyam, Girish Sastry, Amanda Askell, Sandhini Agarwal, Ariel Herbert-Voss, Gretchen Krueger, Tom Henighan, Rewon Child, Aditya Ramesh, Daniel M. Ziegler, Jeffrey Wu, Clemens Winter, Christopher Hesse, Mark Chen, Eric Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess, Jack Clark, Christopher Berner, Sam McCandlish, Alec Radford, Ilya Sutskever, and Dario Amodei. Language models are few-shot learners. *NeurIPS*, 2020.
- [8] Fabian Caba Heilbron, Victor Escorcia, Bernard Ghanem, and Juan Carlos Nieves. Activitynet: A large-scale video benchmark for human activity understanding. In *CVPR*, 2015.
- [9] Joao Carreira and Andrew Zisserman. Quo vadis, action recognition? a new model and the kinetics dataset. In *CVPR*, 2017.
- [10] Yen-Chun Chen, Linjie Li, Licheng Yu, Ahmed El Kholy, Faisal Ahmed, Zhe Gan, Yu Cheng, and Jingjing Liu. Uniter: Universal image-text representation learning. In *ECCV*, 2020.
- [11] Jaemin Cho, Jie Lei, Hao Tan, and Mohit Bansal. Unifying vision-and-language tasks via text generation. In *ICML*, 2021.
- [12] Jinwoo Choi, Gaurav Sharma, Manmohan Chandraker, and Jia-Bin Huang. Unsupervised and semi-supervised domain adaptation for action recognition from drones. In *WACV*, 2020.
- [13] Ego4D Consortium. Egocentric live 4d perception (Ego4D) database: A large-scale first-person video database, supporting research in multi-modal machine perception for daily life activity. <https://sites.google.com/view/ego4d/home>. Accessed: 2022-11-22.
- [14] Dima Damen, Hazel Doughty, Giovanni Maria Farinella, Antonino Furnari, Evangelos Kazakos, Jian Ma, Davide Moltisanti, Jonathan Munro, Toby Perrett, Will Price, and Michael Wray. Rescaling egocentric vision: collection, pipeline and challenges for Epic-Kitchens-100. *IJCV*, 2022.
- [15] Jia Deng, Wei Dong, Richard Socher, Li-Jia Li, Kai Li, and Li Fei-Fei. ImageNet: A large-scale hierarchical image database. In *CVPR*, 2009.
- [16] Karan Desai and Justin Johnson. Virtex: Learning visual representations from textual annotations. In *CVPR*, 2021.
- [17] Jeffrey Donahue, Lisa Anne Hendricks, Sergio Guadarrama, Marcus Rohrbach, Subhashini Venugopalan, Kate Saenko, and Trevor Darrell. Long-term recurrent convolutional networks for visual recognition and description. In *CVPR*, 2015.
- [18] Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, Jakob Uszkoreit, and Neil Houlsby. An image is worth 16x16 words: Transformers for image recognition at scale. In *ICLR*, 2020.
- [19] Marzieh Fadaee, Arianna Bisazza, and Christof Monz. Data augmentation for low-resource neural machine translation. In *ACL*, 2017.
- [20] Steven Y Feng, Varun Gangal, Jason Wei, Sarath Chandar, Soroush Vosoughi, Teruko Mitamura, and Eduard Hovy. A survey of data augmentation approaches for nlp. In *ACL Findings*, 2021.
- [21] Andrea Frome, Greg S Corrado, Jon Shlens, Samy Bengio, Jeff Dean, Marc’Aurelio Ranzato, and Tomas Mikolov. Devise: A deep visual-semantic embedding model. *NeurIPS*, 2013.
- [22] Rohit Girdhar, Mannat Singh, Nikhila Ravi, Laurens van der Maaten, Armand Joulin, and Ishan Misra. Omnivore: A single model for many visual modalities. In *CVPR*, 2022.
- [23] Ramsri Goutham Golla. High-quality sentence paraphraser using transformers in nlp. <https://huggingface.co/ramsrigouthamg/t5-large-paraphraser-diverse-high-quality>. Accessed: 2022-06-01.
- [24] Kristen Grauman, Andrew Westbury, Eugene Byrne, Zachary Chavis, Antonino Furnari, Rohit Girdhar, Jackson Hamburger, Hao Jiang, Miao Liu, Xingyu Liu, Miguel Martin, Tushar Nagarajan, Ilija Radosavovic, Santhosh Kumar Ramakrishnan, Fiona Ryan, Jayant Sharma, Michael Wray, Mengmeng Xu, Eric Zhongcong Xu, Chen Zhao, Siddhant Bansal, Dhruv Batra, Vincent Cartillier, Sean Crane, Tien Do, Morrie Doulaty, Akshay Erapalli, Christoph Feichtenhofer, Adriano Fragomeni, Qichen Fu, Abrahm Gebreselasie, Cristina Gonzalez, James Hillis, Xuhua Huang, Yifei Huang, Wenqi Jia, Weslie Khoo, Jachym Kolar, Satwik Kotur, Anurag Kumar, Federico Landini, Chao Li, Yanghao Li, Zhenqiang Li, Karttikeya Mangalam, Raghava Modhugu, Jonathan Munro, Tullie Murrell, Takumi Nishiyasu, Will Price, Paola Ruiz Puentes, Merey Ramazanova, Leda Sari, Kiran Somasundaram, Audrey Southerland, Yusuke Sugano, Ruijie Tao, Minh Vo, Yuchen Wang, Xindi Wu, Takuma Yagi, Ziwei Zhao, Yunyi Zhu, Pablo Arbelaez, David Crandall, Dima Damen, Giovanni Maria Farinella, Christian Fuegen, Bernard Ghanem, Vamsi Krishna Ithapu, C. V. Jawahar, Hanbyul Joo, Kris Kitani, Haizhou Li, Richard Newcombe,

- Aude Oliva, Hyun Soo Park, James M. Rehg, Yoichi Sato, Jianbo Shi, Mike Zheng Shou, Antonio Torralba, Lorenzo Torresani, Mingfei Yan, and Jitendra Malik. Ego4d: Around the world in 3,000 hours of egocentric video. In *CVPR*, 2022.
- [25] Tengda Han, Weidi Xie, and Andrew Zisserman. Temporal alignment networks for long-term video. In *CVPR*, 2022.
- [26] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *CVPR*, 2016.
- [27] Sepp Hochreiter and Jürgen Schmidhuber. Long short-term memory. *Neural computation*, 1997.
- [28] Jordan Hoffmann, Sebastian Borgeaud, Arthur Mensch, Elena Buchatskaya, Trevor Cai, Eliza Rutherford, Diego de Las Casas, Lisa Anne Hendricks, Johannes Welbl, Aidan Clark, Tom Hennigan, Eric Noland, Katie Millican, George van den Driessche, Bogdan Damoc, Aurelia Guy, Simon Osindero, Karen Simonyan, Erich Elsen, Jack W. Rae, Oriol Vinyals, and Laurent Sifre. Training compute-optimal large language models. *arXiv preprint arXiv:2203.15556*, 2022.
- [29] Ari Holtzman, Jan Buys, Li Du, Maxwell Forbes, and Yejin Choi. The curious case of neural text degeneration. In *ICLR*, 2020.
- [30] Xiaowei Hu, Zhe Gan, Jianfeng Wang, Zhengyuan Yang, Zicheng Liu, Yumao Lu, and Lijuan Wang. Scaling up vision-language pre-training for image captioning. In *CVPR*, 2022.
- [31] Chao Jia, Yinfei Yang, Ye Xia, Yi-Ting Chen, Zarana Parekh, Hieu Pham, Quoc Le, Yun-Hsuan Sung, Zhen Li, and Tom Duerig. Scaling up visual and vision-language representation learning with noisy text supervision. In *ICML*, 2021.
- [32] Chen Ju, Tengda Han, Kunhao Zheng, Ya Zhang, and Weidi Xie. Prompting visual-language models for efficient video understanding. In *ECCV*, 2022.
- [33] Evangelos Kazakos, Jaesung Huh, Arsha Nagrani, Andrew Zisserman, and Dima Damen. With a little help from my temporal context: Multimodal egocentric action recognition. In *BMVC*, 2021.
- [34] Dan Kondratyuk, Liangzhe Yuan, Yandong Li, Li Zhang, Mingxing Tan, Matthew Brown, and Boqing Gong. Movinets: Mobile video networks for efficient video recognition. In *CVPR*, 2021.
- [35] Hildegard Kuehne, Hueihan Jhuang, Estíbaliz Garrote, Tomaso Poggio, and Thomas Serre. Hmdb: a large video database for human motion recognition. In *ICCV*, 2011.
- [36] Jie Lei, Linjie Li, Luowei Zhou, Zhe Gan, Tamara L Berg, Mohit Bansal, and Jingjing Liu. Less is more: Clipbert for video-and-language learning via sparse sampling. In *CVPR*, 2021.
- [37] Yin Li, Miao Liu, and James M Rehg. In the eye of beholder: Joint learning of gaze and actions in first person video. In *ECCV*, 2018.
- [38] Yanghao Li, Tushar Nagarajan, Bo Xiong, and Kristen Grauman. Ego-exo: Transferring visual representations from third-person to first-person videos. In *CVPR*, 2021.
- [39] Kevin Qinghong Lin, Alex Jinpeng Wang, Mattia Soldan, Michael Wray, Rui Yan, Eric Zhongcong Xu, Difei Gao, Rongcheng Tu, Wenzhe Zhao, Weijie Kong, Hongfa Cai Chengfei, Wang, Dima Damen, Bernard Ghanem, Wei Liu, and Mike Zheng Shou. Egocentric video-language pre-training. In *NeurIPS*, 2022.
- [40] Ziyi Lin, Shijie Geng, Renrui Zhang, Peng Gao, Gerard de Melo, Xiaogang Wang, Jifeng Dai, Yu Qiao, and Hongsheng Li. Frozen clip models are efficient video learners. In *ECCV*, 2022.
- [41] Jiasen Lu, Dhruv Batra, Devi Parikh, and Stefan Lee. Vilbert: Pretraining task-agnostic visiolinguistic representations for vision-and-language tasks. *NeurIPS*, 2019.
- [42] Huaishao Luo, Lei Ji, Botian Shi, Haoyang Huang, Nan Duan, Tianrui Li, Jason Li, Taroon Bharti, and Ming Zhou. Univl: A unified video and language pre-training model for multimodal understanding and generation. *arXiv preprint arXiv:2002.06353*, 2020.
- [43] Maluuba. nlg-eval. <https://github.com/Maluuba/nlg-eval>. Accessed: 2022-06-01.
- [44] Antoine Miech, Jean-Baptiste Alayrac, Lucas Smaira, Ivan Laptev, Josef Sivic, and Andrew Zisserman. End-to-end learning of visual representations from uncurated instructional videos. In *CVPR*, 2020.
- [45] Antoine Miech, Dimitri Zhukov, Jean-Baptiste Alayrac, Makarand Tapaswi, Ivan Laptev, and Josef Sivic. Howto100m: Learning a text-video embedding by watching hundred million narrated video clips. In *ICCV*, 2019.
- [46] Arsha Nagrani, Paul Hongsuck Seo, Bryan Seybold, Anja Hauth, Santiago Manen, Chen Sun, and Cordelia Schmid. Learning audio-video modalities from image captions. In *ECCV*, 2022.
- [47] Bolin Ni, Houwen Peng, Minghao Chen, Songyang Zhang, Gaofeng Meng, Jianlong Fu, Shiming Xiang, and Haibin Ling. Expanding language-image pretrained models for general video recognition. In *ECCV*, 2022.
- [48] Aaron van den Oord, Yazhe Li, and Oriol Vinyals. Representation learning with contrastive predictive coding. *arXiv preprint arXiv:1807.03748*, 2018.
- [49] Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, Gretchen Krueger, and Ilya Sutskever. Learning transferable visual models from natural language supervision. In *ICML*, 2021.
- [50] Alec Radford, Jeffrey Wu, Rewon Child, David Luan, Dario Amodei, and Ilya Sutskever. Language models are unsupervised multitask learners. *OpenAI blog*, 2019.
- [51] Colin Raffel, Noam Shazeer, Adam Roberts, Katherine Lee, Sharan Narang, Michael Matena, Yanqi Zhou, Wei Li, and Peter J. Liu. Exploring the limits of transfer learning with a unified text-to-text transformer. *JMLR*, 2020.
- [52] Aditya Ramesh, Prafulla Dhariwal, Alex Nichol, Casey Chu, and Mark Chen. Hierarchical text-conditional image generation with clip latents. *arXiv preprint arXiv:2204.06125*, 2022.
- [53] Anna Rohrbach, Atousa Torabi, Marcus Rohrbach, Niket Tandon, Christopher Pal, Hugo Larochelle, Aaron Courville, and Bernt Schiele. Movie description. *IJCV*, 2017.
- [54] Christoph Schuhmann, Romain Beaumont, Cade W Gordon, Ross Wightman, Theo Coombes, Aarush Katta, Clayton Mullis, Patrick Schramowski, Srivatsa R Kundurthy, Katherine Crowson, Ludwig Schmidt, Robert Kaczmarczyk, and Jenia Jitsev. Laion-5b: An open large-scale dataset for training next generation image-text models. In *NeurIPS D&B*,

- 2022.
- [55] Rico Sennrich, Barry Haddow, and Alexandra Birch. Neural machine translation of rare words with subword units. In *ACL*, 2016.
  - [56] Paul Hongsuck Seo, Arsha Nagrani, Anurag Arnab, and Cordelia Schmid. End-to-end generative pretraining for multimodal video captioning. In *CVPR*, 2022.
  - [57] Gunnar A Sigurdsson, Abhinav Gupta, Cordelia Schmid, Ali Farhadi, and Karteek Alahari. Actor and observer: Joint modeling of first and third-person videos. In *CVPR*, 2018.
  - [58] Gunnar A Sigurdsson, Abhinav Gupta, Cordelia Schmid, Ali Farhadi, and Karteek Alahari. Charades-ego: A large-scale dataset of paired third and first person videos. *arXiv preprint arXiv:1804.09626*, 2018.
  - [59] Kihyuk Sohn. Improved deep metric learning with multi-class n-pair loss objective. *NeurIPS*, 2016.
  - [60] Khurram Soomro, Amir Roshan Zamir, and Mubarak Shah. Ucf101: A dataset of 101 human actions classes from videos in the wild. *arXiv preprint arXiv:1212.0402*, 2012.
  - [61] Andreas Steiner, Alexander Kolesnikov, Xiaohua Zhai, Ross Wightman, Jakob Uszkoreit, and Lucas Beyer. How to train your vit? data, augmentation, and regularization in vision transformers. *TMLR*, 2022.
  - [62] Weijie Su, Xizhou Zhu, Yue Cao, Bin Li, Lewei Lu, Furu Wei, and Jifeng Dai. Vi-bert: Pre-training of generic visual-linguistic representations. In *ICLR*, 2020.
  - [63] Swathikiran Sudhakaran, Sergio Escalera, and Oswald Lanz. Lsta: Long short-term attention for egocentric action recognition. In *CVPR*, 2019.
  - [64] Chen Sun, Austin Myers, Carl Vondrick, Kevin Murphy, and Cordelia Schmid. Videobert: A joint model for video and language representation learning. In *ICCV*, 2019.
  - [65] Maria Tsimpoukelli, Jacob L Menick, Serkan Cabi, SM Eslami, Oriol Vinyals, and Felix Hill. Multimodal few-shot learning with frozen language models. In *NeurIPS*, 2021.
  - [66] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. In *NeurIPS*, 2017.
  - [67] Ashwin K Vijayakumar, Michael Cogswell, Ramprasad R Selvaraju, Qing Sun, Stefan Lee, David Crandall, and Dhruv Batra. Diverse beam search: Decoding diverse solutions from neural sequence models. *arXiv preprint arXiv:1610.02424*, 2016.
  - [68] Oriol Vinyals, Alexander Toshev, Samy Bengio, and Dumitru Erhan. Show and tell: A neural image caption generator. In *CVPR*, 2015.
  - [69] Wenhui Wang, Hangbo Bao, Li Dong, Johan Bjorck, Zhiliang Peng, Qiang Liu, Kriti Aggarwal, Owais Khan Mohammed, Saksham Singhal, Subhojit Som, et al. Image as a foreign language: Beit pretraining for all vision and vision-language tasks. *arXiv preprint arXiv:2208.10442*, 2022.
  - [70] William Yang Wang and Diyi Yang. That’s so annoying!!!: A lexical and frame-semantic embedding based data augmentation approach to automatic categorization of annoying behaviors using# petpeeve tweets. In *EMNLP*, 2015.
  - [71] Xiaohan Wang, Linchao Zhu, Heng Wang, and Yi Yang. Interactive prototype learning for egocentric action recognition. In *ICCV*, 2021.
  - [72] Jason Wei and Kai Zou. Eda: Easy data augmentation techniques for boosting performance on text classification tasks. In *EMNLP*, 2019.
  - [73] Jason Weston, Samy Bengio, and Nicolas Usunier. Large scale image annotation: learning to rank with joint word-image embeddings. *Machine learning*, 2010.
  - [74] John Wieting and Kevin Gimpel. Paranzmt-50m: Pushing the limits of paraphrastic sentence embeddings with millions of machine translations. In *ACL*, 2018.
  - [75] Michael Wray, Diane Larlus, Gabriela Csurka, and Dima Damen. Fine-grained action retrieval through multiple parts-of-speech embeddings. In *ICCV*, 2019.
  - [76] Chao-Yuan Wu, Yanghao Li, Karttikeya Mangalam, Haoqi Fan, Bo Xiong, Jitendra Malik, and Christoph Feichtenhofer. Memvit: Memory-augmented multiscale vision transformer for efficient long-term video recognition. In *CVPR*, 2022.
  - [77] Hu Xu, Gargi Ghosh, Po-Yao Huang, Dmytro Okhonko, Armen Aghajanyan, Florian Metze, Luke Zettlemoyer, and Christoph Feichtenhofer. Videoclip: Contrastive pre-training for zero-shot video-text understanding. In *EMNLP*, 2021.
  - [78] Hongwei Xue, Yuchong Sun, Bei Liu, Jianlong Fu, Ruihua Song, Houqiang Li, and Jiebo Luo. Clip-vip: Adapting pre-trained image-text model to video-language representation alignment. In *ICLR*, 2023.
  - [79] Shen Yan, Xuehan Xiong, Anurag Arnab, Zhichao Lu, Mi Zhang, Chen Sun, and Cordelia Schmid. Multiview transformers for video recognition. In *CVPR*, 2022.
  - [80] Jiahui Yu, Zirui Wang, Vijay Vasudevan, Legg Yeung, Mojtaba Seyedhosseini, and Yonghui Wu. CoCa: Contrastive captioners are image-text foundation models. *TMLR*, 2022.
  - [81] Lu Yuan, Dongdong Chen, Yi-Ling Chen, Noel Codella, Xiyang Dai, Jianfeng Gao, Houdong Hu, Xuedong Huang, Boxin Li, Chunyuan Li, Ce Liu, Mengchen Liu, Zicheng Liu, Yumao Lu, Yu Shi, Lijuan Wang, Jianfeng Wang, Bin Xiao, Zhen Xiao, Jianwei Yang, Michael Zeng, Luowei Zhou, and Pengchuan Zhang. Florence: A new foundation model for computer vision. *arXiv preprint arXiv:2111.11432*, 2021.
  - [82] Rowan Zellers, Ximing Lu, Jack Hessel, Youngjae Yu, Jae Sung Park, Jize Cao, Ali Farhadi, and Yejin Choi. Merlot: Multimodal neural script knowledge models. *NeurIPS*, 2021.
  - [83] Andy Zeng, Adrian Wong, Stefan Welker, Krzysztof Choromanski, Federico Tombari, Aveek Purohit, Michael Ryoo, Vikas Sindhwani, Johnny Lee, Vincent Vanhoucke, and Pete Florence. Socratic models: Composing zero-shot multimodal reasoning with language. In *ICLR*, 2023.
  - [84] Xiaohua Zhai, Alexander Kolesnikov, Neil Houlsby, and Lucas Beyer. Scaling vision transformers. In *CVPR*, 2022.
  - [85] Xiang Zhang, Junbo Zhao, and Yann LeCun. Character-level convolutional networks for text classification. In *NeurIPS*, 2015.
  - [86] Luowei Zhou, Chenliang Xu, and Jason J Corso. Towards automatic learning of procedures from web instructional videos. In *AAAI*, 2018.
  - [87] Xingyi Zhou, Rohit Girdhar, Armand Joulin, Philipp Krähenbühl, and Ishan Misra. Detecting twenty-thousand classes using image-level supervision. In *ECCV*, 2022.
  - [88] Linchao Zhu and Yi Yang. Actbert: Learning global-local video-text representations. In *CVPR*, 2020.
  - [89] Xizhou Zhu, Jinguo Zhu, Hao Li, Xiaoshi Wu, Hongsheng Li, Xiaohua Wang, and Jifeng Dai. Uni-perceiver: Pre-

training unified architecture for generic perception for zero-shot and few-shot tasks. In *CVPR*, 2022.