

Semantic Video Transformer for Robust Action Recognition

1st Keval Doshi

Department of Electrical Engineering
University of South Florida
Tampa, Florida, USA
kevaldoshi@usf.edu

2nd Yasin Yilmaz

Department of Electrical Engineering
University of South Florida
Tampa, Florida, USA
yasiny@usf.edu

Abstract—Video action recognition has attracted significant research attention over the past several years. Although adversarial effects and robustness in image classification models have been heavily investigated, robustness of action recognition models to natural or adversarial perturbations remain largely unexplored. Moreover, even though transformer based approaches have shown great promise on various vision tasks, they have yet to be evaluated in terms of their robustness. To this end, we propose a Semantic Video Transformer for Action Recognition (SeViTAR), which maps visual features obtained by a video transformer to a more robust visual-semantic representation. We extensively evaluate the proposed approach on the ROSE Challenge dataset, and outperform all baselines with a significant margin.

Index Terms—robustness, action recognition, video transformer, semantic mapping

I. INTRODUCTION

Due to the availability of extensively annotated large datasets as well as the development of improved deep neural network (DNN) architectures, several visual recognition tasks, such as image classification and video action recognition, have made significant progress in recent years. Although they have achieved considerable success, DNNs have been discovered to be vulnerable to adversarial attacks [1]–[4]. An extensive body of research has demonstrated that introducing perturbations into an image or video can cause a given DNN prediction to be incorrect. Specifically, it was recently shown that convolutional neural network (CNN) based architectures for video action recognition are susceptible to such adversarial attacks and added perturbations can significantly affect their overall performance. Moreover, natural perturbations to image or video can happen such as camera shaking due to wind, poor video quality due to the Internet connection problems, etc. Machine learning models should be robust to both natural and adversarial perturbations. While robust image recognition models have been extensively studied, robustness in video action recognition still remains largely unexplored.

In the existing research, video action recognition is characterized as a supervised classification problem, in which a model is trained on a collection of known classes and then used to classify a set of unknown test videos. Most existing

approaches rely on extracting visual feature representations using DNN architectures and then classifying them using one-hot encoding. In this paper, we argue that using the inherent semantic information from the class labels can lead to a more robust representation. For example, classes such as *skiing* and *slalom skiing* might have similar visual feature representations, and thus not be discriminative enough to be classified correctly. Hence, we propose leveraging the semantic information of the class labels to generate a more robust feature representation.

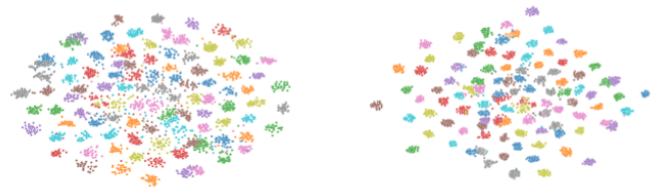


Fig. 1. t-SNE visualization of the features extracted from the UCF-101 dataset by I3D (left) and the proposed transformer (right). Each point represents a video and various classes are represented with different colors. It is seen that the features learned by the proposed transformer model (right) are more separable than the I3D features (left). Best viewed in color.

Several existing video understanding works primarily leverage 3D convolutional models for extracting spatiotemporal feature representations from videos, such as C3D [8] or I3D [9]. However, convolutional kernels are incapable of capturing spatiotemporal correlations that span a large number of time instances. Self-attention based transformer models, on the other hand, have relatively larger receptive fields [10] and can process longer video sequences, capturing long-range dependencies more effectively. In Fig. 1, we show the t-SNE representation of the extracted visual features using the proposed transformer based approach as compared to the I3D model. We see that the transformer model learns more semantically separable features as compared to I3D.

Motivated by these observations, we propose leveraging self-attention architectures, in particular a spatiotemporal transformer model, for extracting visual feature representation from videos. The self-attention approach, in contrast to convolutional kernels, can capture long-range dependencies and is permutation invariant while being computationally efficient

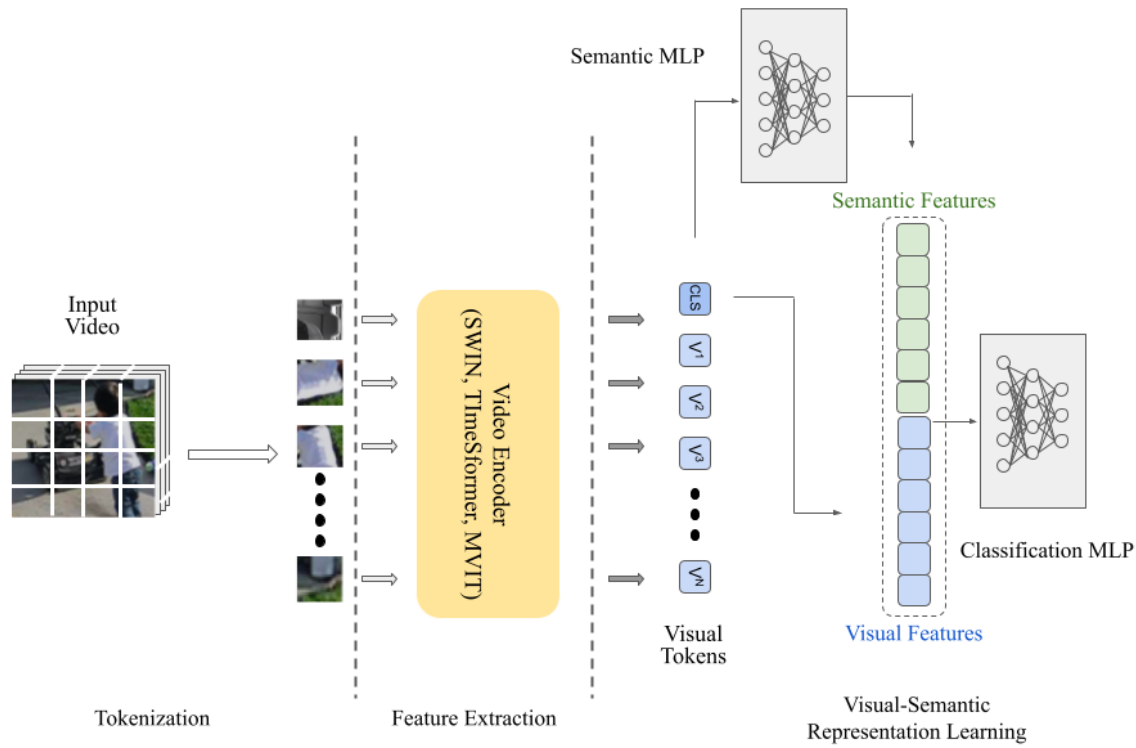


Fig. 2. The proposed semantic video transformer architecture, SeViTAR, for robust action recognition. For feature extraction, any video encoder can be used, such as the recent video transformers, Swin [5], TimeSformer [6], MViT [7].

during training and inference. With the help of a semantic embedding mechanism, our semantic transformer approach is able to learn spatiotemporal features that are more robust to input perturbations compared to the convolutional models. Our contributions can be summarized as follows:

- *A novel semantic video transformer architecture.* We propose a novel architecture that maps the spatiotemporal video features learned by a transformer to a semantic embedding with the help of class labels.
- *A robust transformer based approach for video action recognition.* We show that the semantic embedding together with the spatiotemporal transformer output provide a more robust feature representation.
- *An extensive evaluation of the proposed approach on several benchmark datasets.* The proposed method outperforms existing benchmark approaches by a wide margin on the ROSE Challenge dataset.

II. RELATED WORKS

Video action recognition has been extensively studied over the past several years [9], [11]–[14]. Specifically, 3D-Convolutional models such as C3D [8] and I3D [9] extended 2D-CNN to 3D-CNNs kernels. Recently, transformer based approaches such as the TimeSformer architecture [6] and VidTr [10] have shown promising results on video recognition tasks. While earlier approaches use hand-crafted semantic features [15], recent works have primarily use Word2Vec [16]

for generating the semantic embeddings from the class labels. However, such approaches are prone to suffer from the domain shift problem, which occurs when a model trained on the seen semantic labels is unable to generalize well to the unseen class labels [17].

Despite the fact that adversarial attacks on images have been studied for several years since [2], the vulnerabilities of video recognition models have only recently come to light [1], [18]–[20]. However, several recent approaches address the vulnerabilities of existing deep neural network based models to perturbations or adversarial effects. Specifically, the most common perturbation pattern is sending queries to the target model and estimating gradients to generate adversarial data. PatchAttack (V-BAD) [21] is the first proposed black-box video attack. They proposed a method to generate perturbations for each frame of a video, and then updated their perturbations with queries. Heuristic Attack [22] is another black-box adversarial attack which uses a query-based attack strategy to heuristically select the key frames to be attacked.

III. PROPOSED APPROACH

In this section, we present the proposed approach for robust video action recognition. We introduce a novel transformer model, *SeViTAR* (*Semantic Video Transformer for Action Recognition*), which leverages spatiotemporal self-attention and semantic embedding to learn feature representations that are more robust to video perturbations compared to the benchmark convolutional architectures.



Fig. 3. Different types of perturbations from the ROSE Challenge dataset.

A. Problem Definition

Traditionally, video action recognition has been defined as a supervised classification problem, where given a training set of videos X^s and labels S from seen classes

$$\{(x_1^s, s_1), \dots, (x_N^s, s_N)\},$$

we aim to accurately classify a set of videos

$$X^u = \{(x_1^u), \dots, (x_M^u)\}$$

from unknown classes $U = \{u_1, \dots, u_M\}$, where N and M are the number of training and testing videos respectively.

B. Semantic Video Transformer

Given the promising performance of recently proposed video transformers, such as TimeSformer [6] and VidTr [10], we extract visual features from videos using a spatiotemporal transformer model. Similar to the TimeSformer and VidTr architectures, we leverage the divided space-time attention mechanism. Specifically, applying spatial and temporal attention at the same time requires $\mathcal{O}(n^2)$ complexity with respect to the sequence length, which quickly becomes inefficient. Hence, as shown in Fig. 2, in each encoding block, we first apply temporal attention followed by spatial attention, both of which use the standard qkv attention scheme proposed in [23]. The entire transformer consists of A parallel self-attention heads with L sequential encoding blocks in each head. The visual feature extraction procedure is as follows.

First, a clip y of F frames and size $H \times W$ is sampled from the input video x . We proceed by converting the entire video clip y to a sequence of 2D patches of size $P \times P$, given by $e_t^p \in \mathbb{R}^{3 \times P^2}$. Here $p = 1, \dots, N$ represents the spatial position of each patch (i.e., patch index), $t = 1, \dots, F$ denotes the temporal location of each frame and 3 is the number of color channels. Finally, we flatten each patch e_t^p into $v_t^p \in \mathbb{R}^{3P^2}$ and linearly map it to a score vector using a trainable linear projection $E \in \mathbb{R}^{q \times 3P^2}$:

$$z_t^p(0) = Ev_t^p + \mu_t^p, \quad (1)$$

where $\mu_t^p \in \mathbb{R}^q$ is a latent vector learned to encode the spatiotemporal position of each (p, t) pair. Following the NLP

transformer BERT [24], we add a latent vector $z_0^0(0) \in \mathbb{R}^q$ for an additional fictional patch that will be trained to represent the video's score vector through time and spatial self-attention. The input to the model is $\{z_0^0(0), z_t^p(0)\}_{p,t}$. Each block l receives $\{z_0^0(l-1), z_t^p(l-1)\}$ and produces $\{z_0^0(l), z_t^p(l)\}$. We then use the output of the last step $z_0^0(L)$ as the video's visual feature representation (shown as CLS token in Fig. 2).

Next, for a more robust representation, we augment the visual features with semantic features. Specifically, we obtain a semantic embedding of the input video x by mapping visual features $z_0^0(L)$ to a semantic space through a multilayer perceptron (MLP). We first train the semantic MLP by minimizing the mean squared error (MSE) between the output semantic feature vector $\hat{s}$ and the Sent2Vec embedding s_{vec} of the video class label, s . Finally, the semantic feature vector $\hat{s}$ is concatenated with the visual feature vector $z_0^0(L)$ to form the visual-semantic feature representation $[\hat{s}, z_0^0(L)]$. We finally train the classification MLP on these visual-semantic representations to classify the input videos using the cross-entropy loss.

C. Implementation Details

The proposed SeViTAR model is built upon the SlowFast [25] and TimeSformer packages [6] as the video encoder. We preprocess the input videos by resizing the shorter side to 256 pixels and then a random crop is applied to create a 224×224 ($H \times W$) video snippet. Following the Vision Transformer (ViT) [26] model, we use a patch size of 16×16 , which results in a total of $N = 196$ patches in a frame. During each encoding block for each patch, the size of the score vectors ($z_t^p(l)$) at each encoding block is $q = 768$, and the number of self-attention heads is set as $A = 12$. The number of input frames is set as 8. We extracted the features using 4 NVIDIA A6000 GPUs with a batch size of 8. The loss function is minimized via synchronized SGD with a learning rate of 0.002. To extract semantic embeddings, we use the Sent2Vec algorithm proposed in [27].

TABLE I

VIDEO ACTION RECOGNITION PERFORMANCE ON THE PERTURBED DATASETS OF THE ROSE ROBUSTNESS CHALLENGE. THE PROPOSED SEMANTIC VIDEO TRANSFORMER METHOD OUTPERFORMS THE BENCHMARK METHODS BY A WIDE MARGIN. *ONLY 80K VIDEOS FROM KINETICS-400P ARE USED FOR TRAINING DUE TO TIME AND COMPUTATIONAL CONSTRAINTS.

Method	UCF-101P	HMDB-51P	Kinetics-400P	UCF-101PMini	HMDB-51PMini
SlowFast [28]	-	-	55.48	-	-
I3D [9]	76.9	-	-	-	-
TimeSformer [6]	88.5	70.48	54.26	-	-
SeViTAR (Ours)	89.35	78.92	70.49*	90.54	80.98

IV. EXPERIMENTS

A. Datasets

The UCF-101, HMDB-51, and the Kinetics datasets are three publicly available benchmark datasets that have been commonly used in the majority of recent studies to evaluate the action recognition performance. The UCF-101 dataset contains 13,320 videos from 101 classes, with the majority of the videos focusing on five different types of actions. In the HMDB-51 dataset, there are 6767 videos divided into 51 classes based on everyday human actions. The Kinetics-400 dataset is one of the most comprehensive action recognition datasets available with a total of over 250K videos sourced from YouTube belonging to 400 different classes. To evaluate the performance of the proposed approach in terms of robustness, we use the ROSE Challenge dataset¹, which consists of perturbations such as spatial, temporal, camera-related, and compression corruptions added to the benchmark datasets. A sample set of the perturbations are shown in Fig. 3. The introduced perturbations make it significantly more difficult to recognize activities. Moreover, the intensity of each perturbation is different for each video.

B. Results

In Table I, the performance of the proposed approach is compared with existing state-of-the-art approaches such as SlowFast [28], I3D [9] and the TimeSformer [6]. We used the I3D and TimeSformer models on the ROSE Challenge dataset to obtain benchmark results. The SlowFast results are provided by the ROSE Challenge organizers. It is clearly seen that the proposed SeViTAR model considerably outperforms all benchmark results. Specifically, we see an improvement of 8.44% on the HMDB-51P and 15.01% on the Kinetics-400P datasets as compared to the next best results. These results demonstrate the robustness of the proposed approach as compared to existing models. In particular, the improvement over the TimeSformer results show the contribution of the semantic embedding part.

It should be noted that the performance of the SeViTAR approach on the Kinetics dataset leverages only 80K training videos due to time and computational constraints. It is highly anticipated that the accuracy will further improve when trained on the entire Kinetics dataset. The performance of our method on the mini sets released for the ROSE Challenge are also provided in Table I.

¹<https://rosecvpr22.github.io>

V. ACKNOWLEDGEMENT

This work was supported by the U.S. National Science Foundation (NSF) under Grant #2040572.

VI. CONCLUSION

In this work, we introduce a semantic video transformer for robust video action recognition, called SeViTAR. Specifically, our goal is to highlight the robustness of the visual-semantic embeddings extracted using the proposed SeViTAR model as compared to the existing 3D-CNN and transformer based visual models. Through experiments on the ROSE Robustness Challenge datasets, we show that the proposed approach significantly outperforms the existing approaches by a wide margin.

REFERENCES

- [1] Roi Pony, Itay Naeh, and Shie Mannor, "Over-the-air adversarial flickering attacks against video recognition networks," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2021, pp. 515–524.
- [2] Ian J Goodfellow, Jonathon Shlens, and Christian Szegedy, "Explaining and harnessing adversarial examples," *arXiv preprint arXiv:1412.6572*, 2014.
- [3] Christian Szegedy, Wojciech Zaremba, Ilya Sutskever, Joan Bruna, Dumitru Erhan, Ian Goodfellow, and Rob Fergus, "Intriguing properties of neural networks," *arXiv preprint arXiv:1312.6199*, 2013.
- [4] Nicolas Papernot, Patrick McDaniel, Somesh Jha, Matt Fredrikson, Z Berkay Celik, and Ananthram Swami, "The limitations of deep learning in adversarial settings," in *2016 IEEE European symposium on security and privacy (EuroS&P)*. IEEE, 2016, pp. 372–387.
- [5] Ze Liu, Jia Ning, Yue Cao, Yixuan Wei, Zheng Zhang, Stephen Lin, and Han Hu, "Video swin transformer," in *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, 2022, pp. 3202–3211.
- [6] Gedas Bertasius, Heng Wang, and Lorenzo Torresani, "Is space-time attention all you need for video understanding?," *arXiv preprint arXiv:2102.05095*, 2021.
- [7] Haoqi Fan, Bo Xiong, Karttikeya Mangalam, Yanghao Li, Zhicheng Yan, Jitendra Malik, and Christoph Feichtenhofer, "Multiscale vision transformers," in *Proceedings of the IEEE/CVF International Conference on Computer Vision*, 2021, pp. 6824–6835.
- [8] Du Tran, Lubomir Bourdev, Rob Fergus, Lorenzo Torresani, and Manohar Paluri, "Learning spatiotemporal features with 3d convolutional networks," in *Proceedings of the IEEE international conference on computer vision*, 2015, pp. 4489–4497.
- [9] Joao Carreira and Andrew Zisserman, "Quo vadis, action recognition? a new model and the kinetics dataset," in *proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2017, pp. 6299–6308.
- [10] Yanyi Zhang, Xinyu Li, Chunhui Liu, Bing Shuai, Yi Zhu, Biagio Brattoli, Hao Chen, Ivan Marsic, and Joseph Tighe, "Vidtr: Video transformer without convolutions," in *Proceedings of the IEEE/CVF International Conference on Computer Vision*, 2021, pp. 13577–13587.

- [11] Christoph Feichtenhofer, Axel Pinz, and Andrew Zisserman, "Convolutional two-stream network fusion for video action recognition," in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2016, pp. 1933–1941.
- [12] Christoph Feichtenhofer, Axel Pinz, and Richard P Wildes, "Spatiotemporal multiplier networks for video action recognition," in *Proceedings of the IEEE conference on computer vision and pattern recognition*, 2017, pp. 4768–4777.
- [13] Yunbo Wang, Mingsheng Long, Jianmin Wang, and Philip S Yu, "Spatiotemporal pyramid network for video action recognition," in *Proceedings of the IEEE conference on Computer Vision and Pattern Recognition*, 2017, pp. 1529–1538.
- [14] Basura Fernando, Efstratios Gavves, Jose M Oramas, Amir Ghodrati, and Tinne Tuytelaars, "Modeling video evolution for action recognition," in *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*, 2015, pp. 5378–5387.
- [15] Haroon Idrees, Amir R Zamir, Yu-Gang Jiang, Alex Gorban, Ivan Laptev, Rahul Sukthankar, and Mubarak Shah, "The thumos challenge on action recognition for videos "in the wild"," *Computer Vision and Image Understanding*, vol. 155, pp. 1–23, 2017.
- [16] Tomas Mikolov, Kai Chen, Greg Corrado, and Jeffrey Dean, "Efficient estimation of word representations in vector space," *arXiv preprint arXiv:1301.3781*, 2013.
- [17] Meera Hahn, Andrew Silva, and James M Rehg, "Action2vec: A crossmodal embedding approach to action learning," *arXiv preprint arXiv:1901.00484*, 2019.
- [18] Huanqian Yan, Xingxing Wei, and Bo Li, "Sparse black-box video attack with reinforcement learning," *arXiv preprint arXiv:2001.03754*, 2020.
- [19] Hu Zhang, Linchao Zhu, Yi Zhu, and Yi Yang, "Motion-excited sampler: Video adversarial attack with sparked prior," in *European Conference on Computer Vision*. Springer, 2020, pp. 240–256.
- [20] Shasha Li, Abhishek Aich, Shitong Zhu, Salman Asif, Chengyu Song, Amit Roy-Chowdhury, and Srikanth Krishnamurthy, "Adversarial attacks on black box video classifiers: Leveraging the power of geometric transformations," *Advances in Neural Information Processing Systems*, vol. 34, 2021.
- [21] Linxi Jiang, Xingjun Ma, Shaoxiang Chen, James Bailey, and Yu-Gang Jiang, "Black-box adversarial attacks on video recognition models," in *Proceedings of the 27th ACM International Conference on Multimedia*, 2019, pp. 864–872.
- [22] Zhipeng Wei, Jingjing Chen, Xingxing Wei, Linxi Jiang, Tat-Seng Chua, Fengfeng Zhou, and Yu-Gang Jiang, "Heuristic black-box adversarial attacks on video recognition models," in *Proceedings of the AAAI Conference on Artificial Intelligence*, 2020, vol. 34, pp. 12338–12345.
- [23] Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin, "Attention is all you need," in *Advances in neural information processing systems*, 2017, pp. 5998–6008.
- [24] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova, "Bert: Pre-training of deep bidirectional transformers for language understanding," *arXiv preprint arXiv:1810.04805*, 2018.
- [25] Yunpeng Chen, Haoqi Fan, Bing Xu, Zhicheng Yan, Yannis Kalantidis, Marcus Rohrbach, Shuicheng Yan, and Jiashi Feng, "Drop an octave: Reducing spatial redundancy in convolutional neural networks with octave convolution," in *Proceedings of the IEEE/CVF International Conference on Computer Vision*, 2019, pp. 3435–3444.
- [26] Alexey Dosovitskiy, Lucas Beyer, Alexander Kolesnikov, Dirk Weissenborn, Xiaohua Zhai, Thomas Unterthiner, Mostafa Dehghani, Matthias Minderer, Georg Heigold, Sylvain Gelly, et al., "An image is worth 16x16 words: Transformers for image recognition at scale," *arXiv preprint arXiv:2010.11929*, 2020.
- [27] Matteo Pagliardini, Prakhar Gupta, and Martin Jaggi, "Unsupervised Learning of Sentence Embeddings using Compositional n-Gram Features," in *NAACL 2018 - Conference of the North American Chapter of the Association for Computational Linguistics*, 2018.
- [28] Haoqi Fan, Yanghao Li, Bo Xiong, Wan-Yen Lo, and Christoph Feichtenhofer, "Pyslowfast," <https://github.com/facebookresearch/slowfast>, 2020.