

Supervised Sentence Representation with Interactive Loss Decay

Anonymous ACL submission

Abstract

Contrastive learning has achieved remarkable results in sentence representation, but its semantic representation remains independent in the process of training and inference, and could not pay attention to the interactive information of sentence pairs. This paper proposes InterCSE, a interactive contrastive learning method for sentence embedding, which not only focuses on the semantic similarity sorting of sentence pairs, but also increases the interactive information as a supplement for sentence representation. Meanwhile, we propose a loss decaying strategy to balance embedding similarity and interactive information objectives. We evaluate the performance of InterCSE on standard semantic textual similarity (STS) tasks, and experiments show that our model using $BERT_{base}$ and $BERT_{large}$ achieve 82.11% and 82.88% spearman’s correlation, 0.54% and 0.43% improvement compared to SimCSE respectively. We also conduct experiments that adding interactive network based on Sentence-Transformers, which get 85.18%(+0.88% $BERT_{base}$) and 85.69%(+1.07% $RoBERTa_{base}$) spearman’s correlation on STS Benchmark. Hence, adding interactive features to the traditional siamese network performs very well, and achieves new state-of-the-art performance on sentence representation tasks.

1 Introduction

Sentence representation learning is a vital component of natural language processing tasks (Cer et al., 2017). The rapid development of sentence representation technology has made a wide range of downstream tasks more intelligent, especially information retrieval and text clustering.

Recently, the pre-trained language model has become the cornerstone of natural language processing technology, such as BERT(Devlin et al., 2019; Liu et al., 2019), GPT(Radford et al., 2018, 2019;

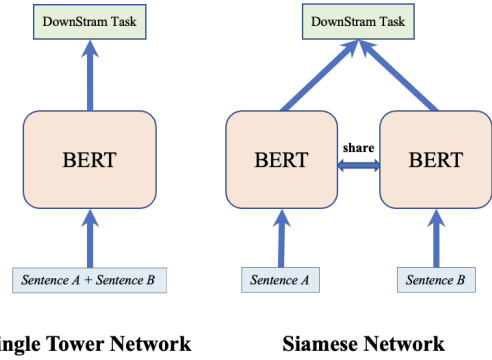


Figure 1: The architecture of single tower network(left) and siamese network(right).

Brown et al., 2020), ERNIE(Sun et al., 2019a,b, 2021), which greatly affects the development of various downstream tasks.

Many sentence representation approaches are transformed into point-wise classification tasks(Reimers and Gurevych, 2019), and there is a gap between the training optimization objectives and good sentence representation. Sentence representation could not get high performance directly with pre-trained language model because of anisotropic phenomenon (Gao et al., 2019), but contrastive learning play the role of a bridge. Comparative learning applies the sorting method with InfoNCE(van den Oord et al., 2019) contrastive loss function, which is more suitable for the optimization goal of sentence representation.

The sentence representation model(Gao et al., 2022; Yan et al., 2021) of contrastive learning combined with pre-trained model is a siamese network, which encodes all sentences independently. The structure of the siamese network makes it possible to quickly produce the vector representation of the sentences and calculate the similarity during training and prediction, so as to complete the retrieval or sentence clustering task in massive data. How-

Network Type	Model	Spearman
Siamese Network	SBERT-base ♡	84.7
	SBERT-large ♡	84.5
	SimCSE-BERT-base ♠	84.3
	SimCSE-BERT-large(reproduce)	85.4
	SROBERTa-base ♡	84.9
	SROBERTa-large ♡	85.0
	SimCSE-RoBERTa-base ♠	85.8
	SimCSE-RoBERTa-large ♠	86.7
Single Tower Network	BERT-base ♦	85.8
	BERT-large ♦	86.5
	RoBERTa-base ◇	87.2
	RoBERTa-large ◇	88.1

Table 1: Comparison of singel tower network and siamese network on STS-Benchmark(Cer et al., 2017) test set with spearman’s correlation. ♡: results from Sentence-Transformers(Reimers and Gurevych, 2019), ♠: results from SimCSE(Gao et al., 2022), ♦ : results from BERT (Devlin et al., 2019), ◇: the performance of RoBERTa-base and RoBERTa-large in single tower network are reproduce through fairseq(Ott et al., 2019). SimCSE models are trained on NLI datasets(Bowman et al., 2015), and the other models are trained on STS-Benchmark train set. All results are rounded to one decimal place for comparison.

ever, the structure of independent encoding makes the sentences pair lose the interactive information, which reduces the accuracy rate.

There are obvious differences between single tower network and siamese network. The network structure is depicted in Figure 1. The single tower network is like the original BERT(Devlin et al., 2019) model structure. After concating the sentence pairs, the semantic features of the sentence pairs are extracted and input to the downstream classification or regression network. The siamese network(Koch et al., 2015) has two encoders that encode each sentence individually where the two encoders share model parameters. After encoding, it can perform similarity calculation or put into the downstream network using sentence embedding, such as Sentence-Transformers(Reimers and Gurevych, 2019).

The performance of the single-tower network and the siamese network on the STS Benchmark test dataset is shown in Table 1. On the whole, the single tower network has a good improvement in spearman’s correlation index compared with the siamese network. We consider that when the single tower network encodes a sentence pair, the sentence is not an independent individual. It will refer to its counterparts for encoding and use the attention mechanism to extract features, which makes the single tower network in the sentence pair similarity task has a higher correlation coefficient. Recently, most sentence representation tasks use

siamese networks. Although it could bring computational advantages on massive data, it inevitably reduces the accuracy of correlation.

In order to take full advantage of the accuracy advantages of single tower networks and the inference speed advantages of siamese networks, this paper proposes a multi-task contrastive learning method. On the basis of SimCSE, we added a single tower network as a supplement to form a framework for multi-task learning. This method could fully obtain the interaction information between sentence pairs during the process of training sentence representation.

Our contributions can be summarized as follows:

1. We propose an effective framework of multi-task contrastive learning for sentence representation tasks which could introduce sentence interactive semantic information.
2. We design a loss function for the proposed framework. When the interactive network introduce information increment, try to minimize the hurt to the original sentence representation.
3. Experiments show that our approach achieves new state-of-the-art performance on STS tasks.
4. We also introduce sentences interactive network based on Sentence-Transformers, and

get a quite outstanding performance on STS benchmark task.

2 Related Work

2.1 Language Model

Recently, transformer(Vaswani et al., 2017) structure shines in the field of deep learning and lays the foundation for large models. The attention mechanism is good at capturing the semantic relationship between sequences. Using transformer’s encoder, many language models have been born, such as GPT(Radford et al., 2018), BERT(Devlin et al., 2019), RoBERTa(Liu et al., 2019), XLNet(Yang et al., 2019), Ernie(Sun et al., 2019a), etc. These models mainly have some differences in masking mechanism, pre-training data and methods. Among them, the BERT model mainly masks 15% words in the sentence and predicts the origin words as a unsupervised pre-training task. After obtaining the pre-trained model, NLP downstream tasks only need to complete fine-tuning on the model to achieve high performances, including sentence classification, sequence tagging, question answering, machine reading and comprehension and sentence-pair classification or regression.

2.2 Contrastive Learning and Sentence Representation

Contrastive learning was first proposed in the field of computer vision. It mainly wants to optimize the similarity of sample pairs in a representation space, and make similar sample pairs gather and dissimilar sample pairs stay away. After the sample representation is obtained using siamese network embedding, the sample similarity is calculated to maximize the similarity of the positive samples. Contrastive learning could take a cross-entropy objective with in-batch-negatives. This optimization scenario is very suitable for unsupervised scenarios. Usually, positive samples can be obtained through simple data enhancement, and a large number of negative samples can be obtained through negative sampling.

In terms of sentence representation, the idea of contrastive learning continues to be used, leading to many research results, such as ConSERT(Yan et al., 2021), SimCSE(Gao et al., 2022) and ES-imCSE(Wu et al., 2022). ConSERT proposes four different data augmentation strategies to generate views for contrastive learning, including adversarial attack, token shuffling, cutoff and dropout.

SimCSE first describes an unsupervised approach, which takes an input sentence and predicts itself in a contrastive objective, with only standard dropout used as noise. Only using dropout as data augmentation becomes a popular method for contrastive learning. ESImCSE introduces two modifications based on SimCSE. It applies a simple repetition operation to modify the input sentence, and then passes the input sentence and its modified counterpart to the pre-trained Transformer encoder, respectively, to get the positive pair. And it introduces a momentum contrast, enlarging the number of negative pairs.

2.3 Multi-Task Learning

Multi-Task Learning (MTL) is an important research topic in machine learning, which aims to learn multiple tasks simultaneously. In MTL, multiple tasks are divided into multiple learning units, and a learner can learn multiple tasks by making multiple small models.

One of the important research directions in MTL is to develop efficient algorithms and theoretical models. In recent years, many works have been carried out in this direction, including deep neural networks, attention mechanism, joint attention, graph attention, and so on. MTL has been successfully used in natural language processing applications, including text classification(Liu et al., 2017), machine translation(Luong et al., 2016), sequence labeling(Rei, 2017) and sentence representations(Ahmad et al., 2018). These algorithms and theoretical models can help MTL learn multiple tasks with better performance.

3 InterCSE

In this section, we present InterCSE (introducing **I**nteractive semantic information in **C**ontrastive **S**entences **E**mboding) for sentence representation task. Firstly, we present the overall framework of our approach. Then, we introduce the training dataset for our model. Finally, we talk about the combination strategies of loss function for multiple objectives.

3.1 Framework

The main idea of our approach is to introduce interactive information while encoding a sentence. Hence, a single tower network which we call interactive network in InterCSE is added on the basis of SimCSE(Gao et al., 2022) as shown in Figure 3.

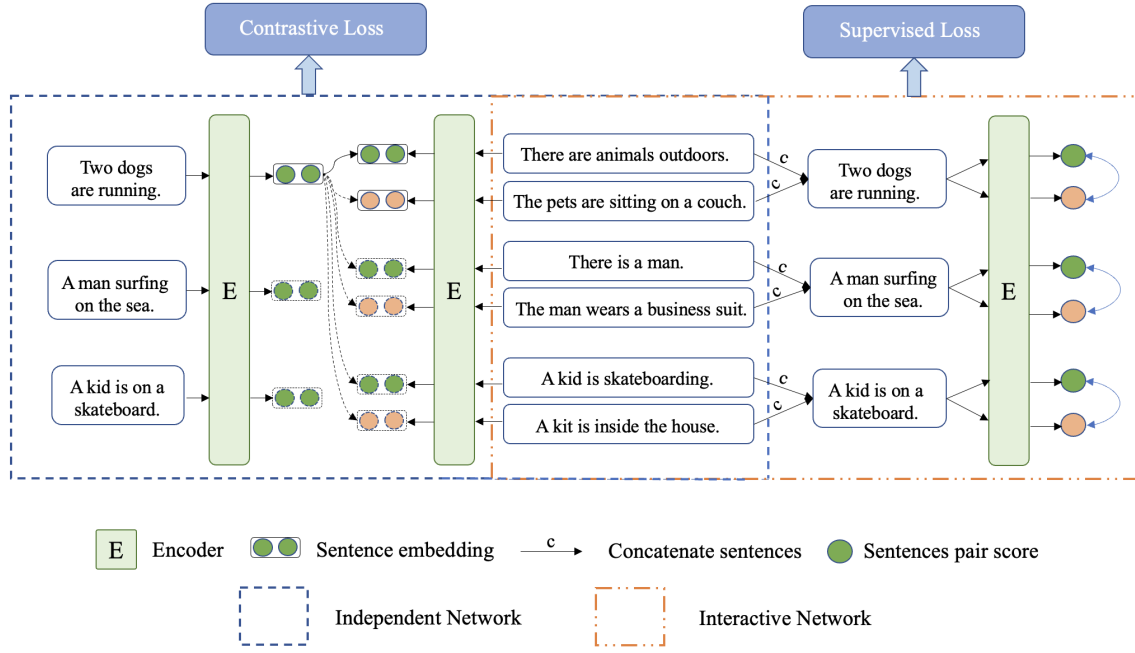


Figure 2: The architecture of InterCSE. Independent network focuses on the representation of sentence, and interactive network perceives sentences interactive semantic. Independent network encodes sentences individually, while interactive network concatenates two sentences and predicts the degree of semantic relevance. Three encoders share same parameters when training.

There are two major components in our framework: independent network and interactive network. During the training process, the interactive network completes the interactive feature extraction of sentence pairs. While inference, only independent network works and generates sentence embedding. All the transformer encoders share same parameters.

Independent Network The goal of the independent network is to quickly produce semantic representation of single sentence without relying on other sentences. The independent network is composed of a transformer encoder and cosine network. Transformer encoder is a siamese network based on BERT-like model, and cosine network is generally cosine similarity calculation. The input is exactly two sentence, such as sentence A and sentence A^+ (or sentence A^-)¹, and independent network encodes the sentence pairs independently by $[CLS]$ representation, then calculates the cosine similarity. There are a high semantic score between sentence A and sentence A^+ , and a low semantic score between sentence A and sentence A^- .

Interactive Network The goal of the interac-

¹For convenience of description, sentence A refers *Two dogs are running.*, sentence A^+ refers *There are animals outdoors.*, and sentence A^- refers *The pets sitting on a couch.* in Figure 3.

tive network is to obtain non-independent sentence semantic information through the attention interaction of words between sentence pairs. Non-independent sentences semantic information is an important correlation signal, which can keenly perceive the semantics of sentence to details. The interactive network consists of an transformer encoder and a feed forward network. Transformer encoder is a BERT-like model, and the input is the concatenation of two sentence and $[SEP]$ token. The $[CLS]$ embedding of transformer encoder is input into the feed forward network to evaluate the similarity of sentence pairs or classification. For example, concatenation of sentence A and sentence A^+ is positive, concatenation of sentence A and sentence A^- is negative.

3.2 Training Dataset

The model we proposed is suitable for supervised tasks, and the in-batch-negative sampling method of the independent network requires discrete label data, so we introduce natural language inference (NLI) datasets to train our model, including the SNLI(Bowman et al., 2015) and MNLI(Williams et al., 2018) datasets. NLI datasets consist of high-quality pairs, and given a premise, human annotators generate three types of sentences: entailment(is absolutely true), neutral(might be true),

Type of NLI datasets	Numbers
Entailment	314k
Neutral	314k
Contradiction	314k
Entailment+Contradiction	270k

Table 2: Statistics of different type of SNLI+MNLI datasets.

and contradiction(is definitely false).

Following the work of SimCSE(Gao et al., 2022), we adopt the hard negative strategy, using entailment and contradiction to represent positive and negative samples, respectively. Therefore, a triplet is generated, (x_i, x_i^+, x_i^-) , where x_i is the premise, x_i^+ and x_i^- are entailment and contradiction hypotheses. Statistics of training datasets are shown in Table 2.

3.3 Loss Expression

The model we proposed is a multi-task model, including two modules of interactive network and independent network. The modeling objectives of the two modules are different, and the corresponding loss functions are also different. The following describes the loss functions corresponding to the two modules in detail.

For the independent encoding module, the input is a mini-batch of data, we take a cross-entropy objective with in-batch negative sampling. We also use hard negative in our model, and the contrastive loss $loss_{cl}$ can be expressed as:

$$-\log \frac{e^{\text{sim}(h_i, h_i^+)/\tau}}{\sum_{j=1}^N (e^{\text{sim}(h_i, h_i^+)/\tau} + e^{\text{sim}(h_i, h_i^-)/\tau})} \quad (1)$$

where $\text{sim}(\cdot)$ indicates cosine similarity function, τ controls the temperature, h_i , h_i^+ and h_i^- is the embedding of x_i , x_i^+ and x_i^- respectively.

For the interaction module, the input is a mini-batch data where there are N pairs (x, x^+) and (x, x^-) , and the label of (x, x^+) and (x, x^-) are 1 and 0 respectively. We consider two kind of loss functions to perceive the interactive semantic information.

The first is classification loss for all sentences in a mini-batch, and we use cross entropy loss $loss_{ce}$ to express as:

$$-\sum_{i=1}^{2N} (y_i \log(p_i) + (1 - y_i) \log(1 - p_i)) \quad (2)$$

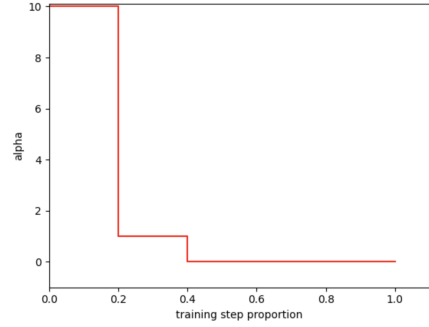


Figure 3: α piecewise constant decaying when training.

where y_i is the label, p_i is the predicted score of feed forward network.

The second is margin ranking loss for the ranking relation between positive and negative samples in a mini-batch. Margin ranking loss $loss_{mr}$ could enhance the pairwise distinction in semantic space, and it can be express as:

$$\sum_{i=1}^N \max(0, \text{margin} - (p_i^+ - p_i^-)) \quad (3)$$

where p_i^+ and p_i^- is the predicted score of feed forward network of (x_i, x_i^+) and (x_i, x_i^-) , margin controls the boundary.

With classification loss and margin ranking loss, we get supervised loss for interactive network as:

$$loss_{sl} = loss_{ce} + loss_{mr} \quad (4)$$

Multi-task learning needs to integrate different loss functions. We propose the weighted decaying methods to regularize three different loss functions:

While capturing the interactive semantic information, we need to minimize the hurt to siamese encoder. We propose the loss decaying methods to regularize two different loss functions as:

$$loss = loss_{cl} + \alpha * loss_{sl} \quad (5)$$

where α is piecewise constant decaying when training. For example, we evenly divide the training process into 5 stages, each stage corresponds to a different α . When the value list of α is $[10, 1, 0.1, 0.01, 0.001]$, the decaying is shown on Figure 4.

4 Experiments

Our approach is mainly proposed for supervised tasks, and we conducted multiple experiments on Semantic Textual Similarity (STS) task to verify the effectiveness of this approach.

4.1 Setups

Datasets Following previous works(Reimers and Gurevych, 2019; Gao et al., 2022; Yan et al., 2021), we evaluate our approach on 7 STS tasks: STS 2012-2016(Agirre et al., 2012, 2013, 2014, 2015, 2016), STS Benchmark(Cer et al., 2017) and SICK-Relatedness(Marelli et al., 2014). A pair of sentences in those datasets has a gold score between 0 and 5 to indicate their semantic similarity. The higher the score, the higher the similarity of sentences pair.

Since the gold score label of the STS datasets is not suitable for the training process of the independent network, we introduce the SNLI(Bowman et al., 2015) and MNLI(Williams et al., 2018) (NLI) datasets as supervised data to train our model. The details of NLI datasets have described in Section 3.2.

Evaluation When evaluating the trained model, we first obtain the representation of sentences by $[CLS]$ token embeddings, then we report the spearman correlation between the cosine similarity scores of sentence representations and the human-annotated gold scores. When calculating spearman correlation, we merge all sentences together in a STS task, and calculate the mean of spearman correlation.

4.2 Training Details

Our implementation is based on the SimCSE. We start from pre-trained checkpoints of BERT-base(uncased) and BERT-large (uncased), take the $[CLS]$ representation as the sentence embedding, and train model on the combination of MNLI and SNLI datasets on 4 Tesla V100 gpus for 4 epochs. Hyper-parameters τ and $margin$ are set to 0.04 and 0.5. Since α is used to learn interactive information and could not hurt performance of sentence embedding, the value is chose from $[10, 1, 0.1, 0.01, 0.001]$ following with the training process.

4.3 Main Results

We compare our approach to previous state-of-the-art sentence embedding methods on STS tasks, including InferSent(Conneau et al., 2017), Universal Sentence Encoder(Cer et al., 2018), Sentence-BERT(Reimers and Gurevych, 2019), ConSERT(Yan et al., 2021), SimCSE(Gao et al., 2022) and PromCSE(Jiang et al., 2022a).

Table 5 shows the evaluation results on 7 STS

tasks. InterCSE can substantially improve results on all the datasets, greatly outperforming the previous state-of-the-art models. Specifically, InterCSE-BERT(base) improves the averaged spearman’s correlation from 81.57% to 82.11% compared to SimCSE-BERT(base). InterCSE-BERT(large) achieve 82.88% spearman’s correlation, 0.43% improvement compared to SimCSE-BERT(large). Our models achieve new state-of-the-art performance on STS tasks.

5 Inter-Sentence-Transformers

In this section, for demonstrating the generality of interactive information on sentence representation tasks, we propose to add interactive network to the Sentence-Transformers , and conduct some experiments to verify the performance on the STS-Benchmark dataset.

5.1 Model and Loss Function

Our proposed model is called Inter-Sentence-Transformers . It also consists of independent network and interactive network.

The independent network follows Sentence-Transformers with siamese transformers, and uses the mean of all tokens embedding as sentences representation. The predicted score produced by cosine-similarity on sentences representation is called *indep_score*.

The interactive network extracts the $[CLS]$ token representation as input of a layer of regression network which is made up of fully-connected layer and sigmoid function. The predicted score produced by sigmoid function is called *inter_score*.

Interactive and independent network both use mean-square error as loss function, and we combine these to generate a final loss with piecewise-constant decay as shown in equation (6) to (8).

$$L_{indep} = \frac{1}{N} \sum_{i=1}^N (y_i - indep_score_i)^2 \quad (6)$$

$$L_{inter} = \frac{1}{N} \sum_{i=1}^N (y_i - inter_score_i)^2 \quad (7)$$

$$Loss = L_{indep} + \alpha * L_{inter} \quad (8)$$

where y_i is the true score of sentences pair, α is decaying as training.

5.2 Experiment Results

We conduct experiments on the STS-Benchmark dataset that use STS-Benchmark training dataset to

Model	STS12	STS13	STS14	STS15	STS16	STS-B	SICK-R	Avg
InferSent-Glove ♣	52.86	66.75	62.15	72.77	66.87	68.03	65.65	65.01
Universal Sentence Encoder ♣	64.49	67.80	64.61	76.83	73.18	74.92	76.69	71.22
SBERT(base) ♣	70.97	76.53	73.19	79.09	74.30	77.03	72.91	74.89
SBERT-flow(base) ♠	69.78	77.27	74.35	82.01	77.46	79.12	76.21	76.60
SBERT-whitening(base) ♠	69.65	77.57	74.66	82.27	78.39	79.52	76.91	77.00
ConSERT-joint(base) ◇	70.53	79.96	74.85	81.45	76.72	78.82	77.53	77.12
SimCSE-BERT(base) ♠	75.30	84.67	80.19	85.40	80.82	84.25	80.39	81.57
PromCSE-BERT(base) △	75.58	84.33	79.67	85.79	81.24	84.25	80.79	81.81
InterCSE-BERT(base) ♥	75.83	85.05	80.82	86.00	81.07	84.86	81.16	82.11
SBERT(large) ♣	72.27	78.46	74.90	80.99	76.25	79.23	73.75	76.55
SimCSE-BERT(large) ◆	76.15	86.29	80.80	86.27	81.29	85.39	80.98	82.45
InterCSE-BERT(large) ♥	76.59	86.69	81.83	86.32	81.85	85.78	81.12	82.88

Table 3: Sentence embedding performance on STS tasks (Spearman’s correlation, "all" setting). ♥: results of our approach, ♣: results from sentence-transformer(Reimers and Gurevych, 2019), ◇: results from ConSERT(Yan et al., 2021), ♠: results from SimCSE(Gao et al., 2022), △: results from PromCSE(Jiang et al., 2022b), ◆: results from our reproduced.

train our model and evaluate on test dataset with spearman’s correlation. The label of dataset is a number between 0 and 5, so we divide it by 5.0 as the two sentences semantic similarity.

The piecewise constant value α is chose from [10, 1, 0.1, 0.01, 0.001] in BERT-like base models and [1, 0.1, 0.01, 0.001, 0.0001] in BERT-like large models. Training epochs is 32, and the other parameters are the same as Sentence-Transformer.

The final performance is shown in the table 4. Inter-Sentence-Transformers get a state-of-the-art performances. It improves 0.88% and 1.07% based on BERT-base and RoBERTa-base respectively compared with Sentence-Transformers. And it slight increases based on BERT-large and RoBERTa-large. Inter-Sentence-Transformers treats sentence representation as a regression task. Experiment results show that our interactive network plays an effective role for sentence representation.

6 Ablation Experiment

In this section, we conduct further analyses to understand the role of interactive network in InterCSE, including training strategy, loss combination and loss decay.

Training Strategy We think that there are three training strategies that can help the model to consider the interaction information of sentence pairs while capturing sentence independent semantics:

- Single-Siamese*, we firstly train single tower network, and then train siamese network;

Model	Spearman
SBERT-STSB-base ♣	84.30
SBERT-STSB-large ♣	84.28
SRoBERTa-STSB-base ♣	84.62
SRoBERTa-STSB-large ♣	84.41
Inter-SBERT-STSB-base ◇	85.18
Inter-SBERT-STSB-large ◇	84.91
Inter-SRoBERTa-STSB-base ◇	85.69
Inter-SRoBERTa-STSB-large ◇	84.82

Table 4: Evaluation on the STS Benchmark test set with spearman’s correlation. ♣: reproduced on Sentence-Transformers(Reimers and Gurevych, 2019); ◇: adding interactive network based on Sentence-Transformers. All results are trained on STS Benchmark train set.

Training Strategy	STS-B Spearman
<i>Single-Siamese</i>	37.74
<i>Siamese-Single</i>	84.76
<i>InterCSE</i>	86.25

Table 5: Evaluation on the STS Benchmark dev set with spearman’s correlation of training strategies, and the initial checkpoint is BERT-base.

- Siamese-Single*, we firstly train siamese network, and then train single tower network;
- InterCSE*, we use joint training strategy with loss decay.

The performance of three training strategy are shown as Table 5, and the joint training strategy has a obvious advantage.

Loss Combination From the experimental results of InterCSE and Inter-Sentence-Transformers,

Model	STS-B	STS Avg.
<i>Loss</i>		
w/o $loss_{ce}$	86.19	82.02
w/o $loss_{mr}$	86.16	81.97
constant α	86.08	81.75
square penalty	86.09	81.76
decay α	86.25	82.11

Table 6: Evaluation on the STS Benchmark dev set with spearman’s correlation of loss function.

we can see the effectiveness of the interaction network for improving the overall performance. At the same time, in the section *Training Strategy*, the benefits of multi-task methods for training two sub-networks at the same time are also confirmed. And there are still many methods how to combine the loss of multi-tasking. Therefore, we have carried out a variety of combination strategies:

- (a) remove $loss_{ce}$ from equation (4),
- (b) remove $loss_{mr}$ from equation (4),
- (c) set α as a constant value, and we use 1.0,
- (d) add square penalty as below,

$$loss = loss_{cl}^2 + loss_{sl}^2 \quad (9)$$

- (e) use loss decay method.

The performance of different loss is shown in Table 6. Removing $loss_{ce}$ or $loss_{mr}$ lose some semantic supervision signal which lower overall performance, and loss decay get an awesome result than constant α and square penalty.

Loss Decay In our approach, we use a dynamically changing loss to express the multi-task learning goal, which does not affect the sentence semantic representation while taking more semantic information. Hence, we propose an attenuation strategy for the loss of the interactive network, as shown in the equation (5), and α gradually decreases with the training process going on. The value of α represents the degree of interactive information introduced. The larger the α , the more sufficient the interactive information introduced, but it will also affect the sentence representation of the original siamese network. The results of different descent methods we compared are shown in Table 7. When α decay list is [10, 1, 0.1, 0.01, 0.001], our model get best performance. In the early stage of training, the interactive network has a large loss, which makes the model perceive the interactive

α decay list	STS-B Spearman
[100, 10, 1, 0.1, 0.01]	86.19
[10, 1, 0.1, 0.01, 0.001]	86.25
[1, 0.1, 0.01, 0.001, 0.0001]	86.09
[0.1, 0.01, 0.001, 0.0001, 0.00001]	86.02

Table 7: Evaluation on the STS Benchmark dev set with spearman’s correlation of α decay list.

information and update the model parameters by back-propagation. In the later stage of training, the interaction network is weakened, and the model fully learns the semantic embedding of the sentence.

7 Conclusion

In this paper, we propose multiple task method joint contrastive learning for sentence representation which is termed InterCSE. Experiments shows that InterCSE achieves considerable performance on 7 semantic text similarity tasks. Through InterCSE uses a simple framework, it makes a perfect combination of the advantages of single tower network and siamese network by decaying loss function. When perceiving fine-grained word information, our approach try to minimize damage to sentence semantics. Meanwhile, the performance of Inter-Sentence-Transformers is improved a lot than Sentence-Transformers, it proves that the interactive network can bring a good positive gain again. In the future, we could focus on designing reinforcement learning method with human feedback, which help to enhance the sentence embedding.

Limitations

One limitation of our work is that we can not experiment on unsupervised setting, though the interactive network of InterCSE and Inter-Sentence-Transformers needs labeled sentence pairs. In the unsupervised task of sentence representation, how to extract interaction information is still a tough challenge.

References

- Eneko Agirre, Carmen Banea, Claire Cardie, Daniel Cer, Mona Diab, Aitor Gonzalez-Agirre, Weiwei Guo, Iñigo Lopez-Gazpio, Montse Maritxalar, Rada Mihalcea, German Rigau, Larraitz Uria, and Janyce Wiebe. 2015. [SemEval-2015 task 2: Semantic textual similarity, English, Spanish and pilot on interpretability](#). In *Proceedings of the 9th International Workshop on Semantic Evaluation (SemEval 2015)*,

553	pages 252–263, Denver, Colorado. Association for	
554	Computational Linguistics.	
555	Eneko Agirre, Carmen Banea, Claire Cardie, Daniel Cer,	
556	Mona Diab, Aitor Gonzalez-Agirre, Weiwei Guo,	
557	Rada Mihalcea, German Rigau, and Janyce Wiebe.	
558	2014. SemEval-2014 task 10: Multilingual semantic	
559	textual similarity . In <i>Proceedings of the 8th Interna-</i>	
560	<i>tional Workshop on Semantic Evaluation (SemEval</i>	
561	<i>2014)</i> , pages 81–91, Dublin, Ireland. Association for	
562	Computational Linguistics.	
563	Eneko Agirre, Carmen Banea, Daniel Cer, Mona Diab,	
564	Aitor Gonzalez-Agirre, Rada Mihalcea, German	
565	Rigau, and Janyce Wiebe. 2016. SemEval-2016	
566	task 1: Semantic textual similarity, monolingual	
567	and cross-lingual evaluation . In <i>Proceedings of the</i>	
568	<i>10th International Workshop on Semantic Evaluation</i>	
569	<i>(SemEval-2016)</i> , pages 497–511, San Diego, Californ-	
570	ia. Association for Computational Linguistics.	
571	Eneko Agirre, Daniel Cer, Mona Diab, and Aitor	
572	Gonzalez-Agirre. 2012. SemEval-2012 task 6: A	
573	pilot on semantic textual similarity . In <i>*SEM 2012:</i>	
574	<i>The First Joint Conference on Lexical and Computa-</i>	
575	<i>tional Semantics – Volume 1: Proceedings of the</i>	
576	<i>main conference and the shared task, and Volume</i>	
577	<i>2: Proceedings of the Sixth International Workshop</i>	
578	<i>on Semantic Evaluation (SemEval 2012)</i> , pages 385–	
579	393, Montréal, Canada. Association for Computa-	
580	tional Linguistics.	
581	Eneko Agirre, Daniel Cer, Mona Diab, Aitor Gonzalez-	
582	Agirre, and Weiwei Guo. 2013. *SEM 2013 shared	
583	task: Semantic textual similarity . In <i>Second Joint</i>	
584	<i>Conference on Lexical and Computational Semantics</i>	
585	<i>(*SEM), Volume 1: Proceedings of the Main Confer-</i>	
586	<i>ence and the Shared Task: Semantic Textual Similar-</i>	
587	<i>ity</i> , pages 32–43, Atlanta, Georgia, USA. Association	
588	for Computational Linguistics.	
589	Wasi Uddin Ahmad, Xueying Bai, Zhechao Huang,	
590	Chao Jiang, Nanyun Peng, and Kai-Wei Chang. 2018.	
591	Multi-task learning for universal sentence embed-	
592	dings: A thorough evaluation using transfer and aux-	
593	iliary tasks .	
594	Samuel R. Bowman, Gabor Angeli, Christopher Potts,	
595	and Christopher D. Manning. 2015. A large anno-	
596	tated corpus for learning natural language inference .	
597	In <i>Proceedings of the 2015 Conference on Empiri-</i>	
598	<i>cal Methods in Natural Language Processing</i> , pages	
599	632–642, Lisbon, Portugal. Association for Computa-	
600	tional Linguistics.	
601	Tom B. Brown, Benjamin Mann, Nick Ryder, Melanie	
602	Subbiah, Jared Kaplan, Prafulla Dhariwal, Arvind	
603	Neelakantan, Pranav Shyam, Girish Sastry, Amanda	
604	Askell, Sandhini Agarwal, Ariel Herbert-Voss,	
605	Gretchen Krueger, Tom Henighan, Rewon Child,	
606	Aditya Ramesh, Daniel M. Ziegler, Jeffrey Wu,	
607	Clemens Winter, Christopher Hesse, Mark Chen, Eric	
608	Sigler, Mateusz Litwin, Scott Gray, Benjamin Chess,	
609	Jack Clark, Christopher Berner, Sam McCandlish,	
610	Alec Radford, Ilya Sutskever, and Dario Amodei.	
611	2020. Language models are few-shot learners .	
	Daniel Cer, Mona Diab, Eneko Agirre, Iñigo Lopez-	612
	Gazpio, and Lucia Specia. 2017. SemEval-2017	613
	task 1: Semantic textual similarity multilingual and	614
	crosslingual focused evaluation . In <i>Proceedings</i>	615
	<i>of the 11th International Workshop on Semantic</i>	616
	<i>Evaluation (SemEval-2017)</i> , pages 1–14, Vancouver,	617
	Canada. Association for Computational Linguistics.	618
	Daniel Cer, Yinfei Yang, Sheng-yi Kong, Nan Hua,	619
	Nicole Limtiaco, Rhomni St. John, Noah Constant,	620
	Mario Guajardo-Cespedes, Steve Yuan, Chris Tar,	621
	Brian Strope, and Ray Kurzweil. 2018. Universal	622
	sentence encoder for English . In <i>Proceedings of</i>	623
	<i>the 2018 Conference on Empirical Methods in Nat-</i>	624
	<i>ural Language Processing: System Demonstrations</i> ,	625
	pages 169–174, Brussels, Belgium. Association for	626
	Computational Linguistics.	627
	Alexis Conneau, Douwe Kiela, Holger Schwenk, Loïc	628
	Barraut, and Antoine Bordes. 2017. Supervised	629
	learning of universal sentence representations from	630
	natural language inference data . In <i>Proceedings of</i>	631
	<i>the 2017 Conference on Empirical Methods in Nat-</i>	632
	<i>ural Language Processing</i> , pages 670–680, Copen-	633
	hagen, Denmark. Association for Computational Lin-	634
	guistics.	635
	Jacob Devlin, Ming-Wei Chang, Kenton Lee, and	636
	Kristina Toutanova. 2019. Bert: Pre-training of deep	637
	bidirectional transformers for language understand-	638
	ing .	639
	Jun Gao, Di He, Xu Tan, Tao Qin, Liwei Wang, and Tie-	640
	Yan Liu. 2019. Representation degeneration problem	641
	in training natural language generation models .	642
	Tianyu Gao, Xingcheng Yao, and Danqi Chen. 2022.	643
	Simcse: Simple contrastive learning of sentence em-	644
	beddings .	645
	Yuxin Jiang, Linhan Zhang, and Wei Wang. 2022a. Im-	646
	proved universal sentence embeddings with prompt-	647
	based contrastive learning and energy-based learning .	648
	Yuxin Jiang, Linhan Zhang, and Wei Wang. 2022b. Im-	649
	proved universal sentence embeddings with prompt-	650
	based contrastive learning and energy-based learning .	651
	In <i>Findings of the Association for Computational</i>	652
	<i>Linguistics: EMNLP 2022</i> , pages 3021–3035, Abu	653
	Dhabi, United Arab Emirates. Association for Com-	654
	putational Linguistics.	655
	Gregory Koch, Richard Zemel, and Ruslan Salakhut-	656
	dinov. 2015. Siamese neural networks for one-shot	657
	image recognition.	658
	Pengfei Liu, Xipeng Qiu, and Xuanjing Huang. 2017.	659
	Adversarial multi-task learning for text classification .	660
	In <i>Proceedings of the 55th Annual Meeting of the</i>	661
	<i>Association for Computational Linguistics (Volume</i>	662
	<i>1: Long Papers)</i> , pages 1–10, Vancouver, Canada.	663
	Association for Computational Linguistics.	664
	Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Man-	665
	dar Joshi, Danqi Chen, Omer Levy, Mike Lewis,	666
	Luke Zettlemoyer, and Veselin Stoyanov. 2019.	667

668	Roberta: A robustly optimized bert pretraining approach.	Adina Williams, Nikita Nangia, and Samuel R. Bowman. 2018. A broad-coverage challenge corpus for sentence understanding through inference.	723
669			724
670	Minh-Thang Luong, Quoc V. Le, Ilya Sutskever, Oriol Vinyals, and Lukasz Kaiser. 2016. Multi-task sequence to sequence learning.	Xing Wu, Chaochen Gao, Liangjun Zang, Jizhong Han, Zhongyuan Wang, and Songlin Hu. 2022. ESIM-CSE: Enhanced sample building method for contrastive learning of unsupervised sentence embedding.	725
671			726
672			727
673	Marco Marelli, Stefano Menini, Marco Baroni, Luisa Bentivogli, Raffaella Bernardi, and Roberto Zamparelli. 2014. A SICK cure for the evaluation of compositional distributional semantic models. In <i>Proceedings of the Ninth International Conference on Language Resources and Evaluation (LREC'14)</i> , pages 216–223, Reykjavik, Iceland. European Language Resources Association (ELRA).		728
674			729
675			730
676			731
677			732
678			733
679		Yuanmeng Yan, Rumei Li, Sirui Wang, Fuzheng Zhang, Wei Wu, and Weiran Xu. 2021. Consert: A contrastive framework for self-supervised sentence representation transfer.	734
680			735
681	Myle Ott, Sergey Edunov, Alexei Baevski, Angela Fan, Sam Gross, Nathan Ng, David Grangier, and Michael Auli. 2019. fairseq: A fast, extensible toolkit for sequence modeling. In <i>Proceedings of NAACL-HLT 2019: Demonstrations</i> .		736
682			737
683			
684			738
685			739
686	Alec Radford, Karthik Narasimhan, Tim Salimans, and Ilya Sutskever. 2018. Improving language understanding by generative pre-training.		740
687			741
688			742
689			743
690	Alec Radford, Jeff Wu, Rewon Child, David Luan, Dario Amodei, and Ilya Sutskever. 2019. Language models are unsupervised multitask learners.		744
691			
692	Marek Rei. 2017. Semi-supervised multitask learning for sequence labeling. In <i>Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)</i> , pages 2121–2130, Vancouver, Canada. Association for Computational Linguistics.		
693			
694			
695			
696			
697			
698	Nils Reimers and Iryna Gurevych. 2019. Sentence-bert: Sentence embeddings using siamese bert-networks.		
699			
700	Yu Sun, Shuohuan Wang, Shikun Feng, Siyu Ding, Chao Pang, Junyuan Shang, Jiaxiang Liu, Xuyi Chen, Yanbin Zhao, Yuxiang Lu, Weixin Liu, Zhihua Wu, Weibao Gong, Jianzhong Liang, Zhizhou Shang, Peng Sun, Wei Liu, Xuan Ouyang, Dianhai Yu, Hao Tian, Hua Wu, and Haifeng Wang. 2021. Ernie 3.0: Large-scale knowledge enhanced pre-training for language understanding and generation.		
701			
702			
703			
704			
705			
706			
707			
708	Yu Sun, Shuohuan Wang, Yukun Li, Shikun Feng, Xuyi Chen, Han Zhang, Xin Tian, Danxiang Zhu, Hao Tian, and Hua Wu. 2019a. Ernie: Enhanced representation through knowledge integration.		
709			
710			
711			
712	Yu Sun, Shuohuan Wang, Yukun Li, Shikun Feng, Hao Tian, Hua Wu, and Haifeng Wang. 2019b. Ernie 2.0: A continual pre-training framework for language understanding.		
713			
714			
715			
716	Aaron van den Oord, Yazhe Li, and Oriol Vinyals. 2019. Representation learning with contrastive predictive coding.		
717			
718			
719	Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Lukasz Kaiser, and Illia Polosukhin. 2017. Attention is all you need.		
720			
721			
722			