
Chirality in Action: Time-Aware Video Representation Learning by Latent Straightening

Piyush Bagad Andrew Zisserman

VGG, Dept. of Engineering Science, University of Oxford

Abstract

Our objective is to develop compact video representations that are sensitive to visual change over time. To measure such time-sensitivity, we introduce a new task: chiral action recognition, where one needs to distinguish between a pair of temporally opposite actions, such as “opening vs. closing a door”, “approaching vs. moving away from something”, “folding vs. unfolding paper”, *etc.* Such actions (i) occur frequently in everyday life, (ii) require understanding of simple visual change over time (in object state, size, spatial position, count . . .), and (iii) are known to be poorly represented by many video embeddings. Our goal is to build time aware video representations which offer linear separability between these chiral pairs. To that end, we propose a self-supervised adaptation recipe to inject time-sensitivity into a sequence of frozen image features. Our model is based on an auto-encoder with a latent space with inductive bias inspired by *perceptual straightening*. We show that this results in a compact but time-sensitive video representation for the proposed task across three datasets: Something-Something, EPIC-Kitchens, and Charade. Our method (i) outperforms much larger video models pre-trained on large-scale video datasets, and (ii) leads to an improvement in classification performance on standard benchmarks when combined with these existing models.

1 Introduction

The ever-increasing scale of video content on the Internet demands efficient and compact descriptors that can be readily used for classification, ranking and search. The goodness of video descriptors (or video representations) is largely measured in terms of action recognition on standard benchmarks such as Kinetics-400 [11], UCF101 [75] and Something-Something [28] to name a few. While action recognition performance provides a reliable single measure for the representation, more insight is obtained by establishing how well a given video descriptor encodes various aspects of the video such as objects, scene context, motion and temporal dynamics. We can coarsely categorize these aspects into: *static* properties (objects, scene, *etc.*) and *dynamic* properties (motion, visual change, *etc.*). It is well established that, apart from Something-Something, most contemporary video benchmarks tend to focus more on static properties [10, 35, 44, 49, 108]. While there has been an effort to shift the focus to evaluating dynamic [28, 49, 59, 73], properties of actions are still entangled with static understanding without a clear definition of dynamics. In this work, our objective is to study a specific time-sensitive property: understanding how well video descriptors encode *visual change* in a video.

What do we mean by “understanding” visual change? Consider an example action pair: “a person climbing up a ladder” vs. “a person climbing down a ladder”. In this case, the vertical position of the person changes over time and it is temporally opposite in the two actions. An ideal video descriptor should encode this change and use it to distinguish between such action pairs. We call such action pairs ‘*chiral*’, and the task of distinguishing between them ‘*chiral action recognition*’. Where can we find such actions? These are quite common in everyday life, and humans effortlessly recognize them.

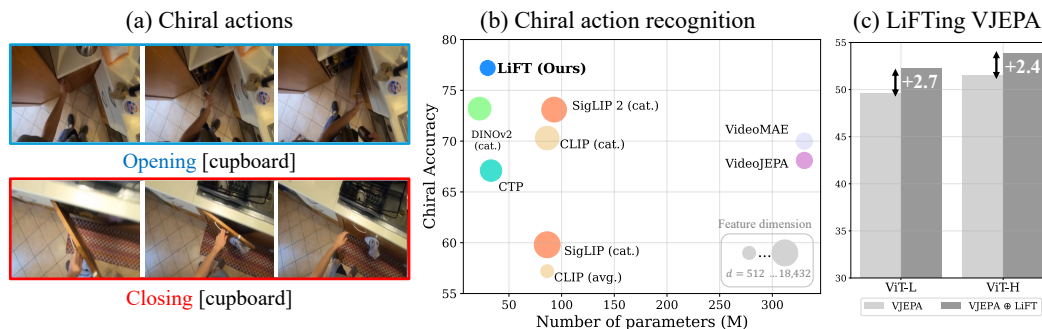


Figure 1: (a) We introduce *chiral actions*: temporally opposite action pairs to test time-awareness of video descriptors. We build a meta-dataset of chiral actions by mining SSv2, EPIC and Charades. To build time-aware descriptors, we propose LiFT (Linearized Feature Trajectories) that disentangles a sequence of DINOv2 features into a *static* and *dynamic* descriptor. (b) LiFT outperforms contemporary video and image models in recognizing chiral actions across the three datasets. Moreover, LiFT descriptors can be plugged into existing models (VideoJEPa or VideoMAE) to improve action recognition on standard benchmarks. (c) shows a linear probe of LiFT with VideoJEPa on SSv2.

In this work, we mine chiral pairs from three existing datasets (SSv2 [28], EPIC-Kitchens [17], and Charades [74]), to set up a chiral evaluation benchmark.

Turning to the descriptors, most existing video representations are obtained by two classes of methods, either (i) self-supervised video embeddings – natively multi-frame models, trained on millions of videos [5, 79, 89], or (ii) per-frame image models adapted for video data that are usually trained for specific datasets [57, 64]. We borrow from both lines of work and propose a self-supervised adaptation recipe that yields general video descriptors that outperform much larger models [5, 89] for chiral action recognition. Specifically, we hypothesize that a representation will be time sensitive if the per-frame features form a smooth trajectory in latent space. Inspired by Perceptual Straightening [32], we operationalize this by learning a model that maps per frame features from a strong image model to ordered points on lines in latent space. We show that the two vectors representing these high-dimensional lines yield time-aware video descriptors. We call our model *LiFT* for Linearized Feature Trajectories. Qualitatively, we show that LiFT learns compact video descriptors that encode a smooth, continuous approximations of the feature trajectories. Quantitatively, we show that LiFT descriptors are time-aware: they can distinguish between chiral action pairs across three datasets without specialized fine-tuning.

While LiFT descriptors achieve strong results on chiral action recognition, can they be more generally useful, say by combining with other video models? In this spirit, we evaluate linear and attentive probes with LiFT descriptors combined with video models such as VideoJEPa [5] on four standard action recognition datasets: Kinetics-400 [11], UCF-101 [75], HMDB-51 [42] and SSv2 [28]. We show that the combination of LiFT and a given video model always outperforms solely using the video model, across different video models, across all four benchmarks. This demonstrates that the time-sensitivity in LiFT preserves information that is complementary to standard video models, which in turn helps *lift* performance on action recognition benchmarks. In summary, our contributions are:

1. We propose a new task called chiral action recognition which requires discounting static context and accounting for the dynamic change in a video. We formulate a meta-dataset from three action recognition datasets to benchmark this task.
2. We propose LiFT: a self-supervised recipe to adapt DINOv2 features into a compact, time-sensitive, and general video descriptor. LiFT outperforms much larger video models (e.g., $10\times$ bigger and trained on over $6\times$ samples) by over 7% on the proposed chiral benchmark.
3. We demonstrate that LiFT encodes time-sensitive information that is complementary to contemporary video models. We show that combining LiFT descriptors with video models such as VideoJEPa [5], VideoMAE [79] and InternVideo2.5 [89] lifts performance with linear as well as attentive probes as compared to only using these models.

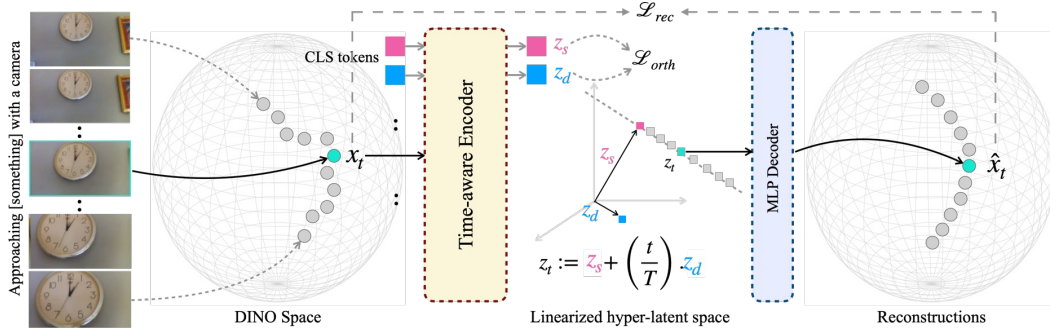


Figure 2: **Linearized Feature Trajectories (LiFT)**. We propose LiFT as a simple adaptation of image features to obtain time-aware video descriptors. First, we encode each frame independently with a DINOv2 backbone. Then, we pass them through a Transformer encoder with two learnable tokens: $\mathbf{z}_s$ (static) and $\mathbf{z}_d$ (dynamic). Next, inspired by Perceptual Straightening [32], we enforce linearity in the latent space which enables reconstruction of the trajectory with only the two learnable tokens. We do not show position encodings and projection layers for brevity. The network is trained with the usual reconstruction loss and an orthogonality regularization between the static and dynamic tokens. Once the model is trained, an input video is represented by concatenation of $\mathbf{z}_s, \mathbf{z}_d$.

2 LiFT: Video Representation by Linearized Feature Trajectories

Our objective is to learn a single time-aware descriptor vector for a given video. We want to avoid training large video models [5, 79] from scratch given an academic compute budget. In such a scenario, usually Parameter-Efficient Fine-Tuning (PEFT) is employed to adapt image models for video data [57, 58]. However, PEFT is dataset-specific supervised tuning while our goal is to obtain a more general video descriptor while still adapting an image model without any label supervision.

Considering these desiderata, we propose a simple recipe based on adapting a strong image model with reconstruction in the latent space. Our central hypothesis is that for videos that depict a visual change, the per-frame features lie on smooth trajectories that encode such change over time. While these trajectories tend to be non-linear, we can map them to a latent space in which they are parametrized by a line. This is loosely inspired by the Perceptual Straightening Hypothesis [32]: humans perceive image sequences that are non-linear in pixel space as straight lines in the perceptual space. Thus, a video can be represented simply by the vectors that define this line in the latent space. Owing to this linear formulation, we call the model *LiFT*: Linearized Feature Trajectories.

The schematic diagram of LiFT is shown in Fig. 2. Formally, consider a video as a sequence of images $\{I_t\}_{t=1}^T, I_t \in \mathbb{R}^{C \times H \times W}, \forall t$. First, we encode each frame independently using an image model Φ . We only retain the global CLS token per frame.

$$\mathbf{x}_t = \Phi(I_t) \in \mathbb{R}^D, \forall t. \quad (1)$$

Then, we train an autoencoder network to reconstruct $\{\mathbf{x}_t\}_1^T$ while learning meaningful video descriptors in its latent space. Now, we describe the Encoder-Decoder network and how we train it.

Encoder. The encoder takes in sequence of frame features and outputs a descriptor for the sequence. First, we project the feature sequence to a potentially lower-dimensional space.

$$\mathbf{e}_t = P_{\downarrow}(\mathbf{x}_t) \in \mathbb{R}^d, \forall t. \quad (2)$$

Sinusoidal position encoding is added to encode the frame index. Then, a Transformer Encoder takes in this sequence $\{\mathbf{e}_t\}_1^T$ along with two learnable CLS-like tokens: $\mathbf{e}_s$ and $\mathbf{e}_d$ that are used to encode *static* and *dynamic* information respectively in the feature sequence.

$$\mathbf{z}_s, \mathbf{z}_d = \text{TransformerEncoder}(\mathbf{e}_s, \mathbf{e}_d, \{\mathbf{e}_t\}) \quad (3)$$

We collect these tokens output by the Transformer, denoted by $\mathbf{z}_s$ and $\mathbf{z}_d \in \mathbb{R}^d$. The overall video descriptor is given by their concatenation $\mathbf{z} \in \mathbb{R}^{2d}$.

Decoder. The decoder takes in $\mathbf{z}_s, \mathbf{z}_d$ and time index t and outputs feature vector as time t . In the latent space, we enforce a linearity constraint defined by $\mathbf{z}_s, \mathbf{z}_d$ as shown in Fig. 2.

$$\mathbf{z}_t := \mathbf{z}_s + \left(\frac{t}{T}\right) \cdot \mathbf{z}_d \in \mathbb{R}^d, \forall t \quad (4)$$

The decoder is a two-layer MLP that takes in $\mathbf{z}_t$ and outputs the reconstructed feature at time t .

$$\hat{\mathbf{x}}_t = \text{MLPDecoder}(\mathbf{z}_t) \in \mathbb{R}^D, \forall t. \quad (5)$$

Training objective. We train the network with the usual reconstruction loss and a regularizer that encourages orthogonality between the static and dynamic latent vectors.

$$\mathcal{L} := \mathcal{L}_{\text{rec}} + \lambda \mathcal{L}_{\text{orth}} = \sum_{t=1}^T \|\mathbf{x}_t - \hat{\mathbf{x}}_t\|_2^2 + \lambda \cdot \text{cos-sim} \left(\frac{\mathbf{z}_s}{\|\mathbf{z}_s\|_2}, \frac{\mathbf{z}_d}{\|\mathbf{z}_d\|_2} \right) \quad (6)$$

Note that this is unsupervised. Once trained, we discard the Decoder and use the Encoder to get a LiFT video descriptor $(\mathbf{z}_s, \mathbf{z}_d)$.

Time-awareness of LiFT descriptors. To gain an insight into what LiFT learns, first, we visualize the joint tSNE embeddings of the original and reconstructed trajectories on sample videos in Fig. 3(a). We also visualize the tSNE embeddings of an example action pair in Fig. 3(b). These illustrate that: (a) LiFT outputs a smooth, continuous approximation of the original trajectories evident from the tSNE plot. In a sense, LiFT captures the “arc of change” depicted in the video; and, (b) Thanks to the simple linearization, LiFT is compelled to learn compact descriptors that can distinguish between temporally opposite actions such as “opening vs closing a door”. Note that the linearization is only applied in the latent space, thus, LiFT’s reconstructed trajectories can be non-linear.

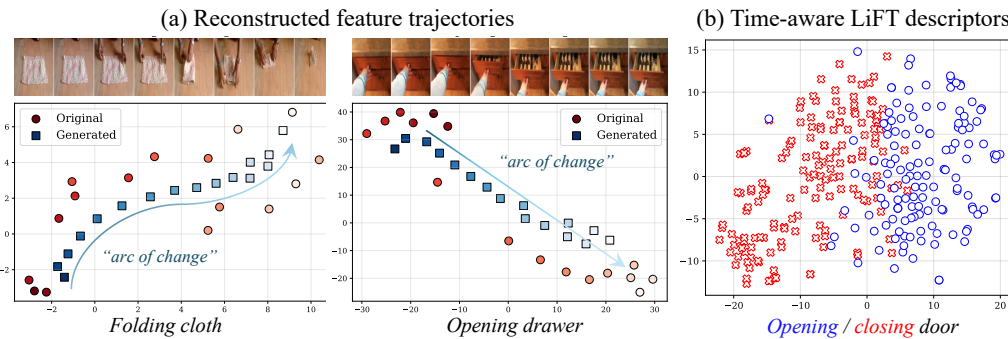


Figure 3: **Qualitative analysis.** (a) LiFT reconstructs a smooth, continuous approximation of the original feature trajectories roughly encoding the arc of visual change. (b) LiFT descriptors are time-aware: LiFT distinguishes between temporally opposite actions such as “opening/ closing door”.

Implementation details. We use DINOv2 (ViT-S/14) [56] with registers [18] as the base image feature extractor with dimension $D=384$. We only use the CLS token output for each frame. We linearly sample $T=16$ frames from each video and compute the features ahead of training. The input feature sequence is first projected to the space of the Encoder $\mathbb{R}^D \rightarrow \mathbb{R}^d$; we choose $d=384$. The Encoder is a standard Transformer [14] with 4 layers and 8 attention heads each. The Encoder uses sinusoidal position encoding [14] to encode the frame index. The outputs from the two CLS tokens, $\mathbf{z}_s, \mathbf{z}_d \in \mathbb{R}^d$ are then projected to a space where we impose the linearity constraint: $\mathbb{R}^d \rightarrow \mathbb{R}^d$. The Decoder is an MLP with 2 hidden layers each followed by a GeLU activation [34] and LayerNorm [2]. Overall, the model has 8.7M trainable parameters beyond the 22M parameters of the frozen DINOv2 encoder. We provide more details on the architecture in the Supplemental. We also provide ablations varying d and the amount of training data in the Supplemental.

Training. We train LiFT on Kinetics-400 [11] which has about 240K videos. Since the image encoder is frozen, we pre-compute features which makes the training very efficient. The model is trained for 500 epochs with a batch size of 128 with Adam optimizer [41] with a learning rate of 0.001 and a LRPlateau scheduler.



Figure 4: **Example chiral pairs** from each of the three datasets. To distinguish the actions in each pair, one needs to discount the spatial context and account for *what* is changing over time and *how*.

3 Chirality in Action (CiA) Dataset

To quantitatively measure time-awareness of LiFT (or other video descriptors), we propose a new task, namely, *chiral action recognition*. The study of time has covered aspects such as arrow of time [60, 91], order of frames [53, 102], recognizing temporally fine-grained actions [73, 104], or space-time tracking [96]. Prior work has developed specialized benchmarks and models for such tasks. In contrast, we want to measure time-awareness of a general video descriptor in recognizing simple everyday actions. It is well-established that existing action recognition benchmarks are biased to spatial understanding [35, 72]. Thus, we narrow down our focus on what we call *chiral actions*.

Chiral actions. In daily life, we often perform actions such as “closing/opening a door”, “folding/unfolding a cloth” or “getting in/out of a car”. Our proposal is that a good video model should distinguish between such temporally opposite actions. Loosely inspired by the notion of *visual chirality* [50], we call such action pairs as *chiral*, *i.e.*, pairs that are approximately mirror reflections along the time dimension. Consider the examples shown in Fig. 4. In distinguishing between these actions, one needs to discount the spatial context and account for the visual change over time.

Note that unlike in the study of arrow of time [91], we do not artificially reverse time-arrow but in a sense, our chiral actions have a naturally opposite arrow of time. The notion of chiral actions is also closely related to *reversible* actions studied in [62]. However, our work is complementary since we can use the methods in [62] to identify chiral actions and then evaluate video models on them. Finally, chiral actions, as we define them, are similar to *nearly symmetric actions* introduced in concurrent work by Ponbavagathi and Roitberg [61]. However, we build a meta-dataset of a more general mix of datasets that includes a richer set of actions, has both exo- and ego-centric videos and is larger in size.

Constructing CiA dataset. We build a meta-dataset out of chiral subsets of popular action recognition datasets. We identify three datasets to build a benchmark for chiral action recognition: Something-Something (SSv2) [28], EPIC-Kitchens (EPIC) [17], and Charades [74]. These datasets come with action recognition labels with separate verb and noun annotations. For each dataset, we build chiral pairs from the provided labels as follows.

1. We pass the list of action verbs to ChatGPT and ask it to find antonym pairs. We manually verify the output to remove pairs that have hallucinated verbs or those that are not visually antonymous.
2. For each verb pair, we group similar nouns together. For example, for verb pair “opening” vs “closing”, nouns such as “door/cupboard/drawer” that represent visually similar actions are grouped. This group is represented by the triplet (“opening”, “closing”, “[door]”). Likewise, (“opening”, “closing”, “[box]”) represents a separate chiral group where objects such as “tiffin/box/parcel”, *etc* are grouped together. Thus, each chiral group is a triplet consisting of a pair of opposite verbs and the associated noun.
3. For each chiral group, we split the videos into train and test sets following the split defined in the original dataset.

Some basic numbers for each dataset are provided in Table 1 and visual examples are shown in Fig. 4.

Base dataset	Chiral groups	Avg videos/group	Example chiral group
Something-Something (SSv2) [28]	16	852.8	Folding / Unfolding [something]
EPIC-Kitchens (EPIC) [17]	66	412.2	Opening / Closing [door]
Charades [74]	28	768.4	Taking / Putting a [laptop]

Table 1: **Numbers for CiA meta-dataset.** We mine chiral action pairs in three existing action recognition datasets to build our benchmark. Visual examples are shown in Fig. 4.

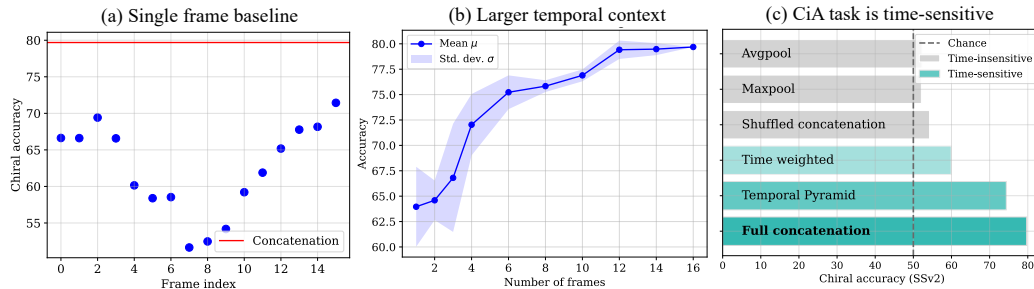


Figure 5: **Time-sensitivity of chiral action recognition.** On SSv2 (Chiral), we run simple baselines with DINOv2 features to test how time-sensitive the task is. In (a), (b), the video descriptor is obtained by concatenating per frame features. (a) Single frames chosen at the ends tend to do well but lag far behind using all frames. (b) Increasing number of frames provides consistent benefit while saturating around 16 frames. (c) **Time insensitive pooling** (e.g., mean) of features is noticeably worse than **time-sensitive pooling** (e.g., full concatenation or time-weighting).

Evaluation protocol. We measure time-sensitivity as a test of linear separability of chiral actions. For a given video model and a chiral group, we extract video representations for each video from the two antonym classes. We train a linear classifier in the feature space on the train set. We repeat this for each chiral group and report the average accuracy on test set across all groups. There are several reasons to choose this evaluation protocol: (i) We want the evaluation method to be as simple as possible so that the true strength of the representation is measured without confounding with the strength of the evaluation model (*i.e.*, in this case, a linear probe); (ii) Since there can be several chiral groups (e.g., $K=16$ in SSv2), it is computationally convenient to train simple linear probes rather than a more complicated and compute-heavy model for each group; (iii) Evaluating frozen features is much more common and practical [5, 12] as models get larger and evaluations need to be faster.

Properties of chiral action recognition. To analyze the time-sensitivity of chiral action recognition, we consider simple pooling baselines for DINOv2 features. The simplicity of the model enables us to study the task while discounting model strength. We show the results on SSv2 in Fig. 5. This shows that (a) A single frame chosen at either ends of the video yields a decent baseline as the end frames depict either the start or the end state of an action which is sometimes informative. However, the performance of such single frame baseline significantly lags using more frames as shown in (b). Increasing the temporal context with more frames consistently improves performance. (c) Finally, we compare **time-insensitive pooling** methods (e.g., average) with **time-sensitive** ones (e.g., full concatenation). We find that time-sensitive pooling does significantly better. Overall, this analysis highlights that the task at hand (chiral action recognition) benefits from time-aware ordering of many frames which establishes that it does not suffer as much from a single- or static-frame bias [35, 72].

4 Experiments

In this section, we first present results on our proposed chiral action recognition task. Then, we explore more general action recognition tasks where our descriptor is useful. In particular, in Section 4.1, we show that LiFT outperforms much larger video models in distinguishing between temporally opposite actions without specialized fine-tuning; and in Section 4.2, we show that plugging the LiFT descriptor with other general video models lifts their performance on standard action recognition benchmarks.

4.1 Chiral action recognition

Experimental details. We follow the evaluation protocol detailed in Section 3. For each chiral group, we compute video descriptors as described in Section 2 and train a linear classifier. We compare

against two sets of baselines. (i) image models such as the image encoder in CLIP. (ii) Video models trained with self- or language supervision on large-scale video datasets. Typically, we sample $T=16$ frames linearly from the video and compute a single descriptor. For image models, we concatenate the per-frame features to obtain the video descriptor. For video models based on R(2+1)D architecture (e.g., TCLR [19]), we sample $T' = 8$ clips over the span of the video and concatenate clip features to obtain a single descriptor. For Transformer-based video models, if the model uses a CLS token, we treat that as descriptor or average pool all the output tokens unless stated otherwise.

Main result. From the results shown in Table 2, we observe that the proposed LiFT features achieve the best performance on SSv2, EPIC and Charade, while being compact ($d = 768$). Notably, LiFT beats much heavier video models such as VideoMAE, VideoJEPa and InternVideo2.5. Interestingly, naively concatenating image features with a strong model (DINOv2, SigLIP2) generally performs better or at par with much heavier video models (e.g., VideoMAE) reinforcing that the complete sequence of image features does retain rich information for chiral action recognition. However, naturally, naively concatenating yields very bulky descriptors which may be impractical, say, in indexing a database with millions of videos. Less surprisingly, Transformer-based video models outperform ResNet (R(2+1)D) based models.

Finally, we find that careful feature pooling can make a notable difference. For example, InternVideo2.5/VideoMAE/VideoJEPa output a sequence of tokens per frame. Average pooling over space and concatenating over time does much better than average pooling over space and time. Nevertheless, average pooling output tokens (e.g., from VideoMAE) is still reasonable and time sensitive because (i) it does not have an explicit CLS token as a single video embedding, (ii) it comprises of 3D space-time tokens enhanced by temporal position encoding; different frames interact with each other in every layer of the transformer.

Model	Architecture	Pooling	D ↓	Chiral Accuracy ↑		
				SSv2	EPIC	Charades
Chance	-	-	-	50.0	50.0	50.0
<i>Image models with naive concatenation</i>						
CLIP [65]	ViT-B/16	Average	512	53.5	63.4	54.7
CLIP [65]	ViT-B/16	Concat.	8192	71.6	71.6	67.7
BLIP2 [45]	ViT-B/16 + Q-Former	Concat.	12288	73.2	70.3	67.3
SigLIP [103]	ViT-B/16	Concat.	18432	57.9	66.2	55.2
SigLIP 2 [80]	ViT-B/16	Concat.	12288	76.8	74.7	67.8
DINOv2 [56]	ViT-S/14	Concat.	6144	79.7	74.1	65.8
<i>Image models adapted for video (trained on Kinetics-400)</i>						
ST-Adapter [57]	ViT-B/16	Learned	768	50.5	63.7	54.4
DiST [64]	ViT-B/16	Learned	512	52.1	59.9	55.9
<i>Video models</i>						
Tubelet Contrast [77]	R(2+1)D	Concat.	4096	64.6	62.8	58.9
TCLR [19]	R(2+1)D	Concat.	4096	67.9	62.5	58.8
CTP [86]	R(2+1)D	Concat.	4096	78.8	64.4	58.0
VideoMAE [79]	ViT-L/16x16x2	Average	1024	80.3	70.5	59.1
VideoMAEv2 [87]	ViT-B/16x16x2	Average	768	65.3	67.5	55.5
SIGMA [69]	ViT-B/16x16x2	Average	768	66.5	69.1	56.1
MME [76]	ViT-B/16x16x2	Average	768	78.4	70.8	57.5
VideoJEPa [5]	ViT-L/16x16x2	Average	1024	80.4	67.4	56.4
InternVideo 2.5 [89]	InternViT-6B	Average	4096	55.8	66.1	55.4
VideoMAE [79]	ViT-L/16x16x2	Time concat.	8192	85.7	75.0	66.1
VideoJEPa [5]	ViT-L/16x16x2	Time concat.	8192	85.4	70.8	57.1
InternVideo 2.5 [89]	InternViT-6B	Time concat.	32768	80.0	70.9	62.8
LiFT (Ours)	ViT-S/14	Learned	768	86.6	75.5	69.5

Table 2: **Results on chiral action recognition.** (1) Our method (LiFT) of efficiently adapting sequential information in DINOv2 features has the best performance on the chiral splits of all three datasets. (2) On average, sequence of image features contain stronger discriminative information for chiral actions in comparison to native video models.

Ablation study. We run ablation over key design choices such as the image encoder in LiFT. Results of ablation study on SSv2 are shown in Table 3. We note that self-supervised image features such as iBOT [107] and DINOv2 [56] outperform language-supervised models such as CLIP [65] or SigLIP2 [80]. We hypothesize that since DINOv2/iBOT are better at capturing spatial details, their feature trajectories of a video tend to better capture smooth visual change compared to language-supervised models. From the last two rows, we also establish that using the orthogonal loss does provide a small benefit in chiral action recognition.

What kinds of change are easier to understand? We categorize the visual change involved in each chiral action pair. For example, in "moving towards/away from the camera", the object size or depth changes or in "taking/putting one of many objects on table", the object count changes. We average the performance across all chiral pairs that depict a given kind of visual change across all three datasets and report in Table 4. We find that LiFT features shine in distinguishing chiral actions that involve change in object state or count but struggle in those with change in position along x -axis.

Image encoder	Architecture	SSv2 (Chiral)
CLIP	ViT-B/16	75.9
SigLIP2	ViT-B/16	77.8
BLIP2	ViT-B/16 + QF	75.7
iBOT	ViT-B/16	80.3
DINOv2	ViT-B/14	85.4
DINOv2	ViT-L/14	85.9
LiFT w/o $\mathcal{L}_{\text{orth}}$	ViT-S/14	85.9
LiFT	ViT-S/14	86.6

Table 3: **Ablation on image encoders.** (i) self-supervised image features (*e.g.*, iBOT/DINOv2) outperform language-supervised features (*e.g.*, CLIP), (ii) with DINOv2, features out of larger models do not necessarily show improvement, (iii) using orthogonality loss helps.

Change type	VMAE	VJEPa	DINO	LiFT
Dist. bet. objects	70.8	87.5	83.3	87.5
Object count	64.2	62.4	69.5	72.4
Object size/depth	96.8	96.8	92.2	96.8
Object state	72.9	66.3	75.9	80.7
Spatial position $\leftrightarrow$	96.3	96.1	75.7	75.2
Spatial position $\updownarrow$	91.5	89.7	79.7	93.6

Table 4: **Performance across kinds of change.** Color **green** denotes performance at par or better than competing models and **red** denotes worse. LiFT features shine in distinguishing chiral actions that involve change in object state or count but struggle in those with change in position along x -axis. DINO here denotes naive concatenation of DINO features.

4.2 LiFTing video models on standard benchmarks

While LiFT outperforms much heavier video models in recognizing chiral actions, can it help improve performance on standard action recognition? We conduct an extensive linear probe evaluation across four standard datasets: Kinetics-400 (K400) [85], UCF-101 [75], HMDB-51 [42] and SSv2 [28]. The experimental details are provided in the Supplemental. As shown in Table 5a, while LiFT by itself does not beat top video models, concatenating LiFT with such video models consistently lifts their performance. This indicates that LiFT descriptors have complementary information. Furthermore, in Table 5b, we ablate over kinds of probes and model sizes used. We consistently observe a benefit with LiFT. Interestingly, with VideoJEPa as well as VideoMAE, ViT-L combined with LiFT even outperforms a scaled up ViT-H. Thus, overall, an adapter such as LiFT when combined with standard video models can provide strong video representations useful for classification, retrieval and search.

4.3 Investigating performance difference between vertical vs. horizontal motion

As noted in Table 4, there is a substantial difference between the ability of LiFT's to distinguish vertical translation motion ($\up$ vs. $\down$) and horizontal motion ($\rightarrow$ vs. $\leftarrow$). We uncover two reasons for this. First, the base model is better at encoding changes along vertical axis compared to horizontal axis. This is shown in Table 4 by comparing LiFT with a baseline with concatenated DINOv2 features over time. The latter is better at encoding $\updownarrow$ changes compared to $\leftrightarrow$ changes. Second, and more substantially, this difference is caused by the anisotropic sensitivity of DINOv2 feature sequence, *i.e.*, DINOv2 features vary less with horizontal movement vs vertical movement. To quantify this, we generate synthetic sequences with a checkerboard background and a colored disc that moves either horizontally or vertically at a constant rate. As the measure, we compute the variance over time and then average it across the feature dimensions. This time variance in vertical shift sequences is about 25% higher than that in horizontal shift sequences. Finally, we demonstrate three tricks to reduce this gap and improve overall performance: (i) scaling up image resolution, (ii) using a better base

Model	K400	UCF	HMDB	SSv2	Arch.	Probe	Base	Base $\oplus$ LiFT	Δ
Chance	0.25	0.99	1.96	0.58	<i>VideoJEPA</i>				
LiFT	55.4	86.6	65.2	30.8	ViT-L	Non-lin.	51.7	54.2	+2.5
VJEPA	59.8 [†]	91.3	76.1	49.6 [†]	ViT-L	Attentive	65.9 [†]	66.9	+1.0
VJEPA $\oplus$ LiFT	63.7	92.6	78.0	52.3	ViT-L	Linear	49.6 [†]	52.3	+2.7
Δ	+3.9	+1.3	+1.9	+2.7	ViT-H	Linear	51.5	53.9	+2.4
VideoMAE	55.0	83.6	66.5	38.3	<i>VideoMAE</i>				
VideoMAE $\oplus$ LiFT	63.6	88.8	72.6	46.3	ViT-L	Non-lin.	43.9	50.1	+6.2
Δ	+8.6	+5.2	+6.1	+6.0	ViT-L	Attentive	61.5	63.7	+2.2
InternVid2.5	62.8	88.2	71.9	23.4	ViT-S	Linear	19.4	37.3	+17
InternVid2.5 $\oplus$ LiFT	65.9	90.3	75.3	35.9	ViT-B	Linear	25.6	41.1	+16
Δ	+3.1	+2.1	+3.4	+11.5	ViT-L	Linear	38.3	46.3	+6.0
					ViT-H	Linear	40.0	46.9	+6.9

(a) LiFT combined with video models lifts their performance results across four action recognition benchmarks.

(b) LiFT consistently improves performance of video models across model scale and probes.

Table 5: **Results on standard action recognition datasets.** LiFT improves probing accuracies with standard video models across datasets and model sizes. [†]Note: The numbers for VideoJEPA are obtained with our experimental setup. We could not precisely reproduce the numbers reported in the paper [5] even using their codebase. We have reached out to the authors for clarification.

model (WebSSL [22]), and (iii) using test-time rotation augmentation. More details are provided in the Appendix C.4.

5 Related Work

Human perception of videos and time. Psychologists have tried to understand how humans perceive visual change in videos (*e.g.*, motion) for a long time [83, 90]. More recently, Hénaff et al. [32] present a remarkable finding: visual system in humans and macaques transforms complex pixel dynamics in videos into straighter temporal trajectories [32, 33]. Straighter trajectories make predictions easier and predictions are a fundamental part of human perception. Inspired by this insight, we learn an auto-encoder on image feature trajectories with linearity baked in the latent space. Although there is some prior work inspired by Perceptual Straightening [27, 31, 55], we apply it to a sequence of image features and show that it leads to more general time-aware representations.

Time-aware video representations. Temporally pooling image sequences has been a classical way of representing videos. Carefully crafted pooling in pixel space [8], in motion/flow space [9] and in embedding spaces [23] have been devised. Since the prominence of deep learning on videos, time has been creatively used as a source of self-supervision: space-time jigsaw [40], time arrow [62, 91], time order [26, 94, 97], speed [6], tracking [36, 85], contrasting temporal views [19, 63, 68], cycle consistency in time [21] or explicitly modeling temporal dynamics [15, 37, 105]. Modern video encoders are based on Transformers [1, 5, 7, 51, 79, 87]. Data-efficiency [77, 79] and time-sensitivity [20, 69, 78, 98] of video models continue to be active areas of research [71]. In this work, we investigate time-sensitivity of existing models through chiral action recognition and propose a simple recipe to embed videos based on summarizing trajectories of image features. Concurrent to our work, Xue et al. [95] propose a reinforcement learning-based training strategy to instill arrow of time awareness in video LLMs. Our chiral actions are related to the arrow of time, but we do not artificially reverse the arrow of time in a video; instead we aim to distinguish actions that are naturally opposite along the arrow of time.

Action recognition benchmarks. Early datasets for action recognition in videos include UCF-101 [75], HMDB-51 [42], Sports-1M [38] and Kinetics [11]. Transformers [84] prompted the rise of multimodal video datasets with text [4, 52], audio [13, 25] and 3D [29, 81]. LLMs led to the rise of instruction-tuning datasets [46, 106] and benchmarks [24, 47, 59] for videos. However, the community has repeatedly discovered that a majority of these do not actually test for time; a single frame or an unordered set of frames would suffice to recognize the action in the video [10, 16, 35, 44, 49, 108]. SSv2 [28], Diving-48 [49] introduced temporally sensitive actions while other datasets evaluate specific aspects: causal/counterfactual reasoning [59, 93, 99], compositionality [30, 101], concept-binding [39, 70], temporal prepositions [3] and verbs [54, 72]. In this work, we propose chiral action recognition that evaluates video features in discriminating temporally opposite actions. Our definition

of chirality is related to that of equivariant actions in Price and Damen [62] but their aim was more to discover actions invariant/equivariant to time flipping. It is also related to the subset of action antonyms in Cores et al. [16]. Chirality is also related to *nearly symmetric actions* in concurrent work by Ponbagavathi and Roitberg [61]. However, unlike [61], we propose a more general video embedding model trained in an unsupervised manner.

Efficient adaptation of image models to videos. Given the computational cost of training video models from scratch, Parameter Efficient Fine-Tuning (PEFT) methods to adapt image models for videos have emerged [48, 57, 64, 67, 82]. Since we use frozen DINOv2 features, our work also adapts image model for video recognition. However, PEFT methods are usually trained separately for each downstream dataset and generally used in a supervised learning setup. Our method is more generally applicable. It is trained in an unsupervised manner on Kinetics-400 and the resulting video embeddings are shown to be applicable for chiral action recognition across three datasets.

6 Discussion and Conclusion

In an effort to develop time-sensitive video descriptors, we proposed Linearized Feature Trajectories (LiFT): a simple recipe to adapt DINO per-frame features with an auto-encoder with an inductive bias inspired by Perceptual Straightening [32]. As a measure of time-sensitivity, we introduce chiral action recognition to distinguish between temporally opposite actions such as “opening vs. closing a door”. We created the CiA meta-dataset with chiral pairs mined from three public datasets: SSv2 [28], EPIC [17], and Charades [74]. On CiA, we show that LiFT outperforms much heavier video models including VideoJEPa [5] and VideoMAE [79] while being compact. Furthermore, we show that the time-sensitive LiFT descriptors contain information that is complementary to standard video models. For example, LiFT when combined with VideoJEPa lifts performance across four action recognition benchmarks: Kinetics [11], UCF [75], HMDB [42] and SSv2 [28].

Future work. Since we only use per frame CLS tokens, LiFT likely misses out on some spatial details, especially horizontal translation as shown in Table 4. Investigating ways of mitigating this, *e.g.*, using a sequence of dense feature maps, is an open avenue for future research. Furthermore, since our recipe is self-supervised, combining it with other compute-heavy self-supervised pre-training paradigms such as Masked Modeling [5, 79] or Autoregression [66, 92] should be interesting avenues to imbue more time-sensitivity into these representations.

Societal impact. Our work improves video representations for temporal actions thus impacting video applications, *e.g.*, indexing a massive video database or text to video retrieval. A less direct impact can be on improving domestic robots that frequently need to perform temporal actions like opening/closing a door. Our work also has a positive impact in terms of energy efficiency: it is much more efficient in training and inference compared to larger video models. Finally, our work inherits the risks of any video models in being applied to NSFW video content.

Acknowledgments. This research is funded by the EPSRC Programme Grant VisualAI EP/T028572/1, and a Royal Society Research Professorship RSRP\R\241003. We thank Ashish Thandavan for support with infrastructure, and Sindhu Hegde and Makarand Tapaswi for useful discussions. We also thank the anonymous reviewers that helped improve this work.

References

- [1] Anurag Arnab, Mostafa Dehghani, Georg Heigold, Chen Sun, Mario Lučić, and Cordelia Schmid. Vivit: A Video Vision Transformer. In *International Conference on Computer Vision (ICCV)*, 2021. 9
- [2] Jimmy Lei Ba, Jamie Ryan Kiros, and Geoffrey E Hinton. Layer Normalization. *arXiv preprint arXiv:1607.06450*, 2016. 4
- [3] Piyush Bagad, Makarand Tapaswi, and Cees GM Snoek. Test of Time: Instilling Video-Language Models with a Sense of Time. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2023. 9
- [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 *International Conference on Computer Vision (ICCV)*, 2021. 9
- [5] Adrien Bardes, Quentin Garrido, Jean Ponce, Xinlei Chen, Michael Rabbat, Yann LeCun, Mido Assran, and Nicolas Ballas. Revisiting Feature Prediction for Learning Visual Representations from Video. *Transactions on Machine Learning Research*, 2024. 2, 3, 6, 7, 9, 10, 27, 28
- [6] Sagie Benaim, Ariel Ephrat, Oran Lang, Inbar Mosseri, William T Freeman, Michael Rubinstein, Michal Irani, and Tali Dekel. Speednet: Learning the Speediness in Videos. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2020. 9
- [7] Gedas Bertasius, Heng Wang, and Lorenzo Torresani. Is Space-Time Attention All You Need for Video Understanding? In *International Conference on Machine Learning (ICML)*, 2021. 9
- [8] Hakan Bilen, Basura Fernando, Efstratios Gavves, Andrea Vedaldi, and Stephen Gould. Dynamic Image Networks for Action Recognition. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2016. 9
- [9] Aaron F. Bobick and James W. Davis. The Recognition of Human Movement Using Temporal Templates. *IEEE Transactions on Pattern Analysis and Machine Intelligence (TPAMI)*, 2001. 9
- [10] Shyamal Buch, Cristóbal Eyzaguirre, Adrien Gaidon, Jiajun Wu, Li Fei-Fei, and Juan Carlos Niebles. Revisiting the "Video" in Video-Language Understanding. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2022. 1, 9
- [11] Joao Carreira and Andrew Zisserman. Quo Vadis, Action Recognition? A New Model and the Kinetics Dataset. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2017. 1, 2, 4, 9, 10
- [12] João Carreira, Dilara Gokay, Michael King, Chuhan Zhang, Ignacio Rocco, Aravindh Mahendran, Thomas Albert Keck, Joseph Heyward, Skanda Koppula, Etienne Pot, et al. Scaling 4D Representations. *arXiv preprint arXiv:2412.15212*, 2024. 6
- [13] Honglie Chen, Weidi Xie, Andrea Vedaldi, and Andrew Zisserman. Vggsound: A Large-scale Audio-visual Dataset. In *International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, 2020. 9
- [14] Junlin Chen, Chengcheng Xu, Yangfan Xu, Jian Yang, Jun Li, and Zhiping Shi. Flatten: Video Action Recognition is an Image Classification Task. *arXiv preprint arXiv:2408.09220*, 2024. 4
- [15] Siyi Chen, Minkyu Choi, Zesen Zhao, Kuan Han, Qing Qu, and Zhongming Liu. Unfolding Videos Dynamics Via Taylor Expansion. *arXiv preprint arXiv:2409.02371*, 2024. 9
- [16] Daniel Cores, Michael Dorkenwald, Manuel Mucientes, Cees GM Snoek, and Yuki M Asano. Tvbench: Redesigning video-language evaluation. *Arxiv*, 2024. 9, 10
- [17] Dima Damen, Hazel Doughty, Giovanni Maria Farinella, Sanja Fidler, Antonino Furnari, Evangelos Kazakos, Davide Moltisanti, Jonathan Munro, Toby Perrett, Will Price, et al. Scaling Egocentric Vision: The EPIC-Kitchens Dataset. In *European Conference on Computer Vision (ECCV)*, 2018. 2, 5, 6, 10
- [18] Timothée Darcet, Maxime Oquab, Julien Mairal, and Piotr Bojanowski. Vision Transformers Need Registers. *arXiv preprint arXiv:2309.16588*, 2023. 4
- [19] Ishan Dave, Rohit Gupta, Mamshad Nayeem Rizve, and Mubarak Shah. TCLR: Temporal Contrastive Learning for Video Representation. *Computer Vision and Image Understanding*, 2022. 7, 9, 27
- [20] Ishan Rajendrakumar Dave, Mamshad Nayeem Rizve, Chen Chen, and Mubarak Shah. Timebalance: Temporally-invariant and Temporally-distinctive Video Representations for Semi-supervised Action Recognition. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2023. 9

- [21] Debidatta Dwibedi, Yusuf Aytar, Jonathan Tompson, Pierre Sermanet, and Andrew Zisserman. Temporal Cycle-Consistency Learning. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2019. 9
- [22] David Fan, Shengbang Tong, Jiachen Zhu, Koustuv Sinha, Zhuang Liu, Xinlei Chen, Michael Rabbat, Nicolas Ballas, Yann LeCun, Amir Bar, et al. Scaling language-free visual representation learning. *arXiv preprint arXiv:2504.01017*, 2025. 9
- [23] Basura Fernando, Efstratios Gavves, José Oramas, Amir Ghodrati, and Tinne Tuytelaars. Rank Pooling for Action Recognition. *IEEE Transactions on Pattern Analysis and Machine Intelligence (TPAMI)*, 2016. 9
- [24] Chaoyou Fu, Yuhan Dai, Yondong Luo, Lei Li, Shuhuai Ren, Renrui Zhang, Zihan Wang, Chenyu Zhou, Yunhang Shen, Mengdan Zhang, et al. Video-MME: The First-Ever Comprehensive Evaluation Benchmark of Multi-modal LLMs in Video Analysis. *arXiv preprint arXiv:2405.21075*, 2024. 9
- [25] Jort F Gemmeke, Daniel PW Ellis, Dylan Freedman, Aren Jansen, Wade Lawrence, R Channing Moore, Manoj Plakal, and Marvin Ritter. Audio Set: An Ontology and Human-Labeled Dataset for Audio Events. In *International Conference on Acoustics, Speech and Signal Processing (ICASSP)*, 2017. 9
- [26] Amir Ghodrati, Efstratios Gavves, and Cees GM Snoek. Video Time: Properties, Encoders and Evaluation. *arXiv preprint arXiv:1807.06980*, 2018. 9
- [27] Ross Goroshin, Michael F Mathieu, and Yann LeCun. Learning to Linearize Under Uncertainty. *Advances in neural information processing systems*, 2015. 9
- [28] Raghav Goyal, Samira Ebrahimi Kahou, Vincent Michalski, Joanna Materzynska, Susanne Westphal, Heuna Kim, Valentin Haenel, Ingo Fruend, Peter Yianilos, Moritz Mueller-Freitag, et al. The "Something Something" Video Database for Learning and Evaluating Visual Common Sense. In *International Conference on Computer Vision (ICCV)*, 2017. 1, 2, 5, 6, 8, 9, 10
- [29] Kristen Grauman, Andrew Westbury, Eugene Byrne, Zachary Chavis, Antonino Furnari, Rohit Girdhar, Jackson Hamburger, Hao Jiang, Miao Liu, Xingyu Liu, et al. Ego4d: Around the World in 3,000 Hours of Egocentric Video. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2022. 9
- [30] Madeleine Grunde-McLaughlin, Ranjay Krishna, and Maneesh Agrawala. Agqa: A Benchmark for Compositional Spatio-temporal Reasoning. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2021. 9
- [31] Anne Harrington, Vasha DuTell, Ayush Tewari, Mark Hamilton, Simon Stent, Ruth Rosenholtz, and William T Freeman. Exploring Perceptual Straightness in Learned Visual Representations. In *International Conference on Learning Representations (ICLR)*, 2022. 9
- [32] Olivier J Hénaff, Robbe LT Goris, and Eero P Simoncelli. Perceptual Straightening of Natural Videos. *Nature neuroscience*, 2019. 2, 3, 9, 10
- [33] Olivier J Hénaff, Yoon Bai, Julie A Charlton, Ian Nauhaus, Eero P Simoncelli, and Robbe LT Goris. Primary Visual Cortex Straightens Natural Video Trajectories. *Nature communications*, 2021. 9
- [34] Dan Hendrycks and Kevin Gimpel. Gaussian Error Linear Units (GELUs). *arXiv preprint arXiv:1606.08415*, 2016. 4
- [35] De-An Huang, Vignesh Ramanathan, Dhruv Mahajan, Lorenzo Torresani, Manohar Paluri, Li Fei-Fei, and Juan Carlos Niebles. What Makes a Video a Video: Analyzing Temporal Information in Video Understanding Models and Datasets. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2018. 1, 5, 6, 9, 29
- [36] Allan Jabri, Andrew Owens, and Alexei Efros. Space-Time Correspondence as a Contrastive Random Walk. *Advances in Neural Information Processing Systems (NeurIPS)*, 2020. 9
- [37] Dinesh Jayaraman and Kristen Grauman. Slow and Steady Feature Analysis: Higher Order Temporal Coherence in Video. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2016. 9
- [38] Andrej Karpathy, George Toderici, Sanketh Shetty, Thomas Leung, Rahul Sukthankar, and Li Fei-Fei. Large-scale Video Classification with Convolutional Neural Networks. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2014. 9
- [39] Zeeshan Khan, CV Jawahar, and Makarand Tapaswi. Grounded Video Situation Recognition. *Advances in Neural Information Processing Systems (NeurIPS)*, 2022. 9

- [40] Dahun Kim, Donghyeon Cho, and In So Kweon. Self-supervised Video Representation Learning with Space-Time Cubic Puzzles. In *Association for the Advancement of Artificial Intelligence (AAAI)*, 2019. 9, 27
- [41] Diederik P Kingma and Jimmy Ba. Adam: A Method for Stochastic Optimization. *arXiv preprint arXiv:1412.6980*, 2014. 4
- [42] Hildegard Kuehne, Hueihan Jhuang, Estíbaliz Garrote, Tomaso Poggio, and Thomas Serre. HMDB: A Large Video Database for Human Motion Recognition. In *International Conference on Computer Vision (ICCV)*, 2011. 2, 8, 9, 10
- [43] Hyeonmin Lee, Jin-Young Kim, Kyungjune Baek, Jihwan Kim, Hyojun Go, Seongsu Ha, Seokjin Han, Jiho Jang, Raehyuk Jung, Daewoo Kim, et al. TWLV-I: Analysis and Insights from Holistic Evaluation on Video Foundation Models. *arXiv preprint arXiv:2408.11318*, 2024. 27
- [44] Jie Lei, Tamara L Berg, and Mohit Bansal. Revealing Single Frame Bias for Video-and-Language Learning. *arXiv:2206.03428*, 2022. 1, 9
- [45] Junnan Li, Dongxu Li, Silvio Savarese, and Steven Hoi. Blip-2: Bootstrapping Language-Image Pre-training with Frozen Image Encoders and Large Language Models. In *International Conference on Machine Learning (ICML)*, 2023. 7
- [46] KunChang Li, Yinan He, Yi Wang, Yizhuo Li, Wenhai Wang, Ping Luo, Yali Wang, Limin Wang, and Yu Qiao. Videochat: Chat-centric Video Understanding. *arXiv preprint arXiv:2305.06355*, 2023. 9
- [47] Kunchang Li, Yali Wang, Yinan He, Yizhuo Li, Yi Wang, Yi Liu, Zun Wang, Jilan Xu, Guo Chen, Ping Luo, et al. Mvbench: A Comprehensive Multi-modal Video Understanding Benchmark. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2024. 9
- [48] Xinhao Li, Yuhan Zhu, and Limin Wang. Zero2v: Zero-cost Adaptation of Pre-trained Transformers from Image to Video. In *European Conference on Computer Vision (ECCV)*, 2024. 10
- [49] Yingwei Li, Yi Li, and Nuno Vasconcelos. Resound: Towards Action Recognition Without Representation Bias. In *European Conference on Computer Vision (ECCV)*, 2018. 1, 9
- [50] Zhiqiu Lin, Jin Sun, Abe Davis, and Noah Snavely. Visual Chirality. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2020. 5
- [51] Ze Liu, Jia Ning, Yue Cao, Yixuan Wei, Zheng Zhang, Stephen Lin, and Han Hu. Video Swin Transformer. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2022. 9
- [52] 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 *International Conference on Computer Vision (ICCV)*, 2019. 9
- [53] Ishan Misra, C Lawrence Zitnick, and Martial Hebert. Shuffle and Learn: Unsupervised Learning Using Temporal Order Verification. In *European Conference on Computer Vision (ECCV)*, 2016. 5
- [54] Liliane Momeni, Mathilde Caron, Arsha Nagrani, Andrew Zisserman, and Cordelia Schmid. Verbs in Action: Improving Verb Understanding in Video-Language Models. In *International Conference on Computer Vision (ICCV)*, 2023. 9
- [55] Julie Xueyan Niu, Cristina Savin, and Eero Simoncelli. Learning Predictable and Robust Neural Representations by Straightening Image Sequences. *Advances in Neural Information Processing Systems*, 2024. 9
- [56] Maxime Oquab, Timothée Darcet, Théo Moutakanni, Huy Vo, Marc Szafraniec, Vasil Khalidov, Pierre Fernandez, Daniel Haziza, Francisco Massa, Alaaeldin El-Nouby, et al. Dinov2: Learning Robust Visual Features Without Supervision. *arXiv preprint arXiv:2304.07193*, 2023. 4, 7, 8
- [57] Junting Pan, Ziyi Lin, Xiatian Zhu, Jing Shao, and Hongsheng Li. ST-Adapter: Parameter-Efficient Image-to-Video Transfer Learning. *Advances in Neural Information Processing Systems*, 2022. 2, 3, 7, 10
- [58] Jungin Park, Jiyoung Lee, and Kwanghoon Sohn. Dual-path Adaptation from Image to Video Transformers. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2023. 3

- [59] Viorica Patraucean, Lucas Smaira, Ankush Gupta, Adria Recasens, Larisa Markeeva, Dylan Banarse, Skanda Koppula, Mateusz Malinowski, Yi Yang, Carl Doersch, et al. Perception Test: A Diagnostic Benchmark for Multimodal Video Models. *Advances in Neural Information Processing Systems (NeurIPS)*, 2024. 1, 9
- [60] Lyndsey C Pickup, Zheng Pan, Donglai Wei, YiChang Shih, Changshui Zhang, Andrew Zisserman, Bernhard Scholkopf, and William T Freeman. Seeing the Arrow of Time. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2014. 5
- [61] Thinesh Thiyakesan Ponbagavathi and Alina Roitberg. Order Matters: On Parameter-Efficient Image-to-Video Probing for Recognizing Nearly Symmetric Actions. *arXiv preprint arXiv:2503.24298*, 2025. 5, 10
- [62] Will Price and Dima Damen. Retro-Actions: Learning ‘close’ by Time-reversing ‘open’ videos. In *International Conference on Computer Vision Workshops (ICCVW)*, 2019. 5, 9, 10
- [63] Rui Qian, Tianjian Meng, Boqing Gong, Ming-Hsuan Yang, Huisheng Wang, Serge Belongie, and Yin Cui. Spatiotemporal Contrastive Video Representation Learning. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2021. 9
- [64] Zhiwu Qing, Shiwei Zhang, Ziyuan Huang, Yingya Zhang, Changxin Gao, Deli Zhao, and Nong Sang. Disentangling Spatial and Temporal Learning for Efficient Image-to-Video Transfer Learning. In *International Conference on Computer Vision (ICCV)*, 2023. 2, 7, 10
- [65] 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 *International Conference on Machine Learning (ICML)*, 2021. 7, 8
- [66] Jathushan Rajasegaran, Ilija Radosavovic, Rahul Ravishankar, Yossi Gandelsman, Christoph Feichtenhofer, and Jitendra Malik. An Empirical Study of Autoregressive Pre-training from Videos. *arXiv preprint arXiv:2501.05453*, 2025. 10
- [67] Hanoona Rasheed, Muhammad Uzair Khattak, Muhammad Maaz, Salman Khan, and Fahad Shahbaz Khan. Fine-tuned CLIP Models are Efficient Video Learners. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2023. 10
- [68] Adria Recasens, Pauline Luc, Jean-Baptiste Alayrac, Luyu Wang, Florian Strub, Corentin Tallec, Mateusz Malinowski, Viorica Pătrăucean, Florent Althé, Michal Valko, et al. Broaden Your Views for Self-supervised Video Learning. In *International Conference on Computer Vision (ICCV)*, 2021. 9
- [69] Mohammadreza Salehi, Michael Dorkenwald, Fida Mohammad Thoker, Efstratios Gavves, Cees GM Snoek, and Yuki M Asano. SIGMA: Sinkhorn-Guided Masked Video Modeling. In *European Conference on Computer Vision (ECCV)*, 2024. 7, 9
- [70] Darshana Saravanan, Darshan Singh, Varun Gupta, Zeeshan Khan, Vineet Gandhi, and Makarand Tapaswi. VELOCITY: Can Video-Language Models Bind Semantic Concepts Through Time? *arXiv preprint arXiv:2406.10889*, 2024. 9
- [71] Madeline C Schiappa, Yogesh S Rawat, and Mubarak Shah. Self-supervised Learning for Videos: A Survey. *ACM Computing Surveys*, 2023. 9
- [72] Laura Sevilla-Lara, Shengxin Zha, Zhicheng Yan, Vedanuj Goswami, Matt Feiszli, and Lorenzo Torresani. Only Time Can Tell: Discovering Temporal Data for Temporal Modeling. In *Winter Conference on Applications of Computer Vision (WACV)*, 2021. 5, 6, 9, 29
- [73] Dian Shao, Yue Zhao, Bo Dai, and Dahua Lin. Finegym: A Hierarchical Video Dataset for Fine-grained Action Understanding. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2020. 1, 5
- [74] Gunnar A Sigurdsson, Gül Varol, Xiaolong Wang, Ali Farhadi, Ivan Laptev, and Abhinav Gupta. Hollywood in Homes: Crowdsourcing Data Collection for Activity Understanding. In *European Conference on Computer Vision (ECCV)*, 2016. 2, 5, 6, 10
- [75] K Soomro. UCF101: A Dataset of 101 Human Actions Classes from Videos in the Wild. *arXiv:1212.0402*, 2012. 1, 2, 8, 9, 10
- [76] Xinyu Sun, Peihao Chen, Liangwei Chen, Changhao Li, Thomas H Li, Minghui Tan, and Chuang Gan. Masked Motion Encoding for Self-Supervised Video Representation Learning. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2023. 7

- [77] Fida Mohammad Thoker, Hazel Doughty, and Cees GM Snoek. Tubelet-contrastive Self-supervision for Video-efficient Generalization. In *International Conference on Computer Vision (ICCV)*, 2023. 7, 9, 29
- [78] Fida Mohammad Thoker, Letian Jiang, Chen Zhao, and Bernard Ghanem. SMILE: Infusing Spatial and Motion Semantics in Masked Video Learning. *arXiv preprint arXiv:2504.00527*, 2025. 9
- [79] Zhan Tong, Yibing Song, Jue Wang, and Limin Wang. VideoMAE: Masked Autoencoders are Data-Efficient Learners for Self-Supervised Video Pre-Training. *Advances in Neural Information Processing Systems (NeurIPS)*, 2022. 2, 3, 7, 9, 10, 27
- [80] Michael Tschannen, Alexey Gritsenko, Xiao Wang, Muhammad Ferjad Naeem, Ibrahim Alabdulmohsin, Nikhil Parthasarathy, Talfan Evans, Lucas Beyer, Ye Xia, Basil Mustafa, et al. Siglip 2: Multilingual Vision-language Encoders with Improved Semantic Understanding, Localization, and Dense Features. *arXiv preprint arXiv:2502.14786*, 2025. 7, 8
- [81] Vadim Tschernezki, Ahmad Darkhalil, Zhifan Zhu, David Fouhey, Iro Larina, Diane Larlus, Dima Damen, and Andrea Vedaldi. EPIC Fields: Marrying 3D Geometry and Video Understanding. In *Advances in Neural Information Processing Systems (NeurIPS)*, 2023. 9
- [82] Shuyuan Tu, Qi Dai, Zuxuan Wu, Zhi-Qi Cheng, Han Hu, and Yu-Gang Jiang. Implicit Temporal Modeling with Learnable Alignment for Video Recognition. In *International Conference on Computer Vision (ICCV)*, 2023. 10
- [83] Jan PH Van Santen and George Sperling. Temporal Covariance Model of Human Motion Perception. *Journal of the Optical Society of America A*, 1984. 9
- [84] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Lukasz Kaiser, and Illia Polosukhin. Attention Is All You Need. *Advances in Neural Information Processing Systems (NeurIPS)*, 2017. 9
- [85] Shashanka Venkataramanan, Mamshad Nayeem Rizve, João Carreira, Yuki M Asano, and Yannis Avrithis. Is ImageNet Worth 1 Video? Learning Strong Image Encoders from 1 Long Unlabelled Video. *arXiv preprint arXiv:2310.08584*, 2023. 8, 9
- [86] Guangting Wang, Yizhou Zhou, Chong Luo, Wenxuan Xie, Wenjun Zeng, and Zhiwei Xiong. Unsupervised Visual Representation Learning by Tracking Patches in Video. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2021. 7
- [87] Limin Wang, Bingkun Huang, Zhiyu Zhao, Zhan Tong, Yinan He, Yi Wang, Yali Wang, and Yu Qiao. VideoMAE v2: Scaling Video Masked Autoencoders with Dual Masking. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2023. 7, 9, 27
- [88] Yi Wang, Kunchang Li, Xinhao Li, Jiashuo Yu, Yinan He, Guo Chen, Baoqi Pei, Rongkun Zheng, Zun Wang, Yansong Shi, et al. Internvideo2: Scaling Foundation Models for Multimodal Video Understanding. In *European Conference on Computer Vision (ECCV)*, 2024. 27
- [89] Yi Wang, Xinhao Li, Ziang Yan, Yinan He, Jiashuo Yu, Xiangyu Zeng, Chenting Wang, Changlian Ma, Haian Huang, Jianfei Gao, et al. InternVideo2. 5: Empowering Video MLLMs with Long and Rich Context Modeling. *arXiv preprint arXiv:2501.12386*, 2025. 2, 7, 27
- [90] Andrew B Watson and Albert J Ahumada Jr. Model of Human Visual-motion Sensing. *Journal of the optical Society of America A*, 1985. 9
- [91] Donglai Wei, Joseph J Lim, Andrew Zisserman, and William T Freeman. Learning and Using the Arrow of Time. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2018. 5, 9
- [92] Dirk Weissenborn, Oscar Täckström, and Jakob Uszkoreit. Scaling Autoregressive Video Models. *arXiv preprint arXiv:1906.02634*, 2019. 10
- [93] Junbin Xiao, Xindi Shang, Angela Yao, and Tat-Seng Chua. Next-qa: Next Phase of Question-Answering to Explaining Temporal Actions. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2021. 9
- [94] Dejing Xu, Jun Xiao, Zhou Zhao, Jian Shao, Di Xie, and Yueting Zhuang. Self-supervised Spatiotemporal Learning Via Video Clip Order Prediction. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2019. 9
- [95] Zihui Xue, Mi Luo, and Kristen Grauman. Seeing the Arrow of Time in Large Multimodal Models. *arXiv preprint arXiv:2506.03340*, 2025. 9

- [96] Bin Yan, Houwen Peng, Jianlong Fu, Dong Wang, and Huchuan Lu. Learning Spatio-temporal Transformer for Visual Tracking. In *International Conference on Computer Vision (ICCV)*, 2021. 5
- [97] Charig Yang, Weidi Xie, and Andrew Zisserman. Made to Order: Discovering Monotonic Temporal Changes Via Self-supervised Video Ordering. *arXiv preprint arXiv:2404.16828*, 2024. 9
- [98] Di Yang, Yaohui Wang, Quan Kong, Antitza Dantcheva, Lorenzo Garattoni, Gianpiero Francesca, and François Brémond. Self-supervised Video Representation Learning Via Latent Time Navigation. In *Association for the Advancement of Artificial Intelligence (AAAI)*, 2023. 9
- [99] Kexin Yi, Chuang Gan, Yunzhu Li, Pushmeet Kohli, Jiajun Wu, Antonio Torralba, and Joshua B Tenenbaum. Clevrer: Collision Events for Video Representation and Reasoning. *arXiv preprint arXiv:1910.01442*, 2019. 9
- [100] Xueyang Yu, Xinlei Chen, and Yossi Gandelsman. Learning Video Representations Without Natural Videos. *arXiv preprint arXiv:2410.24213*, 2024. 29
- [101] Zhou Yu, Lixiang Zheng, Zhou Zhao, Fei Wu, Jianping Fan, Kui Ren, and Jun Yu. ANetQA: A Large-scale Benchmark for Fine-grained Compositional Reasoning Over Untrimmed Videos. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2023. 9
- [102] Sukmin Yun, Jaehyung Kim, Dongyoon Han, Hwanjun Song, Jung-Woo Ha, and Jinwoo Shin. Time is Matter: Temporal Self-supervision for Video Transformers. *arXiv preprint arXiv:2207.09067*, 2022. 5
- [103] Xiaohua Zhai, Basil Mustafa, Alexander Kolesnikov, and Lucas Beyer. Sigmoid Loss for Language Image Pre-Training. In *International Conference on Computer Vision (ICCV)*, 2023. 7
- [104] Chuhan Zhang, Ankush Gupta, and Andrew Zisserman. Temporal Query Networks for Fine-Grained Video Understanding. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2021. 5
- [105] Heng Zhang, Daqing Liu, Qi Zheng, and Bing Su. Modeling Video as Stochastic Processes for Fine-grained Video Representation Learning. In *Conference on Computer Vision and Pattern Recognition (CVPR)*, 2023. 9
- [106] Yuanhan Zhang, Jinming Wu, Wei Li, Bo Li, Zejun Ma, Ziwei Liu, and Chunyuan Li. LLaVA-Video: Video Instruction Tuning With Synthetic Data. *Transactions on Machine Learning Research*, 2025. 9
- [107] Jinghao Zhou, Chen Wei, Huiyu Wang, Wei Shen, Cihang Xie, Alan Yuille, and Tao Kong. ibot: Image BERT Pre-training with Online Tokenizer. *arXiv preprint arXiv:2111.07832*, 2021. 8
- [108] Orr Zohar, Xiaohan Wang, Yann Dubois, Nikhil Mehta, Tong Xiao, Philippe Hansen-Estruch, Licheng Yu, Xiaofang Wang, Felix Juefei-Xu, Ning Zhang, et al. Apollo: An Exploration of Video Understanding in Large Multimodal Models. *arXiv preprint arXiv:2412.10360*, 2024. 1, 9

NeurIPS Paper Checklist

1. Claims

Question: Do the main claims made in the abstract and introduction accurately reflect the paper's contributions and scope?

Answer: [Yes]

Justification: The abstract and introduction reflect the main contribution which is (i) LiFT; an adaptation of image features to obtain demonstrably time-aware video descriptors (Section 2); (ii) presenting a benchmark to probe time-awareness of video descriptors (Section 3); and (iii) utility of LiFT in improving existing video models for action recognition (Section 4).

Guidelines:

- The answer NA means that the abstract and introduction do not include the claims made in the paper.
- The abstract and/or introduction should clearly state the claims made, including the contributions made in the paper and important assumptions and limitations. A No or NA answer to this question will not be perceived well by the reviewers.
- The claims made should match theoretical and experimental results, and reflect how much the results can be expected to generalize to other settings.
- It is fine to include aspirational goals as motivation as long as it is clear that these goals are not attained by the paper.

2. Limitations

Question: Does the paper discuss the limitations of the work performed by the authors?

Answer: [Yes]

Justification: The scope of the work and some limitations are discussed in Section 6, *e.g.*, in the lack of sensitivity of LiFT to certain changes.

Guidelines:

- The answer NA means that the paper has no limitation while the answer No means that the paper has limitations, but those are not discussed in the paper.
- The authors are encouraged to create a separate "Limitations" section in their paper.
- The paper should point out any strong assumptions and how robust the results are to violations of these assumptions (e.g., independence assumptions, noiseless settings, model well-specification, asymptotic approximations only holding locally). The authors should reflect on how these assumptions might be violated in practice and what the implications would be.
- The authors should reflect on the scope of the claims made, e.g., if the approach was only tested on a few datasets or with a few runs. In general, empirical results often depend on implicit assumptions, which should be articulated.
- The authors should reflect on the factors that influence the performance of the approach. For example, a facial recognition algorithm may perform poorly when image resolution is low or images are taken in low lighting. Or a speech-to-text system might not be used reliably to provide closed captions for online lectures because it fails to handle technical jargon.
- The authors should discuss the computational efficiency of the proposed algorithms and how they scale with dataset size.
- If applicable, the authors should discuss possible limitations of their approach to address problems of privacy and fairness.
- While the authors might fear that complete honesty about limitations might be used by reviewers as grounds for rejection, a worse outcome might be that reviewers discover limitations that aren't acknowledged in the paper. The authors should use their best judgment and recognize that individual actions in favor of transparency play an important role in developing norms that preserve the integrity of the community. Reviewers will be specifically instructed to not penalize honesty concerning limitations.

3. Theory assumptions and proofs

Question: For each theoretical result, does the paper provide the full set of assumptions and a complete (and correct) proof?

Answer: [NA]

Justification: This paper does not include any theoretical results. The formulation of the LiFT method is empirical and uses standard notation and terminology.

Guidelines:

- The answer NA means that the paper does not include theoretical results.
- All the theorems, formulas, and proofs in the paper should be numbered and cross-referenced.
- All assumptions should be clearly stated or referenced in the statement of any theorems.
- The proofs can either appear in the main paper or the supplemental material, but if they appear in the supplemental material, the authors are encouraged to provide a short proof sketch to provide intuition.
- Inversely, any informal proof provided in the core of the paper should be complemented by formal proofs provided in appendix or supplemental material.
- Theorems and Lemmas that the proof relies upon should be properly referenced.

4. Experimental result reproducibility

Question: Does the paper fully disclose all the information needed to reproduce the main experimental results of the paper to the extent that it affects the main claims and/or conclusions of the paper (regardless of whether the code and data are provided or not)?

Answer: [Yes]

Justification: We provide all the details to construct the CiA dataset based on existing public datasets in Section 3. We also provide implementation details in Section 2 to reproduce LiFT video descriptor. In Section 4, we provide details for each of the main experiments conducted. The model is trained and evaluated on public datasets.

Guidelines:

- The answer NA means that the paper does not include experiments.
- If the paper includes experiments, a No answer to this question will not be perceived well by the reviewers: Making the paper reproducible is important, regardless of whether the code and data are provided or not.
- If the contribution is a dataset and/or model, the authors should describe the steps taken to make their results reproducible or verifiable.
- Depending on the contribution, reproducibility can be accomplished in various ways. For example, if the contribution is a novel architecture, describing the architecture fully might suffice, or if the contribution is a specific model and empirical evaluation, it may be necessary to either make it possible for others to replicate the model with the same dataset, or provide access to the model. In general, releasing code and data is often one good way to accomplish this, but reproducibility can also be provided via detailed instructions for how to replicate the results, access to a hosted model (e.g., in the case of a large language model), releasing of a model checkpoint, or other means that are appropriate to the research performed.
- While NeurIPS does not require releasing code, the conference does require all submissions to provide some reasonable avenue for reproducibility, which may depend on the nature of the contribution. For example
 - (a) If the contribution is primarily a new algorithm, the paper should make it clear how to reproduce that algorithm.
 - (b) If the contribution is primarily a new model architecture, the paper should describe the architecture clearly and fully.
 - (c) If the contribution is a new model (e.g., a large language model), then there should either be a way to access this model for reproducing the results or a way to reproduce the model (e.g., with an open-source dataset or instructions for how to construct the dataset).

- (d) We recognize that reproducibility may be tricky in some cases, in which case authors are welcome to describe the particular way they provide for reproducibility. In the case of closed-source models, it may be that access to the model is limited in some way (e.g., to registered users), but it should be possible for other researchers to have some path to reproducing or verifying the results.

5. Open access to data and code

Question: Does the paper provide open access to the data and code, with sufficient instructions to faithfully reproduce the main experimental results, as described in supplemental material?

Answer: [Yes]

Justification: The proposed CiA meta-dataset is based on three public datasets with instructions for construction provided. CiA meta-dataset will be publicly released. We will also provide it along with the supplementary material. Likewise, the code and trained model will be released as open-source.

Guidelines:

- The answer NA means that paper does not include experiments requiring code.
- Please see the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- While we encourage the release of code and data, we understand that this might not be possible, so “No” is an acceptable answer. Papers cannot be rejected simply for not including code, unless this is central to the contribution (e.g., for a new open-source benchmark).
- The instructions should contain the exact command and environment needed to run to reproduce the results. See the NeurIPS code and data submission guidelines (<https://nips.cc/public/guides/CodeSubmissionPolicy>) for more details.
- The authors should provide instructions on data access and preparation, including how to access the raw data, preprocessed data, intermediate data, and generated data, etc.
- The authors should provide scripts to reproduce all experimental results for the new proposed method and baselines. If only a subset of experiments are reproducible, they should state which ones are omitted from the script and why.
- At submission time, to preserve anonymity, the authors should release anonymized versions (if applicable).
- Providing as much information as possible in supplemental material (appended to the paper) is recommended, but including URLs to data and code is permitted.

6. Experimental setting/details

Question: Does the paper specify all the training and test details (e.g., data splits, hyper-parameters, how they were chosen, type of optimizer, etc.) necessary to understand the results?

Answer: [Yes]

Justification: The paper uses data splits based on the original public datasets. All the details and hyper-parameters are stated in Section 2, Section 4 and supplementary material.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The experimental setting should be presented in the core of the paper to a level of detail that is necessary to appreciate the results and make sense of them.
- The full details can be provided either with the code, in appendix, or as supplemental material.

7. Experiment statistical significance

Question: Does the paper report error bars suitably and correctly defined or other appropriate information about the statistical significance of the experiments?

Answer: [No]

Justification: While we do not report error bars in the main paper, we will report error bars for the key experiments in the Supplemental.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The authors should answer "Yes" if the results are accompanied by error bars, confidence intervals, or statistical significance tests, at least for the experiments that support the main claims of the paper.
- The factors of variability that the error bars are capturing should be clearly stated (for example, train/test split, initialization, random drawing of some parameter, or overall run with given experimental conditions).
- The method for calculating the error bars should be explained (closed form formula, call to a library function, bootstrap, etc.)
- The assumptions made should be given (e.g., Normally distributed errors).
- It should be clear whether the error bar is the standard deviation or the standard error of the mean.
- It is OK to report 1-sigma error bars, but one should state it. The authors should preferably report a 2-sigma error bar than state that they have a 96% CI, if the hypothesis of Normality of errors is not verified.
- For asymmetric distributions, the authors should be careful not to show in tables or figures symmetric error bars that would yield results that are out of range (e.g. negative error rates).
- If error bars are reported in tables or plots, The authors should explain in the text how they were calculated and reference the corresponding figures or tables in the text.

8. Experiments compute resources

Question: For each experiment, does the paper provide sufficient information on the computer resources (type of compute workers, memory, time of execution) needed to reproduce the experiments?

Answer: [Yes]

Justification: We will state the compute resources used in the Supplemental. Briefly, we use various consumer-grade GPUs (e.g., Nvidia RTX A4000, Tesla P40, Quadro RTX 8000, NVIDIA RTX A6000). Since per-frame features are pre-computed, training and evaluation is very efficient.

Guidelines:

- The answer NA means that the paper does not include experiments.
- The paper should indicate the type of compute workers CPU or GPU, internal cluster, or cloud provider, including relevant memory and storage.
- The paper should provide the amount of compute required for each of the individual experimental runs as well as estimate the total compute.
- The paper should disclose whether the full research project required more compute than the experiments reported in the paper (e.g., preliminary or failed experiments that didn't make it into the paper).

9. Code of ethics

Question: Does the research conducted in the paper conform, in every respect, with the NeurIPS Code of Ethics <https://neurips.cc/public/EthicsGuidelines>?

Answer: [Yes]

Justification: This research does not use human subjects. The data used is sourced from publicly available datasets for action recognition.

Guidelines:

- The answer NA means that the authors have not reviewed the NeurIPS Code of Ethics.
- If the authors answer No, they should explain the special circumstances that require a deviation from the Code of Ethics.

- The authors should make sure to preserve anonymity (e.g., if there is a special consideration due to laws or regulations in their jurisdiction).

10. **Broader impacts**

Question: Does the paper discuss both potential positive societal impacts and negative societal impacts of the work performed?

Answer: [Yes]

Justification: This work concerns foundational research on probing and improving representations of video data. While applications of human action recognition can be applied for surveillance, there is nothing substantial that this research adds to already available video recognition methods for such use cases.

Guidelines:

- The answer NA means that there is no societal impact of the work performed.
- If the authors answer NA or No, they should explain why their work has no societal impact or why the paper does not address societal impact.
- Examples of negative societal impacts include potential malicious or unintended uses (e.g., disinformation, generating fake profiles, surveillance), fairness considerations (e.g., deployment of technologies that could make decisions that unfairly impact specific groups), privacy considerations, and security considerations.
- The conference expects that many papers will be foundational research and not tied to particular applications, let alone deployments. However, if there is a direct path to any negative applications, the authors should point it out. For example, it is legitimate to point out that an improvement in the quality of generative models could be used to generate deepfakes for disinformation. On the other hand, it is not needed to point out that a generic algorithm for optimizing neural networks could enable people to train models that generate Deepfakes faster.
- The authors should consider possible harms that could arise when the technology is being used as intended and functioning correctly, harms that could arise when the technology is being used as intended but gives incorrect results, and harms following from (intentional or unintentional) misuse of the technology.
- If there are negative societal impacts, the authors could also discuss possible mitigation strategies (e.g., gated release of models, providing defenses in addition to attacks, mechanisms for monitoring misuse, mechanisms to monitor how a system learns from feedback over time, improving the efficiency and accessibility of ML).

11. **Safeguards**

Question: Does the paper describe safeguards that have been put in place for responsible release of data or models that have a high risk for misuse (e.g., pretrained language models, image generators, or scraped datasets)?

Answer: [NA]

Justification: The data used in this research is sourced from publicly available action recognition datasets.

Guidelines:

- The answer NA means that the paper poses no such risks.
- Released models that have a high risk for misuse or dual-use should be released with necessary safeguards to allow for controlled use of the model, for example by requiring that users adhere to usage guidelines or restrictions to access the model or implementing safety filters.
- Datasets that have been scraped from the Internet could pose safety risks. The authors should describe how they avoided releasing unsafe images.
- We recognize that providing effective safeguards is challenging, and many papers do not require this, but we encourage authors to take this into account and make a best faith effort.

12. **Licenses for existing assets**

Question: Are the creators or original owners of assets (e.g., code, data, models), used in the paper, properly credited and are the license and terms of use explicitly mentioned and properly respected?

Answer: [No]

Justification: The paper cites all the assets (datasets, models) used in this research. However, the specific licenses and terms of use for these assets are not explicitly mentioned in the paper text or appendix.

Guidelines:

- The answer NA means that the paper does not use existing assets.
- The authors should cite the original paper that produced the code package or dataset.
- The authors should state which version of the asset is used and, if possible, include a URL.
- The name of the license (e.g., CC-BY 4.0) should be included for each asset.
- For scraped data from a particular source (e.g., website), the copyright and terms of service of that source should be provided.
- If assets are released, the license, copyright information, and terms of use in the package should be provided. For popular datasets, paperswithcode.com/datasets has curated licenses for some datasets. Their licensing guide can help determine the license of a dataset.
- For existing datasets that are re-packaged, both the original license and the license of the derived asset (if it has changed) should be provided.
- If this information is not available online, the authors are encouraged to reach out to the asset's creators.

13. New assets

Question: Are new assets introduced in the paper well documented and is the documentation provided alongside the assets?

Answer: [No]

Justification: We do not release the code and data at the time of submission. But we shall do it in the near future with all required documentation. We will provide the CiA meta-dataset along with the supplementary material.

Guidelines:

- The answer NA means that the paper does not release new assets.
- Researchers should communicate the details of the dataset/code/model as part of their submissions via structured templates. This includes details about training, license, limitations, etc.
- The paper should discuss whether and how consent was obtained from people whose asset is used.
- At submission time, remember to anonymize your assets (if applicable). You can either create an anonymized URL or include an anonymized zip file.

14. Crowdsourcing and research with human subjects

Question: For crowdsourcing experiments and research with human subjects, does the paper include the full text of instructions given to participants and screenshots, if applicable, as well as details about compensation (if any)?

Answer: [NA]

Justification: This paper does not use any crowdsourcing experiments or research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Including this information in the supplemental material is fine, but if the main contribution of the paper involves human subjects, then as much detail as possible should be included in the main paper.

- According to the NeurIPS Code of Ethics, workers involved in data collection, curation, or other labor should be paid at least the minimum wage in the country of the data collector.

15. Institutional review board (IRB) approvals or equivalent for research with human subjects

Question: Does the paper describe potential risks incurred by study participants, whether such risks were disclosed to the subjects, and whether Institutional Review Board (IRB) approvals (or an equivalent approval/review based on the requirements of your country or institution) were obtained?

Answer: [NA]

Justification: This paper does not use any crowdsourcing experiments or research with human subjects.

Guidelines:

- The answer NA means that the paper does not involve crowdsourcing nor research with human subjects.
- Depending on the country in which research is conducted, IRB approval (or equivalent) may be required for any human subjects research. If you obtained IRB approval, you should clearly state this in the paper.
- We recognize that the procedures for this may vary significantly between institutions and locations, and we expect authors to adhere to the NeurIPS Code of Ethics and the guidelines for their institution.
- For initial submissions, do not include any information that would break anonymity (if applicable), such as the institution conducting the review.

16. Declaration of LLM usage

Question: Does the paper describe the usage of LLMs if it is an important, original, or non-standard component of the core methods in this research? Note that if the LLM is used only for writing, editing, or formatting purposes and does not impact the core methodology, scientific rigorousness, or originality of the research, declaration is not required.

Answer: [Yes]

Justification: This paper only uses LLMs for minor data verification (*e.g.*, to gather temporally opposite actions from a list of action verbs). While it is not most fundamental part, we describe it in Section 3.

Guidelines:

- The answer NA means that the core method development in this research does not involve LLMs as any important, original, or non-standard components.
- Please refer to our LLM policy (<https://neurips.cc/Conferences/2025/LLM>) for what should or should not be described.

A Dataset: Chirality in Action (CiA)

Metadata and examples. In Table 6, we show the chiral groups constructed in SSv2. Similarly, we construct 66 chiral groups in EPIC and 28 groups in Charades. In Table 7, we show the number of videos in the chiral splits of each of the three datasets. We attach the chiral splits for all three datasets as part of the Supplemental. We also provide a single CSV file that includes the chiral groups for all three combined data sets. We also show more examples of chiral pairs from each of the three datasets in Fig. 6, Fig. 7 and Fig. 8. In general, since SSv2 has single canonical actions, it is a cleaner test bed for chiral action recognition. EPIC and Charades usually have a more cluttered visual context where cues for chiral recognition are more subtle.

Verb →	Verb ←	Noun (object)
Pulling [something] from left to right	Pulling [something] from right to left	['something']
Pushing [something] from left to right	Pushing [something] from right to left	['something']
Turning the camera left while filming [...]	Turning the camera right while filming [...]	['something']
Approaching [something] with your camera	Moving away from [something] with your camera	['something']
Closing [something]	Opening [something]	['object']
Closing [something]	Opening [something]	['door']
Closing [something]	Opening [something]	['bottle']
Closing [something]	Opening [something]	['book']
Closing [something]	Opening [something]	['purse']
Closing [something]	Opening [something]	['drawer']
Moving [...] and [...] away from each other	Moving [...] and [...] closer to each other	['something']
Moving [something] away from the camera	Moving [something] towards the camera	['something']
Moving [something] down	Moving [something] up	['something']
Putting [something similar to other things ...]	Taking [one of many similar things on the table]	['something']
Turning the camera downwards while filming [...]	Turning the camera upwards while filming [...]	['something']
Folding [something]	Unfolding [something]	['something']

Table 6: **Chiral groups in SSv2.** We construct 16 chiral groups in SSv2 by identifying temporally opposite verbs. Note that “opening vs. closing” is split across different objects representing entirely different actions. Noun “[something]” denotes a placeholder which can include any appropriate object that fits with the action verb.

Dataset	Chiral groups	Total videos		Avg. videos per chiral group		Avg. duration (s)
		Train	Validation	Train	Validation	
SSv2	16	12216	1430	763.5	89.4	3.6
EPIC	66	24101	3108	365.1	47.1	1.6
Charades	28	16018	5498	572.1	196.4	8.6

Table 7: **CiA dataset size.** For each of the constituent datasets, we show the total number of videos in the proposed chiral split and also the average number of videos per chiral group. Note that we train one linear probe for each chiral group.

Time-sensitivity of CiA. In Fig. 9 and Fig. 10, we repeat the experiments to check time-sensitivity (Fig 5 in the main paper) of the CiA benchmark on all three datasets. Our inferences about time-sensitivity hold for all three datasets.

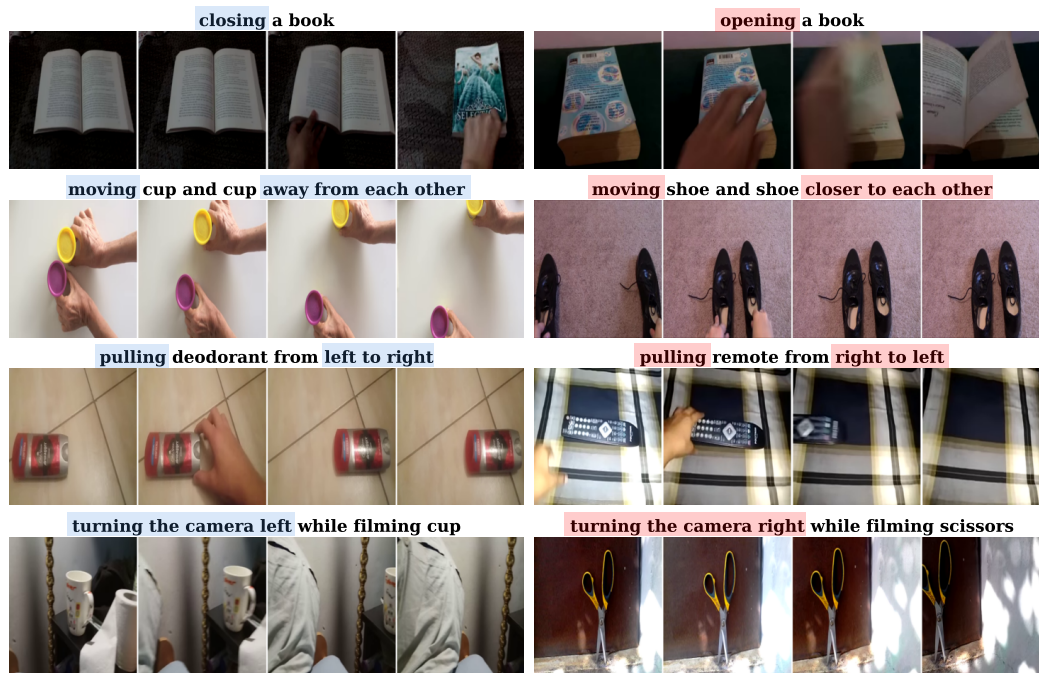


Figure 6: **CiA samples from SSv2.** More examples of chiral pairs from the train set of SSv2. The positive direction actions are marked in blue while the negative direction ones are marker in red.



Figure 7: **CiA samples from EPIC.** More examples of chiral pairs from the train set of EPIC. The positive direction actions are marked in blue while the negative direction ones are marker in red.



Figure 8: **CiA samples from Charades.** More examples of chiral pairs from the train set of Charades. The positive direction actions are marked in blue while the negative direction ones are marker in red.

B Model: Linearized Feature Trajectories (LiFT)

Architecture. A detailed sketch of the architecture is provided in Fig. 11. In LiFT, the Encoder takes in a sequence of features $\{\mathbf{x}_t\}_t$ and outputs two descriptor tokens $\mathbf{z}_s, \mathbf{z}_d$. First, a linear layer projection is applied $\mathbb{R}^D \rightarrow \mathbb{R}^d$. Then, Sinusoidal position encoding is added representing frame index t . The two CLS tokens $\mathbf{e}_s, \mathbf{e}_d$ are initialized randomly. Then, the CLS tokens along with the sequence tokens are passed through a Transformer with $L=4$ blocks and $H=8$ heads each. Each block has a multi-head self-attention (MHSA) layer followed by a FFN layer. Both the layers are preceded by LayerNorm layers. Then, the outputs for the two CLS tokens are projected with a linear layer ($\mathbb{R}^d \rightarrow \mathbb{R}^d$) followed by LayerNorm. This gives the latent descriptors $\mathbf{z}_s$ and $\mathbf{z}_d$.

The decoder takes in $\mathbf{z}_s, \mathbf{z}_d, t$ and outputs $\hat{\mathbf{x}}_t$. First, we construct an intermediate representation for the frame at index t using our linearity constraint in the latent space.

$$\mathbf{z}_t = \mathbf{z}_s + (t/T) \cdot \mathbf{z}_d \quad (7)$$

Then, this is passed to an MLP network with two hidden layers each followed by GeLU activation and LayerNorm. The first hidden layer maps $\mathbb{R}^d \rightarrow \mathbb{R}^{2d}$ and the second layer maps $\mathbb{R}^{2d} \rightarrow \mathbb{R}^{2d}$. This is followed by a linear projection ($\mathbb{R}^{2d} \rightarrow \mathbb{R}^D$) back to the DINOv2 space.

Compute resources. In order to train LiFT, we first compute and store feature vectors for DINOv2 ViT-S/14. This feature computation is run on 4 NVIDIA RTX A4000 16GB GPUs in parallel. It takes about 12 GPU hours to compute features for 250K videos in Kinetics-400. Once features are computed, LiFT is trained on a single consumer-grade GPU (*e.g.*, NVIDIA RTX A4000, Tesla P40, Quadro RTX 8000, NVIDIA RTX A6000). A single training run takes about 15 GPU hours.

Computational cost comparison to video models. LiFT has a total of 30M parameters (22M frozen DINOv2 parameters and 8M trainable parameters) trained on 240K videos. In contrast, native video models, *e.g.*, Video JEPa (smallest variant is ViT-L) has about $10\times$ the parameters (305M) and training samples (2M). Furthermore, during inference, for a clip with $T=16$ frames with resolution 224×224 on the same hardware, LiFT takes 2.3 ms/clip while VJEPa takes 8.5 ms/clip. Thus, overall, LiFT provides a compact, fast, and complementary video embedding when compared to native video models.

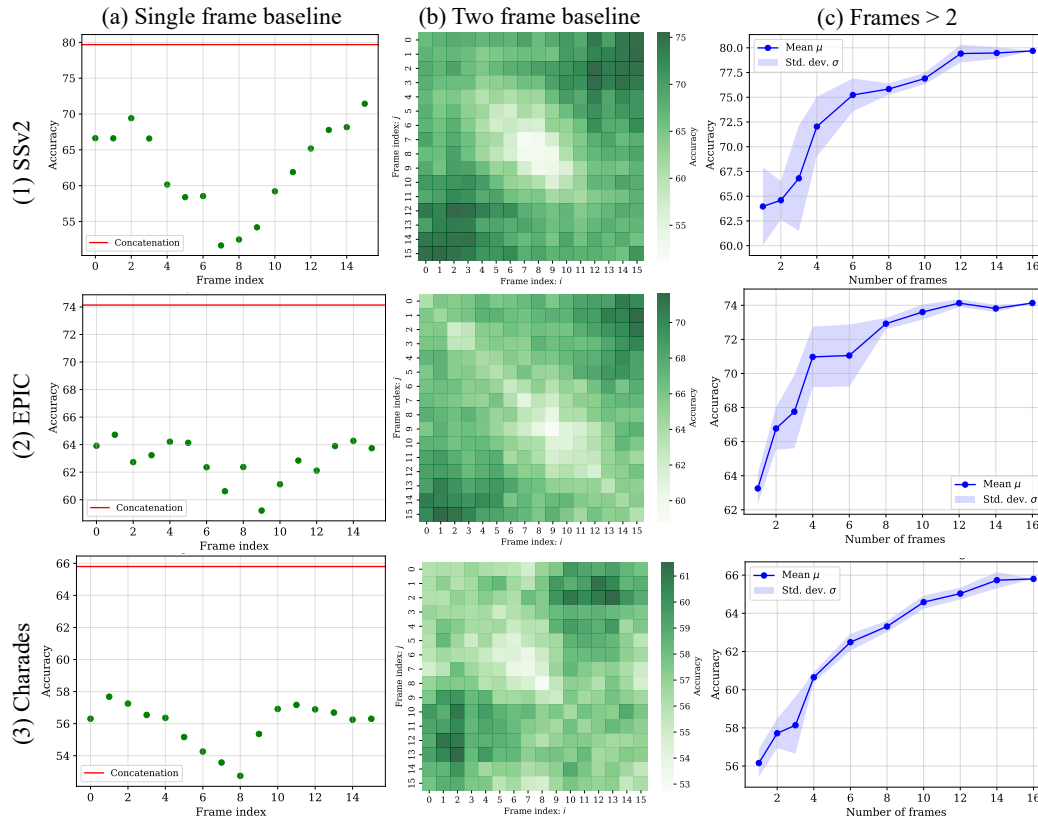


Figure 9: **Time-sensitivity of CiA: Part 1.** We repeat the experiment shown in Fig 5 (a)/(b) of the main paper for all datasets. Rows represent datasets while columns represent different properties of the task. (a) A single-frame baseline tends to do well on frames at the end of the video sequence since those usually encode either the start or end state of the action. (b) A two-frame baseline usually does best if the frames are picked at the two ends of the video. (c) As more frames are considered in the context, accuracy on chiral action recognition improves. Overall, these demonstrate that chiral action recognition is time-sensitive: it benefits predictably from more frames, especially at the ends.

C Experiments

C.1 Setup details

Details for chiral action recognition. To benchmark a given video model for chiral action recognition, we require a single descriptor vector for a video. There are two important details here: (i) *input processing pipeline*: Different methods differ in the way they sample frames, apply cropping operations, etc. (ii) *pooling*: existing methods [5, 19, 79] usually only represent short clips (sequence of frames with a fixed stride), so we need a way of pooling clip-level descriptors into a video-level descriptor. For (i), we follow the data pipeline for each model as provided. For (ii), depending on the method, we either average pool per-clip representations following [43] (e.g., for VideoMAEv2 [87]) or concatenate them (e.g., for 3D ResNet methods like TCLR [19]), or we hand-craft a pooling mechanism (e.g., averaging spatial tokens for each frame and concatenating across time for InternVideo2.5 [88]). Investigating a general pooling method that gives more time-aware descriptors is an avenue we leave for future work. For image-based model, we sample T frames linearly and simply concatenate per-frame features to represent the video.

Details for standard action recognition. For the experiments with probing video models [40, 79, 89] with LiFT, we sample a single clip of $T=16$ frames with a stride of $s=4$, resize the short side and center crop to (224, 224). Since VideoMAE, VideoJEPA and InternVideo2.5 all produce a sequence of space-time tokens without any global CLS token, we compute the average of all tokens to represent

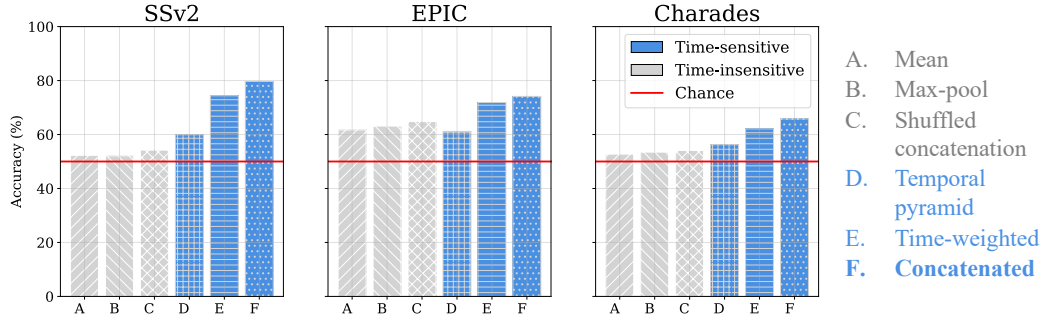


Figure 10: **Time-sensitivity of CiA: Part 2.** We repeat the experiment shown in Fig 5 (c) of the main paper across all three datasets. We show that time-insensitive pooling of per-frame features (e.g., average pooling) leads to much worse performance than with time-sensitive pooling (e.g., concatenation) on chiral action recognition. Note that all the pooling methods considered are non-parametric. This demonstrates the time-sensitivity of chiral action recognition since incorporating the time order of frames substantially improves performance.

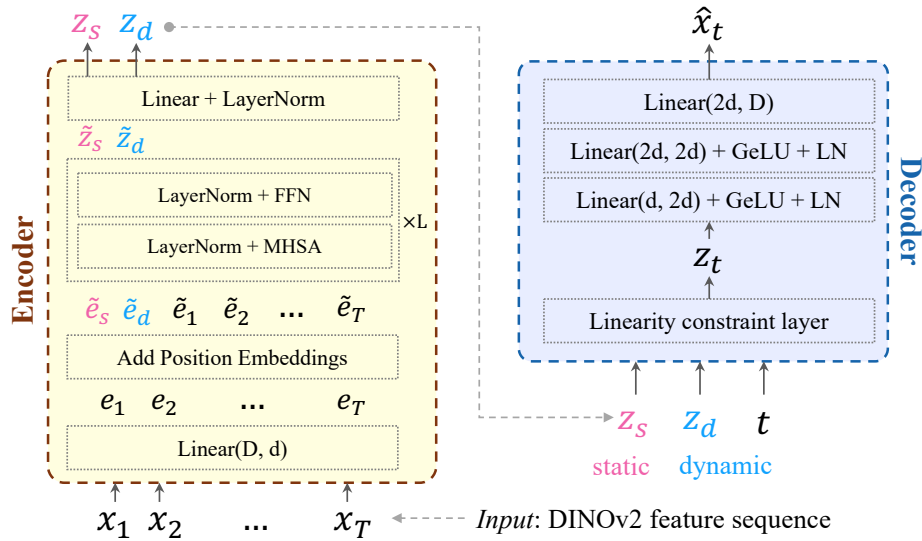


Figure 11: **LiFT architecture details.** The encoder takes in a sequence of DINOv2 features and outputs a video descriptor disentangled into static and dynamic vectors. The decoder reconstructs the feature sequence with linearity baked in the latent space.

the video in case of linear/non-linear probing. Then, we concatenate the LiFT descriptor with this descriptor and train a classifier head on top to output the action class. In case of *linear probe*, the classifier head is a linear layer. In case of *non-linear probe*, it is a two layer MLP with 512 hidden dimensions with ReLU non-linearity and Dropout of 0.1. In case of an attentive probe, following [5], we train a single attention layer with a learnable query to pool the space-time tokens into a single descriptor. Then, LiFT is concatenated with the query vector and a linear classifier layer is added on top of the concatenation. We train the probe for 100 epochs using Adam optimizer with learning rate of $1e^{-5}$ and LRplateau scheduler.

C.2 Additional Ablations

Varying the latent dimension d . In Tab. 3(a), we vary the latent dimension of the Encoder in LiFT. While the number of parameters increases with d , we find that with $d = 384$, LiFT achieves the best performance while still being compact and containing only 8.7M parameters. Note that we do not account for the fixed DINOv2 parameters (22.1M) in this experiment.

Varying the amount of training data. In Tab. Appendix C.2 (b), we vary the amount of training data used to train LiFT with the unsupervised reconstruction loss. We fix the latent dimension to be $d = 384$ and note that the model capacity is fixed. For each row, we run the experiment with three different random seeds and report the average and standard deviation in accuracy. Surprisingly, even with 10% of the data, LiFT gets to 83.3% accuracy. With increasing data samples, the mean accuracy increases marginally. We hypothesize that this is due to two reasons. (i) The model size remains fixed (8.7M) and may not have the capacity to significantly benefit from more samples, (ii) since Kinetics-400 is known to be biased to static understanding (single frame or unordered set of frames [35, 72]), for videos with little visual change, reconstructing the per-frame feature trajectories may not have sufficient signal to inform LiFT. This experiment raises some interesting research questions. Is it possible to achieve time-sensitive video representations by training on selected, *temporally hard* samples only? Is using synthetic data with hand-crafted temporal patterns sufficient [77, 100]? We leave these questions for future work.

Latent dim d	Accuracy	Parameters (M)
192	85.9	2.3
256	85.8	4.0
384	86.6	8.7
512	84.9	15.3

(a) Varying the latent dimension d of Encoder.

Data frac.	Accuracy
0.1	83.3 $\pm$ 0.1
0.2	84.4 $\pm$ 0.2
0.4	85.9 $\pm$ 0.4
0.6	85.6 $\pm$ 0.6
0.8	85.8 $\pm$ 0.8
1.0	86.2 $\pm$ 0.4

(b) Varying the % train data.

Table 8: **Ablations.** Both ablations are conducted on the chiral subset of SSv2. In (a) we vary the latent dimension of the LiFT encoder. We find the best performance with $d = 384$. In (b), we vary the amount of training data (Kinetics-400) used to adapt LiFT. The given % is uniformly randomly chosen from the entire dataset. Surprisingly, even with 10% of the data, LiFT gets to 83% accuracy. We hypothesize that at fixed model capacity, scaling up to more samples gives diminishing returns.

Error bars. To compute error bars, we train LiFT on Kinetics-400 with five different random seeds. The rest of the training configuration is kept constant across all runs. Then, we evaluate the trained models on our main task: chiral action recognition as described in the main paper across the three datasets, SSv2, Charades and EPIC-Kitchens. We report the mean and standard deviation in accuracy in Table 9. The table illustrates these results, highlighting the consistency of the model’s performance.

Dataset	Accuracy (%)
SSv2	86.1 $\pm$ 0.3
EPIC	76.5 $\pm$ 0.8
Charades	70.3 $\pm$ 0.6

Table 9: **Error bars for LiFT.** Mean accuracy across five random seeds. LiFT remains fairly stable and the error bars emphasize the difference between LiFT and other video models.

C.3 Qualitative results

In Fig. 12, we show more examples with tSNE embeddings of the LiFT reconstructed feature trajectories. In most cases, LiFT reconstructs a smoother, continuous approximation of the original trajectory. Note that the original trajectory points seem more scattered than they actually are because tSNE optimizes for local neighborhood distances, which are dominated by closeness of points in the reconstructed trajectory. In case of Fig. 12(f), we observe a divergence between the true and reconstructed trajectories. In this case, the model likely fails to capture the (subtle) visual change which likely causes $\mathbf{z}_d$ to be inaccurate leading to the discrepancy.

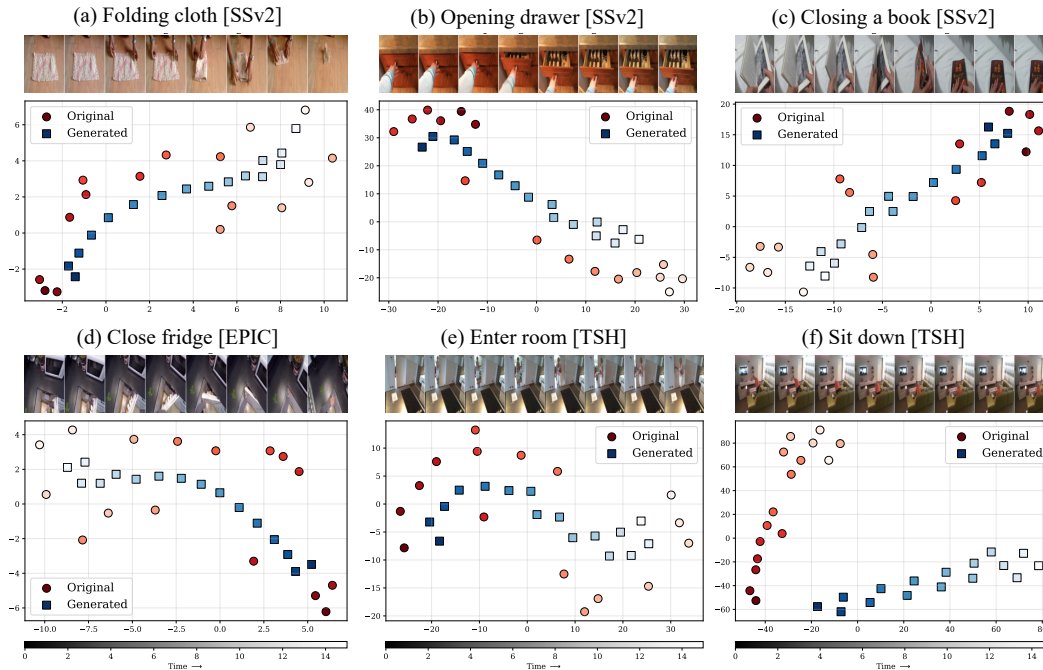


Figure 12: **More qualitative samples of reconstructed features.** We show tSNE embeddings of original and reconstructed features for six videos. The red circles represent original features while blue squares represent reconstructions. Gradient of the color encodes the frame index (time). In general, LiFT tends to output a smooth, continuous approximation of the original feature trajectories in DINO space. (f) is an example failure case where the static token seems reasonable but the direction token is inaccurately predicted causing the direction of original and reconstructed trajectories to differ.

C.4 The curious case of horizontal motion

Based on Table 4, it seems that correctly encoding horizontal motion (distinguish between something moving $\rightarrow$ and $\leftarrow$) is much harder than encoding vertical motion. Notably, we find that the base model itself (DINOv2-ViT-S/14) struggles with this kind of horizontal motion. Ideally, a change in horizontal spatial position (“moving left to right” vs “moving right to left”) should result in dynamic tokens pointing in opposite directions. But this is conditioned on the base model reliably encoding the horizontal spatial position of an object at a given time. Our experiments in Table 10 confirm that the base model itself (DINOv2) does not accurately encode the horizontal spatial position of an object.

Change type	VideoMAE	VJEPa	DINOv2 (concat.)	LiFT
Distance between objects	70.8	87.5	83.3	87.5
Object count	64.2	62.4	69.5	72.4
Object size/depth	96.8	96.8	92.2	96.8
Object state	72.9	66.3	75.9	80.7
Spatial position $\leftrightarrow$	96.3	96.1	75.7	75.2
Spatial position $\updownarrow$	91.5	89.7	79.7	93.6
Average	82.1	83.1	79.4	84.4

Table 10: **LiFT is comparable or superior to much larger video models for all types of visual changes except horizontal shift.** On horizontal shift (e.g., “Pulling something from left to right vs. right to left”), LiFT is worse than these video models. As evident from the DINOv2 (concat.) column, we confirm that this is because the concatenated base DINOv2 features do not encode such motion as well as the video models.

Change type	LiFT	(1.) w/ 224 → 448	(2.) w/ WebSSL	(3.) w/ TTR
Distance between objects	87.5	95.8	91.7	91.7
Object count	72.4	73.7	73.9	73.2
Object size/depth	96.8	92.9	92.2	96.9
Object state	80.7	80.3	80.7	79.2
Spatial position ↔	75.2	82.4	79.7	77.9
Spatial position ↕	93.6	94.4	92.8	92.4
Average	84.4	86.6	85.2	85.2

Table 11: **We show three directions that improve the performance on encoding horizontal shift motion.** TTR denotes test-time rotation augmentation.

Furthermore, we dug deeper into analyzing the DINOv2 feature sequences for horizontal vs vertical shift samples. We hypothesize that the root cause of the difference in performance of DINOv2 concat. (and consequentially, LiFT) on horizontal vs vertical shifts is due to anisotropic sensitivity of DINO feature sequence, i.e., DINO features vary less with horizontal movement vs vertical movement. Below, we explain our experimental setup and the observations.

To have a perfectly controlled test setting, we generate $N=2000$ synthetic sequences with a checkerboard background and a colored disc that moves either horizontally (from left end of the image to right) or vertically (from top to bottom) at a constant rate. An example sequence of horizontal motion is shown in Fig. 13(a). We compute the DINOv2 feature vector for each frame in the sequence. To measure the variation over time, we compute the variance over time and then average it across the feature dimensions. We call this Time Variance (TV). We compute the TV for each sequence and then average it over all sequences for horizontal (or vertical) shifts.

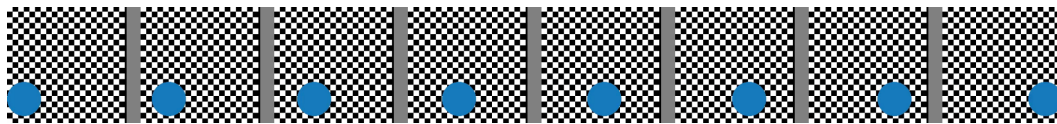
We find that mean Time Variance in vertical shift sequences is about 25% higher than that in horizontal shift sequences. This supports our hypothesis about inherent anisotropic sensitivity of DINOv2 features in case of horizontal or vertical shift motion. Even qualitatively, as shown in Fig. 13(b), vertical motion sequences of DINOv2 concatenated embeddings seem more separable compared to horizontal motion sequences.

It is worth asking why this difference is observed in horizontal vs vertical motion. Is it something to do with the DINO’s training procedure (e.g., cropping mechanism) or position encodings in DINO or something else? Likewise, this connects to how we remedy this (e.g., by training DINO with rotated images?). All these questions require more time and deeper investigation and we defer them to future work. However, we do offer three directions that improve the performance on horizontal motion.

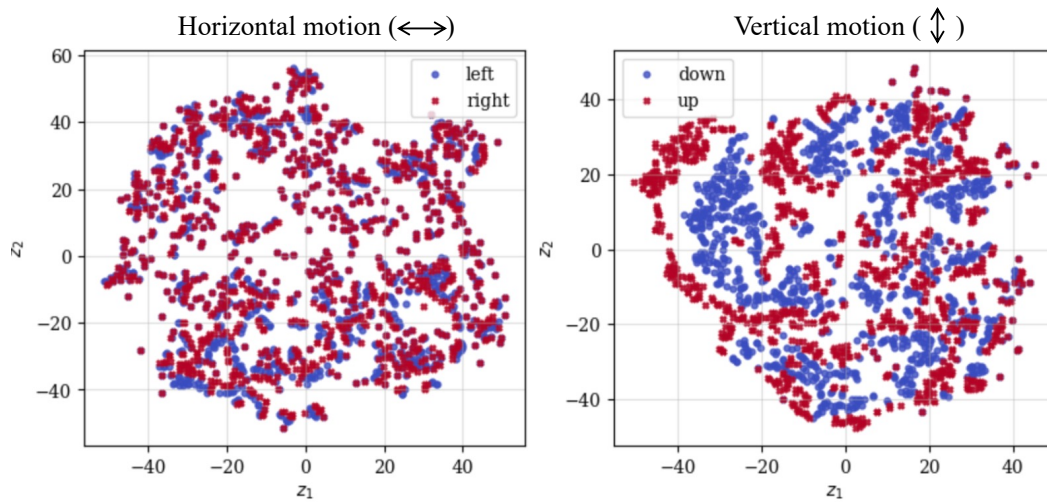
Possible mitigation. We highlight three promising directions to fix this. In the following, we show the resulting improvements in Table 11, and then explain the rationale for each direction below.

- Scaling up the image resolution at test-time: we hypothesize that encoding of fine-grained information such as the spatial positions of objects should improve with image resolution. This provides a +7.2 point improvement in the spatial position while improving the average across all types of changes by 2.2%.
- Improving the base model: an inherently better model should encode spatial positions better. We use WebSSL [2] which is a scaled up DINO-like image model trained on 2B samples. It yields a boost of +4.5% on horizontal shift.
- Using image rotations as a form of test-time recovery: interestingly, we note that encoding of position along the vertical axis is better than that along the horizontal axis. We exploit this fact and concatenate embeddings of videos rotated by $\pi/2, \pi, 2\pi/3$ with that of the upright video.

There is still a gap of 14% between best LiFT model and VideoMAE on horizontal shift. There is more work to do here but we re-iterate that LiFT is stronger than the video encoders for all other kinds of visual change.



(a) Example synthetic video showing horizontal motion



(b) tSNE projections of DINOv2 (concatenated over time) features.

Figure 13: **Synthetic sequence experiments for horizontal vs vertical motion.** Concatenated DINOv2 features are more separable in distinguishing two kinds of vertical motion sequences as compared to horizontal motion sequences.