

Video-adverb retrieval with compositional adverb-action embeddings

Thomas Hummel¹

thomas.hummel@uni-tuebingen.de

Otniel-Bogdan Mercea¹

otniel-bogdan.mercea@uni-tuebingen.de

A. Sophia Koepke¹

a-sophia.koepke@uni-tuebingen.de

Zeynep Akata^{1,2}

zeynep.akata@uni-tuebingen.de

¹ University of Tübingen

Germany

² MPI for Intelligent Systems

Germany

Abstract

Retrieving adverbs that describe an action in a video poses a crucial step towards fine-grained video understanding. We propose a framework for video-to-adverb retrieval (and vice versa) that aligns video embeddings with their matching compositional adverb-action text embedding in a joint embedding space. The compositional adverb-action text embedding is learned using a residual gating mechanism, along with a novel training objective consisting of triplet losses and a regression target. Our method achieves state-of-the-art performance on five recent benchmarks for video-adverb retrieval. Furthermore, we introduce dataset splits to benchmark video-adverb retrieval for unseen adverb-action compositions on subsets of the MSR-VTT Adverbs and ActivityNet Adverbs datasets. Our proposed framework outperforms all prior works for the generalisation task of retrieving adverbs from videos for unseen adverb-action compositions. Code and dataset splits are available at <https://hummelth.github.io/ReGaDa/>.

1 Introduction

Fine-grained video understanding is concerned with the detailed analysis of video content beyond action recognition. This is relevant for improving and potentially accelerating video search and retrieval. While there has been significant progress in action retrieval and recognition in videos [2, 28, 39, 43], the fine-grained understanding of actions remains challenging. In particular, it can be useful to perceive how an action is performed in order to better understand the action itself [12, 13, 30]. For instance, in addition to recognising the action *cutting*, it is useful to understand details about the execution of an action, *e.g. cutting slowly*. Specifically, we consider the bidirectional video-adverb retrieval task where we retrieve adverbs that match an action in a video and vice versa.

For bidirectional video-adverb retrieval, adverbs and action words can be combined in a compositional manner. The same adverb can describe multiple actions, such as *cutting slowly* or *dancing slowly*. The compositional nature of the adverb-action pairings can also

be exploited when learning adverb-action representations. Our proposed REGADA framework for video-adverb retrieval uses a **residual gating** mechanism to compose **adverb-action** (REGADA) representations for retrieval.

At its core, our framework learns to align adverb representations and video representations in a shared embedding space using a novel training objective which consists of a direct regression loss between the adverb and video representations and triplet losses. To obtain the adverb representation, the adverb and action are jointly embedded using a residual gating mechanism, which we adapted to the video-adverb retrieval task from [46]. It models the composition as a transformation of the adverb embedding based on the action by using a gate and a residual mechanism. The gate facilitates the preservation of meaningful information from the adverb embeddings based on the adverb-action composition. Our final composition is learned as a residual combination on top of the gated adverb embeddings. This allows our composed embeddings to be in the same “feature space” as the original adverb embeddings. Similar to previous works for this task, our model assumes knowledge of the ground-truth action class to perform video-adverb retrieval.

The compositional adverb-action embeddings and our proposed training objective prove beneficial for the retrieval performance, specifically for the retrieval of unseen adverb-action compositions. REGADA obtains state-of-the-art results on the five video-adverb retrieval benchmarks HowTo100M Adverbs [13, 27], VATEX Adverbs [12, 48], ActivityNet Adverbs [1, 12], MSR-VTT Adverbs [12, 54], and Adverbs in Recipes [27, 30]. Furthermore, we propose two additional splits for benchmarking the retrieval of unseen adverb-action compositions on the ActivityNet Adverbs and MSR-VTT Adverbs datasets.

To summarise, we make the following contributions: 1) Our proposed method for video-adverb retrieval uses a text encoder based on a gated residual mechanism and a novel training objective. 2) We evaluate REGADA on the challenging unseen video-adverb retrieval task and introduce new benchmark splits, compliant with zero-shot learning principles, for the retrieval of unseen adverb-action compositions based on the ActivityNet Adverbs and MSR-VTT Adverbs datasets. 3) Our framework outperforms prior work for both the seen and the unseen adverb-action composition retrieval tasks.

2 Related work

Fine-grained action understanding in video retrieval. Early works for video understanding extended retrieval approaches for images to videos, by temporally aggregating frames in a video [11, 37, 44, 53]. With the availability of large video-text datasets [3, 5, 20, 27, 36, 43, 54, 56], different methods focused on sentence disambiguation [9, 51], self-supervision [11, 41, 57], weakly supervised learning [27, 28, 39], multiple embedding experts [14, 23, 26], or the use of large pre-trained models [21, 24, 38, 52]. Video-action retrieval specifically aims at retrieving videos based on an action, *e.g.* using a verb to describe the same [16, 50]. Moreover, [10, 13, 51, 55, 58] use nouns in addition to verbs for video-text retrieval. In a more general setting, [6] recently proposed to use a large language model to generate modified captions to improve verb understanding in video-language models. Different to these methods, we focus on adverbs in the video-adverb retrieval task.

Video-adverb retrieval. The video-adverb retrieval task was introduced by [13] along with the HowTo100M Adverbs dataset. [13] learns a shared representation between videos and adverbs, modelling adverb information as learned linear transformations on action class label word embeddings, similar to [63] for object attributes. Unlike [13], we choose to utilise se-

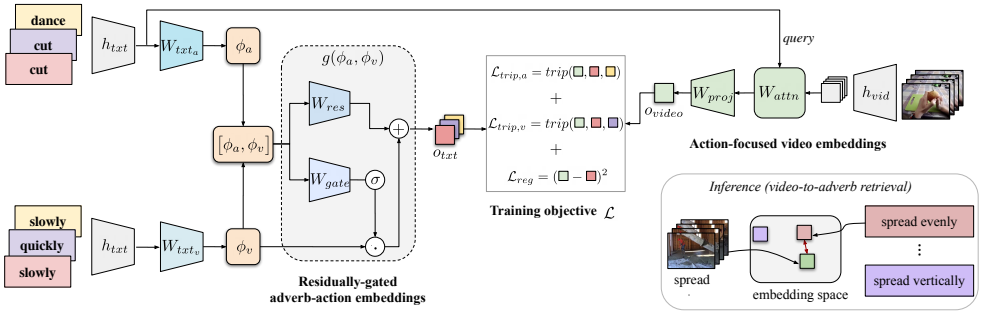


Figure 1: **Overview of our REGADA framework for video-adverb retrieval.** Our framework composes adverb-action embeddings with a gated residual between the adverbs ϕ_v and the concatenated action and adverb embeddings $[\phi_a, \phi_v]$. The training objective $\mathcal{L}$ aligns the learned text and video representations in a joint embedding space. For test time inference, outputs are obtained based on similarity in the embedding space.

mantic information from adverb embeddings in addition to action embeddings for modelling adverb-action compositions. [12] extends [13] to the low-data regime with pseudo-labelling. The recently proposed [60] tackles the task either as a classification or regression problem. Its video encoder builds on [13] with an additional projection following the attention while keeping the text representations frozen. The classification variant is trained with a cross-entropy loss for adverb classification, while the regression variant uses a regression target describing the change an adverb induced in an action embedding. Different from [60], we aim at learning the adverb-action representations and the video representations in a shared embedding space. Formulating the task as an alignment problem in a shared embedding space combined with compositional adverb-action representations significantly boosts the performance for video-adverb retrieval.

Learning with object attributes. Approaches for learning object-attribute pairs from images can be broadly categorized into classification [22, 25, 29, 52, 53] and retrieval approaches [6, 8, 18, 55, 46, 47, 49]. Our adverb-action compositions are most closely related to [46], which proposed a residual gating mechanism for learning compositional image-text embeddings. This mechanism proved particularly useful for retrieving images using both an image and a text query, the text describing a desired modification of the query image. We adapt a similar residual gating mechanism for learning compositional adverb-action embeddings by aligning the composition with action-focused video embeddings.

3 REGADA framework for video-adverb retrieval

In this section, we provide details about our proposed REGADA framework for video-adverb retrieval which is visualised in Fig. 1. We first describe the video-adverb retrieval task, and then provide details about our framework. Finally, we detail our training objective and the inference procedure for retrieval.

Task setting and dataset. The adverb-to-video retrieval task aims at retrieving matching videos from a pool of videos for a given adverb. Similarly, for the video-to-adverb retrieval task, given a video, the aim is to retrieve the adverb that best describes the action depicted

in the video from a pool of pre-set adverbs. We denote a dataset with N samples, A action classes and V adverb classes by $\mathcal{D} = \{\mathcal{X}_{[i]}, y_{[i]}\}_{i=1}^N$, consisting of video data $\mathcal{X}_{[i]}$, and ground-truth action and adverb labels $y_{[i]} = \{a_{[i]}, v_{[i]}\}$ with one-hot encodings for the action $a_{[i]} \in \mathbb{R}^A$ and adverb $v_{[i]} \in \mathbb{R}^V$. We define the sets of possible actions and adverbs as $\mathcal{A}$ and $\mathcal{V}$. The set of all possible adverb-action combinations is $\mathcal{C} = \mathcal{V} \times \mathcal{A}$.

Our REGADA framework learns to align video and adverb-action representations in a joint embedding space. It generates compositional textual representations for adverb-action pairs using a text encoder. Additionally, the visual information is processed in a video encoder to obtain visual representations that contain information about the adverb associated with a given action. In the following, we describe how we obtain class label embeddings for the actions and adverbs, and how the video and text encoders process the video features and class label embeddings.

Residually-gated adverb-action embeddings. We obtain word embeddings for the action $a \in \mathcal{A}$ and for the adverb $v \in \mathcal{V}$ from a pre-trained language encoder h_{txt} , giving $\theta_v = h_{\text{txt}}(v)$, and $\theta_a = h_{\text{txt}}(a)$ with $\theta_a, \theta_v \in \mathbb{R}^{d_\theta}$. We then apply two linear maps $W_{\text{txt}_a}, W_{\text{txt}_v} : \mathbb{R}^{d_\theta} \rightarrow \mathbb{R}^{d_{\text{dim}}}$, such that $\phi_a = W_{\text{txt}_a}(\theta_a)$ and $\phi_v = W_{\text{txt}_v}(\theta_v)$. The action and adverb embeddings are then further processed jointly in our text encoder. Additionally, the action word embedding θ_a serves as a query vector in the video encoder’s attention for generating an action-focused video embedding.

Our text encoder uses a residual gating mechanism which is based on [46]. Given ϕ_a and ϕ_v as inputs, the output of the text encoder is defined as:

$$o_{\text{txt}_j} = g(\phi_a, \phi_v) = \omega_g * \sigma(W_{\text{gate}}(\phi_a, \phi_v)) \odot \phi_v + \omega_r * W^{\text{res}}(\phi_a, \phi_v), \quad (1)$$

where $j \in \{1, \dots, V\}$, ω_g, ω_r are learnable scalar weights for balancing the gating mechanism and the residual, $\odot$ is an element-wise product, and σ the sigmoid function. W_{res} and W_{gate} are modelled using MLPs with N_r and N_g layers respectively. For those, the input consisting of adverb and action embeddings, is first passed through a concatenation operator and batch normalisation [47] is applied. The subsequent layers consist of a linear map followed by dropout [48] with probability drop_g and a Leaky ReLU [49]. The final layer is a linear projection to $\mathbb{R}^{d_{\text{dim}}}$.

We tackle video-adverb retrieval by aligning text and videos in a learned shared embedding space. Our residual gating mechanism models the composition as a transformation of the adverb embedding based on the action. The gating mechanism thereby allows to retain information from adverbs when actions do not provide useful semantic information.

Action-focused video embeddings. A pre-trained video classification network h_{vid} is used to extract a sequence of visual features $\mathbf{x}_{[i]} = \{x_1, \dots, x_t, \dots, x_T\}_i$, where $\mathbf{x}_{[i]} = h_{\text{vid}}(\mathcal{X}_{[i]})$ and $x_t \in \mathbb{R}^{d_x}$. We use T to denote the number of temporal segments in a video clip.

Given a sequence of video features $\mathbf{x}_{[i]}$ and its associated action word embedding $\theta_{a_{[i]}}$ (for easier readability, we omit the subscripts $_{[i]}$), we obtain action-focused video embeddings using a similar mechanism as the one proposed in [43]. The video embeddings are obtained using weak action-level ground-truth in the multi-head attention mechanism [45]. The action word embedding θ_a serves as the query in the attention to focus on parts of the video that are relevant to the given action, and ignore the temporal segments that may be relevant to other actions.

For the multi-head attention, we map the video features $\{x_t\}_{t \in [1, T]}$ to keys and values using linear maps $W_k : \mathbb{R}^{d_x} \rightarrow \mathbb{R}^{d_{\text{head}_x} H_x}$, $W_v : \mathbb{R}^{d_x} \rightarrow \mathbb{R}^{d_{\text{head}_x} H_x}$ with H_x heads and a dimension of d_{head_x} per head. We also map the action word embeddings θ_a to queries with $W_q : \mathbb{R}^{d_\theta} \rightarrow$

$\mathbb{R}^{d_{head_x} H_x}$. For each attention head j , we have

$$p_{attn}^j = g_{attn}^{DL} \left(\text{softmax} \left(\frac{W_q^j(\theta_a)^T W_k^j(\mathbf{x})}{\sqrt{d_{head_x}}} \right) \right) W_v^j(\mathbf{x}), \quad (2)$$

where g_{attn}^{DL} denotes dropout with probability $drop_{attn}$.

We apply a linear mapping $W_{attn} : \mathbb{R}^{d_{head_x} H_x} \rightarrow \mathbb{R}^{d_{dim}}$ to aggregate the per-head attention giving the output video embedding $o_{attn} = W_{attn}([p_{attn}^1, \dots, p_{attn}^H])$. The final output is obtained with an MLP, $W_{proj} : \mathbb{R}^{d_{dim}} \rightarrow \mathbb{R}^{d_{dim}}$,

$$o_{video} = W_{proj}(o_{attn}), \quad (3)$$

where each of the N_{proj} layers of W_{proj} consists of a linear layer $W_{proj}^l : \mathbb{R}^{d_{dim}} \rightarrow \mathbb{R}^{d_{dim}}$, layer normalisation [9] g_{proj}^{LN} , ReLU [34] g_{proj}^{ReLU} , and dropout g_{proj}^{DL} with probability $drop_{proj}$.

Training objectives. Our REGADA framework is trained with triplet losses based on [14] and with a direct regression loss between the video and text embeddings. We consider the triplet loss function $trip(a, p, n) = \max(0, \|a - p\|_2 - \|a - n\|_2 + \mu)$, with the anchor embedding a , the embeddings for the positive and negative samples p and n , and the margin μ . The **action triplet loss** encourages the alignment of the video representation o_{video} and text embeddings with the matching action as opposed to a sampled negative action $\phi_{\bar{a}}$. For this, we use the video embedding o_{video} as the anchor, the text embedding with ground truth action ϕ_a and adverb ϕ_v as the positive sample, and the text embedding of the same adverb but different action $\phi_{\bar{a}_i}$ as a negative:

$$\mathcal{L}_{trip,a} = \frac{1}{n} \sum_{i=1}^n trip(o_{video_i}, g(\phi_{a_i}, \phi_{v_i}), g(\phi_{\bar{a}_i}, \phi_{v_i})) \quad \text{for } \phi_{\bar{a}_i} \neq \phi_{a_i}. \quad (4)$$

We use an **adverb triplet loss** to push text embeddings containing the adverb antonym $\phi_{\bar{v}}$ away from the ground-truth text embedding:

$$\mathcal{L}_{trip,v} = \frac{1}{n} \sum_{i=1}^n trip(o_{video_i}, g(\phi_{a_i}, \phi_{v_i}), g(\phi_{a_i}, \phi_{\bar{v}_i})). \quad (5)$$

By restricting the negative samples for adverbs to their antonyms, the loss does not punish potential ambiguities of actions in videos (e.g. a drawer being opened slowly can at the same time be opened partially but not quickly). Our **regression loss** directly minimises the distance between the output video and text embeddings:

$$\mathcal{L}_{reg} = \frac{1}{n} \sum_{i=1}^n (o_{video_i} - g(\phi_{a_i}, \phi_{v_i}))^2. \quad (6)$$

The final loss is computed as the weighted sum of the above losses according to

$$\mathcal{L} = \lambda_a * \mathcal{L}_{trip,a} + \lambda_v * \mathcal{L}_{trip,v} + \lambda_{reg} * \mathcal{L}_{reg}, \quad (7)$$

with hyperparameters $\lambda_a, \lambda_v, \lambda_{reg} \in \mathbb{R}$.

Retrieving adverbs and videos (inference). Similar to [14], we evaluate our method on adverb-to-video and video-to-adverb retrieval given the ground-truth action a . For video-to-adverb retrieval, given a video $\mathbf{x}$ and action query a , we embed the video to obtain o_{video} , and

we obtain embeddings for j adverb-action combinations o_{txt_j} for $j \in \{1, \dots, V\}$. Using the cosine similarity metric we rank all the text embeddings o_{txt_j} by their similarity to the query video embedding o_{video} and we consider the highest-ranked pair as the retrieved adverb.

For adverb-to-video retrieval, given an adverb v and action a that are embedded to o_{txt} , we define the set of test videos containing action a as Γ . We rank all video embeddings o_{video_j} for videos in Γ using the similarity computed between each o_{video_j} and o_{txt} and select the video which is closest to o_{txt} .

4 Video-adverb retrieval benchmarks

In this section, we provide details about the datasets used in our experiments. In particular, we use five datasets for video-adverb retrieval. Furthermore, we propose two new dataset splits for the task of retrieving adverbs from videos for unseen adverb-action compositions.

Video-adverb retrieval datasets. HowTo100M Adverbs [13] consists of 5,824 video clips with annotations for 6 adverbs and 72 actions. In the following, we refer to HowTo100M Adverbs as **HowTo100M**. The recently proposed **Adverbs in Recipes** dataset has 10 adverbs, 48 actions and 7,003 videos. VATEX Adverbs [12] dataset has, with 34 adverbs and 135 actions, the largest variety of annotated adverbs and actions, consisting of 14,617 videos. We refer to VATEX Adverbs as **VATEX**. ActivityNet Adverbs [12] consists of 3,099 videos with 20 adverbs and 114 actions. We refer to it as **ActivityNet**. MSR-VTT Adverbs [12] is made up of 1,824 videos with 18 adverbs and 106 actions. In the following, we call this dataset **MSR-VTT**.

Unseen adverb-action compositions splits. We strive to explore the ability to recognise adverbs for novel adverb-action combinations. [12] proposed a dataset split for unseen compositions at test time for the VATEX dataset. Using the available videos in VATEX from [50], we replicate this split for the S3D video and text features used in this work, by omitting unavailable videos. We additionally propose new splits for unseen compositions on the ActivityNet and MSR-VTT datasets. We exclude HowTo100M Adverbs and Adverbs in Recipes, as both are subsets of HowTo100M which was used for pre-training the text and S3D video model. Hence, this would not comply with zero-shot learning principles.

To create splits for ActivityNet and MSR-VTT, we follow the protocol in [12]: We first split the set of possible adverb-action compositions into two non-overlapping sets, so that all adverbs and all actions are present in both sets, but individual compositions are only contained in one of the sets. We additionally constrain the compositions for each set so that for a given adverb-action composition, its antonym-action composition is assigned to the same set. We assign the videos from one of the sets to the training set and split the videos of the other half into two different sets, assigning half of the instances in each composition to the test set and the other to an unlabelled set (which is used to train [12] with pseudo-labelling). Table 1 shows details about the replicated split for VATEX, and for our proposed splits based on ActivityNet and MSR-VTT (full details are provided in the supplementary material).

Dataset	# tr (s)	# t (s)	# tr (p)	# t (p)
VATEX	6603	3293	319	316
MSR-VTT	987	454	225	225
ActivityNet	1490	848	635	543

Table 1: Statistics of the proposed dataset splits for the retrieval of unseen adverb-action compositions on the MSR-VTT and ActivityNet datasets. (tr: train, t: test, s: video samples, p: adverb-action pairs)

5 Experiments

In this section, we provide details about the baselines, implementation details, and evaluation metrics used in this work. Video-adverb retrieval results on five benchmarks are presented in Section 5.1, and we provide model ablation studies in Section 5.2. In Section 5.4, we investigate the transfer to unseen adverb-action compositions during inference.

Baselines. We report results for the **Prior** and **S3D pre-trained** baselines from [60]. **Prior** does not require any training but it uses the data distribution and adverb frequency for scoring. **S3D pre-trained** is also training-free and uses the similarity between frozen video and text representations from the S3D backbone jointly trained on video and text. **TIRG** [46] employs a similar residual gating mechanism as REGADA for image-text retrieval. To adapt it to the video domain, we use the same video encoder as our method. Different from REGADA, it models the composition as a transformation of the action embedding and uses a classification-based training objective. We also compare our framework to **Action Modifier** [43] and to the recently proposed **AC** frameworks [60]. AC tackles the task either as a classification (AC_{CLS}) or regression (AC_{REG}) problem.

Implementation details. We use the video and text features provided by [60] which were extracted using a frozen S3D model that was jointly pre-trained on video-text pairs from HowTo100M [27]. Here, $d_x = 1024$, T is the length of the video in seconds, and $d_\theta = 512$. REGADA uses an internal embedding dimension $d_{dim} = 400$. We use $N_g = 2$, except for HowTo100M and Adverbs in Recipes where $N_g = 3$ and $N_g = 4$ respectively. Additionally, we set $N_r = 2$ except for Adverbs in Recipes where we use $N_r = 3$. The dropout probability in the residual gating mechanism is $drop_g = 0.6$ for all datasets but Adverbs in Recipes and HowTo100M where we use $drop_g = 0.7$. The loss hyperparameters are chosen as $\lambda_a = 1$ for all datasets and $\lambda_v = 2.0$ for all datasets, except for $\lambda_v = 1.5$ on Adverbs in Recipes. Furthermore, we use a $\lambda_{reg} = 1.0$ for all dataset except for HowTo100M where $\lambda_{reg} = 1.5$. We train with a batch size of 512, and employ the Adam [49] optimizer with $\beta_1 = 0.9$, $\beta_2 = 0.999$, and weight decay 10^{-5} . Our method is trained for 2000 epochs using a learning rate of 10^{-5} for all datasets with the exception of HowTo100M where we use $3 * 10^{-5}$. We follow [60], and train all baselines for 1000 epochs using a learning rate of 10^{-4} . We conduct all experiments on a single Nvidia 2080-Ti GPU.

Evaluation metrics. We follow [60], and report mean Average Precision (mAP) scores for adverb-to-video-retrieval, in particular **mAP M** (“adverb-to-video (all)” in [43]) and **mAP W**. mAP M is computed by ranking videos that contain the same ground-truth action according to their similarity to the adverb-action text embedding. For mAP W, the class scores are reweighed according to their support size in the test set. For video-to-adverb retrieval, we report binary antonym accuracy **Acc-A**. This is equivalent to ranking adverb-action embeddings according to their similarity to the embedded video and calculating the mAP by restricting the set of adverbs to the target adverb and its antonym (“video-to-adverb (antonym)” in [43]). Similar to [60], we report the best metrics independently. This means that models corresponding to each result may originate from different epochs.

5.1 Comparison with the state of the art

In Table 2, we present adverb-to-video retrieval and video-to-adverb retrieval results with our REGADA framework on five benchmark datasets. It can be observed that REGADA outperforms the baselines across all datasets. In particular, we see more significant improvements of our framework over the prior methods for the adverb-to-video retrieval metrics

	HowTo100M [■]			Adverbs in Recipes [■]			ActivityNet [■]			MSR-VTT [■]			VATEX [■]		
	mAP W	mAP M	Acc-A	mAP W	mAP M	Acc-A	mAP W	mAP M	Acc-A	mAP W	mAP M	Acc-A	mAP W	mAP M	Acc-A
Priors	0.446	0.354	0.786	0.491	0.263	0.854	0.217	0.159	0.745	0.308	0.152	0.723	0.216	0.086	0.752
S3D pre-tr.	0.339	0.238	0.560	0.389	0.173	0.735	0.118	0.070	0.560	0.194	0.075	0.603	0.122	0.038	0.586
TIRG [■]	0.441	0.476	0.721	0.485	0.228	0.835	0.186	0.111	0.709	0.297	0.113	0.700	0.195	0.065	0.735
Act. M. [■]	0.406	0.372	0.796	0.509	0.251	0.857	0.184	0.125	0.753	0.233	0.127	0.731	0.139	0.059	0.751
AC _{CLS} [†] [■]	0.562	0.420	0.786	0.606	0.289	0.841	0.130	0.096	0.741	0.305	0.131	0.751	0.283	0.108	0.754
AC _{REG} [†] [■]	0.555	0.423	0.799	0.613	0.244	0.847	0.119	0.079	0.714	0.282	0.114	0.774	0.261	0.086	0.755
REGADA	0.567	0.528	0.817	0.704	0.418	0.874	0.239	0.175	0.771	0.378	0.228	0.786	0.290	0.113	0.817

Table 2: Results for adverb-to-video (mAP W/M) and video-to-adverb retrieval (Acc-A). Higher is better for all metrics. [†] refers to updated results provided by the authors.

(mAP W and mAP M) compared to video-to-adverb retrieval (Acc-A). For instance, on the HowTo100M dataset REGADA outperforms AC_{CLS} for adverb-to-video retrieval with mAP M and mAP W scores of 0.528 and 0.567 compared to 0.420 and 0.562. For the video-to-adverb retrieval measure Acc-A, REGADA obtains a score of 0.817 compared to 0.786 with AC_{REG}.

The most recent and strongest competitor [50] optimises its systems using two different losses. The best results obtained from these two models are reported for each dataset and metric, showing no clear pattern as to which model variant is stronger. Our REGADA framework consistently outperforms both model variants [50] on all metrics and datasets. We hypothesise that our framework’s strong performance can be attributed to its compositional embeddings which is a key element of REGADA.

5.2 Model ablations

This section analyses the impact of using different input text information, losses, and components in the text encoder on the overall video-adverb retrieval performance of REGADA.

Input to the text encoder. The gating mechanism in REGADA represents the composition as a residual on top of the adverb and allows the adverb information to be retained, leveraging the action as auxiliary information. We refer to the adverb as the *main* and the action as the *auxiliary* modality in REGADA. We investigate if a compositional adverb-action word embedding ϕ_{comp} , which directly embeds an adverb-action label pair (e.g. “cut quickly”) with h_{text} , can be used as the main modality instead. Table 3 shows the impact of using different main and auxiliary modalities. REGADA obtains scores of 0.290 and 0.113 for mAP W and mAP M on VATEX compared to 0.245 and 0.078 when using ϕ_a as main modality and ϕ_v as auxiliary. This confirms that capturing information about the adverb is crucial for solving the task. Acc-A is less affected by the type of input information, REGADA obtains 0.817 compared to 0.806 when using ϕ_{comp} as main and ϕ_a as auxiliary modality. Overall, using ϕ_v as main and ϕ_a as auxiliary modality is most effective across datasets.

Losses. In Table 4, we show the impact of our three loss functions, $\mathcal{L}_{trip,a}$, $\mathcal{L}_{trip,v}$, and $\mathcal{L}_{reg}$. On VATEX, REGADA obtains a mAP W and mAP M of 0.290 and 0.113 compared to 0.182 and 0.074 when using only $\mathcal{L}_{reg}$. For Acc-A, REGADA obtains a score of 0.817 compared to 0.756 for $\mathcal{L}_{trip,a} + \mathcal{L}_{trip,v}$. The regression loss $\mathcal{L}_{reg}$ boosts the performance on all datasets significantly. Our novel loss combination gives the best video-adverb retrieval performance by better aligning adverb-action compositions and video representations. Previous work either only used triplet losses [2, 3] or used a fixed textual regression target [50].

Residual gating mechanism in the text encoder. Table 5 analyses the contributions of the components of the residual gating mechanism, such as the residual branch, the sigmoid, and weight sharing between the gated and residual branches. On VATEX, REGADA achieves the best results. Interestingly, sharing weights between the gated and residual branches yields

Text Input		HowTo100M [■]			Adverbs in Recipes [■]			ActivityNet [■]			MSR-VTT [■]			VATEX [■]		
<i>main</i>	<i>auxiliary</i>	mAP W	mAP M	Acc-A	mAP W	mAP M	Acc-A	mAP W	mAP M	Acc-A	mAP W	mAP M	Acc-A	mAP W	mAP M	Acc-A
ϕ_a	ϕ_v	0.485	0.390	0.824	0.436	0.221	0.872	0.225	0.147	0.763	0.336	0.144	0.780	0.245	0.078	0.807
ϕ_{comp}	ϕ_v	0.498	0.454	0.827	0.518	0.322	0.877	0.220	0.150	0.751	0.350	0.144	0.771	0.255	0.084	0.808
ϕ_{comp}	ϕ_a	0.503	0.467	0.830	0.524	0.365	0.881	0.222	0.147	0.758	0.348	0.146	0.763	0.255	0.090	0.806
ϕ_v	ϕ_a	0.567	0.528	0.817	0.704	0.418	0.874	0.239	0.175	0.771	0.378	0.228	0.786	0.290	0.113	0.817

Table 3: Effect of using different types of input information for the text encoder in REGADA.

Loss			HowTo100M 			Adverbs in Recipes 			ActivityNet 			MSR-VTT 			VATEX 		
$\mathcal{L}_{trip,a}$	$\mathcal{L}_{trip,v}$	$\mathcal{L}_{reg}$	mAP W	mAP M	Acc-A	mAP W	mAP M	Acc-A	mAP W	mAP M	Acc-A	mAP W	mAP M	Acc-A	mAP W	mAP M	Acc-A
✓	✗	✗	0.361	0.228	0.697	0.429	0.214	0.836	0.162	0.104	0.582	0.259	0.138	0.714	0.133	0.047	0.677
✗	✓	✗	0.340	0.236	0.740	0.430	0.213	0.846	0.128	0.079	0.664	0.260	0.127	0.737	0.166	0.062	0.743
✗	✗	✓	0.470	0.378	0.743	0.468	0.234	0.839	0.202	0.140	0.729	0.288	0.186	0.743	0.182	0.074	0.700
✓	✓	✗	0.367	0.246	0.755	0.468	0.239	0.851	0.157	0.098	0.674	0.273	0.116	0.737	0.174	0.062	0.756
✓	✓	✓	0.567	0.528	0.817	0.704	0.418	0.874	0.239	0.175	0.771	0.378	0.228	0.786	0.290	0.113	0.817

Table 4: Impact of using different losses to train REGADA. For losses that are not used, the corresponding scalar weight in $\mathcal{L}$ is set to zero.

Components			HowTo100M [■]			Adverbs in Recipes [■]			ActivityNet [■]			MSR-VTT [■]			VATEX [■]		
R	σ	SW	mAP W	mAP M	Acc-A	mAP W	mAP M	Acc-A	mAP W	mAP M	Acc-A	mAP W	mAP M	Acc-A	mAP W	mAP M	Acc-A
✓	✓	✓	0.535	0.433	0.811	0.689	0.404	0.875	0.256	0.190	0.771	0.374	0.182	0.766	0.288	0.109	0.808
✓	✗	✗	0.512	0.496	0.811	0.501	0.269	0.862	0.234	0.171	0.770	0.360	0.194	0.780	0.260	0.098	0.804
✗	✓	✗	0.516	0.477	0.817	0.562	0.296	0.877	0.228	0.169	0.765	0.367	0.161	0.783	0.283	0.111	0.815
✓	✓	✓	0.567	0.528	0.817	0.704	0.418	0.874	0.239	0.175	0.771	0.378	0.228	0.786	0.290	0.113	0.817

Table 5: Impact of different components in the residually-gated text encoder. R: With residual branch W_{res} ; σ : With sigmoid; SW: Sharing weights between W_{res} and W_{gate} .

only slightly weaker results, with a mAP-W score of 0.288 compared to 0.290 with REGADA. For mAP M and Acc-A, REGADA obtains 0.113 and 0.817 compared to 0.111 and 0.815 when not using the residual. While some configurations can achieve better results in selected metrics, REGADA yields consistent state-of-the-art results across all metrics, confirming our model design choices.

5.3 Qualitative Results

We show qualitative results for REGADA on the VATEX dataset in Figure 2. In particular, success cases for REGADA which AC_{REG} retrieved a wrong adverb are shown below in the first and second columns. The third and fourth columns show videos with actions performed forwards/backwards, and upwards/downwards but labelled with only one of the adverbs. This makes both outputs plausible. The right-most column shows an example of a wrongly labelled video for which our model retrieves the correct adverb. This confirms REGADA’s strong generalisation capabilities. In general, we observe that REGADA better captures directional movements or speed than AC_{REG} . It is also superior at disentangling the diverse visual effect of adverbs on different actions (e.g. crawl vs. bend backwards). This can potentially be attributed to the compositional nature of our learned adverb-action representations.

5.4 Generalisation to unseen adverb-action compositions

We additionally evaluate the REGADA framework on video-to-adverb retrieval for unseen adverb-action compositions, i.e. compositions that were not seen during training. We consider the existing VATEX benchmark and our proposed MSR-VTT and ActivityNet splits for this task (see Section 4). Following [■], we report binary antonym classification accu-



Figure 2: Example results for REGADA (Ours) on the VATEX dataset compared to those from AC_{REG} . The two left examples are success cases for our model. The third and fourth example show bidirectionally performed actions that are labelled with only one of the adverbs. The right-most example shows a wrongly labelled video. Full videos are available at: <https://hummelth.github.io/ReGaDa>

racy for video-to-adverb retrieval. We provide additional baseline results with the CLIP [40] model (details for this are provided in the supplementary material). In Table 6, we observe that REGADA significantly outperforms AC_{REG} on VATEX with an accuracy of 61.7 compared to 54.9. On ActivityNet, REGADA obtains a score of 58.4, outperforming [40] with a score of 57.0. This is impressive given that [40] was additionally trained on pseudo-labelled data. CLIP obtains an antonym accuracy of only 54.5 on VATEX, showing a limited fine-grained retrieval capability of CLIP. We provide a further analysis of exploiting different word embeddings for unseen compositions in the supplementary material. Overall, our model yields better results than any prior framework for both seen (c.f. Table 2) and unseen compositions.

Model	VATEX	ActivityNet	MSR-VTT
CLIP [40]	54.5	55.1	57.0
Act. Mod. [40]	53.8	57.0	56.0
AC_{CLS} [40]	54.3	55.1	53.7
AC_{REG} [40]	54.9	53.9	59.0
REGADA	61.7	58.4	61.0

Table 6: Retrieval of unseen adverb-action compositions on the VATEX, ActivityNet and MSR-VTT benchmarks. [40] uses pseudo-labelling.

6 Conclusion

In this work, we proposed a framework for video-adverb retrieval that uses a residual gating mechanism to generate compositional adverb-action representations from adverb and action word embeddings. Along with a novel training objective, our model achieves state-of-the-art results on five video-adverb retrieval benchmarks. Moreover, we introduce two additional dataset splits to benchmark the retrieval of unseen adverb-action compositions. Our proposed framework outperforms all prior works on this task, confirming that our text encoder results in better generalisation abilities.

Acknowledgements: This work was supported by BMBF FKZ: 01IS18039A, DFG: SFB 1233 TP 17 - project number 276693517, by the ERC (853489 - DEXIM), and by EXC number 2064/1 - project number 390727645. The authors thank the International Max Planck Research School for Intelligent Systems (IMPRS-IS) for supporting T. Hummel and O.-B. Mercea. Furthermore, we would like to thank Massimiliano Mancini and Shyamgopal Karthik for helpful discussions and proofreading.

References

- [1] Jean-Baptiste Alayrac, Adria Recasens, Rosalia Schneider, Relja Arandjelović, Jason Ramapuram, Jeffrey De Fauw, Lucas Smaira, Sander Dieleman, and Andrew Zisserman. Self-supervised multimodal versatile networks. *NeurIPS*, 2020.
- [2] Taha Alhersh, Heiner Stuckenschmidt, Atiq Ur Rehman, and Samir Brahim Belhaouari. Learning human activity from visual data using deep learning. *IEEE Access*, 2021.
- [3] Lisa Anne Hendricks, Oliver Wang, Eli Shechtman, Josef Sivic, Trevor Darrell, and Bryan Russell. Localizing moments in video with natural language. In *ICCV*, 2017.
- [4] Jimmy Lei Ba, Jamie Ryan Kiros, and Geoffrey E Hinton. Layer normalization. *arXiv preprint arXiv:1607.06450*, 2016.
- [5] 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.
- [6] Damian Borth, Rongrong Ji, Tao Chen, Thomas Breuel, and Shih-Fu Chang. Large-scale visual sentiment ontology and detectors using adjective noun pairs. In *ACM MM*, 2013.
- [7] Fabian Caba Heilbron, Victor Escorcia, Bernard Ghanem, and Juan Carlos Niebles. Activitynet: A large-scale video benchmark for human activity understanding. In *CVPR*, 2015.
- [8] Chao-Yeh Chen and Kristen Grauman. Inferring analogous attributes. In *CVPR*, 2014.
- [9] Hui Chen, Guiguang Ding, Zijia Lin, Sicheng Zhao, and Jungong Han. Cross-modal image-text retrieval with semantic consistency. In *ACM MM*, 2019.
- [10] Shizhe Chen, Yida Zhao, Qin Jin, and Qi Wu. Fine-grained video-text retrieval with hierarchical graph reasoning. In *CVPR*, 2020.
- [11] Jianfeng Dong, Xirong Li, and Cees GM Snoek. Predicting visual features from text for image and video caption retrieval. In *IEEE Transactions on Multimedia*, 2018.
- [12] Hazel Doughty and Cees GM Snoek. How do you do it? Fine-grained action understanding with pseudo-adverbs. In *CVPR*, 2022.
- [13] Hazel Doughty, Ivan Laptev, Walterio Mayol-Cuevas, and Dima Damen. Action Modifiers: Learning from Adverbs in Instructional Videos. In *CVPR*, 2020.
- [14] Valentin Gabeur, Chen Sun, Karteek Alahari, and Cordelia Schmid. Multi-modal transformer for video retrieval. In *ECCV*, 2020.
- [15] Yuying Ge, Yixiao Ge, Xihui Liu, Dian Li, Ying Shan, Xiaohu Qie, and Ping Luo. Bridging video-text retrieval with multiple choice questions. In *CVPR*, 2022.
- [16] Meera Hahn, Andrew Silva, and James M Rehg. Action2vec: A crossmodal embedding approach to action learning. *arXiv preprint arXiv:1901.00484*, 2019.
- [17] Sergey Ioffe and Christian Szegedy. Batch normalization: Accelerating deep network training by reducing internal covariate shift. In *ICML*, 2015.

- [18] Phillip Isola, Joseph J Lim, and Edward H Adelson. Discovering states and transformations in image collections. In *CVPR*, 2015.
- [19] Diederik P Kingma and Jimmy Ba. Adam: A method for stochastic optimization. *arXiv:1412.6980*, 2014.
- [20] Ranjay Krishna, Kenji Hata, Frederic Ren, Li Fei-Fei, and Juan Carlos Niebles. Dense-captioning events in videos. In *ICCV*, 2017.
- [21] 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.
- [22] Yong-Lu Li, Yue Xu, Xiaohan Mao, and Cewu Lu. Symmetry and group in attribute-object compositions. In *CVPR*, 2020.
- [23] Yang Liu, Samuel Albanie, Arsha Nagrani, and Andrew Zisserman. Use what you have: Video retrieval using representations from collaborative experts. *BMVC*, 2019.
- [24] Huaishao Luo, Lei Ji, Ming Zhong, Yang Chen, Wen Lei, Nan Duan, and Tianrui Li. Clip4clip: An empirical study of clip for end to end video clip retrieval and captioning. *Neurocomputing*, 2022.
- [25] Massimiliano Mancini, Muhammad Ferjad Naeem, Yongqin Xian, and Zeynep Akata. Open world compositional zero-shot learning. In *CVPR*, 2021.
- [26] Antoine Miech, Ivan Laptev, and Josef Sivic. Learning a text-video embedding from incomplete and heterogeneous data. *arXiv preprint arXiv:1804.02516*, 2018.
- [27] 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.
- [28] Antoine Miech, Jean-Baptiste Alayrac, Lucas Smaira, Ivan Laptev, Josef Sivic, and Andrew Zisserman. End-to-end learning of visual representations from uncured instructional videos. In *CVPR*, 2020.
- [29] Ishan Misra, Abhinav Gupta, and Martial Hebert. From red wine to red tomato: Composition with context. In *CVPR*, 2017.
- [30] Davide Moltisanti, Frank Keller, Hakan Bilen, and Laura Sevilla-Lara. Learning action changes by measuring verb-adverb textual relationships. In *CVPR*, 2023.
- [31] Liliane Momeni, Mathilde Caron, Arsha Nagrani, Andrew Zisserman, and Cordelia Schmid. Verbs in action: Improving verb understanding in video-language models. *arXiv preprint arXiv:2304.06708*, 2023.
- [32] Muhammad Ferjad Naeem, Yongqin Xian, Federico Tombari, and Zeynep Akata. Learning graph embeddings for compositional zero-shot learning. In *CVPR*, 2021.
- [33] Tushar Nagarajan and Kristen Grauman. Attributes as operators: factorizing unseen attribute-object compositions. In *ECCV*, 2018.

- [34] Vinod Nair and Geoffrey E Hinton. Rectified linear units improve restricted boltzmann machines. In *ICML*, 2010.
- [35] Zhixiong Nan, Yang Liu, Nanning Zheng, and Song-Chun Zhu. Recognizing unseen attribute-object pair with generative model. In *AAAI*, 2019.
- [36] Andreea-Maria Oncescu, Joao F Henriques, Yang Liu, Andrew Zisserman, and Samuel Albanie. Queryd: A video dataset with high-quality text and audio narrations. In *ICASSP*, 2021.
- [37] Mayu Otani, Yuta Nakashima, Esa Rahtu, Janne Heikkilä, and Naokazu Yokoya. Learning joint representations of videos and sentences with web image search. In *ECCV*, 2016.
- [38] Jae Sung Park, Sheng Shen, Ali Farhadi, Trevor Darrell, Yejin Choi, and Anna Rohrbach. Exposing the limits of video-text models through contrast sets. In *ACL*, 2022.
- [39] Mandela Patrick, Po-Yao Huang, Yuki Asano, Florian Metze, Alexander G Hauptmann, Joao F. Henriques, and Andrea Vedaldi. Support-set bottlenecks for video-text representation learning. In *ICLR*, 2021.
- [40] 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. In *ICML*, 2021.
- [41] Andrew Rouditchenko, Angie Boggust, David Harwath, Brian Chen, Dhiraj Joshi, Samuel Thomas, Kartik Audhkhasi, Hilde Kuehne, Rameswar Panda, Rogerio Feris, et al. AVLnet: Learning audio-visual language representations from instructional videos. In *Interspeech*, 2021.
- [42] Nitish Srivastava, Geoffrey Hinton, Alex Krizhevsky, Ilya Sutskever, and Ruslan Salakhutdinov. Dropout: a simple way to prevent neural networks from overfitting. *JMLR*, 2014.
- [43] Senem Tanberk, Zeynep Hilal Kilimci, Dilek Bilgin Tükel, Mitat Uysal, and Selim Akyokuş. A hybrid deep model using deep learning and dense optical flow approaches for human activity recognition. *IEEE Access*, 2020.
- [44] Atousa Torabi, Niket Tandon, and Leonid Sigal. Learning language-visual embedding for movie understanding with natural-language. *arXiv preprint arXiv:1609.08124*, 2016.
- [45] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. Attention is all you need. *NeurIPS*, 2017.
- [46] Nam Vo, Lu Jiang, Chen Sun, Kevin Murphy, Li-Jia Li, Li Fei-Fei, and James Hays. Composing text and image for image retrieval-an empirical odyssey. In *CVPR*, 2019.
- [47] Xiaoyang Wang and Qiang Ji. A unified probabilistic approach modeling relationships between attributes and objects. In *ICCV*, 2013.

- [48] Xin Wang, Jiawei Wu, Junkun Chen, Lei Li, Yuan-Fang Wang, and William Yang Wang. Vatex: A large-scale, high-quality multilingual dataset for video-and-language research. In *ICCV*, 2019.
- [49] Yang Wang and Greg Mori. A discriminative latent model of object classes and attributes. In *ECCV*, 2010.
- [50] Michael Wray and Dima Damen. Learning visual actions using multiple verb-only labels. In *BMVC*, 2019.
- [51] Michael Wray, Diane Larlus, Gabriela Csurka, and Dima Damen. Fine-grained action retrieval through multiple parts-of-speech embeddings. In *ICCV*, 2019.
- [52] Wenhao Wu, Haipeng Luo, Bo Fang, Jingdong Wang, and Wanli Ouyang. Cap4Video: What can auxiliary captions do for text-video retrieval? In *CVPR*, 2023.
- [53] Bing Xu, Naiyan Wang, Tianqi Chen, and Mu Li. Empirical evaluation of rectified activations in convolutional network. *arXiv preprint arXiv:1505.00853*, 2015.
- [54] Jun Xu, Tao Mei, Ting Yao, and Yong Rui. Msr-vtt: A large video description dataset for bridging video and language. In *CVPR*, 2016.
- [55] Ran Xu, Caiming Xiong, Wei Chen, and Jason Corso. Jointly modeling deep video and compositional text to bridge vision and language in a unified framework. In *AAAI*, 2015.
- [56] Luowei Zhou, Chenliang Xu, and Jason Corso. Towards automatic learning of procedures from web instructional videos. In *AAAI*, 2018.
- [57] Linchao Zhu and Yi Yang. Actbert: Learning global-local video-text representations. In *CVPR*, 2020.
- [58] Dimitri Zhukov, Jean-Baptiste Alayrac, Ramazan Gokberk Cinbis, David Fouhey, Ivan Laptev, and Josef Sivic. Cross-task weakly supervised learning from instructional videos. In *CVPR*, 2019.