

MS-CLIP: MODALITY-SHARED CONTRASTIVE LANGUAGE-IMAGE PRE-TRAINING

Anonymous authors

Paper under double-blind review

ABSTRACT

Large-scale multimodal contrastive pretraining has demonstrated great utility to support high performance in a range of downstream tasks by mapping multiple modalities into a shared embedding space. Typically, this has employed separate encoders for each modality. However, recent work suggest that transformers can support learning across multiple modalities and allow knowledge sharing. Inspired by this, we investigate how to build a modality-shared Contrastive Language-Image Pre-training framework (MS-CLIP). More specifically, we question how many parameters of a transformer model can be shared across modalities during contrastive pre-training, and rigorously study architectural design choices that position the proportion of parameters shared along a spectrum. We observe that a mostly unified encoder for vision and language signals outperforms all other variations that separate more parameters. Additionally, we find that light-weight modality-specific parallel adapter modules further improve performance. Experimental results show that the proposed MS-CLIP outperforms OpenAI CLIP by 13% relatively in zero-shot ImageNet classification (pre-trained on YFCC100M), while simultaneously supporting a reduction of parameters. In addition, our approach outperforms OpenAI CLIP by 1.6 points on a collection of 19 downstream vision tasks. Furthermore, we discover that sharing parameters leads to semantic concepts from different modalities being encoded more closely in the embedding space, facilitating the learning of common semantic structures (e.g., attention patterns) across modalities.

1 INTRODUCTION

Contrastive Language-Image Pre-training (CLIP) has drawn much attention recently in the field of Computer Vision and Natural Language Processing (Jia et al., 2021; Radford et al., 2021), where large-scale image-caption data are leveraged to learn generic vision and language representations through contrastive loss. This allows the learning of open-set visual concepts and imbues the learned visual feature with a robust capability to transfer to diverse vision tasks.

Prior work in this topic often employs separate language and image encoders, despite architectural similarities between the encoders for both modalities. For instance, the original CLIP work (Radford et al., 2021) uses a ViT (Dosovitskiy et al., 2020) based image encoder, and a separate transformer () based language encoder. However, Lu et al. (2021) recently discovered that transformer models pre-trained on language data could generalize well to visual tasks without altering the majority of parameters, suggesting useful patterns and structures may exist across modalities. In addition, shared architectures have been used to achieve state-of-art performance on a variety of vision-language tasks (Zellers et al., 2021; Li et al., 2019; Chen et al., 2019). These observations suggest that a unified encoder for CLIP may potentially be leveraged to realize performance and efficiency gains.

In this paper, we consequently investigate the feasibility of building a modality-shared CLIP (MS-CLIP) architecture, where parameters in vision encoder and text encoder can be shared. Through this framework, we seek answers to the following three questions: (i) In the CLIP training setting, which layers of the encoders for the two modalities should be shared, and which should be modality-specific? (ii) Within each layer, which sub-module should be shared and which should not? (iii) Lastly, what is the impact to performance and efficiency when including lightweight modality-specific auxiliary modules to accommodate specializations in each modality?

In order to answer these questions, we first perform a comprehensive analysis on the impact of varying the degree of sharing of components across different layers. Our results show that in order to maximize performance, the input embedding, layer normalization (LN) (Ba et al., 2016), and output projection should be modality-specific. However, all the remaining components can be shared across vision and text transformers, including the weights in self-attention and feed-forward modules. Sharing all these layers even outperforms more complex strategies where we employ greedy selection of layers or use Neural Architecture Search (NAS) (Dong & Yang, 2019) to search for the optimal weight sharing policy.

Finally, we explore whether introducing lightweight modality-specific components to the shared backbone may yield a better balance between cross-modality modeling and specializations within each modality. Studied designs include: (i) *Early Specialization*. The first layers in vision Transformer and text Transformer are replaced by extra modules that are specialized for each modality, respectively. This includes a set of lightweight cascaded residual convolutional neural networks (CNNs) for vision, and an additional Transformer layer for language. These early layers allow the representations in each modality to lift to higher level patterns before merging, and introduce shift invariance early in the visual branch. (ii) *Efficient Parallel Branch*. For the visual modality, we explore a lightweight multi-scale CNN network, parallel to the main modality-shared branch, and incorporate its multi-scale features to the main branch through depth-wise convolutional adaptors. This parallel branch enables augmenting the main branch with the benefits convolutions can instill from better modeling of spatial relationships.

We pre-train our MS-CLIP on the major public image-caption dataset YFCC100M (Thomee et al., 2016), and rigorously evaluate on 19 downstream datasets that encompass a broad variety of vision tasks. The experimental results demonstrate that MS-CLIP can out-perform original CLIP with fewer parameters on the majority of tasks, including zero-shot recognition, few-shot learning, and linear probing. Moreover, in order to better understand the success of MS-CLIP, we conduct studies on the learned embedding space, namely with a measurement on multi-modal feature fusion degree (Cao et al., 2020) and quantitatively assess to what degree semantic structures (e.g., attention patterns) are shared across modalities. Our results reveal that sharing parameters can pull semantically-similar concepts from different modalities closer and facilitate the learning of common semantic structures (e.g., attention patterns).

The paper is subsequently organized as follows: in Section 2, we cover datasets and describe the shareable modules and modality-specific designs. In Section 3, we present a rigorous study varying amount of parameters shared across modalities and measure the impact to downstream performance and efficiency. In Section 4 we measure the impact of modality-specific designs to performance, and compare to model architectures with the adapters absent. Section 5 covers related work, and Section 6 concludes.

2 METHODS

2.1 SHARABLE MODULES

Following Radford et al. (2021), we use ViT-B/32 as the basic vision encoder, and the transformer encoder as the basic text encoder, as shown in Fig.1a. We adjust the hidden dimension of text transformer from 512 to 768 to match the token width in the vision transformer. The resulted additional baseline method is noted as CLIP (ViT-B/32, T768). After the adjustment, the vast majority of parameters between the two encoders can be shared, such as the attention modules, feedforward modules, and LayerNorm (LN) layers. Modules that cannot be shared include the input embedding layer (where the vision encoder deploys a projection layer to embed image patches, while the text encoder encodes word tokens), and the output projection layer. Both encoders have 12 transformer layers.

2.2 MODALITY-SPECIFIC AUXILIARY MODULE

In this section we describe the two variations of lightweight modality-specific auxiliary modules used in our study.

Early Specialization In the field of multi-modal learning, it is found beneficial to employ different specialized feature extractors for different modalities and unify them together with the same module in latter layers (Castrejon et al., 2016; Hu & Singh, 2021). Motivated by above, we begin the modality-specific design with making only the first layer specialized for visual and text, leaving other layers shared. Concretely, on vision side, we employ a series of convolutional networks with residual connection as our specialization layer, in which the feature resolution is down-sampled and the channel dimension is increased. The detailed configuration is shown in Tab.1. That is inspired by a recent work (Xiao et al., 2021) where they replace the first layer in ViT with several convolution layers. Here we add residual connection to make it more stable for large-scale training. On the language side, since the Transformer has been a de-facto model for language, we keep the Transformer layer and of course, the parameters are not shared.

Efficient Parallel Branch In image representation, multi-scale information has always been essential. However, vanilla vision Transformer (Dosovitskiy et al., 2020) first patchify the image and use a set of patch features of fixed size all along. In recent works that introduce multi-scale into ViT (Liu et al., 2021a; Wu et al., 2021), they gradually reduce the patch size and increase the dimension of channel stage by stage. Nevertheless, if sharing the weight with language Transformer, above methods can not be incorporated because of varied channel dimension. Motivated by Feichtenhofer et al. (2019), we propose to have an auxiliary parallel branch alongside the shared vision Transformer. It consists of one convolution layer and four residual convolution layers, to lower the resolution and widen the channel. Different from plain residual convolution in Early Specialization, here we utilize the bottleneck design in ResNet (He et al., 2016) to save parameters. The main function of parallel branch is to provide multi-scale feature to shared branch. Therefore, we also employ one adapter after each parallel layer to integrate feature in different scales into different layer of shared Transformer. For efficiency, we adopt depth-wise convolutions (DWConv) and point-wise convolution (PWConv) in adapters to adjust the feature map size and depth. The adapter can be formulated as:

$$\begin{aligned} H'_p &= bn(PWConv(DWConv(H_p))) \\ H' &= ln(bn(DWConv(H)) + H'_p) \end{aligned} \quad (1)$$

where H_p is the multi-scale feature in parallel branch and H' is adapter’s output. bn and ln denote batch normalization and layer normalization. It’s noted the *CLS* token is not fused with parallel branch and keeps unchanged. Detailed configuration is shown in Tab.2.

Table 1: Setting of Early Specialization

Module	Dim_In	Dim_Out
3*3 Conv	3	48
Residual 3*3 Conv	48	96
Residual 3*3 Conv	96	192
Residual 3*3 Conv	192	384
Residual 3*3 Conv	384	768
1*1 Conv	768	768
Total # Parameters	4.1M	

Table 2: Setting of Efficient Parallel Branch.

Parallel Module	Adapter Module	Fusion Layer Id
3*3 Conv	16*16 DWConv	2
Bottleneck 3*3 Conv	8*8 DWConv	4
Bottleneck 3*3 Conv	4*4 DWConv	6
Bottleneck 3*3 Conv	2*2 DWConv	8
Bottleneck 3*3 Conv	1*1 DWConv	10
Total # Parameters	3.9M	

3 INVESTIGATING MODALITY SHARING OF VISION/TEXT TRANSFORMER

Here we explore how varying the degree of sharing weights across modalities impacts performance. We use the models mentioned in Sec. 2.1 for initial investigation.

3.1 PRETRAINING DATASET

We use YFCC100M (Thomee et al., 2016) as the pre-training dataset. Following the filtering process in Radford et al. (2021), we only keep image-text pairs where caption is in English. This leaves us around 22 million data pairs. It is noted that YFCC filtered by OPENAI CLIP has 15M image-caption data. Our filtered YFCC has 22M data because we use a slightly different English dictionary to exclude non-English words.

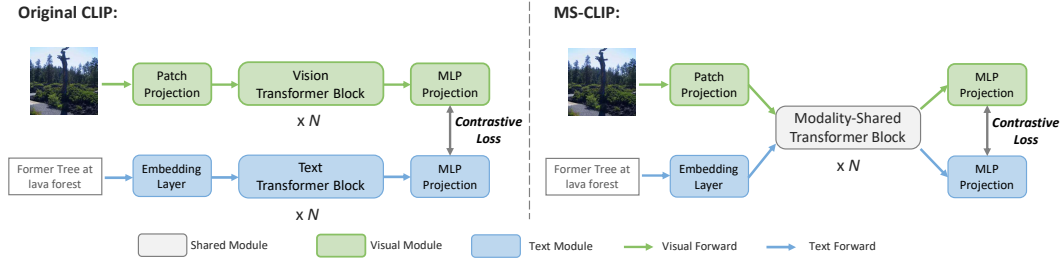


Figure 1: Overview of MS-CLIP, compared with original CLIP.

Table 3: Experimental results of sharing different components in Transformer layer.

Text Width	# Params	Shared Module	Non-Shared Module	Zero-shot Acc(%)
512	150M	-	Attn, FFN, LN1, LN2	32.15
768	209M	-	Attn, FFN, LN1, LN2	31.85
768	125M	Attn, FFN, LN1, LN2	-	28.40
768	125M	Attn, FFN, LN1	LN2	27.57
768	125M	Attn, FFN, LN2	LN1	32.16
768	125M	Attn, FFN	LN1, LN2	32.99

3.2 TRAINING CONFIGURATIONS

Similar to the original CLIP paper Radford et al. (2021), we maintain separate attention masks for image and text: vision transformer allows every patch to attend to others with a bi-directional mask, while text transformer only allows tokens to attend to previous tokens with a uni-directional mask. We train all the models for 32 epochs. The optimizer is Adam with decoupled weight decay regularization (Loshchilov & Hutter, 2017). The learning rate is decayed from $1.6e-3$ to $1.6e-4$, with a cosine scheduler and a warm up at first 5 epochs. We train our models on 16 NVIDIA V100 GPUs with the batch size per GPU to be 256.

3.3 ZERO-SHOT EVALUATION

We use zero-shot accuracy on ImageNet (Deng et al., 2009) validation set as a common evaluation metric. Following CLIP, we use an ensemble of multiple prompts to extract text features as category features.

3.4 INITIAL OBSERVATIONS

1. LNs need to be modality-specific. We mainly examine the shareable modules within each Transformer layer, as the input and output projection layers could not be shared. As shown in Tab.3, the first model variant shares all components, including two LN layers and transformation weights in self-attention module and feedforward module, which results in worse performance compared to CLIP (ViT-B/32) and CLIP (ViT-B/32, T768). Then we make the two LN layers modality-specific, which yields better performance and even surpasses the non-shared version in both zero-shot accuracy and parameter efficiency. Noted that the number of parameters in LNs is almost negligible compared with the transformation weights. The sharing is applied in all 12 layers for simplicity. Our observation echos the finding in FPT (Lu et al., 2021) that only tuning LNs in a mostly-frozen pretrained language model yield satisfactory performance on vision tasks.

2. Less is more: Sharing all layers is better than some. We further study which layer should be modality-specific and which should be modality-shared. We conduct experiments on sharing last N layers where N is ranging from 12 to 0. $N = 12$ indicates all layers are shared and $N = 0$ indicates the non-shared baseline CLIP (ViT-B/32, T768). Tab. 4 suggests that sharing all 12 layers performs

Table 4: Results of sharing different layers in Transformer.

Share Last X layers	12	11	10	8	6	4	2	0	NAS-Search
Zero-shot Acc(%)	32.99	31.25	32.21	32.39	32.85	30.91	nan	31.85	30.97
# Parameters	125M	132M	139M	153M	167M	181M	195M	209M	174M

Table 5: Layer-wise NMI scores of models.

Layer	0	1	2	3	4	5	6	7	8	9	10	11	Avg.
CLIP (ViT-B/32, T768)	0.586	0.387	0.265	0.252	0.255	0.241	0.239	0.243	0.235	0.23	0.227	0.185	0.278
MS-CLIP (B/32)	0.589	0.332	0.235	0.211	0.2	0.21	0.2	0.202	0.214	0.197	0.192	0.173	0.246
w/ Early Specialization	0.471	0.348	0.215	0.21	0.218	0.221	0.22	0.213	0.19	0.183	0.179	0.161	0.235
MS-CLIP-S (B/32)	0.519	0.536	0.243	0.216	0.199	0.221	0.19	0.247	0.216	0.215	0.224	0.217	0.270

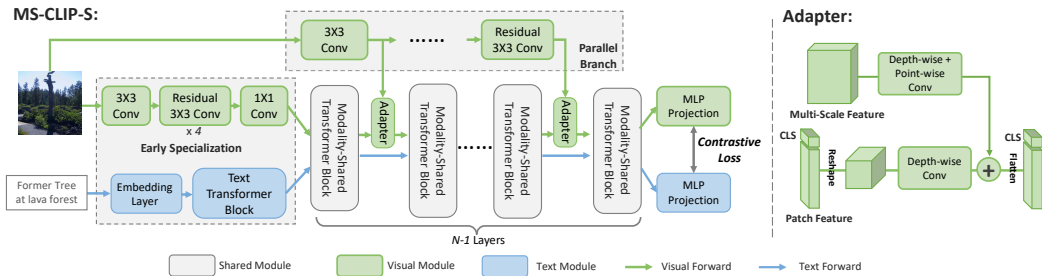


Figure 2: Overview of MS-CLIP-S.

the best while requires the least number of parameters. We name this model MS-CLIP. Additionally, inspired by recent work on Neural Architecture Search (NAS) (Zheng et al., 2021; Dong & Yang, 2019), we train a model that learns a policy to control which layer to (not) share via Gumbel Softmax (Dong & Yang, 2019). Despite its sophistication, it still underperforms MS-CLIP.

3. Shared model exhibits higher multi-modal fusion degree. To probe the multi-modal fusion degree, following Cao et al. (2020), we measure the Normalized Mutual Information (NMI) between visual features and text features at each layer. For each image-caption pair, we use K-means algorithm ($K=2$) to group all feature vectors from the forward pass of visual input and text input into 2 clusters. Then, NMI is applied to measure the difference between the generated clusters and ground-truth clusters. The higher the NMI score is, the easier the visual features and text features can be separated, and the lower the multi-modal fusion degree is.

NMI scores are then used to probe the multi-modal fusion degree of the shared model (MS-CLIP (B/32)) vs. non-shared model (CLIP (ViT-B/32, T768)). Here we choose CLIP (ViT-B/32, T768) instead of CLIP (ViT-B/32) in that the feature dimensions of two modalities have to be the same for clustering. NMI scores of all 12 layers and the average are listed in the first two rows of Tab.5. Shared model has lower NMI scores than original CLIP on almost all the layers and the average, indicating a higher degree of multi-modal fusion.

4 EXPERIMENTS

Given the results in Section 3 demonstrating the robust results from sharing most parameters across modalities, we further explore whether introducing lightweight modality-specific components to the shared backbone may yield a better balance between cross-modality modeling and specializations within each modality. We name the MS-CLIP equipped with both modality-specific designs as MS-CLIP-S, where "S" means supreme. A detailed diagram of MS-CLIP-S is shown in Fig. 2

We first introduce the pre-training setting and details. Then we validate the representation capability of MS-CLIP-S from both zero-shot recognition and linear probing. Finally we analyze MS-CLIP-S both quantitatively and qualitatively.

Table 6: Experimental results of zero-shot recognition on ImageNet validation.

Module Name	# Parameters	Zero-shot Acc(%)
CLIP (ViT-B/32)	150M	32.15
CLIP (ViT-B/32, T768)	209M	31.85
MS-CLIP (B/32)	125M	32.99
w/ Early Specialization	129M	35.18
w/ Parallel Branch	129M	34.18
MS-CLIP-S (B/32)	133M	36.66

4.1 SETUP

Training Details: In addition to training details mentioned in Section 3, the weight decay for non-shared parameters and shared parameters are separately set to 0.05 and 0.2. We found that a higher weight decay for share parameters works better, because shared parameters are updated twice in each iteration, and a higher weight decay can mitigate over-fitting.

Evaluation Datasets: In addition to zero-shot Imagenet mentioned in Section 3, we choose 19 public datasets to prove the representation learning capabilities of MS-CLIP: ImageNet, Food-101 (Bossard et al., 2014), CIFAR-10 (Krizhevsky et al., 2009), CIFAR-100 (Krizhevsky et al., 2009), SUN397 (Xiao et al., 2010), Stanford Cars (Krause et al., 2013), FGVC Aircraft (Maji et al., 2013), Pascal Voc 2007 Classification (Everingham et al.), Describable Texture (dtd) (Cimpoi et al., 2014), Oxford-IIIT Pets (Parkhi et al., 2012), Caltech-101 (Fei-Fei et al., 2004), Oxford Flowers 102 (Nilsback & Zisserman, 2008), MNIST (LeCun et al., 1998), Facial Emotion Recognition (Pantic et al., 2005), STL-10 (Coates et al., 2011), GTSRB (Stallkamp et al., 2012), PatchCamelyon (Veeling et al., 2018), UCF101 (Soomro et al., 2012), Hateful Memes (Kiela et al., 2020). Those datasets cover various visual scenarios, including generic objects, memes, scenes and etc. We take the frozen visual encoder and apply a linear classifier to do logistic regression on top of extracted features. Models are separately trained on each dataset for one epoch and test the accuracy.

Compared Models: We conduct comprehensive experiments with following settings. (1) CLIP (ViT-B/32): The same as Radford et al. (2021), this uses ViT-B32 as visual encoder and Text Transformer as text encoder with width to be 512. (2) CLIP (ViT-B/32, T768): This model sets the width of Text Transformer as 768 to unify the dimension of both encoders. (3) MS-CLIP (B/32): Compared with CLIP (ViT-B/32, T768), this model utilizes the modality-shared transformer blocks to substitute non-shared transformer blocks in visual and text encoders. (4) MS-CLIP (B/32) + Early Specialization: Based on (3), we specialize the first layer of shared visual&text encoders following Sec. 2. (5) MS-CLIP (B/32) + Parallel Branch: Based on (3), we add a parallel branch to shared visual encoder. (6) MS-CLIP-S (B/32): Based on (3), we apply both early specialization and parallel branch to our shared visual&text encoders.

4.2 EXPERIMENTAL RESULTS

Zero-Shot ImageNet: The experimental results are reported in Tab.6. In the first row, we reproduce the CLIP (ViT-B/32) pre-trained on YFCC, following the officially released code. On YFCC, Radford et al. (2021) only reported the result of CLIP (ResNet50), which is 31.3% on zero-shot recognition of ImageNet. It proves that our re-implementation can basically re-produce the results reported. By comparing 1-st row and last row, we find MS-CLIP-S (B/32) can outperform CLIP (ViT-B/32) by 4.5% absolutely and 13.9% relatively in zero-shot recognition accuracy on ImageNet, with less parameters.

Ablation Study: In Tab.6, we further analyze the effect of components in MS-CLIP. By comparing 2-nd row and 3-rd row, it is found that directly increasing the text transformer’s capacity is useless and even a bit harmful. That is also mentioned in Radford et al. (2021). Then comparing 3-rd row and 4-th row, we find that sharing parameters in vision and text transformer improves the

Table 8: Common Semantic Structure distance

Layer	0	1	2	3	4	5	6	7	8	9	10	11	Avg.
CLIP (ViT-B/32)	0.18	0.203	0.227	0.186	0.178	0.164	0.118	0.103	0.106	0.109	0.105	0.074	0.143
MS-CLIP (B/32)	0.175	0.128	0.153	0.132	0.136	0.136	0.106	0.119	0.092	0.106	0.083	0.058	0.113
+ Early Specialization	-	0.107	0.142	0.16	0.12	0.12	0.103	0.103	0.096	0.111	0.11	0.058	0.111
MS-CLIP-S (B/32)	-	0.085	0.162	0.105	0.102	0.103	0.105	0.114	0.093	0.094	0.093	0.061	0.101

performance and even can outperform CLIP (ViT-B/32) by 0.8%. It demonstrates that sharing the parameters enables the visual and text information to benefit and complement each other. Then we evaluate the proposed auxiliary modality-specific modules one by one. The comparison between 5-th row and 4-th row tells that early specialization can bring 2.1% improvement with only 4M parameters increased. On the other hand, from 6-th row and 5-th row, we realize that auxiliary parallel branch on vision can also improve by 1.1%. Those two auxiliary modules can work together to further boost the accuracy to 36.66%.

Table 7: Linear probing results on 19 datasets

Datasets	CLIP (ViT-B/32)	MS-CLIP-S (B/32)	Δ
Food-101	71.3	76.0	+ 4.7
SUN397	68.1	71.7	+ 3.6
Stanford Cars	21.8	27.5	+ 5.7
FGVC Aircraft	31.8	32.9	+ 1.1
Pascal Voc 2007	84.4	86.1	+ 1.7
Describable Texture (dtd)	64.1	69.4	+ 5.3
Oxford-IIIT Pets	61.1	62.1	+ 1.0
Caltech-101	82.8	81.6	- 1.2
Oxford Flowers 102	90.7	93.8	+ 3.1
MNIST	96.5	97.2	+ 0.7
Facial Emotion Recognition	54.9	53.6	- 1.3
STL-10	95.4	95.1	- 0.3
GTSRB	67.1	69.9	+ 2.8
PatchCamelyon	78.3	81.3	+ 3.0
UCF101	72.8	74.6	+ 1.8
CIFAR-10	91.0	87.2	- 3.8
CIFAR-100	71.9	66.7	- 5.2
Hateful Memes	50.6	52.4	+ 1.8
ImageNet	58.5	63.7	+ 5.1
Avg.	69.1	70.7	+ 1.6

Linear Probing: Since we already conduct ablation study under zero-shot recognition, in linear probing, we only compare the CLIP (ViT-B/32) and MS-CLIP-S (B/32). All the results are listed in Tab. 7. Overall, MS-CLIP-S (B/32) outperforms CLIP (ViT-B/32) on 14 out of 19 tasks. The average improvement of 19 tasks in total is 1.6%. The reason behind the improvement of visual encoder might be that, the integration of modality-shared module and modality-specific module enables the visual encoder to benefit from useful language information.

4.3 FURTHER ANALYSIS

NMI Score In Sec. 3, we already explain how to measure NMI score and reports the NMI scores of CLIP (ViT-B/32, T768)

and MS-CLIP (B/32). We further measure the NMI scores of MS-CLIP (B/32) + Early Specialization and MS-CLIP-S (B/32). The result shows that introducing early specialization can further improve the multi-modal fusion degree. But adding parallel branch leads to a decrease of multi-modal fusion degree. That might be due to the integration of modality-specific multi-scale visual features. From the Tab. 6, adding parallel branch indeed improves the transferable representation, which means NMI score may not be a direct indicator of representation quality. In following subsection, we introduce another metric to analyze the knowledge learnt in MS-CLIPs.

Multi-modal Common Semantic Structure To understand why modality-shared Transformer blocks and proposed auxiliary modality-specific modules can improve the representation, we dig deeper into the what our modules have learnt after training. Our hypothesis is that MS-CLIPs should better capture the common semantic structures existing inside concepts in different modalities. To quantitatively measure it, we probe the attention weights during inference and measure the similarity between attentions in visual and attentions in text. To be more specific, the dataset we use is Flick30K-Entity (Plummer et al., 2015), where there are multiple objects in each image grounded

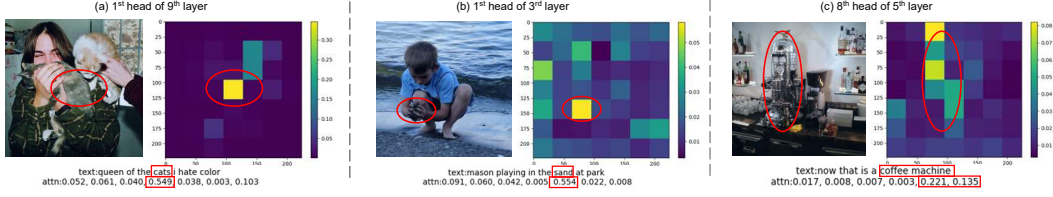


Figure 4: Visualized attention maps of shared attention head.

to corresponding concepts in caption. Given an image, assume there are grounded objects (visual concepts) $\{vc_1, vc_2, \dots, vc_n\}$ and corresponding grounded text concepts $\{tc_1, tc_2, \dots, tc_n\}$, in which tc_i refers to vc_i . In the h -th head of l -th attention layer, we take the raw visual attention map M^{lh} and raw text attention map K^{lh} . In order to get the relationship between concepts, we map the text concept tc_i to its last token t_i , and map the visual concept vc_i to its center patch v_i . Through this mapping, we can treat the attention value between tc_i and tc_j as K_{ij}^{lh} , and attention value between vc_i and vc_j as M_{ij}^{lh} . Then for each concept pairs $\{i, j\}$ in both vision and text, we normalize the attention value over starting concept i with softmax function, and average the normalized attention values over all heads in that attention layer. Further, we compute the $l1$ distance between attention values of the same concept pair in different modalities. Finally, we sum the $l1$ distances of all the concept pairs and treat it as the Common Semantic Structure (CSC) distance of that attention layer. A lower CSC distance means more common attention patterns learnt in Transformer across two modalities. The whole process can be formulated as:

$$dis_{ij}^l = \left| \sum_{h=1}^H \frac{1}{H} softmax_i(M_{ij}^{lh}) - \sum_{h=1}^H \frac{1}{H} softmax_i(K_{ij}^{lh}) \right| \quad (2)$$

$$CSC^l = dis^l = \sum_{i=1}^n \sum_{j=1}^n (dis_{ij}^l) \quad (3)$$

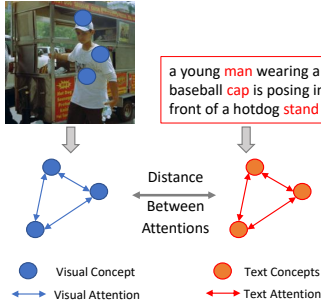


Figure 3: Diagram of computing Common Semantic Structure distance

The layer-wise CSC distance of CLIP (ViT-B/32), MS-CLIP (B/32), MS-CLIP (B/32) + Early Specialization and MS-CLIP-S (B/32) are reported in Tab. 8. It is worth noting we use 10k image-caption pairs from Flickr30k-Entity to compute, which is large enough for getting a stable CSC distance. Since the first layer of MS-CLIP (B/32) + Early Specialization and MS-CLIP-S (B/32) doesn't contain attention module in vision branch, we average the last 11 layers' CSC distance to evaluate it. We can find that both the modality-shared Transformer blocks and proposed auxiliary modality-specific modules can lower the CSC distance and learn more semantic structure similarity of vision and text. It is natural that sharing parameters can enforce the attention to learn more common information. As for proposed modality-specific modules, we suspect that those well designed models can account for the discrepancy of separate modalities and make the remaining shared modules focus more on the common patterns.

Visualization of Shared Attention Head In order to intuitively understand how shared attention module works, we visualize the visual attention patterns and text attention patterns of the same shared attention head during inference. More precisely, for vision, we visualize the attention weights between *CLS* token and all patches. For text, we visualize attention weights between *EOS* token and all other tokens. The reason is that both *CLS* token and *EOS* token will be used as output global feature. The model we use is MS-CLIP-S (B/32). We surprisingly find some heads being able to highlight the same concepts from different modalities. Some samples are visualized in Fig. 3. Take Fig. 3(a) as an example. Given the image and caption respectively as input, the 1st head of 9-th attention layer gives the highest

attention value to the region of "cat" in image and token "cats" in text. It validates that the attention heads in MS-CLIP can learn the co-reference between concepts across vision and language.

5 RELATED WORK

5.1 VISION AND LANGUAGE MODELLING

This work is built on the recent success of learning visual representation from text supervision. VirTex (Desai & Johnson, 2021) proposes to learn visual encoder from image captioning objectives. LocTex (Liu et al., 2021b) introduces localized textual supervision to guide visual representation learning. Both studies are conducted on a relatively small scale. A more recent work CLIP (Radford et al., 2021) demonstrates that generic multimodal pre-training could benefit from extremely large scale training (i.e., a private dataset with 400 million image-caption pairs) and obtain strong zero-shot capability. It adopts a simple but effective contrastive objective that attracts paired image and caption and repels unpaired ones. ALIGN (Jia et al., 2021) has a similar model design except for using EfficentNet (Tan & Le, 2019) as their visual encoder, and is pre-trained on an even larger dataset. Besides recognition task, Gu et al. (2021) distills the learnt CLIP knowledge into object detector to perform zero-shot object detection task. In terms of text prompts, CoOp (Zhou et al., 2021) introduce to model context in prompts with continuous learnable representation to avoid the ad-hoc prompt engineering. Our work focuses on the shareability of transformers in vision and text in large-scale contrastive pre-training and are orthogonal to above mentioned works.

5.2 PARAMETER-SHARING ACROSS MODALITIES

Humans reason over various modalities simultaneously. Sharing modules for multi-modal processing has attracted increasing interests recently from the community. Among them, the one most related to us is VATT (Akbari et al., 2021). VATT introduces a transformer shared by video, text and audio and is pre-trained on a contrastive objective. The proposed model naively reuses the entire network for all modalities and yields results worse than the non-shared counterpart. Lee et al. (2020) proposes to share the parameters of Transformers across both layers and modalities to extremely save parameters. They focuses on video-audio multi-modal downstream task and has an additional multi-modal Transformer for modality fusion. In this work, we focus on the transferable visual representation and zero-shot capability with no fusion part. In multi-task multi-modal learning, Hu & Singh (2021) introduces a shared Transformer decoder to handle multiple tasks. In terms of multimodal fusion, Nagrani et al. (2021) utilizes a set of shared tokens across different modalities to enable the information sharing.

6 CONCLUSION

We propose MS-CLIP, a modality-shared contrastive language-image pre-training approach, where most parameters in vision and text encoders are shared. To explore how many parameters of a transformer model can be shared across modalities, we carefully investigate various architectural design choices through plenty of experiments. In addition, we propose two modality-specific auxiliary designs: Early Specialization and Auxiliary Parallel Branch. Experiments on both zero-shot recognition and linear probing demonstrate the superior of MS-CLIP over CLIP in both effectiveness and parameter efficiency. Finally, we analyze the reasons behind and realize that sharing parameters can map two modalities into a closer embedding space and promote the common semantic structure learning across modalities.

REFERENCES

- Hassan Akbari, Linagzhe Yuan, Rui Qian, Wei-Hong Chuang, Shih-Fu Chang, Yin Cui, and Boqing Gong. Vatt: Transformers for multimodal self-supervised learning from raw video, audio and text. *arXiv preprint arXiv:2104.11178*, 2021.
- Jimmy Lei Ba, Jamie Ryan Kiros, and Geoffrey E Hinton. Layer normalization. *arXiv preprint arXiv:1607.06450*, 2016.

- Lukas Bossard, Matthieu Guillaumin, and Luc Van Gool. Food-101 – mining discriminative components with random forests. In *European Conference on Computer Vision*, 2014.
- Jize Cao, Zhe Gan, Yu Cheng, Licheng Yu, Yen-Chun Chen, and Jingjing Liu. Behind the scene: Revealing the secrets of pre-trained vision-and-language models. In *European Conference on Computer Vision*, pp. 565–580. Springer, 2020.
- Lluís Castrejon, Yusuf Aytar, Carl Vondrick, Hamed Pirsiavash, and Antonio Torralba. Learning aligned cross-modal representations from weakly aligned data. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 2940–2949, 2016.
- Yen-Chun Chen, Linjie Li, Licheng Yu, Ahmed El Kholy, Faisal Ahmed, Zhe Gan, Yu Cheng, and Jingjing Liu. Uniter: Learning universal image-text representations. 2019.
- M. Cimpoi, S. Maji, I. Kokkinos, S. Mohamed, , and A. Vedaldi. Describing textures in the wild. In *Proceedings of the IEEE Conf. on Computer Vision and Pattern Recognition (CVPR)*, 2014.
- Adam Coates, Andrew Ng, and Honglak Lee. An analysis of single-layer networks in unsupervised feature learning. In *Proceedings of the fourteenth international conference on artificial intelligence and statistics*, pp. 215–223. JMLR Workshop and Conference Proceedings, 2011.
- Jia Deng, Wei Dong, Richard Socher, Li-Jia Li, Kai Li, and Li Fei-Fei. Imagenet: A large-scale hierarchical image database. In *2009 IEEE conference on computer vision and pattern recognition*, pp. 248–255. Ieee, 2009.
- Karan Desai and Justin Johnson. Virtex: Learning visual representations from textual annotations. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 11162–11173, 2021.
- Xuanyi Dong and Yi Yang. Searching for a robust neural architecture in four gpu hours. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pp. 1761–1770, 2019.
- 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.
- M. Everingham, L. Van Gool, C. K. I. Williams, J. Winn, and A. Zisserman. The PASCAL Visual Object Classes Challenge 2007 (VOC2007) Results. <http://www.pascal-network.org/challenges/VOC/voc2007/workshop/index.html>.
- Li Fei-Fei, Rob Fergus, and Pietro Perona. Learning generative visual models from few training examples: An incremental bayesian approach tested on 101 object categories. In *2004 conference on computer vision and pattern recognition workshop*, pp. 178–178. IEEE, 2004.
- Christoph Feichtenhofer, Haoqi Fan, Jitendra Malik, and Kaiming He. Slowfast networks for video recognition. In *Proceedings of the IEEE/CVF international conference on computer vision*, pp. 6202–6211, 2019.
- Xiuye Gu, Tsung-Yi Lin, Weicheng Kuo, and Yin Cui. Zero-shot detection via vision and language knowledge distillation. *arXiv preprint arXiv:2104.13921*, 2021.
- Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pp. 770–778, 2016.
- Ronghang Hu and Amanpreet Singh. Unit: Multimodal multitask learning with a unified transformer. *arXiv preprint arXiv:2102.10772*, 2021.
- Chao Jia, Yinfei Yang, Ye Xia, Yi-Ting Chen, Zarana Parekh, Hieu Pham, Quoc V Le, Yunhsuan Sung, Zhen Li, and Tom Duerig. Scaling up visual and vision-language representation learning with noisy text supervision. *arXiv preprint arXiv:2102.05918*, 2021.

- Douwe Kiela, Hamed Firooz, Aravind Mohan, Vedanuj Goswami, Amanpreet Singh, Pratik Ring-shia, and Davide Testuggine. The hateful memes challenge: Detecting hate speech in multimodal memes. *arXiv preprint arXiv:2005.04790*, 2020.
- Jonathan Krause, Michael Stark, Jia Deng, and Li Fei-Fei. 3d object representations for fine-grained categorization. In *Proceedings of the IEEE international conference on computer vision workshops*, pp. 554–561, 2013.
- Alex Krizhevsky, Geoffrey Hinton, et al. Learning multiple layers of features from tiny images. 2009.
- Yann LeCun, Léon Bottou, Yoshua Bengio, and Patrick Haffner. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11):2278–2324, 1998.
- Sangho Lee, Youngjae Yu, Gunhee Kim, Thomas Breuel, Jan Kautz, and Yale Song. Parameter efficient multimodal transformers for video representation learning. *arXiv preprint arXiv:2012.04124*, 2020.
- Liunian Harold Li, Mark Yatskar, Da Yin, Cho-Jui Hsieh, and Kai-Wei Chang. Visualbert: A simple and performant baseline for vision and language. *arXiv preprint arXiv:1908.03557*, 2019.
- Ze Liu, Yutong Lin, Yue Cao, Han Hu, Yixuan Wei, Zheng Zhang, Stephen Lin, and Baining Guo. Swin transformer: Hierarchical vision transformer using shifted windows. *arXiv preprint arXiv:2103.14030*, 2021a.
- Zhijian Liu, Simon Stent, Jie Li, John Gideon, and Song Han. Loctex: Learning data-efficient visual representations from localized textual supervision. *arXiv preprint arXiv:2108.11950*, 2021b.
- Ilya Loshchilov and Frank Hutter. Decoupled weight decay regularization. *arXiv preprint arXiv:1711.05101*, 2017.
- Kevin Lu, Aditya Grover, Pieter Abbeel, and Igor Mordatch. Pretrained transformers as universal computation engines. *arXiv preprint arXiv:2103.05247*, 2021.
- Subhransu Maji, Esa Rahtu, Juho Kannala, Matthew Blaschko, and Andrea Vedaldi. Fine-grained visual classification of aircraft. *arXiv preprint arXiv:1306.5151*, 2013.
- Arsha Nagrani, Shan Yang, Anurag Arnab, Aren Jansen, Cordelia Schmid, and Chen Sun. Attention bottlenecks for multimodal fusion. *arXiv preprint arXiv:2107.00135*, 2021.
- Maria-Elena Nilsback and Andrew Zisserman. Automated flower classification over a large number of classes. In *2008 Sixth Indian Conference on Computer Vision, Graphics & Image Processing*, pp. 722–729. IEEE, 2008.
- Maja Pantic, Michel Valstar, Ron Rademaker, and Ludo Maat. Web-based database for facial expression analysis. In *2005 IEEE international conference on multimedia and Expo*, pp. 5–pp. IEEE, 2005.
- Omkar M Parkhi, Andrea Vedaldi, Andrew Zisserman, and CV Jawahar. Cats and dogs. In *2012 IEEE conference on computer vision and pattern recognition*, pp. 3498–3505. IEEE, 2012.
- Bryan A Plummer, Liwei Wang, Chris M Cervantes, Juan C Caicedo, Julia Hockenmaier, and Svetlana Lazebnik. Flickr30k entities: Collecting region-to-phrase correspondences for richer image-to-sentence models. In *Proceedings of the IEEE international conference on computer vision*, pp. 2641–2649, 2015.
- Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Aspell, Pamela Mishkin, Jack Clark, et al. Learning transferable visual models from natural language supervision. *arXiv preprint arXiv:2103.00020*, 2021.
- Khurram Soomro, Amir Roshan Zamir, and Mubarak Shah. Ucf101: A dataset of 101 human actions classes from videos in the wild. *arXiv preprint arXiv:1212.0402*, 2012.

- Johannes Stalldkamp, Marc Schlipsing, Jan Salmen, and Christian Igel. Man vs. computer: Benchmarking machine learning algorithms for traffic sign recognition. *Neural networks*, 32:323–332, 2012.
- Mingxing Tan and Quoc Le. Efficientnet: Rethinking model scaling for convolutional neural networks. In *International Conference on Machine Learning*, pp. 6105–6114. PMLR, 2019.
- Bart Thomee, David A Shamma, Gerald Friedland, Benjamin Elizalde, Karl Ni, Douglas Poland, Damian Borth, and Li-Jia Li. Yfcc100m: The new data in multimedia research. *Communications of the ACM*, 59(2):64–73, 2016.
- Bastiaan S Veeling, Jasper Linmans, Jim Winkens, Taco Cohen, and Max Welling. Rotation equivariant CNNs for digital pathology. June 2018.
- Haiping Wu, Bin Xiao, Noel Codella, Mengchen Liu, Xiyang Dai, Lu Yuan, and Lei Zhang. Cvt: Introducing convolutions to vision transformers. *arXiv preprint arXiv:2103.15808*, 2021.
- Jianxiong Xiao, James Hays, Krista A Ehinger, Aude Oliva, and Antonio Torralba. Sun database: Large-scale scene recognition from abbey to zoo. In *2010 IEEE computer society conference on computer vision and pattern recognition*, pp. 3485–3492. IEEE, 2010.
- Tete Xiao, Mannat Singh, Eric Mintun, Trevor Darrell, Piotr Dollár, and Ross Girshick. Early convolutions help transformers see better. *arXiv preprint arXiv:2106.14881*, 2021.
- Rowan Zellers, Ximing Lu, Jack Hessel, Youngjae Yu, Jae Sung Park, Jize Cao, Ali Farhadi, and Yejin Choi. Merlot: Multimodal neural script knowledge models. *arXiv preprint arXiv:2106.02636*, 2021.
- Xiawu Zheng, Rongrong Ji, Yuhang Chen, Qiang Wang, Baochang Zhang, Jie Chen, Qixiang Ye, Feiyue Huang, and Yonghong Tian. Migo-nas: Towards fast and generalizable neural architecture search. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 43(9):2936–2952, 2021.
- Kaiyang Zhou, Jingkang Yang, Chen Change Loy, and Ziwei Liu. Learning to prompt for vision-language models. *arXiv preprint arXiv:2109.01134*, 2021.