

HU at SemEval-2024 Task 8A: Can Contrastive Learning Learn Embeddings to Detect Machine-Generated Text?

Shubhashis Roy Dipta

Department of CSEE
University of Maryland, Baltimore County
Baltimore, Maryland, USA
sroydip1@umbc.edu

Sadat Shahriar

University of Houston
Houston, Texas, USA
sshahria@cougarnet.uh.edu

Abstract

This paper describes our system developed for SemEval-2024 Task 8, “Multigenerator, Multidomain, and Multilingual Black-Box Machine-Generated Text Detection” Machine-generated texts have been one of the main concerns due to the use of large language models (LLM) in fake text generation, phishing, cheating in exams, or even plagiarizing copyright materials. A lot of systems have been developed to detect machine-generated text. Nonetheless, the majority of these systems rely on the text-generating model. This limitation is impractical in real-world scenarios, as it’s often impossible to know which specific model the user has used for text generation. In this work, we propose a **single** model based on contrastive learning, which uses **$\approx 40\%$ of the baseline’s parameters** (149M vs. 355M) but shows a comparable performance on the test dataset (**21st out of 137 participants**). Our key finding is that even without an ensemble of multiple models, a single base model can have comparable performance with the help of data augmentation and contrastive learning.¹

1 Introduction

In recent years, Natural Language Processing (NLP) has been totally dependent on Deep Learning rather than statistical machine learning. With multi-task learning (Caruana, 1997), attention-based transformers (Vaswani et al., 2017), and the use of Reinforcement Learning in NLP (Christiano et al., 2017), it has been used in our day-to-day life from mathematical calculations (Yang et al., 2023) to email writing. But with huge help, it has also been used to generate fake news (Zellers et al., 2019), to plagiarize copyright materials (Dehouche, 2021), and also to cheat in exams or assignments (Cotton et al., 2023; Fyfe, 2023). Humans can identify machine-generated text only at the chance level

(Jawahar et al., 2020). There has been a dire need to develop a system to detect machine-generated text.

Though a lot of works (Badaskar et al., 2008; Gehrmann et al., 2019; Zellers et al., 2019; Jawahar et al., 2020; Ippolito et al., 2020; Chakraborty et al., 2023; Pu et al., 2023; Mitchell et al., 2023; He et al., 2023; Guo et al., 2023) have already been deployed for detecting machine-generated text, with the current development of LLMs, most of the systems are failing to find out which one is human-generated vs. machine-generated (mostly due to the improvement of coherency, fluency and usage of real-world dataset (Radford et al., 2019)). In this context, the task “Multigenerator, Multidomain, and Multilingual Black-Box Machine-Generated Text Detection” provides a dataset for training models to classify machine-generated texts. The shared task consists of three sub-tasks: Binary Classification (Machine vs. Human), Multi-class Classification (Which model/human generated this?), and Span Detection (Which part of the text is machine-generated?). A detailed description of the task can be found in the shared task paper (Wang et al., 2024).

In this paper, we describe our final submission on Subtask A (Binary Classification). There were two big challenges of this task: **First**, five Different models have been used to generate the machine-generated text. Zellers et al. (2019) has shown that the best defense for machine-generated text is the model itself that was used for generation. However, in reality, there is a massive surge in large language models (LLMs), each with its own unique style of text generation. The challenge in this particular subtask has heightened due to the utilization of five different LLMs. This complexity demands a versatile, model-agnostic architecture capable of detecting text generated by LLMs in a generalized manner. **Second**, Following the previous challenge, the organizers have employed a different model

¹Our code is publicly available at <https://github.com/dipta007/SemEval24-Task8>

for generating the validation and test datasets compared to those used in the training set. This implies that the text was drawn from a completely distinct distribution. As a result, participants must develop a generalized model capable of performing effectively regardless of the specific model used in the text generation process.

In response to the key challenges, we have investigated the performance of contrastive learning for this particular task. Contrastive learning has been used as a valuable technique across various domains, including Text Embedding (Neelakantan et al., 2022), Document Embedding (Luo et al., 2021), Event Embedding (Roy Dipta et al., 2023), vision (Chen et al., 2020) and Language-Vision learning (Radford et al., 2021). Notably, unlike the majority of submissions in any shared task like competition, Our final submission utilized a **single** model to classify the machine-generated texts rather than an ensemble of multiple models. Hence, our contributions to this paper are as follows,

1. We proposed a novel data augmentation technique, which nearly makes the data X times bigger (X is the number of models used for data augmentation).
2. We propose a single unified model that shows a comparable performance on the test dataset.
3. We have shown that even with a single model, contrastive learning with data augmentation shows a comparable performance, which opens up a door for future exploration.

2 Related Works

In this section, we will provide the prior works that have been done in the realm of machine-generated text detection (§2.1) and contrastive learning (§2.1).

2.1 Machine Generated Text detection

With the progress of LLMs, much prior research has been done to counter-attack the misuse of the LLMs. Before the attention and transformers, Badaskar et al. (2008) has shown how the syntactic and semantic features can help in classifying between human and machine-generated text. Later, Gehrmann et al. (2019) has provided a statistical detection system based on the assumption that the machine samples from the high probability words through max sampling (Gu et al., 2017), k-max sampling (Fan et al., 2018), beam search

(Shao et al., 2017). So, the authors used the probability, rank, and entropy of words as features to classify a machine-generated text. Jawahar et al. (2020) has shown that state-of-the-art LLM can generate texts with human-like fluency and coherence without grammatical or spelling errors. Lastly, Mitchell et al. (2023) have used the change of log-probability between the original text and after random perturbation.

2.2 Contrastive Learning

Contrastive learning was first introduced in the visual domain (Chen et al., 2020). Later, it has been widely used in NLP for representation learning (Xu et al., 2023; Wang and Dou, 2023), event similarity tasks (Gao et al., 2022) and event modeling (Roy Dipta et al., 2023). Inspired by the latter works, we have explored whether contrastive learning can help in machine-generated text detection.

3 System Overview

Our system is divided into three parts: where the first part is data augmentation (described on §3.1), the second part is contrastive learning (described on §3.2), and the last part is the classification head (described on §3.3) over the document embeddings.

3.1 Data Augmentation

The dataset provided in the shared task has text and their corresponding label. However, we need a positive and a (hard) negative pair to use contrastive learning. Our main inspiration for using contrastive learning is that as the texts come from two different entities (machine vs. human), the embedding space should also be different. To facilitate the task, we have used a paraphrase model to generate alternate texts for each text in the dataset. In that way, now, every instance of the dataset has one human/machine-generated text and one machine-generated text. We have utilized the human-generated text as the hard negative ² and the machine-generated text as the soft positive ³.

Another challenge we faced during the paraphrasing of the dataset is that the texts are long. If we give the whole text to the paraphrase model and ask for alternate text, it gives a much shorter text (an issue we observed in the used paraphrase model). In our primary validation, that gives bad

²Hard negatives are the total opposite of the given text

³Soft positives expressed the same idea but might not be the exact one

results due to the loss of information while shortening the text. So, instead of giving the whole text at once, we have split the data by end-of-sentence or newline. Then, each sentence was paraphrased on its own and then joined together again to get the previous structure. The technical details behind generating paraphrases and using them for contrastive learning have been discussed in §4.1 and §4.2, respectively.

3.2 Contrastive Learning

With the availability of an appropriate dataset for contrastive learning, we proceeded to develop our model. Our main assumption was that the embedding of the machine-generated text and human-generated text would exhibit significant differences. A simple overview of the model is shown in the Fig. 1.

The positive and negative data go through the same shared encoder to generate an embedding. This embedding is then used in contrastive learning. We have used the following loss formulation for our contrastive learning:

$$\mathcal{L}_{con} = (1 - y) * \cos(x_1, x_2) + y * \max(0, \cos(x_1, x_2)) \quad (1)$$

Here, x_1 and x_2 are the embeddings of two different pairs, respectively. $\cos(x_1, x_2)$ is the cosine-similarity score between two embeddings. y is $+1$ for positive-positive pairs and -1 otherwise. In our task, y is $+1$ if the data instance contains text from a machine and the other is paraphrased text and -1 if the given text is from a human and the other is paraphrased.

3.3 Classification Loss

In contrastive learning, our primary objective is to acquire meaningful embeddings containing sufficient information for distinguishing between human-generated and machine-generated text. However, we also need to use a classifier model for the downstream task of outputting the actual label. Keeping that in mind, we have used a simple two-linear layer classifier head on top of the embeddings generated by the encoder. During inference time, we used this classifier head to output the labels. We have optimized our model using a simple binary cross-entropy (BCE) loss.

The total loss of our model is defined as,

$$\mathcal{L} = \alpha * \mathcal{L}_{con} + \beta * \mathcal{L}_{cls+} + \gamma * \mathcal{L}_{cls-} \quad (2)$$

Here, $\mathcal{L}_{cls+}$ is the BCE loss of the positive example, and $\mathcal{L}_{cls-}$ is the BCE loss of the negative sample of the data instance. α , β , and γ are hyperparameters that were set to 0.7, 0.8, and 0.1, respectively, based on validation data.

4 Experimental Setup

The following sections are used to describe the technical details behind our data augmentation technique (§4.1), Encoder (§4.2), Classifier Head (§4.3) and Hyperparameters (§4.4).

4.1 Data Augmentation & Pre-processing

We preprocess the raw input, splitting each document into multiple sentences for paraphrasing. After the preprocessing, we got ≈ 3.6 million sentences. Even if we are splitting by new lines or end-of-sentences, we kept exactly the same format during joining, i.e., two new lines rather than 1, to keep most information intact. As the paraphrasing is done on the sentence level rather than the paragraph level, the number of paraphrased sentences is the same as the input sentences (3.6M). So, ideally, we got double the number of training data just by using the data augmentation.

We have tried multiple models from HuggingfaceHub ^{4 5} to generate paraphrase. In our final submission, we have used [Bandel et al. \(2022\)](#)'s model ⁴ for our data augmentation. Use of multiple models or use of prompt-based models ([Achiam et al., 2023](#); [Touvron et al., 2023](#)) for data augmentation has been left out for future exploration due to time and compute constraints. For data split, we use the official train & dev data split. Only train data is used for data augmentation, and the dev data is used to calculate evaluation metrics.

4.2 Pre-trained Encoder

To encode the document, we have used a pre-trained version of longformer-base ([Beltagy et al., 2020](#)) ⁶. The reason behind using this encoder rather than others is, **One**, longformer is good for getting embeddings for long documents because of using global vs. local attention (more details in [Beltagy et al. \(2020\)](#)). **Second**, the pre-trained version was fine-tuned for paraphrase detection, which is kind of similar to our task.

⁴[ibm/qcpg-sentences](#)

⁵[ceshine/t5-paraphrase-paws-msrp-opinosis](#)

⁶[jpwahle/longformer-base-plagiarism-detection](#)

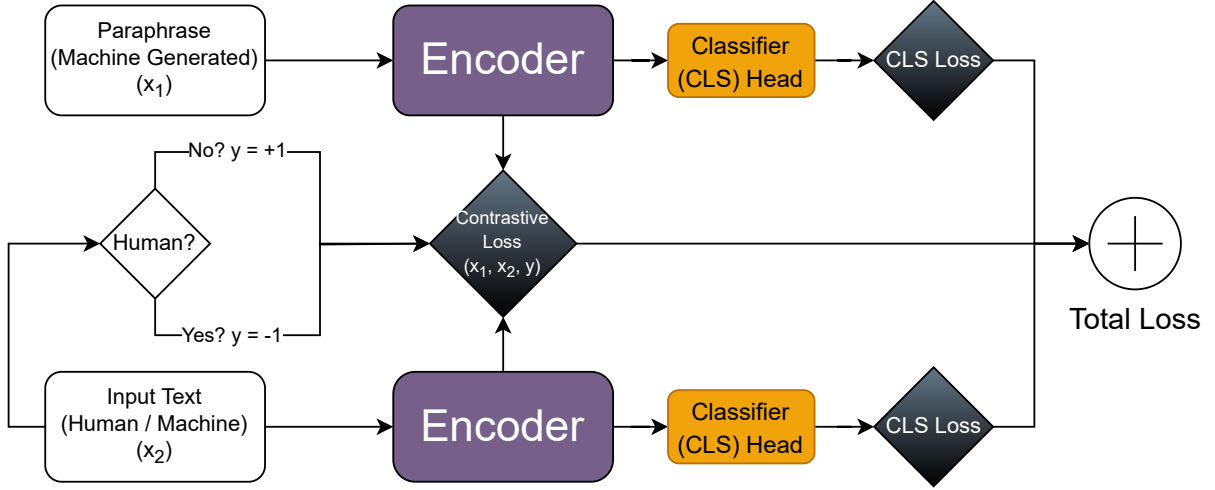


Figure 1: Overview of our model architecture. The same color weights are shared (encoder & classifier head). Diamond boxes represent the loss function, and the plus sign represents the summation of the three losses. The input to the contrastive loss depends on the original label ($y=+1$ if human, else -1).

4.3 Classifier Head

We have used two linear layers for classifier heads with *tanh* activation loss between them. We also have used a dropout layer between them with a probability of 60%. The primary rationale for using a high dropout rate was to enhance the model’s generalization ability and reduce its dependence on the training data.

4.4 Hyperparameters

For training our model, we have used AdamW (Loshchilov and Hutter, 2017) optimizer with a learning rate of $1e-5$. We have used a batch size of 2 with gradient accumulation for 8 steps (effective batch size 16). We have used early stopping on the validation data with patience 10. Maximum document length was set to 4096 as most of the documents are large. We use the PyTorch-lightning⁷ library to run the experiments and Weight & Biases⁸ for logging. All of our experiments are run on NVIDIA Quadro RTX 8000 48GB.

5 Results

In this section, we report our results on subtask A and discuss our analysis. Our evaluation is based on the accuracy metric, but we have provided the micro and macro-f1 for better comparison. All the results are averaged on 3 runs with 3 different random seeds.

⁷<https://lightning.ai/>

⁸<https://wandb.ai/>

5.1 Baseline & Our Model

We use the official baseline provided by the task organizers. They have used RoBERTa-large (Liu et al., 2019) as the encoder and fine-tuned on the train data. Throughout the paper, we refer to this model as *baseline_{rob}*.

We have fine-tuned our model (shown in Fig. 1) on the training dataset. Throughout the paper, we refer to this model as *ours_{con}*.

In the Table 1, we have reported the results on the official test file. *Ours_{con}* is the final submission, and *Ours_{con}+* is the modified version of our final model for more analysis (not official results; used for ablation study - details on §5.2). We can get a comparable result using 60% fewer parameters than the baseline. In the next section, we will see that after hyperparameter tuning, we can get around 5.7% improvement over the baseline. This supports our assumption that using a contrastive learning-based method can help machine-generated text identification.

5.2 Ablation Study

Effect of Maximum Sentence Length: The maximum sentence length is used to tokenize the document. The optimal test score is achieved with a maximum sentence length of 256. This demonstrates that the model can effectively identify machine-generated text even with documents as large as 256 words. This underscores the effectiveness and adaptability of our model’s learning capabilities.

	Max Sen Length	CLS Dropout	Effective Batch Size	Macro-f1	Micro-f1	Accuracy		
$Ours_{con}$	4096	0.6	16	88.81	89.07	89.07		
$baseline_{rob}$	-	-	-	-	-	88.47		
Maximum Sentence Length								
$Ours_{con}+$	128	0.6	16	88.88	89.14	89.14		
$Ours_{con}+$	256			93.30	93.37	93.36		
$Ours_{con}+$	512			88.78	89.04	89.04		
$Ours_{con}+$	1024			90.99	91.13	91.13		
$Ours_{con}+$	2048			91.81	91.93	91.93		
$Ours_{con}$	4096			88.81	89.07	89.07		
Classification Layer Dropout								
$Ours_{con}+$	0	4096	16	92.73	92.81	92.81		
$Ours_{con}+$	0.2			90.16	90.33	90.33		
$Ours_{con}+$	0.4			78.98	80.21	80.21		
$Ours_{con}$	0.6			88.81	89.07	89.07		
$Ours_{con}+$	0.9			82.60	83.31	83.31		
Effective Batch Size								
$Ours_{con}+$	4096			0.6	2	93.80	93.86	93.86
$Ours_{con}+$		4	70.79		73.52	73.52		
$Ours_{con}+$		8	76.82		78.43	78.43		
$Ours_{con}$		16	88.81		89.07	89.07		
$Ours_{con}+$		32	79.72		80.83	80.83		
$Ours_{con}+$		64	90.64		90.80	90.80		
$Ours_{con}+$		128	91.39		91.51	91.51		

Table 1: Macro-f1, Micro-f1, and Accuracy score on the test result. *Ours_{con}* - final submitted model on the shared task, *baseline_{rob}* - official baseline model, and *Ours_{con}+* - modified versions of our final model with more hyperparameter tuning. The **bold** value signifies the best score within a specific section, whereas the underlined value denotes the best score across all sections.

Effect of Classification Dropout: The classification dropout is applied between the two classification layers. Contrary to our initial assumption, the results presented in Table 1 indicate that using a low dropout rate (as low as 0.0) contributes positively to the model’s learning process. This suggests that, even without dropout, the model’s generalization to unseen data (text generated by a new model) is enabled primarily through contrastive learning and data augmentation.

Effects of (Effective) Batch Size: Due to computational constraint, we have used a fixed batch size of 2 and gradient accumulation steps of {1, 2, 4, 8, 16, 32, 64} resulting in an effective batch size of {2, 4, 8, 16, 32, 64, 128}. From the results report on Table 1, we found that using only an effective batch size of 2 yielded superior performance compared to gradient accumulation. Notably, this configuration represents the most optimal result obtained following hyperparameter tuning, positioning us at the 8th rank in the final standings. This suggests that, in this particular context, the benefits of gradient accumulation may be limited compared to simply using a smaller batch size.

6 Conclusion & Future Work

In this work, we introduce our contrastive learning-based system, which shows a comparable performance. We demonstrate that a model with half the parameters and without an ensemble of large models or hand-engineered features can show a comparable performance, which requires more exploration in this field. For future work, the use of recent prompt-based models⁹ can be used for data augmentation. Also, the effect of more advanced contrastive loss, i.e., Triplet loss (Chechik et al., 2010) or InfoNCE loss (Oord et al., 2018), need to be explored.

References

- Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, et al. 2023. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*.
- Sameer Badaskar, Sameer Badaskar, Sameer Badaskar, Sonali Agarwal, Sachin Agarwal, Shilpa Arora, and Shilpa Arora. 2008. [Identifying real or fake articles: Towards better language modeling](#). *International Joint Conference on Natural Language Processing*.
- Elron Bandel, Ranit Aharonov, Michal Shmueli-Scheuer, Ilya Shnayderman, Noam Slonim, and Liat Ein-Dor. 2022. [Quality controlled paraphrase generation](#). In *Proceedings of the 60th Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)*, pages 596–609, Dublin, Ireland. Association for Computational Linguistics.
- Iz Beltagy, Matthew E Peters, and Arman Cohan. 2020. Longformer: The long-document transformer. *arXiv preprint arXiv:2004.05150*.
- Rich Caruana. 1997. Multitask learning. *Machine learning*, 28:41–75.
- Souradip Chakraborty, Amrit Singh Bedi, Sicheng Zhu, Bang An, Dinesh Manocha, and Furong Huang. 2023. [On the possibilities of ai-generated text detection](#). *arXiv.org*.
- Gal Chechik, Varun Sharma, Uri Shalit, and Samy Bengio. 2010. Large scale online learning of image similarity through ranking. *Journal of Machine Learning Research*, 11(3).
- Ting Chen, Simon Kornblith, Mohammad Norouzi, and Geoffrey Hinton. 2020. A simple framework for contrastive learning of visual representations. In *International conference on machine learning*, pages 1597–1607. PMLR.

⁹<https://chat.openai.com/>

- Paul F Christiano, Jan Leike, Tom Brown, Miljan Maric, Shane Legg, and Dario Amodei. 2017. Deep reinforcement learning from human preferences. *Advances in neural information processing systems*, 30.
- Debby RE Cotton, Peter A Cotton, and J Reuben Shipway. 2023. Chatting and cheating: Ensuring academic integrity in the era of chatgpt. *Innovations in Education and Teaching International*, pages 1–12.
- Nassim Dehouche. 2021. Plagiarism in the age of massive generative pre-trained transformers (gpt-3). *Ethics in Science and Environmental Politics*, 21:17–23.
- Angela Fan, Mike Lewis, and Yann Dauphin. 2018. Hierarchical neural story generation. *arXiv preprint arXiv:1805.04833*.
- Paul Fyfe. 2023. How to cheat on your final paper: Assigning ai for student writing. *AI & SOCIETY*, 38(4):1395–1405.
- Jun Gao, Wei Wang, Changlong Yu, Huan Zhao, Wilfred Ng, and Ruifeng Xu. 2022. Improving event representation via simultaneous weakly supervised contrastive learning and clustering. *arXiv preprint arXiv:2203.07633*.
- Sebastian Gehrmann, Sebastian Gehrmann, Hendrik Strobelt, Hendrik Strobelt, Gérard Rushton, Alexander M. Rush, Alexander M. Rush, and Alexander M. Rush. 2019. Gltr: Statistical detection and visualization of generated text. *Annual Meeting of the Association for Computational Linguistics*.
- Jiatao Gu, Kyunghyun Cho, and Victor OK Li. 2017. Trainable greedy decoding for neural machine translation. *arXiv preprint arXiv:1702.02429*.
- Biyang Guo, Xin Zhang, Ziyuan Wang, Minqi Jiang, Jinran Nie, Yuxuan Ding, Jianwei Yue, and Yupeng Wu. 2023. How close is chatgpt to human experts? comparison corpus, evaluation, and detection. *arXiv.org*.
- Xiaotong He, Xinyue Shen, Zeyuan Chen, Michael Backes, and Yang Zhang. 2023. Mgtbench: Benchmarking machine-generated text detection. *arXiv.org*.
- Daphne Ippolito, Daphne Ippolito, Daniel Duckworth, Daniel Duckworth, Chris Callison-Burch, Chris Callison-Burch, Douglas Eck, and Douglas Eck. 2020. Automatic detection of generated text is easiest when humans are fooled. *Annual Meeting of the Association for Computational Linguistics*.
- Ganesh Jawahar, Ganesh Jawahar, Muhammad Abdul-Mageed, Muhammad Abdul-Mageed, Laks V. S. Lakshmanan, and Laks V. S. Lakshmanan. 2020. Automatic detection of machine generated text: A critical survey. *International Conference on Computational Linguistics*.
- Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. 2019. Roberta: A robustly optimized bert pretraining approach. *arXiv preprint arXiv:1907.11692*.
- Ilya Loshchilov and Frank Hutter. 2017. Decoupled weight decay regularization. *arXiv preprint arXiv:1711.05101*.
- Dongsheng Luo, Wei Cheng, Jingchao Ni, Wenchao Yu, Xuchao Zhang, Bo Zong, Yanchi Liu, Zhengzhang Chen, Dongjin Song, Haifeng Chen, et al. 2021. Unsupervised document embedding via contrastive augmentation. *arXiv preprint arXiv:2103.14542*.
- Eric Mitchell, Yoonho Lee, Alexander Khazatsky, Christopher D. Manning, and Chelsea Finn. 2023. Detectgpt: Zero-shot machine-generated text detection using probability curvature. *International Conference on Machine Learning*.
- Arvind Neelakantan, Tao Xu, Raul Puri, Alec Radford, Jesse Michael Han, Jerry Tworek, Qiming Yuan, Nikolas Tezak, Jong Wook Kim, Chris Hallacy, et al. 2022. Text and code embeddings by contrastive pre-training. *arXiv preprint arXiv:2201.10005*.
- Aaron van den Oord, Yazhe Li, and Oriol Vinyals. 2018. Representation learning with contrastive predictive coding. *arXiv preprint arXiv:1807.03748*.
- Jiameng Pu, Zain Sarwar, Sifat Muhammad Abdullah, Abdullah Rehman, Yoonjin Kim, Parantapa Bhat-tacharya, Mobin Javed, and Bimal Viswanath. 2023. Deepfake text detection: Limitations and opportunities. *IEEE Symposium on Security and Privacy*.
- Alec Radford, Jong Wook Kim, Chris Hallacy, Aditya Ramesh, Gabriel Goh, Sandhini Agarwal, Girish Sastry, Amanda Askell, Pamela Mishkin, Jack Clark, et al. 2021. Learning transferable visual models from natural language supervision. In *International conference on machine learning*, pages 8748–8763. PMLR.
- Alec Radford, Jeffrey Wu, Rewon Child, David Luan, Dario Amodei, Ilya Sutskever, et al. 2019. Language models are unsupervised multitask learners. *OpenAI blog*, 1(8):9.
- Shubhashis Roy Dipta, Mehdi Rezaee, and Francis Ferraro. 2023. Semantically-informed hierarchical event modeling. In *Proceedings of the 12th Joint Conference on Lexical and Computational Semantics (*SEM 2023)*, pages 353–369, Toronto, Canada. Association for Computational Linguistics.
- Louis Shao, Stephan Gouws, Denny Britz, Anna Goldie, Brian Strope, and Ray Kurzweil. 2017. Generating high-quality and informative conversation responses with sequence-to-sequence models. *arXiv preprint arXiv:1701.03185*.
- Hugo Touvron, Thibaut Lavril, Gautier Izacard, Xavier Martinet, Marie-Anne Lachaux, Timothée Lacroix, Baptiste Rozière, Naman Goyal, Eric Hambro,

- Faisal Azhar, et al. 2023. Llama: Open and efficient foundation language models. *arXiv preprint arXiv:2302.13971*.
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. Attention is all you need. *Advances in neural information processing systems*, 30.
- Hao Wang and Yong Dou. 2023. Sncse: contrastive learning for unsupervised sentence embedding with soft negative samples. In *International Conference on Intelligent Computing*, pages 419–431. Springer.
- Yuxia Wang, Jonibek Mansurov, Petar Ivanov, Jinyan Su, Artem Shelmanov, Akim Tsvigun, Chenxi Whitehouse, Osama Mohammed Afzal, Tarek Mahmoud, Giovanni Puccetti, Thomas Arnold, Alham Fikri Aji, Nizar Habash, Iryna Gurevych, and Preslav Nakov. 2024. SemEval-2024 task 8: Multigenerator, multidomain, and multilingual black-box machine-generated text detection. In *Proceedings of the 18th International Workshop on Semantic Evaluation, SemEval 2024*, Mexico City, Mexico.
- Lingling Xu, Haoran Xie, Zongxi Li, Fu Lee Wang, Weiming Wang, and Qing Li. 2023. Contrastive learning models for sentence representations. *ACM Transactions on Intelligent Systems and Technology*, 14(4):1–34.
- Zhen Yang, Ming Ding, Qingsong Lv, Zhihuan Jiang, Zehai He, Yuyi Guo, Jinfeng Bai, and Jie Tang. 2023. Gpt can solve mathematical problems without a calculator. *arXiv preprint arXiv:2309.03241*.
- Rowan Zellers, Rowan Zellers, Ari Holtzman, Ari Holtzman, Hannah Rashkin, Hannah Rashkin, Yonatan Bisk, Yonatan Bisk, Ali Farhadi, Ali Farhadi, Franziska Roesner, Franziska Roesner, Yejin Choi, and Yejin Choi. 2019. [Defending against neural fake news](#). *Neural Information Processing Systems*.